

ALBERT HEIJN

The Adoption of Intelligence Amplification in the
Slotting Process: a case study in the data validation
automation of a Dutch Retailer

MASTER THESIS BUSINESS INFORMATION TECHNOLOGY

WIEGER KLOPPENBURG

**UNIVERSITY
OF TWENTE.**



The Adoption of Intelligence Amplification in the Slotting Process: a case study in the data validation automation of a Dutch Retailer

MASTER THESIS

February 2019

Author

Name: G.W. Kloppenburg (Wieger)
Programme: MSc Business Information Technology
Institute: University of Twente
PO Box 217
7500 AE Enschede
The Netherlands
Email: g.w.kloppenburg@student.utwente.nl

Graduation Committee

First supervisor	Dr. Maria-Eugenia Iacob Department Industrial Engineering and Business Information Systems University of Twente, Enschede, The Netherlands m.e.iacob@utwente.nl
Second supervisor	Dr. Marten van Sinderen Faculty of Electrical Engineering, Mathematics and Computer Science University of Twente, Enschede, The Netherlands m.j.vansinderen@utwente.nl
Albert Heijn supervisor	Pieter Meints MSc. Logistics Support Albert Heijn, Zaandam, The Netherlands pieter.meints@ah.nl
Day-to-day supervisor	Ing. Jean Paul Sebastian Piest MSCM Department Industrial Engineering and Business Information Systems University of Twente, Enschede, The Netherlands j.p.s.piest@utwente.nl

Preface

Na tweeënhalf jaar bezig te zijn geweest met het doen van een studie is deze afgerond. Zes maanden geleden ben ik begonnen met afstuderen en het resultaat ligt voor u. Deze heeft een titel mee gekregen genaamd 'The Adoption of Intelligence Amplification in the Slotting Process: a case study in the data validation automation of a Dutch Retailer'. Het onderzoek voor deze thesis is uitgevoerd bij de Albert Heijn, afdeling Logistics Support in Geldermalsen. De thesis is geschreven in het kader van mijn afstuderen aan de opleiding Business Information Technology aan de Universiteit Twente in Enschede. Het was een mooie tijd, veel nieuwe mensen ontmoet en veel nieuwe dingen geleerd. Ik kon dit niet zonder de hulp van een aantal mensen.

Allereerst, Maria en Marten, zonder hun hulp was het niet gelukt. De kritische blik die zij hadden op de thesis en de vele feedback die zij gaven hielpen mij verder om het afstuderen tot een succes te maken. Sebastian, Andrej, voor de wekelijkse en soms wel dagelijkse begeleiding bij het project. Allebei unieke kwaliteiten om me met een frisse blik toch op het juiste spoor te houden, veel bruikbare tips gegeven om de thesis succesvol af te ronden.

Pieter, dank dat jij vertrouwen in mij had om een afstudeerstage te doen bij Albert Heijn en mij op te nemen in je team. Dank voor al je hulp, richting die je gaf en support gedurende de hele reis die afstuderen heet. Naast het schrijven van de thesis heb je me ook een inkijk gegeven in hoe de afdeling Logistics Support in elkaar steekt en wat er allemaal gebeurt hier. Boyd, Berry en Arjan voor de gegeven hulp en inkijk in de verschillende teams van afdeling.

Daarnaast, wil ik een dank uitspreken voor alle andere collega's van de afdeling Logistics Support van Albert Heijn voor het meenemen in alle taken, opdrachten, projecten en alle andere bronnen en informatie die nodig was voor het afronden van de thesis. Zonder alle inzichten, Excel bestanden, enquêtes en interviews was het niet mogelijk om dit onderzoek te doen. Ik wil ze hartelijk bedanken voor hun openheid en eerlijkheid en heb het gevoel dat ik echt onderdeel geworden ben van het team. Met als mooie toekomst dat ik de afdeling mag ondersteunen op gebied van Business en IT. Dank daarvoor!

Ook wil ik al mijn vrienden en familie bedanken voor het geven van genoeg afleiding, inspiratie en geloof dat ik deze thesis succesvol kan afronden. Speciale dank gaat uit naar mijn oom en tante en nicht die voor mij een plekje vrij hadden in hun huis zodat ik niet elke dag op en neer hoefde te gaan naar Enschede. Elke week met veel plezier en dankbaarheid bij jullie gelogeerd en veel waardevolle gesprekken gehad. Dank pa en ma dat jullie altijd voor me klaar stonden in goede en slechte tijden. En als laatste een grote dank naar mijn broers op wie ik altijd kan rekenen.

Wieger Kloppenburg

Enschede, 29 maart 2019

Management Summary

Slotting is one of the most customer-sensitive, labor-intensive, and complex problems encountered within different companies (Jones & Battieste, 2004). It is the allocation of products in the warehouse based on different rules or strategies. There are many different algorithms and methods to calculate the optimal assignment of the location for products in the warehouse. Different policies and processes on how the warehouse should be arranged are designed but there is no uniform way of the allocating products. Currently, this time-consuming process is done manually. This research explores the possibilities for partially automating this process and thus saving time and, eventually, money.

The slotting process contains many tasks where a lot of expertise is involved. And thus, this process is difficult to fully automatize. We therefore introduced the concept of Intelligence Amplification (IA) to the process of slotting. IA is the symbiosis between human and the computer. The decisions should be made under construction of cooperation between humans and machines instead of depending on the predetermined programs, exclusively when it becomes very complex. In other words, the capabilities of the human brain will be extended with the computer power to achieve a system superior to information processing.

The slotting process is analyzed and split up in different tasks. After the assignment of tasks to human and machine, the tasks are categorized on different categories. Most tasks of the slotting process are based on rules. One of those rule-based tasks of slotting is the data validation of products. The numerous consults of experts were held to understand the purpose of the need for validation because it is a time- and energy-consuming task. In respect to that conclusion, there is a need for a reference architecture for rule-based tasks.

With adopting the concept of Intelligence (Burstein, Muivehill, & Deutsch, 1999) Amplification, we designed a reference architecture to handle rule-based tasks. In other words, a solution is created which focusses on a close collaboration between the human and machine to empower the human in the process of rule-based tasks, in this case, the slotting process. The prototype of data validation, designed and developed in this research, proves that the reference architecture works. The results of the testing of the prototype give enough hints and insights for the implementation of the tool. The workers' performance will be enhanced as well as the quality of work in the slotting process.

The contribution to the research field of intelligence amplification is a reference architecture for rule-based tasks in the slotting process with help of the user. This reference architecture could support future research in the application of intelligence amplification in different sectors. Without the help of humans, machines will never be smarter. To ensure that control remains with the human, the machine remains as smart as we make it ourselves.

The contribution to practice lies in the explanation of the concept of IA and the prototype of the validation tool based on the reference architecture. In principle, many rule-based tasks could be automatized with help of the reference architecture. With the prototype we contributed to the knowledge of AH, by demonstrating how data validation can be done. This in combination with people who are not programmers by nature, but who can create rules that deliver results or advice.

Table of contents

Preface.....	iii
Management Summary.....	iv
Table of contents.....	v
Index of figures.....	vii
Index of tables.....	ix
Abbreviations	x
1. Introduction.....	1
1.1. Research background	1
1.2. Problem statement.....	2
1.3. Research goal.....	3
1.4. Research questions.....	3
1.5. Research method.....	4
1.6. Thesis structure	5
1.7. Modeling language and notation	5
2. Literature review.....	7
2.1. Systematic literature review.....	7
2.2. Additional information gathering.....	9
2.3. Research topics – Slotting.....	10
2.4. Research topics – Intelligence Amplification.....	15
3. Current situation	26
3.1. Introduction	26
3.2. High-level process supply chain.....	27
3.3. Slotting process	29
3.4. Current architecture	34
3.5. Current situation performance indices.....	39
3.6. Conclusion	46
4. Future situation.....	47
4.1. IA Framework	47
4.2. Goals	50
4.3. Platform.....	51
5. Design and Implementation.....	52
5.1. Introduction	52
5.2. Reference Architecture.....	53
5.3. Design and development of application.....	59

6. Validation	69
6.1. Validation criteria	69
6.2. Results	69
6.3. Conclusion	72
7. Conclusion.....	74
7.1. Introduction	74
7.2. Literature	74
7.3. Current Situation	75
7.4. Reference Architecture.....	76
7.5. Prototype and validation	77
7.6. Final conclusion	77
7.7. Contribution	78
7.8. Limitations	78
7.9. Recommendations.....	79
References.....	80
Appendices.....	87
A – Taxonomy proposed by Sheridan and Verplank (1978).....	87
B - Survey ‘Slotting’	88
C - Explanation of each slotting task	89
D - Feedback.....	90
E - User Interface.....	92
F - Subprocesses of the slotting process	93
G – List of errors occurring the slotting process	95

Index of figures

Figure 1: Design Science Research Methodology of Peffers et al. (2007) in combination of Wolfswinkel et al. (2013).....	4
Figure 2: Five-staged Method by Wolfswinkel et al. (2013).....	7
Figure 3: Literature Selection Process.....	8
Figure 4: Selection process for slotting articles.....	10
Figure 5: Selection Process: Intelligence Amplification.....	16
Figure 6: General description of an adaptive storage policy.....	23
Figure 7: Interactions between agents (Tsamis et al., 2015)	23
Figure 8: Swimlane model for the process: 'New product in assortment'	27
Figure 9: Slotting process	30
Figure 10: Subprocess: 'Slot product'	30
Figure 11: Subprocess: 'Pre-process data'	31
Figure 12: Subprocess: 'Check package structure of product'	32
Figure 13: Subprocess: 'Check prognosis of SDF'	32
Figure 14: Subprocess: 'Check ideal location'	33
Figure 15: Subprocess: 'Determine the type of location'	33
Figure 16: Business view of the current architecture	35
Figure 17: Application view of the current architecture	36
Figure 18: Technology view of the current architecture.....	36
Figure 19: Total view of current architecture	38
Figure 20: Percentage of free locations per dc	43
Figure 21: Number of movements, new products, and sequence movements.....	44
Figure 22: IA Framework.....	47
Figure 23: Task decomposition of the slotting process.....	48
Figure 24: Role allocation for information processing behaviors (skill, rule, knowledge, and expertise) and the relationship to uncertainty	49
Figure 25: Slotting platform	51
Figure 26: Example of an algorithm	52
Figure 27: The Black box model	52
Figure 28: The White box 'plus' model.....	52
Figure 29: Relationship between the goal and the driver, the assessment and the outcome	54
Figure 30: The goal definition process	54
Figure 31: Business layer of the reference architecture, the rule creation process	55
Figure 32: Business layer of the reference architecture, the goal validation process	55
Figure 33: Application layer of the reference architecture.....	56
Figure 34: Technology layer of the reference architecture.....	56
Figure 35: Total view of the reference architecture	57
Figure 36: Entity relationship model of the application.....	58
Figure 37: Trello.....	60
Figure 38: Data validation process	61
Figure 39: Business layer of the validation tool, the rule creation process	62
Figure 40: Business layer of the validation tool, the goal validation process	62
Figure 41: Application layer of the reference architecture, the goal validation process.....	63
Figure 42: Technology layer of the validation tool.....	63
Figure 43: Total overview of all the layers of the validation tool.....	64

Figure 44: Technology Readiness Levels summarize.....	72
Figure 45: Slotting process	75
Figure 46: Survey – Slotting.....	88
Figure 47: Screenshot ‘create/update’ rule	92
Figure 48: Screenshot ‘overview rules’	92
Figure 49: Subprocess: ‘Unlink product’	93
Figure 50: Subprocess: ‘Mail to Replenishment for the error/ if the error’	93
Figure 51: Subprocess: ‘Check other product in DC’	94

Index of tables

Table 1: DSRM Phase and research questions per chapter	5
Table 2: Concept framework Wolfswinkel et al. (2013) based on Webster and Watson (2002)	9
Table 3: Search query: Slotting.....	10
Table 4: Literature results for Slotting	11
Table 5: Search query: Intelligence Amplification.....	15
Table 6: Literature results of the IA concept.....	17
Table 7: Locations of all DCs of Albert Heijn	26
Table 8: Results for question three of the survey 'Slotting'	41
Table 9: Results for question four of the survey 'Slotting'	42
Table 10: Number of errors per week per category.....	45
Table 11: Number of unforeseen items per week per category	46
Table 12: Assigning decomposed tasks to the human and the AI.....	49
Table 13: Sprint 1 user stories.....	66
Table 14: Sprint 2 user stories.....	68
Table 15: Taxonomy proposed by Sheridan and Verplank (1978)	87

Abbreviations

AH	Albert Heijn
DC	Distribution Center
DSRM	Design Science Research Methodology
IA	Intelligence Amplification
IVFC	Interface Vivaldi Fulfillment Center
KPI	Key Performance Indicator
LDC	National Distribution Center (Landelijk Distributie Center)
LS	Logistics Support
NASA	Nieuw Artikel Systeem Albert Heijn
RDC	Regional Distribution Center (Regionaal Distributie Center)
SDF	Stored Demand Forecast
SDRS	Store Demand Reporting Service
SFC	Shared Fresh Center
SKU	Stock Keeping Unit
SLAP	Stored Location Assignment Problem
SWK	Shared Warehouse Cheese DC
WMS	Warehouse Management Systems

1. Introduction

1.1. Research background

Today's warehouses must execute more and smaller transactions, handle, and store more products, offer more product and service customization, and need to provide more value-added services while having less time to process orders and with less margin for error (Petersen et al., 2005). A warehouse has typically five basic functions: receiving, sorting, storing, order picking and delivering. In order to Van den Berg and Zijm (1999), the product storing and order picking are the most resource-consuming activities. Effective warehouse management can reduce storage costs by reducing inventory levels while maintaining defined service levels. This is a hard challenge for big organizations (Tompkins, 1998; Keller & Keller, 2014). The product variety increase rapidly, warehouse workers are working under high pressure to improve the speed and reliability of the warehouse process (Pujawan et al., 2017). Slotting can help with the improvement of the warehouse process and the problem of storage allocation in the warehouse and create a more effective warehouse. Slotting is a method to determine the optimum location for products within a warehouse. It is used to reduce the amount of travel time for operators by placing the item that frequently ships together next to each other in the pick-face area (Richards, 2011).

As stated above the role of the warehouse will change extremely in the future (van den Broecke, 2018). It will change from a more physical way for communication to a more automated and internet way of communication between humans. More and more suppliers are in contact with the end-user of the product instead of the middleman like a warehouse. There are suppliers like Amazon which will facilitate all the required steps by itself to ensure that products can be delivered faster to the consumer. This looks like a more pessimistic way of thinking for the normal warehouses or retailers to sell products to their customers. But it will have also the opportunity to be the market leader in your segment. This can only be done when you guarantee the same speed, service, and reliability as a giant like Amazon. The positive side is that retailers have their own warehouses and suppliers do not have their own warehouses. Even with good communication between your company and the customers you can make a difference. You must change to be sure you can survive this world. To ensure that every customer gets his or her product on time, the role of your warehouse is very important. It should ensure that every product is available to be picked when needed. This can be seen as the 'slotting' process and needs to be optimized to be sure that every product is available and can be delivered as fast as possible. You need to be technology-minded. The trend is clear: technology is replacing people in supply chain management. Gartner present eight trends for improving the supply chain by new technology (Petty, 2018). One of them is the use of Artificial Intelligence for the digitization and automation labor-intensive, repetitive tasks and processes such as purchasing, invoicing, and parts of the customer service (DeAngelis, 2018). For the short term is that the executives' focus needs to shift from managing people doing mostly repetitive and transactional tasks to designing and managing information and material flows with a limited set of highly specialized workers or agent systems (Lyll, Mercier, & Gstettner, 2018).

The slotting has a big impact on the output of the warehouse and the number of load carries used in the whole supply chain. When the products in the distribution centre are stored in the wrong order, the stacking of the load carries will be substantially worse. For example, when a heavy product is placed on top of a light-weight product, the light-weight products will be crashed. To ensure that those errors will not prevent, it could imply that there are more load carriers are sent to shop. What will imply more trucks to the shop. To reduce this, it is important to ensure that the products lay in the right place in the distribution center (DC). The slotter has the competence and the responsibility to determine the slotting for every day for every DC of the Albert Heijn (AH).

The correct slotting of products is a complex task. An example of this was experienced recently within one of the DCs of the AH. The slotters allocate products based on their names. This caused them to

place six new soup types together with normal soups. However, the new soups were kept in glass bottles, which are not allowed to be stacked at the bottom of a carrier. The soups had to be reallocated, which resulted in a lot of extra movements than that would have been needed when placed correctly at once.

As already discussed, the human is responsible to determine the allocation of the locations for all products in the specific DC which need to be slotted each day. Without the knowledge of all different persons who do this job, the slotting will be a disaster. Those persons have been doing the job for more than five years and are creative enough to deal with different worst-case scenarios. The adaption for the different changes on the ideal scenario of slotting needs to be high due to the fact of different uncertain variables like weather impact, traffic jams and customer needs for that specific day of the week. Those variables have an impact on the slotting on a daily base. To deal with all the changes there are different rules created by the employees of the department, which will ensure that most of those changes can be handled by the employees.

1.1.1. Ahold Delhaize

Ahold Delhaize is established in 2016 and arose from a merger of Ahold and Delhaize Group in 2016. It is a “world-leading food retailer with 6500 stores worldwide and 375,000 people, serving 50 million satisfied customers a week” (Ahold-Delhaize, 2019c). It is active in the United States, Europe, and Southeast Asia and contains supermarkets, convenience stores, online delivery, pick-up points, hypermarkets, specialty stores and gasoline stations (Ahold-Delhaize, 2019b).

In the Netherlands, Ahold Delhaize has over 2,000 stores and distribution centers (DCs) and is the leading supermarket company, a leader in specialty stores and e-commerce company. The Dutch brands of Ahold Delhaize consist of AH, ah.nl, AH to go, bol.com, Etos and Gall & Gall. These brands resulted in net sales of €12.7 billion in 2018 (Ahold-Delhaize, 2019a).

1.1.2. Albert Heijn

AH is founded in 1887 and has grown to be the largest supermarket chain in the Netherlands. Currently, the brand possesses 1010 stores in the Netherlands and this number is still growing (Ahold-Delhaize, 2017). Among these stores around 40 percent are franchised. AH also has some stores in Belgium and Germany. According to IRI (te Pas, 2018), the total market share of AH drops from 35.4 percent in 2016 to 35.3 percent in 2017. The company's mission stated by its founder AH:

*“Het alledaagse betaalbaar, he bijzondere bereikbaar”
(The everyday payable, the extraordinary reachable)*

The brand has four different types of stores. The first and most familiar one is the neighborhood store. These stores are based on regular grocery shopping and the products are influenced by the neighborhood and the region of the store. The second type is the AH XL, an extra-large supermarket for the bigger grocery trips. It has more choice in products and more parking places. The AH to go is the third type of store AH facilitates. The AH to go concept is created to serve customers while traveling and offer them refreshments for the trip. The last and newest type of store is the online store where you can order your products 24/7. You can either have your groceries delivered at your house, or you can pick-up your groceries at one of the Pick-Up points.

1.2. Problem statement

Slotting is one of the most customer-sensitive, labor-intensive, and complex problems encountered within different companies (Jones & Battieste, 2004). It aims at the allocation of products in the warehouse based on different rules or strategies. Many companies use different rules or strategies.

Factories, for example, use the principle of placing high turnover products as close as possible at the input and output point of the distribution center. While as retail stores need to have a different strategy for placing the product in the warehouse based on assortment groups and high turnover products and other characteristics like weight and height of the products. This to ensure the correct stacking with respect to quality and safety issues. Often this process is carried out manually, costing a lot of time and user effort. Slotting is an inherently difficult task or process to automate because of the external uncertain variables. Every day the slotting will be different, and this requires the soft knowledge skills of the slotters.

Currently, the slotting process at the AH is a manual process, in which the employees use a tool which displays the DC as one whole list of locations and this is sorted on the pick sequence in the specific DC. This process is used to perform slotting for all DCS of AH and will be done by different personnel of AH. Furthermore, the slotting process of the AH includes the allocation of new products, which can contain errors in the data. When those occur the slotting of this product cannot be done, and the product cannot be picked, and thus not be shipped to the shop. Furthermore, clean data is essential for automation (Limoncelli, 2018; Rahm & Do, 2000). Not all the information what is available at the AH is free of errors. This needs an investigation to automate data cleaning with help of the Intelligence Amplification concept.

In this thesis, we aim to create a symbiotic method between the tool and human for slotting and all the subprocesses involved in this process. The computational power of the machine will be used to fast calculate different processes and values. The user can focus on different tasks and use those values for further processes. The machine will take on the repetitive and bothersome tasks to ensure that the human will do more complex tasks and have more time for those tasks. This approach is uncommon in science and gives personnel the possibility to reduce the time needed to perform slotting and the subprocesses.

1.3. Research goal

The goal of this thesis is to evaluate Intelligence Amplification's practical value in improving the amplifying factor in the extension of human capabilities and the system. This implies that the human will focus more on the creative and complex tasks in the process and the systems' tasks will be focused on more on the repetitive tasks. To achieve this research goal, this research needs to find out an approach to implement IA and how to enhance human capabilities.

The approach of applying IA into the whole slotting process, this thesis focuses on the business operations of decision making in the slotting process. Furthermore, all the complementary processes will be sure the most inefficient process is tackled by applying IA. The requirements and goals vary from person to person in the organization. Each process requires a variety of activities and scenarios and will be done differently by the persons. In addition, how to enhance human capabilities and intelligence by IA. The humans in the process need to make a shift to focus more on the complex parts of the process instead of the easy or repetitive tasks. The solution of this thesis focuses on the implementation and creation of a reference architecture for supporting rule-based tasks. This in combination with the knowledge of humans to give importance to different goals to support their tasks.

1.4. Research questions

Based on the research context, problem statement and research goal, the main objective of this research is to find out if intelligence amplification adds value when performing repetitive tasks as seen in the slottings process. Therefore, the research question is formulated to contribute to this objective is as follows:

“How to enhance the workers’ performance and quality of work in the slotting process through the adoption of Intelligence Amplification?”

To aid the research question, the following sub-questions have been formulated:

- What is the current state of the literature of Intelligence Amplification and slotting?
 - What are the applications of IA in enhancing worker's performance or reducing workload?
- How do the current process and all the complementary processes of slotting work in the warehouse of Albert Heijn?
- What is the less optimized or most time-consuming process in the process of slotting?
- Which (sub-) process of slotting benefits most of being amended by the utilization of intelligence amplification?
- How can we design a solution for the most time-consuming process of slotting improved by intelligence amplification?
- How well does the chosen solution meet the original requirements?

1.5. Research method

The design science research methodology (DSRM) of Peffers et al. (2007) is used to guide this thesis.

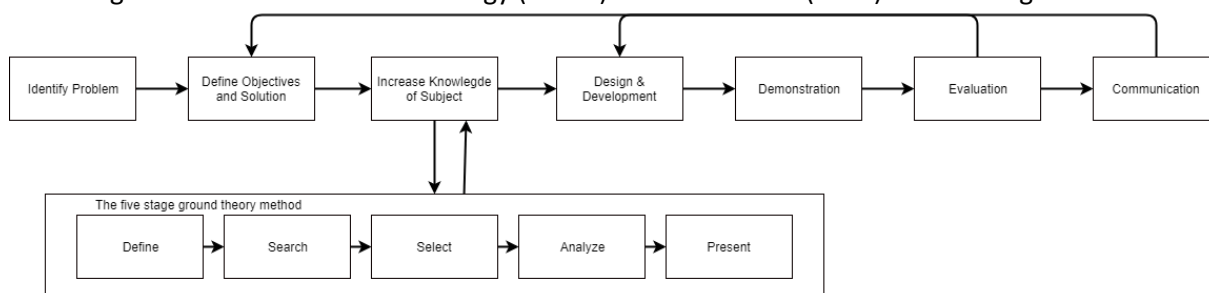


Figure 1: Design Science Research Methodology of Peffers et al. (2007) in combination of Wolfswinkel et al. (2013)

This methodology exists in different phases. Based on the DSRM, the following phases have been distinguished in the current research:

- **Problem identification and motivation:**

During the initial phase, the specific research problem is identified and defined. The research is motivated by showing the importance of the problem.

- **Defining the objectives for a solution:**

In the second phase the objectives for the solution will be defined; which includes a definition of which KPIs are measured, the definition of how the solution is expected to support the problems addressed in phase one.

- **Increase knowledge of the subject:**

In this phase, all the required information and knowledge will be collected. This will be done by the rigorous process of Wolfswinkel et al. (born). Because of the iterative nature of the process of Wolfswinkel et al. the identification of possible new, changed or combinations with regards to the core concepts are more easily identified. All that knowledge will be used to create a solution for the addressed problems in phase one.

- **Design and development:**

This phase involves the design and development of the solution, describing the to-be situation of the slotting process.

- **Demonstration and Evaluation**

In the DSRM, the phases of ‘Demonstration’ and ‘Evaluation’ are two separate phases. In the current research, the phases of demonstration and evaluation are therefore adopted into a single phase. In this phase, the solution is presented to the team and experts and validated by means of structured interviews and reviews of the given solution for the slotting process. This to measure how well the solution has been defined and helps in further developments for automation the slotting process.

- **Communication:**

The current research will be presented in a colloquium that is part of the examination of the master project. In addition, the main results of the current research will be processed into a publication.

1.6. Thesis structure

DSRM	Chapter	(Sub)Questions
Problem Identification and motivation	1. Introduction	
Defining the objectives for a solution	1. Introduction 3. Current Situation 4. Future Situation	
Increase knowledge of the subject	2. Literature review 3. Current Situation	What is the current state of the literature of Intelligence Amplification and slotting? <i>What are the applications of IA in enhancing worker's performance or reducing workload?</i> How do the current process and all the complementary processes of slotting work in the warehouse of Albert Heijn? What is the less optimized or most time-consuming process in the process of slotting?
Design and development	4. Future Situation 5. Design and Implementation	Which (sub-) process of slotting benefits most of being amended by the utilization of intelligence amplification? How can we design a solution for the most time-consuming process of slotting improved by intelligence amplification?
Demonstration and Evaluation	5. Design and Implementation 6. Validation 7. Conclusion	How well does to chosen solution meet the original requirements?
Communication	7. Conclusion	All

Table 1: DSRM Phase and research questions per chapter

1.7. Modeling language and notation

Within this research, different models are presented. For these models, there are two modeling languages used: the ArchiMate[®] language and notation for architectural models, and the BPMN for business and process modeling.

1.7.1. ArchiMate

ArchiMate is an open and independent modeling language for enterprise architecture that is supported by different tool vendors and consulting firms. ArchiMate provides instruments to enable enterprise architects to describe, analyze and visualize the relationships among business domains in an unambiguous way. ArchiMate enables modeling of the architecture domains defined by TOGAF[®], a proven enterprise architecture methodology, and framework used by the world's leading organizations (The Open Group, 2016).

1.7.2. BPMN

In the current research, the Business Process Model and Notation (BPMN) is applied to sketch the different processes. BPMN is a graphical representation for specifying business processes in the business process model. The goal for BPMN is to provide a standard notation readily understandable by all the business stakeholders. These include most of the time business analyst who creates and refine the processes, the technical developers responsible for implementing them, and the business managers who monitor and manage them. BPMN serves as a bridge for the business process design and implementation.

2. Literature review

2.1. Systematic literature review

This section describes the underlying process, method, and structure utilized for the systematic literature review of this research. The focus of the systematic literature review will be on slotting, and intelligence amplification. The topic of machine learning will be explained in a sub-question in the intelligence amplification topic.

There is chosen to perform a literature review in multiple topics to get more knowledge about the different topics. The knowledge is needed to create a good artifact for the pre-processing of data, which implies the validity of the data. Furthermore, with help of machine learning and intelligence amplification the (semi-)automated slotting tool can be made. The slotting process will be executed differently in each warehouse. Each company or warehouse has its own in-house knowledge about the specific warehouse and the products within that warehouse. With the symbiosis between human and machine, the process can be lifted to a higher level, due to the computational speed of the machine and the creativity and in-house knowledge of the human involved in the process. Strengths of humans are creativity (DiBona et al., 2016), self-reflection and the ability to perform a variety of tasks (Barca & Li, 2006). Machines can be designed to perform a specific task, can easily be replaced and can perform non-stop routines (Barca & Li, 2006). The human is needed to document and automate the process (Limoncelli, 2018). With the documentation in the pseudo code, the implementation of the automation and the software will be easier. The automation of different processes can be made easier when all the steps are recognized and be done manually.

To perform a structure literature review, the method of Webster and Watson (2002) is used as a guide. This method can be used to ensure that all the concepts of this thesis are covered by an adequate amount of literature (Nakagawa et al., 2010) and being able to synthesize works of different authors into contributing findings (Webster & Watson, 2002). The first step is to formulate questions that are needed to be answered through the literature review. The next step is to perform the search queries and selection process. Next, a short literature overview is presented. The last step is, with the help of a backward search, to add more and check papers, to get enough knowledge and experience from all the studied literature. A backward search, also known as a reference search, will enable the researcher to examine the sources and identify inconsistencies which are used by the authors of the selected papers.

2.1.1. Literature search process

To be more concrete in the literature search process to answer this question, the refinement process based on a rigorous method by Wolfswinkel et al. (2013) is used. Because of the iterative nature of the process of Wolfswinkel et al. the identification of possible new, changed or combinations with regards to the core concepts like slotting and intelligence amplification are more easily identified. To further structure the systematic literature review such that the research is replicable with the same results, although results obtained by using digital libraries is hard to replicate (Kitchenham & Brereton, 2013). The theory starts with the first of five steps by the definition of the search queries and if needed the research area. Most of this step is basically done in the introduction, motivation and problem statement. The second step of the theory is doing the search. The used engines in this research are Google Scholar and Sciverse Scopus of Elsevier. The method is represented as a scheme in Figure 2: Five-staged Method by Wolfswinkel et al. (2013).

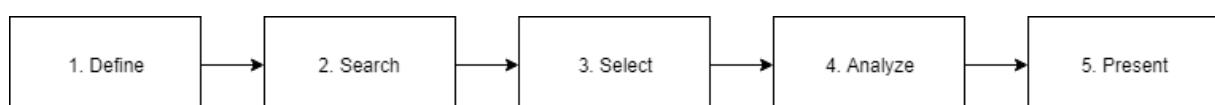


Figure 2: Five-staged Method by Wolfswinkel et al. (2013)

Step four is the selecting of relevant papers for the concept. First, the papers will be received with all the different keywords specified per concept. Second, the duplicates will be removed. After that, the title will be used to limit the papers. If the title deemed interesting for further investigation, the abstract will be analyzed. The abstract will be determinative for the inclusion of exclusion because the abstract will mention the core concepts of the paper. If the abstract is not clear enough but looks interesting, the introduction is used to ensure if the paper is interested enough. The paper will be included in the body of literature if all the above requirements are met. Other guidelines are:

- Sorted by citation count
- Limit to the "Computer Science", "Social Science", "Decision Science", "Engineering" and "Business management and accounting".
- Paper after 2000
- The paper needs to be available. Public or via the external databases to which the UT is entitled.

All steps and rules described above result in the process explained in Figure 3. In this process, forward and backward citation is used. Forward citations are the found citations that are used by the analyzed paper in their own research. Backward citations are the papers that are used by the author of the paper to support the author's own research.

All the phases of the systematic literature review selection process are displayed. The process of searching for literature is continued and repeated until each final concept adequately covered by enough body of literature.

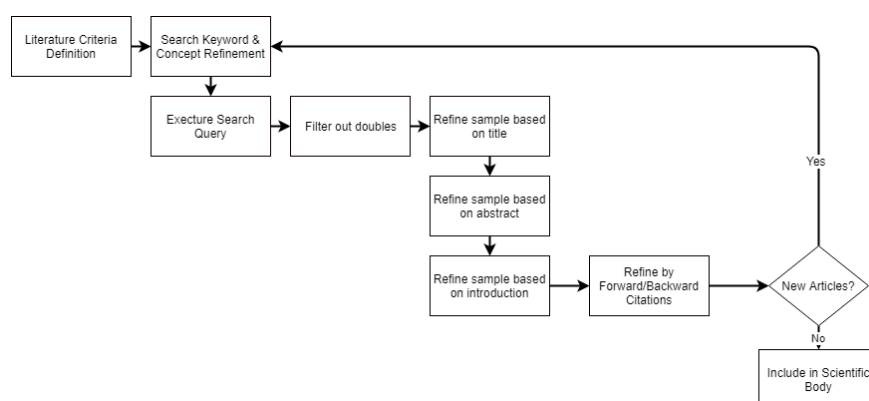


Figure 3: Literature Selection Process

With the use of a strict review protocol, a more rigorous selection of applicable literature will be ensured. The disadvantage of this process is that other possible interesting findings will not be recognized and thus excluded. Highly relevant publications with modest citation counts are left out (Tamm et al., 2011). Keeping on a strict keyword or search word list, synonyms for the concepts could be skipped. This could be tackled with a synonymic search queries search. In this thesis, there are synonyms used for the word warehouse: 'distribution center', 'distribution centre', 'depository', 'depot', and 'storehouse'.

The fourth step of the method is analyzing all the relevant papers which are found. The coding of the papers, by picking one by one random and highlight any finding and insights what seems to be relevant for the concept. Every concept which is relevant for the thesis has its own concept matrix (see Table 2), suggested by Wolfswinkel et al. (2013), which is founded on the work of Webster and Watson (2002). The concept matrix and the process are discussed per concept of the literature review. Step five, the present step, will elaborate all the results in its separate chapter. The concepts are derived from the initial search in the preparation phase of this project and added with some extra concepts found during the literature review.

Articles	Concepts
----------	----------

	A	B	C	...
1		X	X	
2	X	X		
...			X	X

Table 2: Concept framework Wolfswinkel et al. (2013) based on Webster and Watson (2002)

2.2. Additional information gathering

To be sure that all the information is gathered. There are two other methods used to get all the information, exploratory literature review and the empirical research.

2.2.1. Exploratory literature review

An exploratory literature review is to explore the research questions and does not intend to offer final and conclusive solutions to existing problems. This type of literature review is usually conducted to study a problem that has not been clearly defined yet. An exploratory literature review can be used to get a better understanding of the problem (Saunders, Lewis, & Thornhill, 2011).

An exploratory review will help in the determination of the research objective and goal because of the possibility of a too large or ambitious project. This will form the foundation for the initial research questions and rudimentary research goal. A second application of the exploratory review is towards the finding of relevant literature for extending the scientific body for this thesis-research. The topics will possibly span multiple subject areas which could be excluded because of the use of the subject area literature criteria.

The inclusion of an exploratory review can have a great contribution to the finding, scoping of the research and aids in the finding of an applicable development method for the solution of this project. The exploratory literature review is included in the entire review and search process for research scoping and additional material finding purposes. Furthermore, it will be used to observe the information which is available at the Albert Heijn Logistics Support team and maybe other teams.

2.2.2. Empirical research

Empirical research focuses on the gathering of knowledge through observation and experience to describe the relationship between attitude and behavior of a person or group. The empirical cycle consists of five steps (de Groot, 1969). Observation of a phenomenon and inquire concerning its' causes. Induction is the formulation of hypotheses. The deduction is the formulation of experiments that will test the hypotheses. Testing is the procedures by which the hypotheses are tested and data are collected. Evaluation is the interpretation of the data. When those steps did not lead to the proposed result, the hypotheses can be rewritten and do the experiment again (de Groot, 1969). Within this research this discipline is mostly applied by visiting several warehouses, talking with people on the site, talking with people on different departments, but the most important is the fact that we have been part of the team of Logistic Support of Albert Heijn for half a year. One of the contributions of this type of research is the extension of the in-house knowledge of the Albert Heijn. Employees perform certain processes the way that they have learned it at the AH, which cannot be found in the literature. The information will be gathered with semi-structured interviews or by job shadowing the employees in their day-to-day work. Job shadowing is an excellent way to learn more about the experience of the profession (Irby et al., 2016). Another contribution is to learn from the daily routines of the company. By constantly documenting, the future work of automation will be accelerated. That one-shot task that could never happen again, does happen again, and next time it moves faster (Limoncelli, 2018).

2.3. Research topics – Slotting

This section first describes the followed process during literature research. Second, the slotting process is explained. The slotting location assignment problem is discussed to find the more complex problem of Albert Heijn.

2.3.1. Literature search and selection process

Based on the research goals and the research questions, it is necessary to find out what has been studied about slotting by answering the following questions.

- What is slotting?
 - What is the storage location assignment problem (SLAP)?
- Which best practice methods can be used to slot a product in the warehouse?
 - What are the most used techniques?

The knowledge questions can be answered by a literature review. Table 3 shows the different search queries as entered in the websites of *scopus.com* and *scholar.google.com* on 31 August 2018 and the search results of each query.

Search query	Search results
TITLE-ABS-KEY ("slott*" AND ("Warehouse" OR "Distribution Centre" OR "Distribution Center" OR "Depository" OR "Depot" OR "Storehouse"))	60
TITLE-ABS-KEY ("Storage location assignment problem")	35
TITLE-ABS-KEY ("Intelligen*" AND "Slott*" AND ("Warehouse" OR "Distribution Centre" OR "Distribution Center" OR "Depository" OR "Depot" OR "Storehouse"))	1
TITLE-ABS-KEY (("Slotting process" OR "Slotting technique" OR "Slotting strateg*" OR "Slotting method" OR ("Storage location assignment" AND ("method" OR "technique")) AND ("Warehouse" OR "Distribution Centre" OR "Distribution Center" OR "Depository" OR "Depot" OR "Storehouse"))	45
Total	141

Table 3: Search query: Slotting

The initial search leads to 141 usable articles in the literature. Figure 4 illustrates the selection of relevant articles for this concept. First, all the duplicates are removed, after that, there was a selection based on the title and the abstract. If the article looks interesting the introduction was read. Then, all the articles which were not available are removed. Finally, some articles were added with the help of backward or forward citation. The final papers were selected according to the relevance of the content.

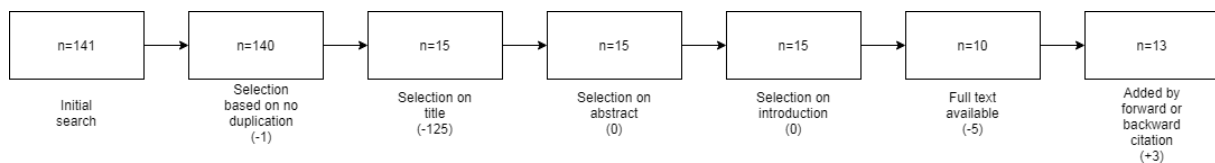


Figure 4: Selection process for slotting articles

Table 4 provides an overview of the resulted literature and the knowledge questions the articles apply to.

	Article	Concepts		
		Slotting	'Storage Location Assignment Problem (SLAP)'	Methods and techniques
1.	H. Brynzer and M.I. Johansson		p. 595-596	
2.	A. Fumi, L. Scarabotti, and M.M. Schiraldi		p. 1-2	

3.	W.H. Hausman, L.B. Schwarz, and S.C. Graves			p. 631-636
4.	Y. Hu, S. Zhang, X. Cheng	p. 143-146		
5.	N.P. Iyer and R. Jayal	p. 326-328		p.328-330
6.	C.C. Jane			p. 59-63
7.	E. Jones and T. Battieste	p. 37-38		
8.	M. Li, X. Chen, and C. Liu			p. 2-4
9.	K.W. Pang and H.L. Chang		p. 4036-4038	p. 4039-4043
10.	C.G. Petersen, C. Siu, and D.R. Heiser	p. 997-1009		
11.	G. Richards	p. 65		
12.	J. Xu, A. Lim, C. Shen, and H. Li			p. 1897-1902
13.	Q. Zu, M.Cao, F.Guo, and Y. Mu			p. 19-20

Table 4: Literature results for Slotting

2.3.2. What is slotting?

Slotting is a method to determine the optimum location for products within a warehouse. It is used to reduce the amount of travel time for operators by placing the item that frequently ships together, next to each other in the pick-face area (Richards, 2011). It refers to the optimal location for the products in the warehouse such that they can be picked in the least time with the maximum accuracy and maximum space utilization (N. P. Iyer & Jayal, 2016). Richards (2011) states that slotting will determine how many and what size of the pick is required for each product line. The fast-moving products are spread efficiently across the warehouse. The needed information to determine the most effective picking system includes:

- dimensions and weight of the product;
- product group;
- the total number of products by category;
- the total number of orders in a period;
- the total number of deliveries;
- mode and the average number of lines per order and line;
- pick-face visits per product;
- typical family groupings;
- items sold together with frequently.

Iyer and Jayal (2016) explained slotting as a process where products and locations are ranked and slotted or grouped according to different attributes. Slotting concerns may differ across warehouses. For example, a promotional product is placed closer to the loading docks, so that they can rapidly be picked and shipped (N. P. Iyer & Jayal, 2016).

Jones and Battieste (2004) mentioned that slotting is crucial for warehouse cost efficiency. The relationship of slotting to order picking is often overlooked. Order picking is time sensitive it is important that clients get orders within the promised time-frame. Optimal slotting usually considers five principles: popularity, similarity, size, characteristics, and space. Most of the warehouse managers focus only on the popularity or the inventory of the product (Jones & Battieste, 2004). 'Popularity' is stocking the most popular products in such ways that it will minimize travel distance and time. So, warehouse managers use this philosophy to move the fast movers close to the outbound location.

In extension to Richards, Petersen et al (2005) describe that slotting is the assignment of items or stock-keeping units (SKUs) to warehouse storage locations. There are different strategies or methods to slot the warehouse:

- Popularity: The number of requests for a given SKU. This will be translated as the number of times that the picker needs to travel to the storage location for the given SKU (Napolitano, 1998).
- Turnover: The total number of SKU shipped during a given period of time, sometimes called the demand for the given SKU.
- Volume: The demand for the SKU multiplied by the volume of the SKU, sometimes called the cube movement of the SKU.
- Pick density: The ratio of popularity for the given SKU to the volume of the SKU. This identifies the SKUs that have the highest pick activity per a given amount of space.
- Cube-per-order-index (COI) (Frazelle, 2001): COI is the ratio of the cube of an SKU to the turnover of the SKU with SKUs ranked in ascending order of the index (Kallina & Lynn, 1976).

According to Harmatuck (1976), the COI strategy is optimal when retrieving a single product by only place it in the warehouse and send it to another shop. Popularity is the most studied strategy in slotting literature (Petersen & Thompson, 2003). When slotting has been completed, the warehouse manager must still decide on which warehouse locations the slotted product will be assigned.

When deciding which is the most appropriate storage strategy for a warehouse, Gu et al. (2007) define the holding efficiency and the access efficiency as the two most important criteria. The holding efficiency corresponds to the capacity that is needed for storing a certain number of items. The access efficiency concerns the ease at which the products can either be retrieved from the pick locations or inserted to the locations.

Slotting in a warehouse can impact the warehouse operations significantly. These operations contribute to the overall cost of supply chain management. Warehouses become more valuable to the company and the supply chain as a whole when they contribute to reduced costs and improved service, flexibility, and responsiveness (Keller et al., 2014). Inefficient slotting can increase the travel times, pick orders and inefficient space utilization, which will lead to increased costs and reduced customer satisfaction (N. P. Iyer & Jayal, 2016). Optimization of the slotting can rapidly increase the sales by moving the right products from the supplier to the stores (Blanchard, 2007).

2.3.2.1. What is the 'Storage Location Assignment Problem'?

An important decision in distribution centers is determining the most efficient assignment of items to locations. This problem is known as 'the storage location assignment problem' (SLAP) and called by practitioners as the problem for inventory slotting (Hu et al., 2017; Jones & Battieste, 2004; Fumi, Scarabotti, & Schiraldi, 2013). Many of those SLAP rules and policies are based on the demand of a specific item. When a product is more popular than another, the product which is more popular will get a bigger space in the distribution center. The demand for the products can change due to the following facts: the introduction of new products, product maturity, or seasonality (Pazour & Carlo, 2015). Typically, when goods are initially stored in the warehouse, the SLAP is solved to optimize the assignments of goods to a storage location. The initial storage location assignment used eventually becomes inefficient as the demand changes. There is some downtime needed to reorganize the whole warehouse by repositioning items according to the solution of the SLAP with the updated demand. This will be done in a frequent way to rearrange the warehouse (Richards, 2011).

There are several researchers who create different algorithms or solutions to tackle this problem more and more. For small-scale problems, the branch and bound algorithms are generally used (Muppani & Adil, 2008). For large-scale problems like at by the Albert Heijn, many intelligent algorithms, like simulated annealing algorithm (Wang et al., 1996) are used in many studies.

2.3.3. Which best practice methods or techniques are used to slot a product in the warehouse?

According to Petersen et al. (2005), there are five different slotting strategies or methods; popularity, turnover, volume, pick density, and COI. Those are explained in the previous section. The warehouse storage location assignment is a widely studied stochastic resource allocation problem (Xu et al., 2008). There are different kinds of warehouses; automated- and manual warehouses. This will imply that there are different kinds of methods or techniques used to create the best slotting in the warehouse. Ciervo (2018) explains which elements and data are crucial to make a strategy for slotting. Historical data, product data, and the location in which the products fits.

Hausman et al. (1976) present three different storage location assignment policies: randomized storage, class-based storage, and dedicated storage. The randomized policy allows products to be stored anywhere in the warehouse pick area. Class-based policy distributes products based on their demand rates and locates the product to different pick areas. Dedicated storage policy each location may be only be used for that specific product.

According to Iyer and Jayal (2016), there are some variants of the slotting process.

- Batch slotting refers to the scheduled creation of a slotting plan for a group of products and locations.
- Re-slotting refers to the renewal of the slotting plans to reflect a business scenario.
- Dynamic slotting refers to the dynamic alteration of the above slotting plans as and when a new item or location is added to the data mix or when orders changed.
- Demand-driven slotting refers to the matching of slotting strategies to demand the signals from the demand-driven supply chain.

Jane (2000) proposes a heuristic method based on historical customers' orders for assigning products into storage zones and constructs a performance index for measuring continuity of each order handled in the pick system. The method aims to allocate all products into proper storage zones to distribute the workload equally to all pickers in the warehouse. In other words, to minimize the difference that might exists between each picker's total number of picks. The heuristic method performs five different steps:

1. Count all the times that product x is demanded in all orders;
2. Compute the pick rate of the product;
3. Sort all products in a descending matter according to the pick rate;
4. Allocate all the products in a sorted manner.
5. Repeat all the following sub-steps until all products are slotted;
 - a. Consider the product which has the largest pick rate among all the unassigned products;
 - b. Compute the summation of all the pick rate of already assigned products in the warehouse;
 - c. Find the best location for the selected product;
 - d. Assign the product to that location.

Pang and Chan (2017) created a storage location assignment algorithm to determine the strategic location to store products to minimize the total travel distance. The method first calculates the weighted travel distance for allocating the popular products. Those products need a closer location to the buffer than less popular products. All those products have a rank and based on this rank the best storage location can be determined. The popularity and association of the products will be determined by analyzing the order history using a data mining algorithm (Pang & Chan, 2017). Furthermore, there are some constraints which need to be considered. There is a constraint that every product must be

allocated to one location in the warehouse. Another one guarantees that all the locations must have a product.

Xu et al. (2008) found a solution for the online warehouse storage assignment problem based on a new designed deterministic approach. They created a framework of the iterative heuristic algorithm. With help of analysis and experiments, the proposed methods can solve the problem very well and they are much more flexible than traditional policies. The framework used the greedy algorithm for the classical scheduling problem, it is the most optimal (Cormen et al., 2009). The greedy algorithm can be explained as the algorithmic paradigm that follows the problem solving heuristic of making the locally optimal choice at each stage with the intent of finding a global optimum (Cormen et al., 2009).

Genetic algorithms (GAs) are proved to be a promising intelligence technology in the engineering field (Goldberg, 1989; Baskar, Subbaraj, & Rao, 2003). To tackle the SLAP of different warehouses, the solution is often given in the Pareto-optimality sense (Li, Chen, & Liu, 2008). In simple words, a solution is Pareto optimal if it cannot result in further improvement of a goal without causing the degradation of another goal (Zitzler & Thiele, 1999). With help of the niche technique, the Pareto-optimum can be reached. The technique is to replace a product-parent by its child when it has a higher offspring than its parent. The algorithm for this technique is to: select two individuals from the parent, carry out crossover and mutation, then calculate the offspring of the two new individuals and compare them with that of the parent. The crossover procedure exists of two different products and takes the dismal labels between them (Li et al., 2008).

Last, Zu et al. (Zu et al., 2011) construct a multi-objective optimization mathematical model to calculate the slotting optimization of the warehouse. The model uses two slotting optimization principles. The first one is the turnover ratio principle. This can be explained as when a good has a high turnover ratio it will be placed in the front of the warehouse, nearby the loading docks. The other one is the weight distribution principle. This can be explained as that the heavier weight products need to be placed before light weighted products; otherwise, the products will be smashed. The process of the hybrid genetic algorithm, which is the base of the model, can be described in the following steps:

1. Encoding; the locations of the warehouse needs to be encoded with a coordinate (row, column, layer)
2. Create an initial population; all the products will be loaded into memory to calculate the difference in weight and its turnover rate.
3. Evaluate and calculate the value of the difference between weight and turnover rate.
4. Generate a random weight coefficient; this to select the best location for this product.

To conclude, there are many different algorithms and methods to calculate the optimized assignment of the location for products in the warehouse. There are different policies on how the warehouse should be arranged but there is no uniform way of doing the slotting in it. There is no uniform process of assigning a location to a product in the warehouse.

2.4. Research topics – Intelligence Amplification

In this section, the Intelligence Amplification (IA) concept is investigated. First, the overview of the papers is given; in the second part, the knowledge questions are answered.

2.4.1. Literature search and selection process

Based on the research questions and research goals, it is fundamental to know what has been studied about the IA by answering the following knowledge questions.

1. What is Intelligence Amplification?
2. What is the role of the human in the IA system?
3. What is the relationship between human and intelligence agents?
4. What are the applications of IA in enhancing worker's performance?
 - a. What are the applications of IA in enhancing slotting or the slotting process?

The concept of “intelligent agent” is described in this thesis by the definition of Wooldridge (2002). “An *intelligent agent* is a computer system that is capable of independent (autonomous) action on behalf of its user or owner in some environment in order to meet its design objectives”. *Multiagent system* is one that consists of several agents, which interact with one another. For a successful interaction, the agents need the ability of cooperation, coordination and negotiation with each other, the same as people can do.

There hasn't been a uniform definition of the IA concept and researchers haven't reached consensus about its values in the practical application. To get enough related and useful articles, the literature search uses more than one query. Table 5 shows all the search queries as entered on the websites of *scopus.com* and *scholar.google.com* on 31 August 2018 and the search results of each query. With the help of the Asterix sign (*) all results for multiple words defining the same or similar words are included.

Search query	Search results
TITLE-ABS-KEY ("Intelligence Amplification" OR "Amplified Intelligence")	20
TITLE-ABS-KEY ("Augment* Intelligence" OR "Intelligence Augmentation")	22
TITLE-ABS-KEY ("*man computer symbiosis" OR "*man machine symbiosis")	97
TITLE-ABS-KEY (("man computer" OR "man machine") AND ("Collaboration" OR "cooperation"))	274
TITLE-ABS-KEY (("man machine symbiosis" OR "man computer symbiosis") AND "performance" AND ("work*" OR "Employee"))	2
TITLE-ABS-KEY (("work*" OR "employee") AND "performance" AND "*man machine collaboration")	11
TITLE-ABS-KEY ("robotic process automation" OR "Artificial Intelligence") AND "reduc* workload")	9
Summary	495

Table 5: Search query: Intelligence Amplification

The initial search resulted in 495 results. Figure 5 illustrates the selection of the relevant articles. First, all duplicates are removed. Second, the articles are selected based on the title. Then, articles are selected based on the abstract. If the abstract looks interested the introduction is read and used to define if it is relevant. After that, all the not accessible articles are removed. In the end, there can be articles added by backward or forward citation. The articles were selected according to the relevance of the content.

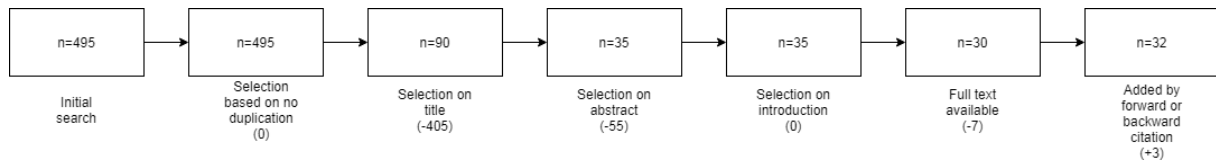


Figure 5: Selection Process: Intelligence Amplification

Table 6 provides the overview of which literature addresses which knowledge questions.

	Article	Concepts				
		Intelligence Amplification	Role of the human	The relationship between human and machine	Application of IA for enhancing worker performance	Application of IA for enhancing the slotting process
1.	Ashby (1956)	p.1-3	p.3-5; p.225-228			
2.	Bechar and Edan (2003)				p.432-435	
3.	Bradshaw et al. (2011)			p.1-5		
4.	Camara et al. (2015)		p.1-3; p.10			
5.	Coetzer et al. (2012)				p.2-4	
6.	Cummings (2014)			p.62-68		
7.	De Greef et al. (2007)				p.439-444	
8.	Deshmukh, McComb, and Wernz (2008)	p.1-5	p.6-9			
9.	Dobrkovic, Liu, Iacob, and Van Hillegersberg (2016)					p. 92-97
10.	Engelbart (1995)	p.3-27				
11.	Farooq and Grudin (2016)		p.27-32	p.27-32		
12.	Ferreira et al. (2014)		p.28-29			
13.	Griffith and Greitzer (2007)		p.42	p.43-49		
14.	Jaccuci et al. (2014)			p.5-13		
15.	Iyer and Jayal (2016)				p. 330	p.330-340
16.	Lange et al. (2014)			p.97-98		
17.	Licklider (1960)	p.4-11	p.4-11			
18.	Padoy and Hager (2011)				p. 5285-5288	
19.	Ren (2016)			p.104-106		
20.	Roy (2004)			p.121-123		
21.	Sankar (2012)	X (movie)				

22.	Smith (1983)	p.267-270	p.267-270			
23.	Tan et al. (2009)			p.152-154	p.153-155	
24.	Tsamis et al. (2015)					p.4-6
25.	Vagia, Transeth, and Fjerdingen (2016)	p.197-199	p.191-196			
26.	Vane and Griffith (2006)	p.1-3	p.3			
27.	Wigness and Gutstein (2018)				p.84-86	
28.	Wilson (2018)				p.19-25	
29.	Xia and Maes (2013)		p.2	p.2	p.2-5	
30.	Yamaba et al. (2008)				p.506-507	
31.	Zhang et al. (2015)				p.200-201; p.211-213	
32.	Zou and Zou (2017)				p.341-342	

Table 6: Literature results of the IA concept

2.4.2. What is Intelligence Amplification?

In the literature, the first person who mentioned Intelligence Amplification (IA) was William Ross Ashby. Ashby (1956) claims that the power of intelligential is equivalent to the power of appropriate selection. Licklider (1960) presents the idea of the symbiosis between human and the computer. He mentioned that decisions should be made under construction of cooperation between humans and machines instead of depending on the predetermined programs, exclusively when it becomes very complex. In other words, the capabilities of the human brain will be extended with the computer power to achieve a system superior to information processing. Licklider (1960) suggests a symbolic partnership between the functions of a human and a computer. So, the time to make the decisions is negligible compared to the time needed to prepare to make the decision and the action after the decision. This will imply that it can be delegated to the computer. Douglas Engelbart (1995) refers also to IA for the goal of the augmentation of the human brain or intellect by adding a higher level of synergistic structuring to their intellectual capabilities.

To conclude, the essence of human' essential role in problem-solving is highlighted by all the researchers. The strong part of humans is that they are flexible and capable of applying different solutions for discovering new patterns, identifying questions, and pose a trait of creativity, which are difficult to replicate for computers (Sankar, 2012). With respect to Licklider, one of the main goals of Intelligence Amplification, or Man-Computer Symbiosis, is to bring the computing machine effectively into the formative parts of technical problems. The other goal is closely related, to bring the computing power of the machine effectively into processes of thinking that must go on in 'real-time'. Furthermore, Licklider mentioned that the power of the computer of decision making is more in the routine tasks instead of the complex ones. For example, relevant actions suggested by both the human and computer, the computer will give the next action that is agreed with the human's initiate judgment based on a solution in the past. Artificial Intelligence focuses more on the replacement of the human. True artificial intelligence which can reason and decide like humans is still decades away (Ertel, 2017). There are quite a few human skills that nothing short of human-intelligence can replicate, like trivial tasks such as picking up items with different shapes and placing them in a basket. This is an extremely complicated task from an AI perspective, again with help of human intelligence this should be made easier (Sankar, 2012).

2.4.3. The role of the human in Intelligence Amplification

According to Griffith and Greitzen (2007), the role of the human in IA is the role to understand the process done by the user of the program. If the user is not able to understand the process, it can never be automated. Furthermore, they contend that the user should be in the superordinate position to overcome the limitations of computers, what can be explained that the human must stay in control over the agent in what they are learning the agents. Otherwise, the human cannot survive, and robots will take over the world.

Camara et al. (2015) state that systems can benefit by receiving different information from humans what is difficult to automatically monitor or analyze (sophisticated sensors), incorporating human input for the decision-making process, or by involving them as system-level effectors to execute adoptions (fallback mechanism). However, human participants are influenced by factors external to the system (training level or fatigue) that affect the likelihood of the success when they perform a task, its duration, or even the willingness to perform the task.

Xia and Maes (2013) describe two influence factors of the human in IA and IA-systems. The first one is the motivation why the human will perform the process. This is required to reach the goal of the entire process. The other one is mood when the human has a bad day, it will take other decisions instead of having a good day.

Ferreira et al. (2014) considered that human intelligence and accuracy are not easy to replace based on results in the past (Bley et al., 2004), especially in human-driven processes which require high precision and flexibility. People can operate in highly uncertain environments and can learn how to react to changes. Extra, disassembly processes operation is highly customized and less automated, making humans and their machines essential. Furthermore, they mentioned that human intelligence, intuition, perception are characteristics that need to be utilized to their maximum to create more efficient production processes.

To conclude, all the researches mentioned the benefits of involving humans in more automating the complex and dynamic problems which occur. The aim of the researches to involve humans is to help overcome the limitation of automation. It focuses more on 'how can we make the system more intelligent', but not on the part to enhance the intelligence and capabilities of the human. Concerning the concept of IA, the human did not only to involve in the system but also it is important to highlight the human central role in solving complex problems. Although, people can argue on this statement and will explain that when making the systems smarter by humans, the human will be smarter by the system because the human is part of the system. The human will not be biologically smarter but the ability to make decisions will be improved when it is part of the system. The previous argue is more a sense what the system can do but it is not implicitly mentioned in the literature.

2.4.4. What is the relationship between human and intelligence agents?

According to Bradshaw et al. (2011) the need to create machines to help the human to stay informed of on-going events happened in their environment. Furthermore, the need for help for alignment and repairment of their perceptions because humans rely on mediated stimuli. Because of the human capability to be very adaptable to changes in the environment, the need for machines is high to see the effect of the positive change after the situation changes. Furthermore, the need for computationally instantiate their models of the world. The other way around, computers need humans to keep them aligned to the context they are made in, machines don't have emotions. Computers cannot detect and adapt to the change of the world, so the human need to keep the computer stable in the highly changing world and need to repair all the ontologies if they need to change. Humans are in control with the computer on the automation decision making and execution of tasks (Smith, 1983).

Furthermore, Bradshaw et al. (2011) mentioned that interdependence is the essence of joint activity. Joint activity can be described as a process, extended in time and space. There is a time when the parties enter the joint activity and a time when it has ended. There are different kinds of joint activity that can be distinguished according to the different types of interdependencies; Co-allocation, cooperation, and collaboration.

Xia and Maes (2013) state that the system and the user can learn from each other achieve co-evolution, especially on decision-making. The actions humans will do shape how humans are. Furthermore, all the technologies the human will create will change the course of human evolution. Many agents that have been developed for interaction with humans play a subordinate role and make team contribution only under direct human supervision (Burstein et al., 1999). To predict human behavior, agents need leadership capabilities, autonomy, and proactively. Humans need to have trust in the agent and accept that they are also team members in their team. The human need to rely on the information what the agent present (Plano & Blasch, 2003).

Jacucci et al. (2014) mentioned the interdependence between humans and machines as telepresence, mixed initiative iterative, affective, and emotional computing, and persuasive technology. They mentioned that the symbiotic interaction can be achieved by combining computation, sensing technology, and the interaction to realize deep perception, awareness, and understanding between humans and computers. The goal is that the symbiosis amplifies the human abilities and knowledge through the co-operation of humans and computer.

Farooq and Grudin (2016) write about the integration between the human and the machine. The integration extends but doesn't replace the interaction. The humans are moving from the relationship 'stimulus-response', where the machine reacts on the command of the human, towards a more integrated world between human and the machine. Here, the machine will assist the human in advance. Integration can create and amplify the risk of software.

Human-engaged computing (HEC) can be described as the synergism between human capacities and technological capacities (Ren, 2016). HEC has three components: engaged humans, engaging computers, and synergized interactions. Humans are engaged when they can develop their inner capacities. Computers are engaged when they can enhance human capabilities. The synergized interaction refers to a state of optimal balance between engaged humans and engaging computers (Ren, 2016). In the best scenario, the humans can design systems that will enable the human the realization of the full potential and in the end will create a better world.

The collaboration of human – agent combines human's flexibility and intelligent agents' efficiency to address the complex requirements (Tan et al., 2009). Intelligent agents perform efficiently in identifying task status, suggesting alternatives, monitoring, processing information, and testing hypotheses (Griffith & Greitzer, 2007). Humans play a fundamental role to creatively think of solutions and visually perceive patterns (Cummings, 2014). Level of automation helps to understand how humans can interact with a complex system, for example in decision making (Roy, 2004). Casini et al. (2015) advocate thinking about how tasks can be best shared between human and intelligent agents and how they can work cooperatively. Furthermore, the question arises how the capabilities of human can be enhanced through the symbiosis between the human and the machine. The human should focus on the more critical tasks, while automation takes charge of making the less critical decision (Lange et al., 2014).

Vagia et al. (2016) summarize different levels of automation. The human need to be involved in the levels 1-5, but they have decreasing access to information and decreasing overrides capabilities in the levels 6-9. If the autonomy level is 10 then there is no human needed (Sheridan & Verplank, 1978). The levels can be seen in Appendix A.

All being said, the relationship between the human and the machine is important in the concept of intelligence amplification. The human is needed to translate their capabilities to the machine to enhance them. The goal of intelligence amplification is to create a symbiosis that will enhance the human abilities and knowledge through the co-operation of humans and computers (Jacucci et al., 2014). The world is moving towards machines that will assist the human more, also based on the capabilities, capacities, and interest of the human itself. Furthermore, to achieve the concept of IA, the right collaboration of human-intelligent agent collaboration, it is necessary to test this concept in the practice, with a case. Furthermore, there is a need to find a solution for the task assignment when implementing IA into problem resolution. At last, there should be a thought about the interaction between the human and the intelligent agent by an interface.

2.4.5. Data cleaning and data quality

Data cleaning deals with the detection and removing of errors and inconsistencies from data to improve the data quality (Rahm & Do, 2000). Those problems are present in single data collections, like files and databases, for example missing data like weight classes in product data. To ensure that the data is accurate and consistent, there is a need for the elimination of duplicate information, add missing values, and validate data (Rahm & Do, 2000). The variety of sources creates the probability that some of the sources contain 'dirty data' is very high. The data which are saved in the data warehouses (files or databases) are used for decision making, so the data need to be correct to avoid wrong conclusions. To clean the data there are several requirements to satisfy. First, the most important one is to detect all the major errors and inconsistencies (Rahm & Do, 2000). After that, the errors could be removed, and the data quality will be improved. This last option is the common approach to improve data, but it will waste data and reduces the power and the output is more biased as before (Barzi & Woodward, 2004). An alternative is to use one of the many methods available for imputing the missing values.

There are many other data cleaning approaches (Rahm & Do, 2000):

- Data analysis: To detect which kinds of errors and inconsistencies are to be removed, a detailed analysis is required. This could be automatic based on different rules or manual by different employees.
- Definition of transformation workflow and mapping rules: Depending on the number of data sources, their degree of heterogeneity and the "dirtiness" of the data, many data transformation, and cleaning steps may have to be executed.
- Verification: The correctness and effectiveness of a transformation workflow and the transformation definitions should be tested and evaluated.
- Transformation: Execution of the transformation steps either by running a workflow for loading and refreshing a data warehouse or during answering queries on multiple sources.
- Backflow of cleaned data: After (single-source) errors are removed, the cleaned data should also replace the dirty data in the original sources to give legacy applications the improved data too and to avoid redoing the cleaning work for future data extractions.

There are two problems with those approaches. The first is the lack of interactivity, the transformation is typically done as batch processes, operating on the whole dataset without any feedback. This can lead to long and frustrating delays during the process. The second problem is the need for much user effort: the transformation and discrepancy detection need significant user effort, making each step of the cleaning process painful and might create more errors (Raman & Hellerstein, 2001). Data cleaning, a labor-intensive and complex process, is estimated that it will accounts for 30%-80% of the development time in a data warehouse project (Shilakes, 1998).

2.4.6. What are the applications of IA in enhancing worker's performance or reducing workload?

In the literature, there are several cases where the computer and human are collaborating to each to reduce the workload complexity and enhance the performance of worker or employees. Those are most of the time done by applying artificial intelligence of the system to the process. The papers will be discussed in this way. The first paper is from 2003 and the last one from this year, 2018.

Bechar and Edan (2003) proposed and developed four different human-robot collaboration levels for target recognition task. The collaboration between the human operator (HO) and robot increases detection by four percent compared to the HO and fourteen percent when compared to a fully autonomous system. In addition, the detection times of the integrated system are reduced. They demonstrate that collaboration will significantly improve the performance which is beyond when humans and computers working alone.

De Greef et al. (2007) creates a framework for the design of cognitive support that augments the capacities of teams and team-members during critical and complex operations. They encourage people to utilize such a method in the designing process of augmented cognition systems. The workload will be lowered by three mitigation triggers, the operator's performance, physiological measurements, and task demands.

Yamaba et al. (2008) tried to obtain proper scheduling rules by means of reinforcement learning. They developed a network-based support system through the man-machine collaboration based on the created 'behavior model' of scheduling activities in decentralized production networks. They used this system to confirm whether proper rules can be learned. Profit sharing was used to obtain rules for selecting the factory where intermediate materials are manufactured. Yamaha et al. confirmed that the rules could converge under profit sharing and worked well and seemed to be rational rules. In summary, the system helps the human to create rules where to schedule the activities, so the performance of the whole supply chain will be better.

Tan et al. (2009) want to improve the man-machine collaboration in cell production by the approach of task modeling. They used six development requirements to benchmark the development. They suggest that the man-machine collaboration is well defined in a structural format by the capability of task analyses method to address the human tasks in a more flexible way. This collaboration has significantly facilitated operation planning in assembly and control levels. With help of the system, the human operator is well guided by the information corresponded to the task components. The information can be used as an evaluation to measure the different working performance of the human operator and system performance as one (Tan et al., 2009). They conclude that there is a need for an evaluation framework to test the symbiosis between the human and the machine.

Padoy and Hager (2011) state that the future of surgery, teleoperated robotic assistants will offer the possibility of performing certain commonly occurring tasks autonomously. They create a Human-Machine Collaborative (HMC) system that can handle portions of surgical tasks autonomously under complete surgeon's control, and other tasks manually. The approach is based on learning from demonstration, the system automatically identifies the completion of a manual subtask, smoothly continuous the next automated task, and returns the control back to the surgeon. The fine motions are performed by the surgeon, while real-time recognition of their termination triggers the automatic execution of previously learned motions. In conclusion, the experiments that the researchers did show that this kind of human-machine collaboration improves the work of the surgeon and because of the automation of simple tasks will reduce the workload of the surgeon.

Coetzer et al. (2012) create a protocol for incorporating a machine into the authentication process which significantly improves the performance of the workforce of human employees for all operating conditions. The protocol focuses on the recognition of flowers and faces. There is an abstraction of the object created which can be modified by a human operator. In addition, the protocol investigates the problem of authentication of static handwritten signatures on cheques in the banking environment. Both the human and the machine will recognize the signatures and the decision of the human and the decision of the system is combined to estimate the performance. They said that the optimal human-machine ensemble outperforms the unaided human workforces for all the cost gradients they researched. Concluded, the protocol for human-machine collaboration will lower the costs for the authentication process instead of doing it all by hand.

Zhang et al. (2015) proposed an adaptive support vector machine-based method to classify operator mental workload (MWL) to detect human operator performance degradation or breakdown. This plays in on the mental part of the work instead of the downgrading of workload complexity or to enhance the performance of employees. With the help of a human-machine system, some signals were recorded to measure the MWL of operators. They mentioned that the role of human operators has been shifted to supervision and decision making in those kinds of systems.

Dobrkovic et al. (2016) propose the intelligence amplification framework that is applicable in typical scheduling problems. The amplification framework is built upon the vision of Licklider (1960). There are several 'intelligence' agents defined for executing different atomic tasks of the scheduling process. Some of them are invoked by the user and some of them run continuously. The proposed framework could be useful in diverse decision-making processes.

Zou and Zou (2017) used the computer and their human abilities to create databases for mapping a variety of ambiguities of language and words. Each word as value-taken can be recorded, not only by computer-assisted enumeration into an electronic dictionary but also by the body of the meaning at all the different levels. This especially for words which have multiple different meanings. This can be used to create a basic development platform that can be built for knowledge production.

Wigness and Gutstein (2018) investigate the impact of group-based label noise on classifier learning and discusses how and why this differs from instance-based label noise. The human is involved in this process by inspecting different data instances of the label and assign a target label. Noisy labels are classified by typos, misidentification or inconsistency by other annotators. With the help of deep learning, the visual classification of labels and the noise of it will be improved. By applying group-based labeling the will drop the complexity of the labeling workload.

The last one of Wilson (2018) at all gives an answer to the question: how immersive virtual environment can relate to an intelligent physical counterpart for allowing more efficient man-machine collaboration. They present an integrated architecture as an interactive and immersive virtual reality environment for telerobotics and telepresence applications. This architecture utilizes a computational server and a multitude of heterogeneous clients which can be virtual reality clients or robotics remote operations' clients. With help of a human, the system is improved in the speed of task completion as well as the task completion accuracy. In conclusion, with help of a human usability study, the system is improved to do the tasks better.

In summary, all the improvements are improvements which are most of the time based on the improvement of the whole system's performance instead of completing tasks, without studying the amplification effects of humans' intelligence. Humans play a critical role in the Intelligence Amplification system, so there is more research needed which will focus more on amplifying human intelligence instead of the system.

2.4.7. What are the applications of Intelligence Amplification in enhancing slotting or the slotting process?

Tsamis et al. (2015) present an adaptive strategy for the storage location assignment. The general description of the way an adaptive policy could be realized in a warehouse environment can be seen in Figure 6.

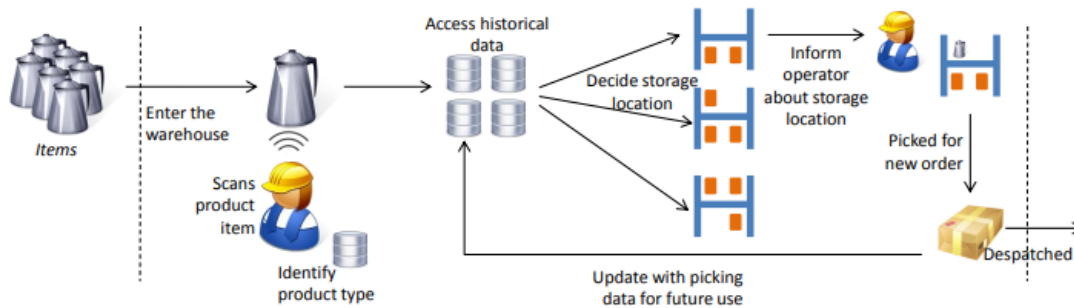


Figure 6: General description of an adaptive storage policy

The product will be automated identified with help of a barcode or RFID-tag. When using a non-adaptive policy, the product will be pre-assigned to a location, but using an adaptive policy, the product will be stored at the best-guessed location based on historical data such as turnover rate, order frequency, and class. To be sure that the decision can be made, the warehouse manager might need to collect information about the available locations. It will inform the operator about the location where it needs to be stored, and the product can be prepared to be dispatched. With help of the product intelligence paradigm (Kärkkäinen et al., 2003; Wong et al., 2002) there can be a suitable solution be found for the deployment of the adaptive policy. Intelligent products are used to identify and track the location in the warehouse (Kärkkäinen et al., 2003). They are also capable of collecting and storing information about themselves or about a similar product (Wong et al., 2002). The communication between them can be used for the assignment of a storage location for themselves after communicate with the same product type elsewhere in the warehouse (Tsamis et al., 2015). The location storage strategy consists of two main elements, the first is the set of interactions between the product agent and three different warehouse agents. The warehouse agents communicate about the locations in the warehouse. Figure 7 shows the interaction between the agents.

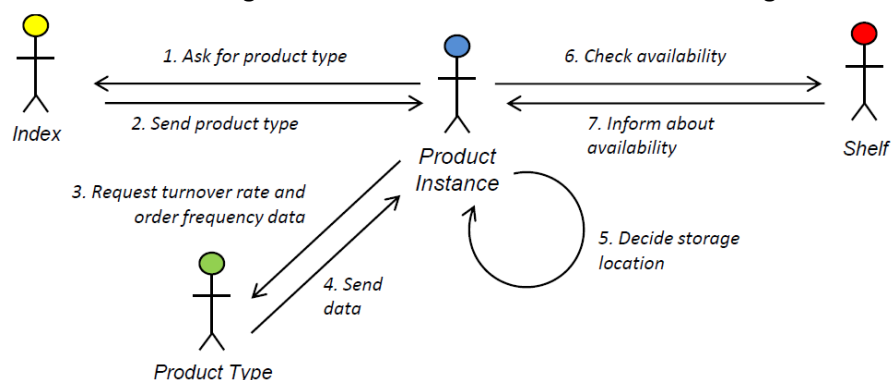


Figure 7: Interactions between agents (Tsamis et al., 2015)

The other element is an algorithm which decides the location for the product. With help of the product intelligence the best location can be found (Tsamis et al., 2015).

Iyer and Jayal (2016) used machine learning for slotting. Machine learning is developed from studies in artificial intelligence of pattern recognition and computational learning theory. “Machine learning is based on algorithms that can learn from data without relying on rules-based programming”

(McKinsey, 2015a). Depending on the nature of the learning model or the feedback which is available, the machine learning tasks are categorized into three classifications (McKinsey, 2015b):

- **Supervised learning:** will be used when an algorithm learns from example data and target responses that can consist of numeric values or string labels in order to predict the correct response when posed with new examples. There is a difference between regression problems and classification problems. The regression targets are numeric values, and a classification target is a qualitative variable (class or tag).
- **Unsupervised learning:** will be used when an algorithm learns from plain examples without the target response, the algorithm needs to learn the data patterns on its own.
- **Reinforcement learning:** will be used when presenting the algorithm with examples that lack labels, as in unsupervised learning. It is often connected to applications for which the algorithm must make decisions, and the decisions bear consequences (in the human world, trial, and error).

There can be added an extra category namely semi-supervised learning. Semi-supervised learning algorithms are trained by a combination of labeled and unlabeled data. This can be useful for a few reasons. First, the process of labeled data is time-consuming. Too much labeling can impose human biases on the model. With help of the unlabeled data during the training process, it will improve the accuracy of the final model (Castle, 2018).

There are many standard algorithms available to apply machine learning concepts to practical problems like slotting (Ray, 2017). Iyer and Jayal (2016) used association rule algorithms to the problem of slotting. An association rule is if/then statements that help to uncover relations between unrelated data (Agrawal et al., 1996). Association rule algorithms determine the association between features used to describe the dataset. Iyer and Jayal (2016) describe the association rule problem in the context of slotting. They used the Apriori algorithm to achieve association rule mining. This algorithm is a groundbreaking algorithm using candidate generation for finding frequent itemsets (Agrawal, Mannila, Srikant, Toivonen, & Verkamo, 1996b). Iyer and Jayal (2016) suggest a supervised learning algorithm to enhance the slotting process. The algorithm learns from past item locations coupled with market basket analysis of the order. They create a new system (Islotting Engine) with several modules to be sure that the automated slotting will operate in a suitable timeframe. The process flow is executed in the following order:

1. Determine Association Rules using MBA; with the help of unsupervised learning, methodology to determine the market basket.
2. Feed the assignment history and MBA data to the Inference Module for filtering; Inference module filters the historic data and MBA data for relevant data based on warehouse constraints.
3. Generate Slotting Rules; machine learning technique to deduce the suitable item-location assignment rules.
4. Generate Slotting Plan; create slotting plan for simulation considering layout constraints.
5. Simulation; Simulate the plan to be sure that the moves are needed.
6. Execution; translate the plan into an export task in the warehouse management system.

Iyer and Jayal (2016) believe that the idea and solution they proposed will improve the picking efficiency in the warehouse. In addition, the pick accuracy and space utilization will be maxed.

To conclude, there are two different cases described for ‘intelligence amplification’ in the slotting process. The word intelligence amplification is put between quotation marks because it is not really intelligence amplification. The solution proposed by Iyer and Payal (2016) will not enhance human intelligence by the computers power and vice versa. The solution made is fully automated and not involving the human. Iyer and Payal (2016)(A. M. Iyer & Jayal, 2016) also mentioned that the future of this slotting application needs to be enhanced by new techniques like Robotics and Intelligent Materials. This can also be in the vein of further research with respect to the application of intelligence amplification. Second, Tsamis et al. (2015) present an adaptive strategy for slotting in the warehouse.

This strategy can be useful for creating an application with help of intelligence amplification to make it smarter and involve the human more in the process. The system needs to take the computation for easy and repetitive tasks and the human for the more complex ones. Furthermore, the other conclusion is that it is hard to find literature about improving slotting by human-computer collaboration but are articles available about the process improvement with help of human or computers in a general way. The improvements which are most of the time based on the improvement of the system in a whole instead of amplifying the human capabilities.

3. Current situation

3.1. Introduction

In the previous chapter, the theoretical framework of the slotting process was established. With this as guidance, the slotting of the Albert Heijn will be explored. Slotting is a method to determine the optimum location for products within a warehouse. It is used to reduce the amount of travel time for operators by placing the item that frequently ships together next to each other in the pick-face area (Richards, 2011). It refers to the optimal location for the products in the warehouse such that they can be picked in the least time with the maximum accuracy and maximum space utilization (N. P. Iyer & Jayal, 2016). An explanation of the slotting process can be found in chapter 2.1.

To get a comprehensive understanding of the slotting process at the Albert Heijn, this chapter will start with a high-level process for the whole supply chain to illustrate how this works when there is a new product arrives. In this high-level process, the need to allocating locations in the warehouse will be made clearer. After that, the current process of slotting is analyzed to give an understanding of how the employees at the AH do this process. Not only the process is made transparent also the hard- and software will be explained with help of an ArchiMate model. Finally, the different KPI's and the weight is discussed to measure the current situation and to find the hiccup to improve, to be sure the right thing will be improved.

The national supply chain of Albert Heijn consists of five distribution centers (DCs) owned and managed by AH. AH owns one national DC (LDC), located in Geldermalsen. Four regional DCs (RDCs), which are in Pijnacker (DCP), Tilburg (DCT), Zaandam (DCZ), and Zwolle (DCO). There are six DCs outsourced to logistics service providers (LSP). Those DCs managed the flow of fresh, non-food, and frozen products like ice creams. Those are located in Bleiswijk, Hoogeveen, Tilburg, Oss, Nieuwegein, and Zeewolde.

Type	#	Operator	Location
National DC (LDC)	1	AH	Geldermalsen
Shared fresh center (SFC)	1	XPO Logistics	Nieuwegein
Regional DC	4	AH	Pijnacker Tilburg Zaandam Zwolle
XPO Oss (Non-food)	1	XPO Logistics	Oss
Shared warehouse Cheese DC (SWK)	1	Bakker Logistiek	Zeewolde
Returns	4	Kuenhe + Nagel	Pijnacker Tilburg Zaandam Zwolle
Frozen products	2	XPO Logistics	Hoogeveen Tilburg
Flowers	1	GIST	Bleiswijk

Table 7: Locations of all DCs of Albert Heijn

3.2. High-level process supply chain

To start with the explanation of the current situation of the process of slotting, it is important to know the need for this process. Without the slotting process, products are not assigned in a location, which will imply that the product is not able to be picked in the order picking process and cannot be bought by the customer. To get a more detailed overview of the different processes or different departments/companies involved in the process, there is a swim lane model created for adding a new product in the assortment, see Figure 8.

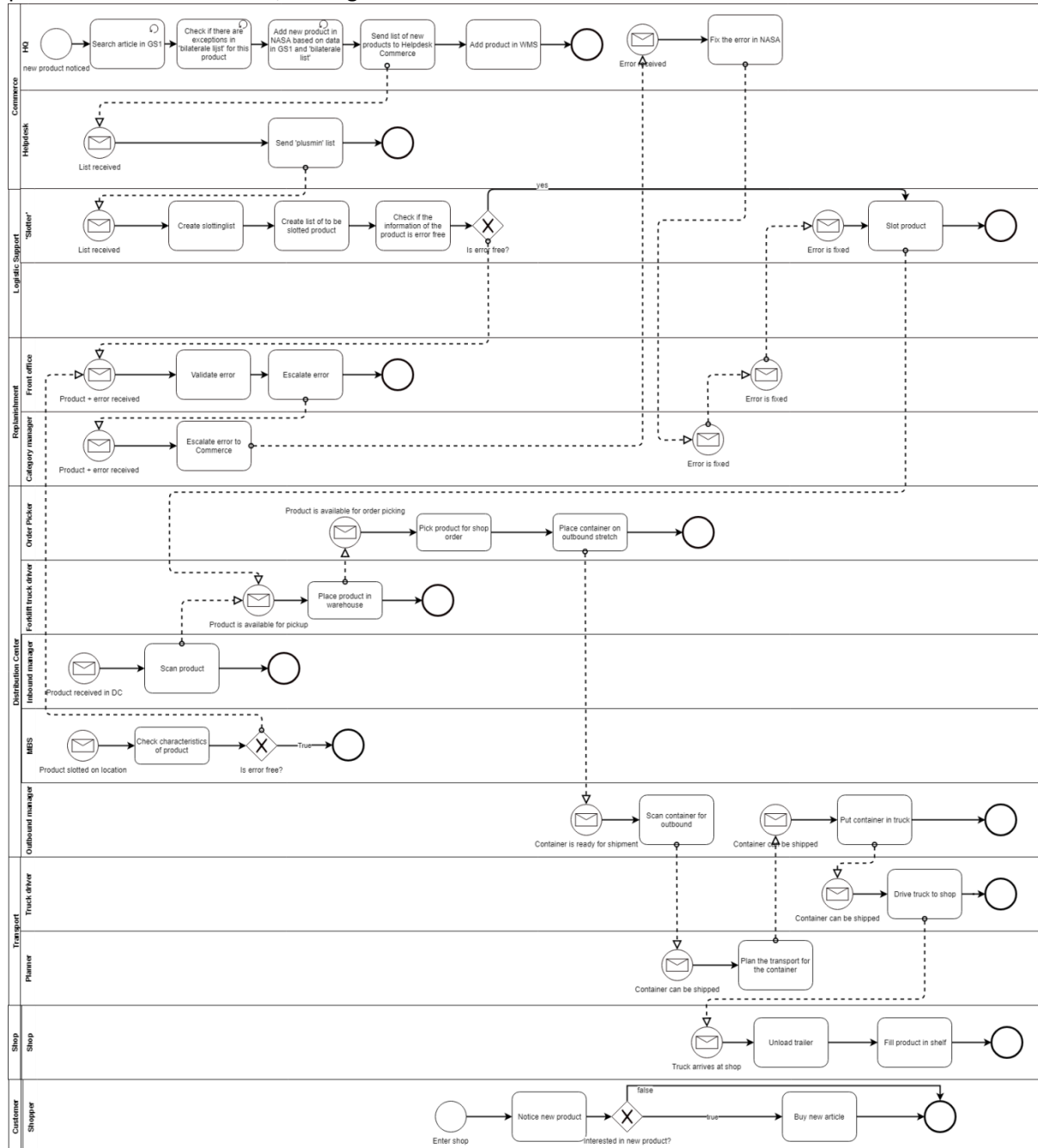


Figure 8: Swimlane model for the process: 'New product in assortment'

3.2.1. Departments and companies

In this process are multiple departments and companies involved to be sure that the new thought product can be bought by the customer. All the different departments will be explained below.

Commerce: The department of commerce introduces the new product of a supplier. They will estimate and guess which revenue this product will be done in the stores. In addition, they will add this product to the database of all the articles.

Logistics Support: The department of logistics support ensures that the processes in the DCs can run smoothly. They do this by, among other things, the following matters:

- Setting up the DC: Where does the 'vakkasten', 'flow racks' and 'pallet' location come?
- Distributing articles within the building and allocate a location for a product(slotting)
- Distribution of articles to stores so that it works out well for production, but also for the stores.
- Ordering roll containers, crates, and rollies.
- Ensure that pick orders are on time in WMS
- Assess whether collide numbers for the coming days/ weeks are feasible for the DCs.

Replenishment: Replenishment is are the supply chain planners of AH. Replenishment is responsible for the demand of the stores and that it is possible to deliver it to them. Furthermore, replenishment ensures that the right stock is in the right place at the right time. They control the flow of the goods. The flow manager is responsible and makes agreements with manufacturers, transporters, internal departments of the organization, etc. to optimally organize the chain as optimally as possible. The planner of replenishment does the operational work. Depending on the expected demand and number of delivery times, they take orders from the suppliers. This is done automatically; in case of errors or deviations, alerts are triggered, and then it will be investigated. Is the error being that big, they adjust the system.

Distribution Center: The distribution centers (DC) are scattered over the Netherlands. In the different DCS, the products are picked on containers or rollies based on the order of that specific store. In the DC, the products of the supplier are placed in the specific allocated location which is specified by the 'slotter'. In the regional dc, there is also a place for the orders which originate from the national dc or the external dc and are cross-docked within the regional warehouse.

Transport: The department of transport is the connecting factor of the supply chain of the AH. For many suppliers in the Netherlands and abroad the transport department delivers the products to the distribution centers.

Shop: There are more than a thousand shops now in the Netherlands and Belgium to sell products among the need of the people lived in those countries. The shop is also responsible for adding the products to the shelves of the store to enable the customer to buy the product.

As can be seen in Figure 8, the process starts when the category manager of that specific category announced a new product. The data will be filled in the NASA-database, by an employee of commerce, which contains all the products of the assortment. This process is error-prone because the data is filled in manually (Rahm & Do, 2000). When all the 'new-to-add' products are added to the database, there will be a list created together with the removed products, which will be sent to the helpdesk of commerce. The helpdesk employee will check the list and will send it to the Logistics Support department of the Albert Heijn.

One of the 'slotters' of the department will add all the products of this list in another list (slot list). In addition, products of SDRS list together with the list of Bakker (for fruits and vegetables) will be added to the slot list. This is base for the products which needs to be slotted or need to be removed in the distribution center. The slotter will check every product if this is correct. A product is correct if it has no data errors. A list of different errors can be found in Appendix G.

If there are errors mentioned this error will be shared with a person of the department of Replenishment. It will land in the inbox of the helpdesk support employee, which will answer the error or escalate it to the category manager. The category manager will answer this error or escalate it to the Department of Commerce if he also does not know the answer. This will investigate the error and will give an answer which will be sent to the category manager. After that to the helpdesk support employee and after that back to the 'slotter'. If this error is fixed than the product can be slotted.

When the list is complete, the next step of the flow will be the slotting process. This process is explained in detail in the next paragraph.

When the product is slotted, the movement order will pop up in the scanner of the forklift truck driver. This will place the product on the assigned location in the warehouse. The product is now ready to be picked up by an order picker which will drop the product on a container. If the order is finished the container is ready for the container to be picked up by a truck driver of a transport company which will drive the container to the store where the product can be sold to the customer. When the container arrives at the store an employee of the store will fill the product on the shelf. After that, the customer can add this product to the basket and can buy this product at the cash desk.

3.3. Slotting process

There are different techniques or strategies of doing slotting in the warehouse as explained in chapter XX of the literature review. In this section, the process at the AH will be explained. This is a combination of slotting with the popularity of the Stock Keeping Units (SKUs) and the volume of the SKU. The volume of the SKU will be got from the stored demand forecast (SDF) of the stores in the country. Also, seasonality will be taken in to account. This is because there are different seasons of products available and all those products need a location in the DC. The weather – is a part of seasonality – will be used to predict the SDF of a product. For example, when there is beautiful weather the soda's and ice cream will be sold more than products like split pea soup.

The slotting process is explained in the different sections below. With help of the BPMN model, the different processes and subprocesses are modeled. The main flow, see Figure 9, starts with the download of the 'slotting' list. This list is created from different source systems which contain the new-in-assortment products and the products which will be out of the assortment. This is the most important source file for the slotters and needs to be clean from errors otherwise products will be wrongly slotted. After the document is downloaded and opened, the slotter will start with unlinking of outgoing products in the warehouse management system (wms) (see subprocess 1). When this is done the slotter will continue with the allocation of location for the new product (see subprocess 2). When this is done the slotter will check if the other products in the dc have the correct location (see subprocess 3). After this, the slotter will create the export files which are used to add the movements and locations changes for wms and upload this to the wms. When this is done the slotter will communicate with the dc to be sure that everything is good. They communicate which locations need to have a facing indication. With help of this indication, order pickers know that the location with the indication is an extension of another location (most of the time the neighbor location). Good to know is that the subprocesses 1-3 can be done at the same time. For example, if there is a new product available and needs a space in the distribution center, but the ideal location is not available, or there are not enough free locations available for that specific aisle. The slotter can choose to first move the other product to a smaller or bigger or other location in the distribution center.

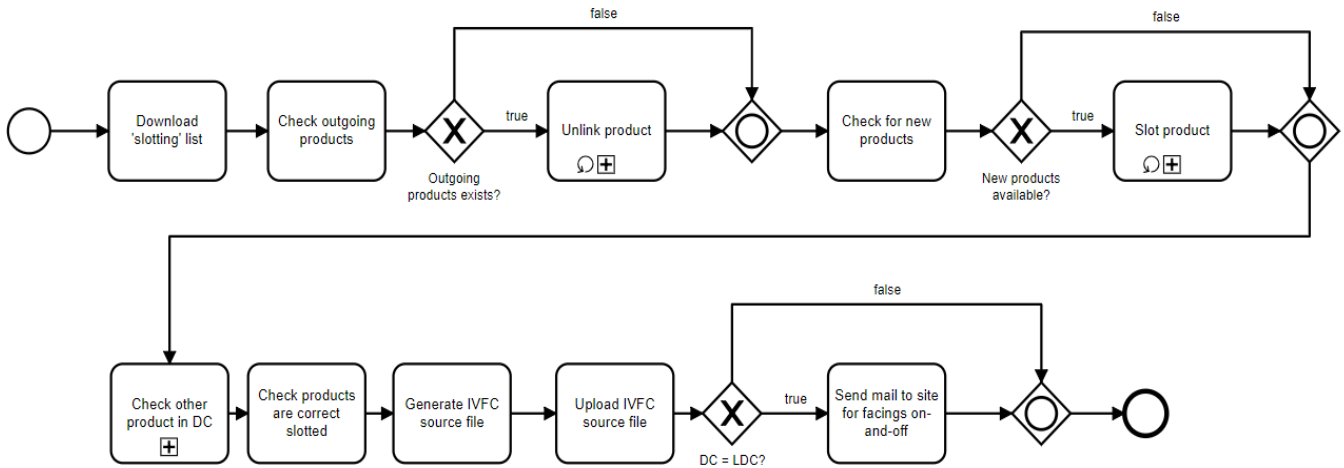


Figure 9: Slotting process

The subprocess of 'allocate product' is explained below, the rest of all the subprocesses stated in Appendix F.

3.3.1. Subprocess: 'allocate product'

As can be seen in the subprocess 'allocate product' or 'slot product' (Figure 10) the first step is to check the product data of the product which will be slotted. This is important to slot the product in the ideal location of that product, this subprocess will be explained in subsection 3.3.1.4. When there are no errors the prognosis of the sdf of that product will be checked. Based on the product characteristics and the sdf, the ideal location will be examined. When the 'ideal location' is free the product will be linked to that location otherwise the product on that specific location will be moved if that is possible (this process will be repeated for that specific product).

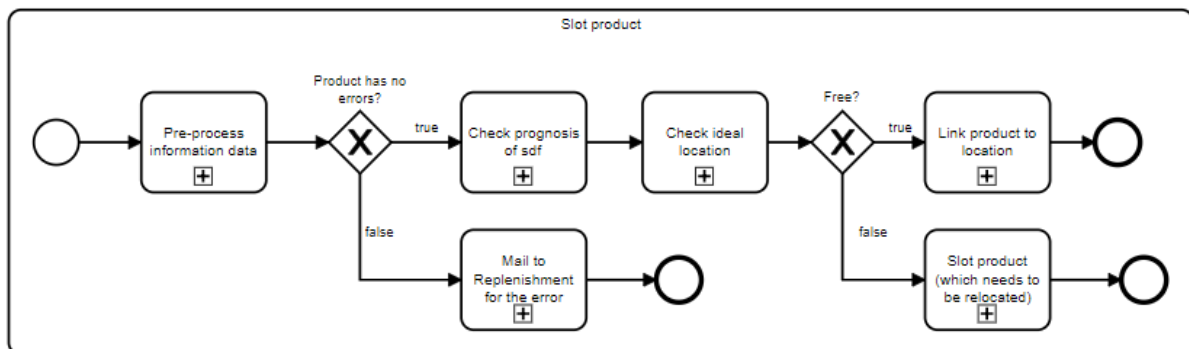


Figure 10: Subprocess: 'Slot product'

3.3.1.1. Subprocess: 'Pre-process data'

This process is a part of the 'slot product' subprocess of slotting. In this process, the data of the product will be checked. First, all the required fields are being checked if they contain a value. After this is checked the product needs to be existing in the NASA-database. Hereafter, the structure of the package of the product is checked, this process will be explained in the next subsection. Furthermore, the dimension of the products is checked. Those dimensions are the height, width, and length of the package of the product. Followed by the check of the weight is not too heavy, heavier than 23 kilograms. Next to check the start dates of the products. The start dates of a new product need to be in the future and need to be filled in before the product can be slotted in the dc. Also, the product needs to be placed in the right distribution center. Otherwise, the products which will arrive in the

national dc will be slotted in the regional dc and vice versa. When the product is short-term, which implies that it will need a location, for example, one week, those need to be placed on a special place in the distribution center and not on a regular spot, this need also is checked. The last two things what will be checked has to deal with the logistics group and product category. The product can be placed in a wrong product category or logistic group and needs to be resolved before it can be slotted.

Note: Every time there is an error occurred this will be sent to the employees of the Replenishment department. This is needed because the employees of replenishment have the right permissions to fix the error in either wms or NASA or can alert the error to a specific department who can fix the error.

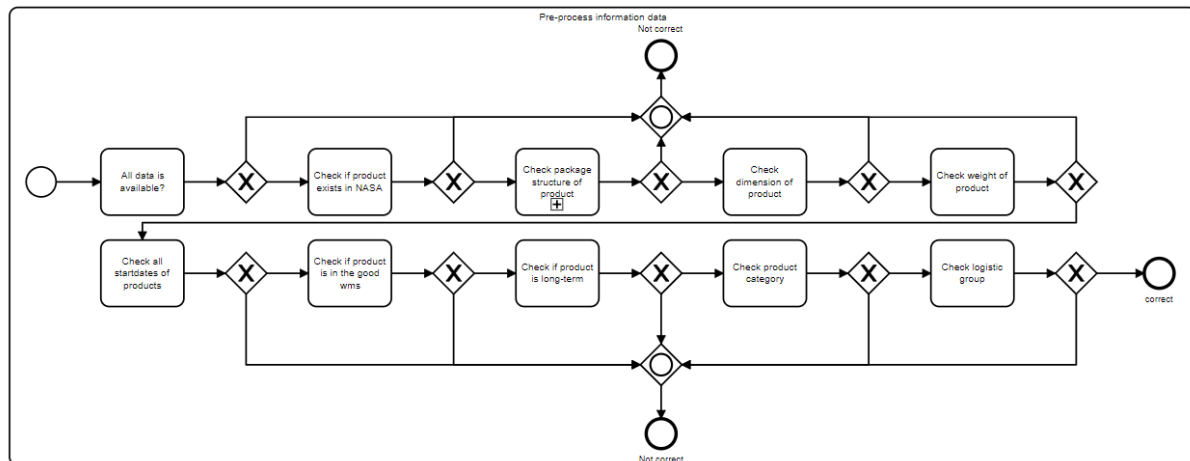


Figure 11: Subprocess: 'Pre-process data'

3.3.1.2. Subprocess: 'check package structure of product'

One subprocess of the pre-process information data process is the checking of the right package structure of each product. In this process, the inbound package of the product is checked in wms. There are several rules drafted to be sure that every product in the dc has the right package to be slotted. Those rules are written down below and will be short explained. Different inbound packages have a different number of rules which needs to be available in the wms of that specific dc.

Those number of rules are as follow:

- There need to be three package rules available if the inbound package contains 'overig'. This package type is in most cases for the boxes and crates.
- When in wms the 'enter lowest um' is checked there is a multiple of three package rules needed.
- When a product has the indication that is needs to be picked out of the box or crate, there is a need for four package rules. Also, when the product is suitable for the kpi machine or when this product will be delivered in a roll container then there are four rules needed.
- Every product which are repack product need to have seven package rules. The implication is that the baspakid need to be 'overig' as explained in bullet one and linked to the pick location.
- Otherwise, the last type is that bullet four is not linked with the pick location the number of rules need to be eight.

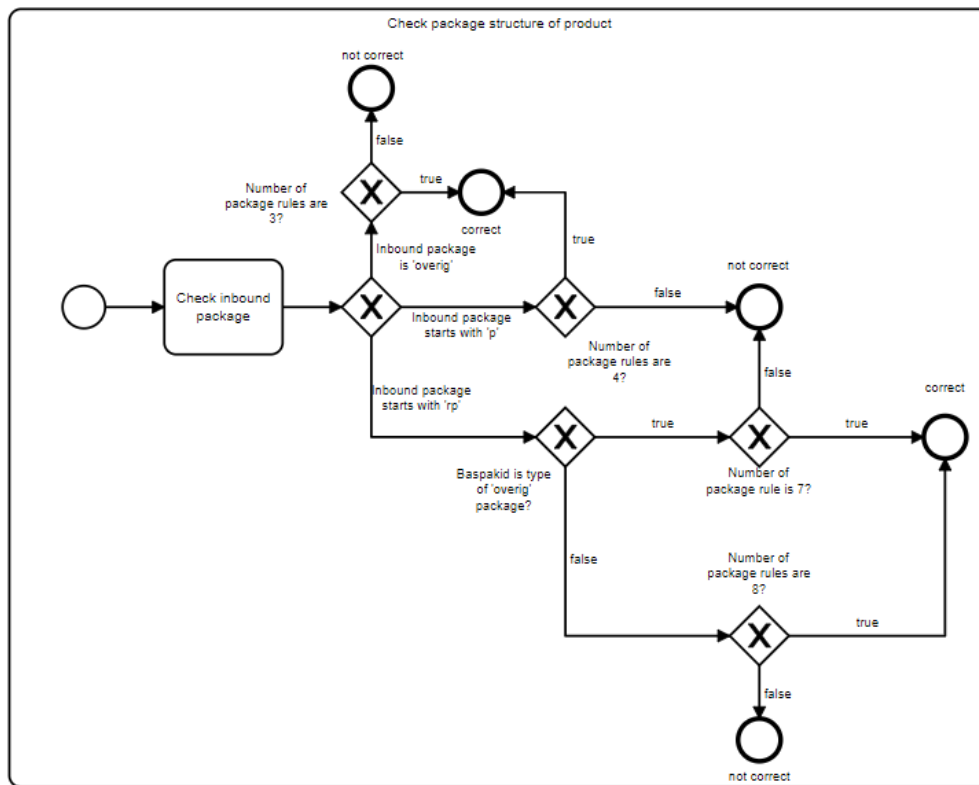


Figure 12: Subprocess: 'Check package structure of product'

3.3.1.3. Subprocess: 'check prognosis of sdf'

When the 'slotter' needs to allocate, remove, enlarge, or reduce the number of locations of the product, the sdf of that product is very important. The location type and the number of locations are based on the SDF of that specific product and will be checked if it is available and how big it is for the next 8 days of the sdf. If the SDF is not available, replenishment is contacted to send the SDF for that product.

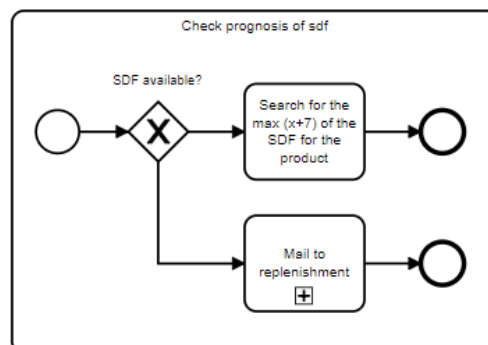


Figure 13: Subprocess: 'Check prognosis of SDF'

3.3.1.4. Subprocess 'check ideal location'

In Figure 14 the process of the ideal location checking is showed. When it is a new product the hall or pick zone needs to be determined. When this is done, the type of location is determined (which is explained in the next section). And when this is done the number of locations is determined, which is in combination with the SDF and the location type. Every location type can contain a maximum of products based on the dimensions of the product.

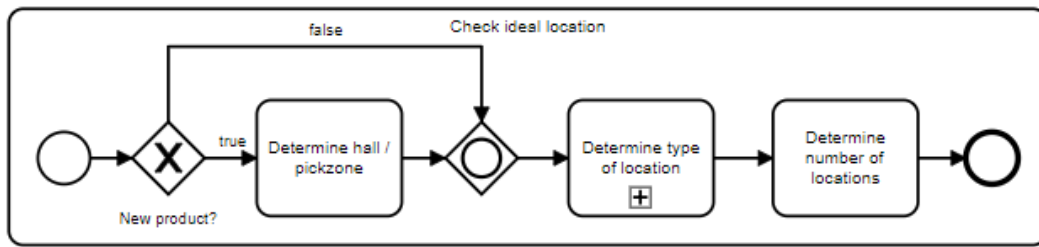


Figure 14: Subprocess: 'Check ideal location'

3.3.1.5. Subprocess 'determine the type of location'

In this process, the type of location is determined. This is based on the characteristics of the package of the product. As can be seen in Figure 15, the different cases will be checked to determine the location of the product and his package. As earlier mentioned, this can only be determined if all the data for this is available.

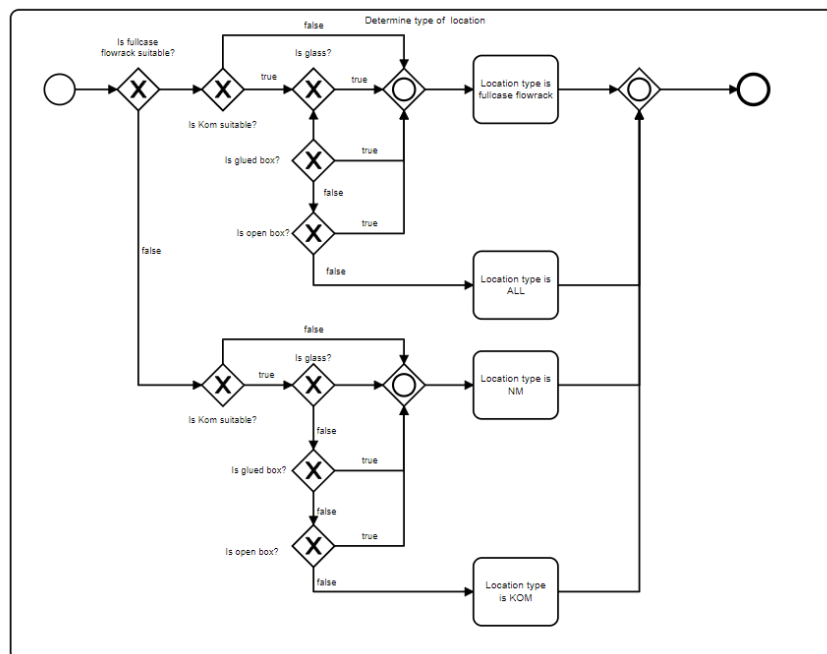


Figure 15: Subprocess: 'Determine the type of location'

3.4. Current architecture

3.4.1. High-over slotting process

The slotting process is already explained in the previous section in detail. The 'slotting tool' is made in Microsoft Access which use the warehouse management system for the locations and products and the assigned location to the product individually. The slotting process can be done in different ways and different process steps. Most of the slotting work will be done in the Access tool, only uploading the file and mailing the facings to the distribution center will not be done in that specific tool. One big challenge for the organization as well as this research is to deal with the outdated software such as the Microsoft Access 2003 and the OS Windows 7. Microsoft Office 2003, where Access is part of, has an end of life support until April of 2014. Windows 7 instead has an end of life support in January 2020 (Microsoft, 2009), and the Access tool won't be working anymore if Albert Heijn decides to migrate to a newer version of the operating system.

3.4.2. Business Layer

The slotting process starts with the data retrieval of different documents. From different sources and systems, also from different departments and supplier will there be data delivered to the department of Logistics Support. As seen in Figure 16 there are different source systems:

- SDRS List: this list contains all the promotional products for the next week and the SDF of that product. This list is generated by another department of the Albert Heijn, replenishment. This will be used to get a better understanding of the products which will be in the action and for which products the current location is too tight.
- Bakker-list: Bakker is the supplier of Albert Heijn which supplies the fruits and vegetables. Those products need to be slotted in the fresh departments of the distribution center. This list contains the changes in different packages or changes in the products.
- 'plus/min' list: this list contains all the products which will be new in the assortment and all the product which will be removed from the assortment of the Albert Heijn.
- Other resources: Google Drive, OneDrive, mail, and telephone; with those devices or tools the different issues are communicated to the department of Logistics Support.

All those lists create a list which is the starting point of the slotting process. An important part is the communication between the dc employee and the slotter. They have close contact to prevent as many mistakes as possible. Also, when there is an error in the data the employees of replenishment will be called or mailed to escalate the issue. Important to report is that errors are mentioned in the slotting, and only when they need to slot that product. The business layer is created with help of the ArchiMate modeling language and can be seen in Figure 16.

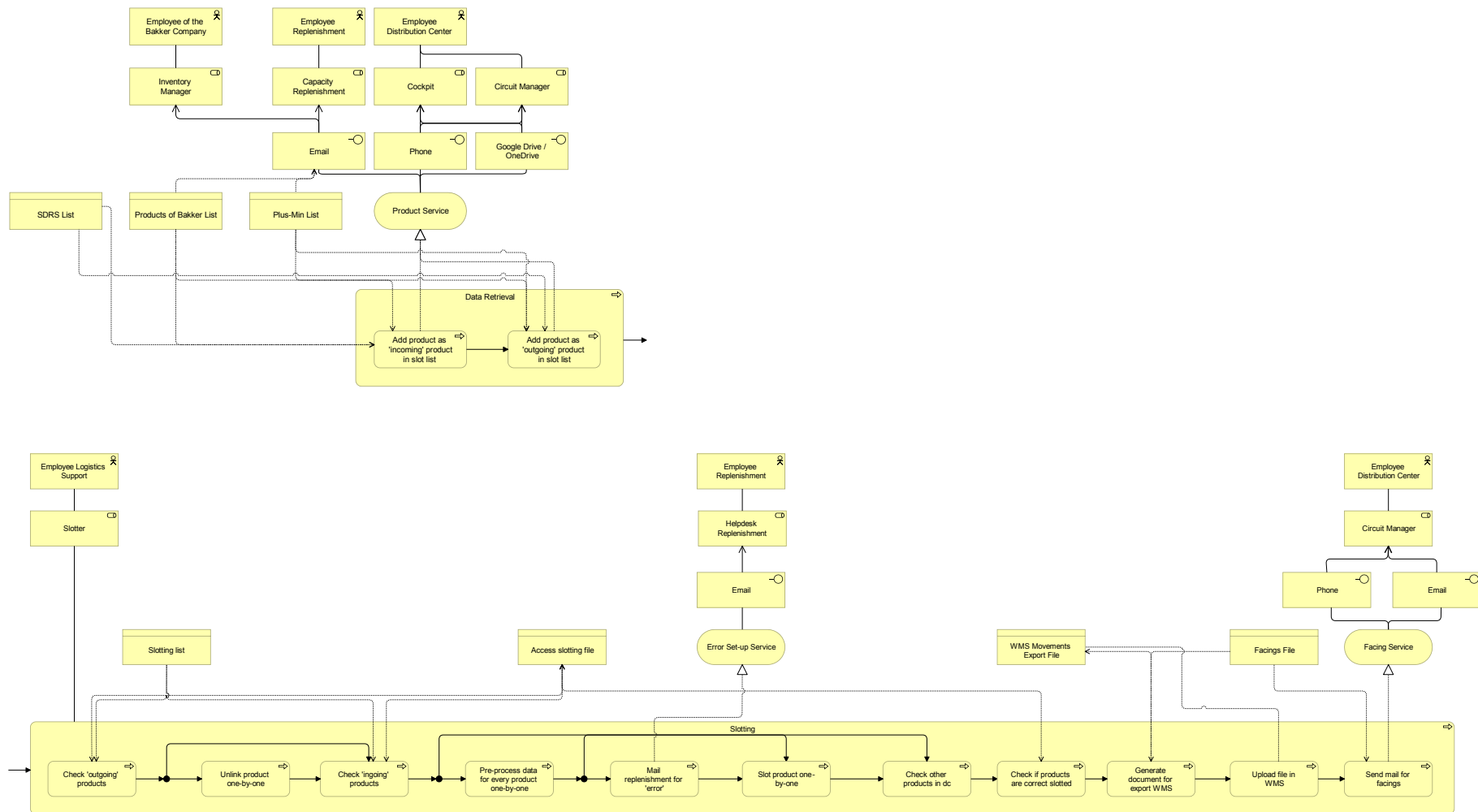


Figure 16: Business view of the current architecture

3.4.3. Application layer

The current application landscape exists in various types of systems. For the current situation and architecture, two systems are relevant as well as two tools, which are displayed in Figure 17. The first type is the replenishment toolbox where different documents are downloaded, like the SDRS list which is explained in the business layer.

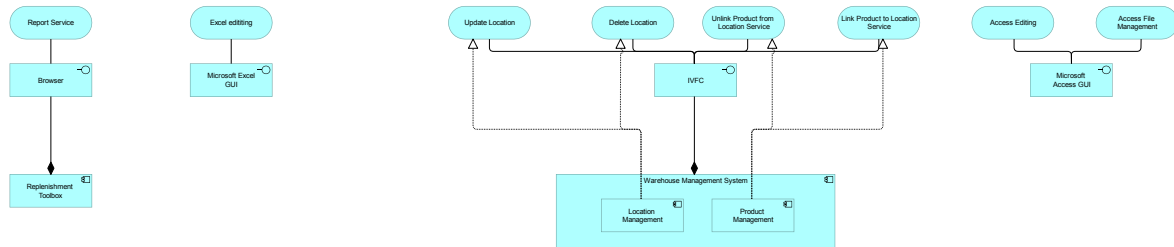


Figure 17: Application view of the current architecture

Next application is the warehouse management system (WMS) of the different dcs. This system contains all information of the specific dc like all products which are in that specific dc and locations in the dc. With help of the interface application IVFC, different information sources can be changed in WMS.

As last, there are different tools or application used of Microsoft Office. Excel is used to create a list and to communicate between different applications. The slotting process will be done in Access. Access shows a long list of locations with the product if it is not empty. When the slotting is done there will be IVFC files created which are used to update all the locations and products in WMS.

3.4.4. Technology layer

In the current architecture, the technology consists of nodes that represent servers and databases, realizing each of the application identified in the previous sector. The technology layer is reflected in Figure 18.

The whole Microsoft Office package Excel and Access included is installed on the computer of the slotter. The IMI WMS is available through a portal where the wms can be downloaded for each specific dc. There is an internet connection needed to download those. In total there are five different wms available in the portal. What can be said is that the access tool, as well as the WMS, are going to be replaced for new tools or systems in the future. There will be an upgrade for the WMS and the access tool are going to be replaced by a new platform, the slotting platform.

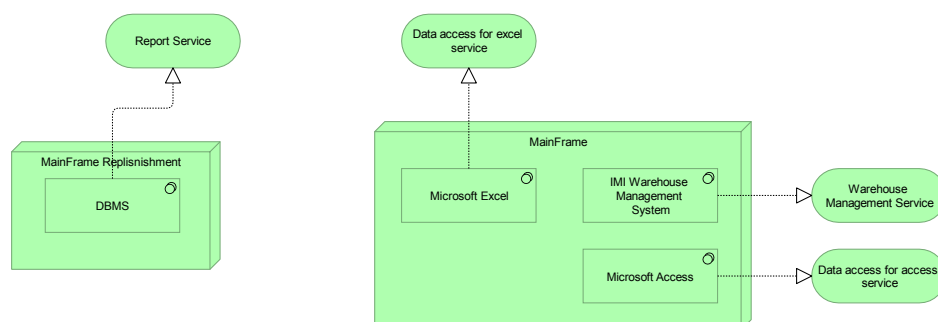


Figure 18: Technology view of the current architecture

The last system what is needed for the slotting process is the toolbox of replenishment where different documents can be downloaded for the slotting. This can be the SDF of products as well as the different collo-numbers available per NASA number.

3.4.5. Overview

To draw a comprehensive overview of the current architecture, the three layers (business, application, and infrastructure) are synthesized in a single architectural overview. This overview is shown in Figure 19. The architecture is structured according to the three layers as identified in the ArchiMate framework (Iacob et al., 2012).

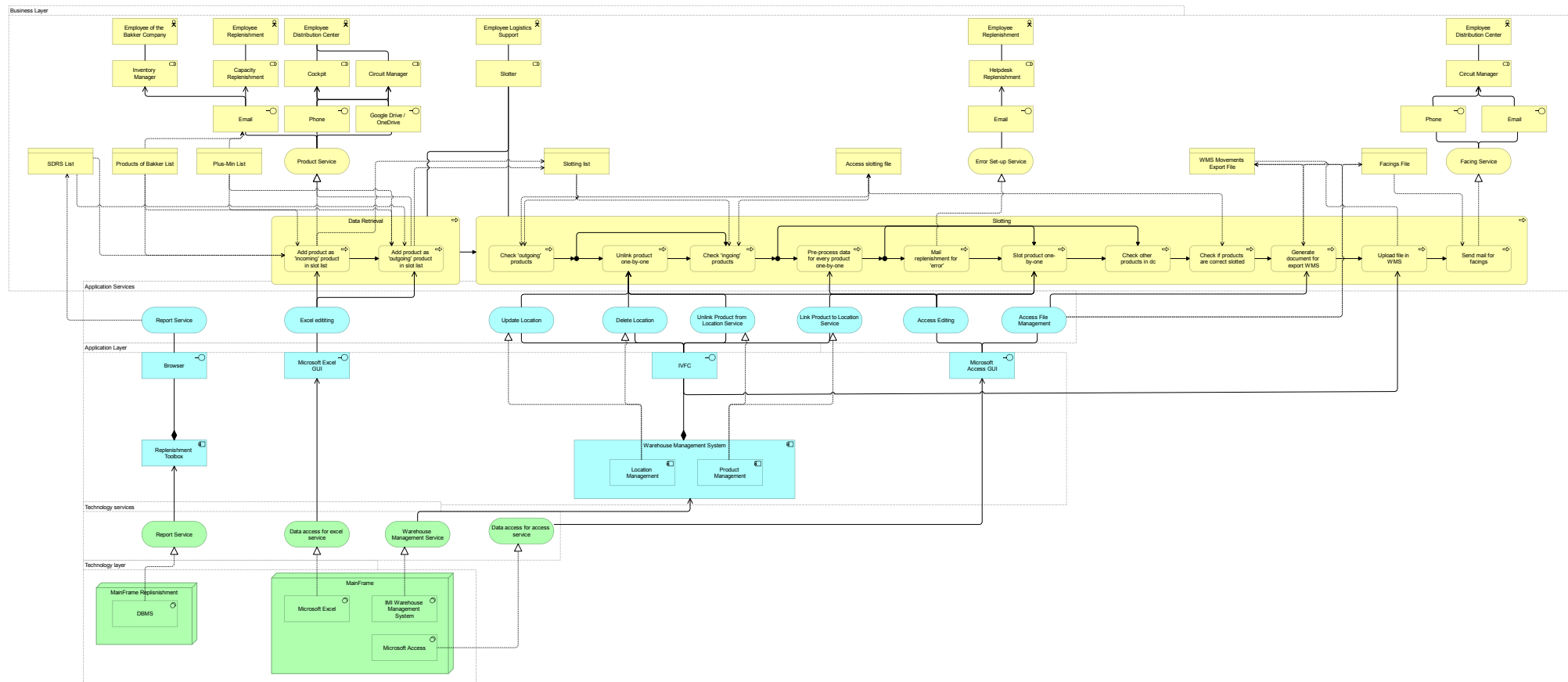


Figure 19: Total view of current architecture

3.5. Current situation performance indices.

To get a better understanding of the current situation there are several performance indicators defined to measure the initial situation of the slotting process. Those indicators are the measurements for the slotting nowadays.

Number of movements

The number of movements will be used as the first performance indicator. At the end of each working day, all the movements and newly arrived articles are written in a report. The movements can be used as a measure for the cost of the distribution center. If there are fewer movements, it will lead to fewer costs.

Number of free locations

The next performance indicator is the number of free locations. The distribution center needs to be sure that there are enough free locations available to be sure all the 'unforeseen' items can be placed when they arrived. Furthermore, due to different other circumstances like the weather, when it will be warm weather the soda will need more locations in the distribution to be sure that order pickers will not block each other when picking that specific product. There is a need for enough free locations in the distribution center.

Number of errors

The number of errors will be used as the third performance indicator. Those will be measured every week from the month of October to the end of week 47. This will be done by reactive note the error if this occurs. In addition, in the survey of the time indication tasks, the number of errors will also be mentioned. Those both results will be checked among each other to be sure nothing is missed.

Number of 'unforeseen items'

The fourth performance indicator is the number of unforeseen items. Those will be measured every week from the month of October to the end of week 47. This will be done by reactive note the error if this occurs. Furthermore, in the survey of the time indication tasks, the number of unforeseen items will also be mentioned. Those both results will be checked among each other to be sure nothing is missed. When there are unforeseen items there is a need for ad hoc slotting to be sure that every product has a location in the dc. In addition, when there is an unforeseen item or error which occurs then this will always be fixed ad-hoc. This does not benefit the plan that was made for that week and will result in more movement and more costs for that week.

Explanation

Most of the above-mentioned numbers give an indication of the slotting performance at that moment. There is a correlation between the time of the ad-hoc slotting moments and the number of free locations, number of movements and the number of 'unforeseen items'. Most of the numbers can be converted into money. This will also be seen in the results of the KPIs of this current situation.

Survey

To measure other difficult to measure indicators, there are different surveys conducted to measure this. To encourage the employees of the department, the easiness for filling in the survey will be high. The survey tool is created in Outsystems. This platform is chosen because the application of this thesis will be created in it and so the first impressions of this platform can be discovered. This platform creates a rapid application. In this tool, the user can fill in the survey anonymously so that they cannot be tracked and be sure that they give not biased answers. In addition, with creating the tool at itself, the freedom to create an intuitive tool is higher than using the standard software for creating surveys. The results of the tool can easily be intertwined with each other and hypotheses can easier be tested.

Furthermore, when creating a tool on itself it can be connected to different tools created in Outsystems or other platforms used by the AH. In the tool, all the answers can be seen publicly and will be shown on the department television to be transparent (which maybe encourage people more to fill in the survey).

Job satisfaction

The first one is to measure the satisfaction and happiness of their job as well as the energy leaks. To measure the job satisfaction of the department, especially the employee with the 'slotter' role, there is chosen to keep a survey under all the employees of the department of Logistics Support. This because when this is only doing under the 'slotters', there will be more bias because they want to tell the pollster what the pollster wants to know. If they do not know that a survey is being held especially for them then the survey will be filled in more honest instead of politically correct answers. The questions are focused more on the overtime if employees can earlier home it is better. And focus on the happiness of tasks which they are doing, when there are tasks which are repetitive (most of those tasks are not funny) can drop the satisfaction of that day.

Time indication of tasks

The second survey is to measure different time indications for the tasks of the process of slotters. Different questions of this survey will give an answer to the time spending part of the process of slotting. In order not to disrupt the work so much, it was decided to create a priority grid where the most time-intensive task is indicated with a 1 and the least with a 12. If there are 12 tasks are defined in the process. If one task is not executed for that day, the answer 'n/a.' is used. This will be used to filter out the tasks which are done be less by slotters. The first solution to the problem will not be focusing on those different tasks. Furthermore, in this, the process steps will be analyzed. It will say that some tasks will be in a different order to be completed. Every 'slotter' will do the process in their own way. It can be said that this need to be uniformed to be sure that everyone does it in the same way to optimize the process more than now. The questions of the survey can be seen in Appendix B.

Validation

After some time, the results of all those performance indicators are been validating by different experts on this process. The slotters themselves will also validate the results and will say if those are correct or not based on their experience in that specific distribution center. In addition, specialist people will have a look at it. Those are people who invent the slotting process at the AH and have enough experience to be sure that the results are logic. At last the team leader of the operational part of the slotting will be asked if the results sound logic.

3.5.1. Results for each KPI

Survey

The survey ran from November 5th until November 23rd. In this time the survey is filled in by all the slotters. The survey can be found in Appendix B. All answers are validated by the slotters and the general conclusion is that all the results sound logical.

Time indication

As can see in Table 8, the results of the third question are showed. The question is about the time indication or priority of the tasks doing by the 'slotter'. The question is separated for the different dc's. The results show independently that checking the slotting of that specific day is the most time-consuming part. What also can be concluded is that the regional dc is mostly the same and different than the national dc. The results show that slotting new products in the dc is a time-consuming part because the ideal location of this product is not empty, and the other product need to be relocated. The unforeseen items have also to deal with the problem of relocating the problem. The need for

empty locations in every aisle is high. This will be more explained in the results of the different number which are measured. In appendix C there is an explanation of each task.

Priority	LDC	DCO	DCT	DCZ	DCP
1	Judge slotting – general	Judge slotting – general	Judge slotting – general	Judge slotting – general	Judge slotting – general
2	New articles	Care for promotion	Care for promotion	Care for promotion	Slotting issues
3	Keep track of the slotting list	New articles	New articles	New articles	Other activities
4	Slotting issues	Keep track of the slotting list	Keep track of the slotting list	Keep track of the slotting list	Care for promotion
5	Other activities	Not announced articles	Other activities	Not announced articles	IVFC (interface to wms)
6	Errors	Other activities	Not announced articles	Disconnect articles	New articles
7	Site visit + preparation	Disconnect articles	Slotting issues	Slotting issues	Keep track of the slotting list
8	Mail to replenishment	IVFC (interface to wms)	Disconnect articles	IVFC (interface to wms)	Errors
9	Care for promotion	Errors	IVFC (interface to wms)	Other activities	Disconnect articles
10	IVFC (interface to wms)	Mail to replenishment	Errors	Mail to replenishment	Mail to replenishment
11	Not announced articles	Slotting issues	Mail to replenishment	Errors	Not announced articles
12	Disconnect articles	Site visit + preparation	Site visit + preparation	Site visit + preparation	Site visit + preparation

Table 8: Results for question three of the survey 'Slotting'

Process steps

Looking at the other difficult to measure indicator, the order of the process. These results are important to get a better understanding of which process step will be done at which time and in which order. The conclusion of this result can be used to maybe reshuffle the process order or make it uniform that all slotters use the same steps. Table 9 represent the result of question four.

One conclusion is that the process is differently done in the LDC than in the other dcs. The LDC slotter first relocates products to a smaller location. The SDF of those products is not big enough to keep them in a big location. The next step is that the slotter will relocate products to a bigger location. In the end, the slotter will allocate a location to newly arrived products for that dc.

Another conclusion is that the process of slotting in RDCs starts with creating a free location with help to unlink products from a location which has no SDF or has no stock anymore. The next step is to allocate new products to the most ideal location in the distribution center.

A strange conclusion can be said that the slotters of the LDC give low priority to checking the product data. This is strange because there are more products in this dc and will arrive more new products than in the others. This will imply that there are more errors in the data and that it will cost more time to control and fix those errors in the LDC. What can be seen in the separate results of the questions is

that the process will be done differently for each day for each dc. Sometimes there is a need for free locations in that dc and another day there are enough free locations to the slot.

Order	LDC	DCO	DCT	DCZ	DCP
1	Tighten number of locations	Disconnect articles	Disconnect articles	Disconnect articles	Disconnect articles
2	Expand the number of locations	Validate product data	Tighten the number of locations	New articles	Tighten the number of locations
3	New articles	New articles	New articles	Tighten the number of locations	New articles
4	Communication	Tighten number of locations	Validate product data	Validate product data	Validate product data
5	Disconnect articles	Communication	Expand the number of locations	Expand the number of locations	Expand the number of locations
6	Validate product data	Expand the number of locations	Communication	Communication	Communication

Table 9: Results for question four of the survey 'Slotting'

There are more important KPIs measured in the survey. Those are combined with another dataset and are explained in the next sections of this chapter. Here are all different numbers measured which are currently used as KPIs at the AH. Those can be separated for the national dc and the regional dcs, this because all the national dc contains other products than the regional DCS. The regional DCS contains almost the same products.

Number of free locations

The number of free locations is measured at the start of the week on Sunday. This because then the week starts at the Albert Heijn and on Monday the start of slotting will begin. If you know how many free spaces you have, you can plan based on this number. Figure 20 shows the number of free locations every week per dc per area of the dc. There will be a difference created between fresh and tenable products. Those numbers are measured from September 2nd until the end of the thesis.

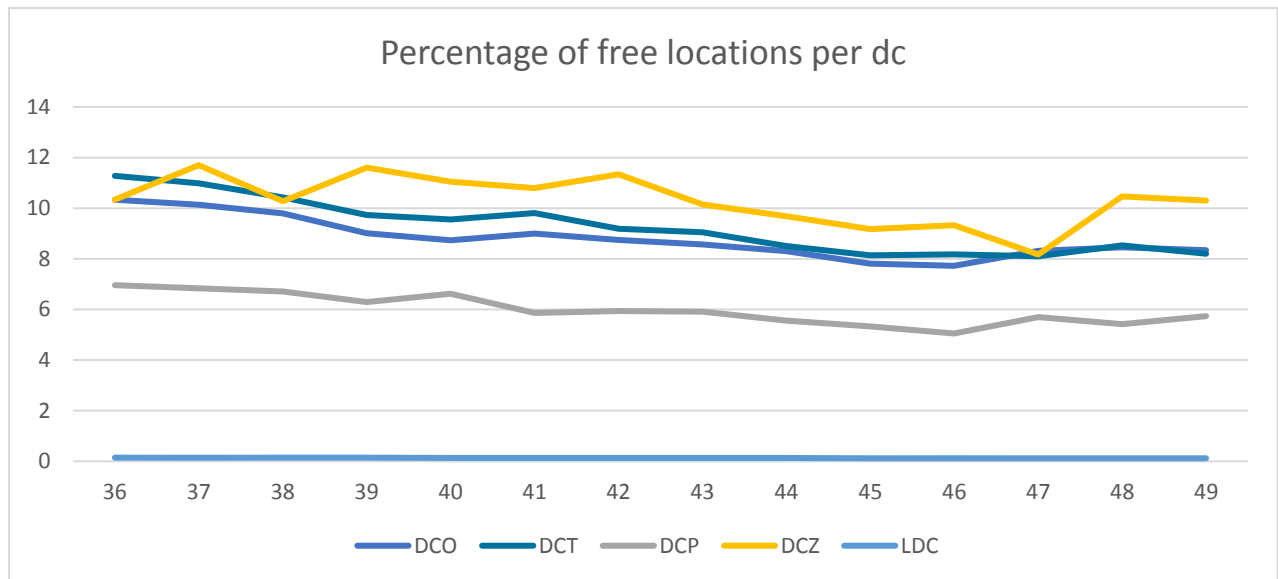


Figure 20: Percentage of free locations per dc

Number of movements and new product movements

The number of movements and new products is being shown in Figure 21. Hereby there is a distinction made between the new products and the movements. The new products stream are also movements but have no from space, the from the location will be the staging lane where the product came from. A movement is from a location in the dc to another location. There will also be sequence movements in the movements. This will be described as a movement to a location what is not free. Hereby, the product which is in the location where the previous product will be moved needs to be free by moving it another free location. This will cost extra time because the location needs to be empty before moving. This will be done over time. The costs of the movements per week per dc will be calculated with the help of 16 movements per hour and one hour will cost 30 euro. The total movements are the number of movements and new products. In Figure 21 the number of movements, newly arrived products, and sequence movements are represented.

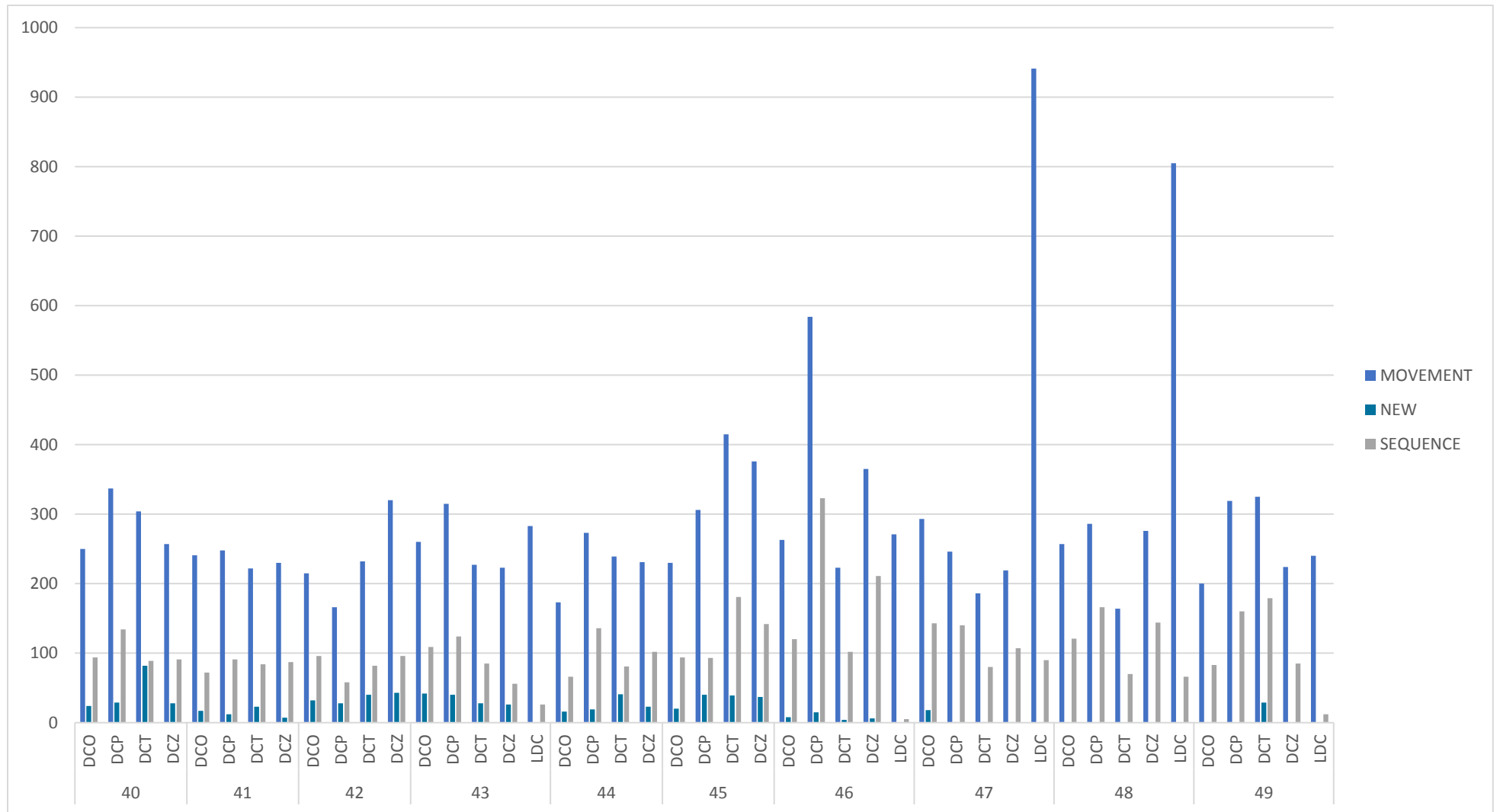


Figure 21: Number of movements, new products, and sequence movements

Number of errors

The number of errors is not separated for every DCS. Those will be handled by one person each day. In Table 10 all those results can be seen. There are different categories defined in the AH for the different errors which can occur. There will be around 158 new products arrived in different dcs. So, on average there will be 8% of the products contain an error. What can be concluded is that there on general many errors in this section. The quality of work is lower due to the errors and much slotting works need to be postponed until the error is fixed by another department. This will result in more ad-hoc slotting, which does not benefit the general plan of the slotting of that dc. What already is mentioned in the problem statement, when there are errors in the data, the slotting week plan cannot be fully executed.

Category	Week 38	Week 39	Week 40	Week 41	Week 42	Week 43	Week 44	Week 45	Week 46	Week 47
Weight	2	2	1	1		1				
OVZE		1	1			1				
Measurements	5	3	2		1	5			1	9
Number of products		2		3		1		1		
Wrong package structure	9	5	4	5	3	2	8			
No package structure	3	2	2	2	2		1	10	6	3
Temperature zone	2			1	1			1		1
Rolly/dolly									1	3
Total	20	15	10	12	7	10	9	12	8	16

Table 10: Number of errors per week per category

The time before the error is recognized is around 2 hours. What also can be concluded is the response time and the costs of every error will be around 10 hours. After the average of 10 hours, the error is fixed and the slotter can slot again. The 10 hours includes also the waiting time of around 8 hours to wait for a response. 1 hour can be taken for problem description and sometimes extra explanation and 1 hour can be taken for the response for the department of Replenishment.

Number of 'unforeseen items'

The number of unforeseen items is not separated for every DCS. Those will be handled by one person each day. In Table 11 all those results can be seen. There are different categories defined in the AH for the different unforeseen items. There will be around 158 new products arrived in different dcs. So, on average there will be 9% of the products not be foreseen and need ad-hoc slotting to be sure that there is a location for the product.

Category	Week 38	Week 39	Week 40	Week 41	Week 42	Week 43	Week 44	Week 45	Week 46	Week 47
Unknown	2		2	2	2	2		1	5	6
Data error				1	1	4	1	10	1	3
Reactivation					3					
BE change		1	2	5	2	2	4			
CKW		2		1	2	3				1
Not received				4	1	2	5	2	5	1

Belgium	1		1							
Total	3	3	5	13	11	13	10	13	11	11

Table 11: Number of unforeseen items per week per category

3.6. Conclusion

This chapter gave an overview of the current situation for slotting in the Albert Heijn. Several aspects of the current situation were described. The elements of the analysis of the current situation are related to the problem of slotting. Furthermore, the problems of wrong slotting or missing information in the master data are described. To conclude this chapter different conclusion can be drawn regarding the analysis of the current situation:

- The most time-consuming part is to check off each product in the distribution has the right locations, has the right number of locations, and will not lead to congestion in the circuit.
- The step of data validation or errors will cost mentally the most effort. In addition, with a solution time of around 10 hours, this is one of the most time-consuming tasks of slotting.
- There arrive around 160 new products in the different distribution centers in the Netherlands.
- Around 9% of the new products are unforeseen, which implies that they are not slotted on beforehand.
- 8% of the products which needs to be slotted contains an error in the data. Which can be that the dimensions of the product are not good, the package structure of the product is wrong, or the product is allocated at the wrong distribution center.
- The number of sequences is on average 33% of the movements. A sequence movement is a movement which can be done when the previous movement in this sequence is done.
- The data will be delivered from different departments in the AH as well as external parties like Bakker. There are four different sources: the plus-min list, the Bakker list, the SDRS list, and the AGF list. Those sources are supplemented by some complaints or ideas in Google Docs, by email, or by telling it through the phone.
- The current slotting process contains different steps to follow to be sure the slotting of the distribution center is the best.

After having analyzed the current situation, it is possible to sketch the future situation of the slotting process in the AH. Therefore, the next chapter describes the future situation.

4. Future situation

As mentioned, Albert Heijn must deal with many new articles, a significant increase in assortment and number of stores. This will lead to an increase in the volume of new articles and current articles that must be processed. In today's slotting the most time-consuming part must deal with the checks if the circuit of the different DCS is correct. Thereby when there is a need to slot new products the data of that needs to be clean from errors. Inevitable this will lead to more errors in the data and more unforeseen articles in the different distribution centers what will lead to more a more time-consuming process to process all the data. The future situation will be the base for the design and development of the solution for this research, this step is related to the design and development phases in the DSRM of Peffers et al. (2007). Without the future situation in mind, there can never be a solution drawn.

Human-in-the-loop

The vision of Licklider will be more applicable in the future of this world as well as at the Albert Heijn. The tightly coupling between the human brains and computing machines, which creates a mutually interdependent relationship, both capable of processing data, and making decisions in a way not even conceivable by each entity separately (Licklider, 1960). The human decision maker and the intelligent agent have both different strength and weaknesses, but the goal is to create a partnership in such a way that they can complement each other. This human-in-the-loop approach will allow the human to stay on top of the process. Which contains making strategic decisions that influence its outcome and need creativity to the change in the market. While the role of the machine is to help with the task by processing large volumes of data, identifying alternatives, simulating outcomes, and automating repetitive operations.

4.1. IA Framework

To find tasks which can be automated the IA framework is used (Dobrkovic et al., 2016). This framework is divided into six different steps whereby the first two steps focus on the identification of the different top-level tasks of the process. The next step is to decompose the top-level tasks into smaller tasks. All those processes are already described in chapter 3. The fourth step must deal with the assignment of tasks between the computer or the human. The last two steps are the design and implementation of different intelligent agents.

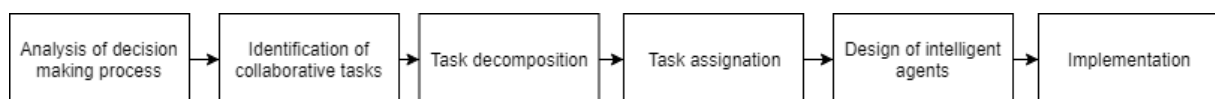


Figure 22: IA Framework

4.1.1. Task assignment

This thesis adopts the hierarchical task analysis (HTA) method to decompose the top-level tasks of the slotters (Annett, 2003) which are received with observation and recording of their work. Using the task list obtained through the help of interviews and hierarchical task analysis, the complexity level is determined for each task. The evaluation is done by all the slotters of the Albert Heijn and with the team lead of that team in a weekly meeting. The tasks which required external knowledge, a creative way of thinking, or contain one or multiple strategic components are assigned to the human.

The complexity of the computational tasks will make that the agents will have a different response to it. If a task has a high level of computations this will be automatically handled by the machine and so this will be given to it. Those tasks need no creativity of the human and are very straight-forward. If the tasks involve some decision, the agent will strive towards the optimal solution with help of the strategy which is set by the human. If there is no solution available or possible or the input is not clear the agent should make the best guess and should ask the human for approval. When the human gives his positive respond, the agent will learn from this response and proceed to execute this action. Alternatively, the task will be left to the human to process.

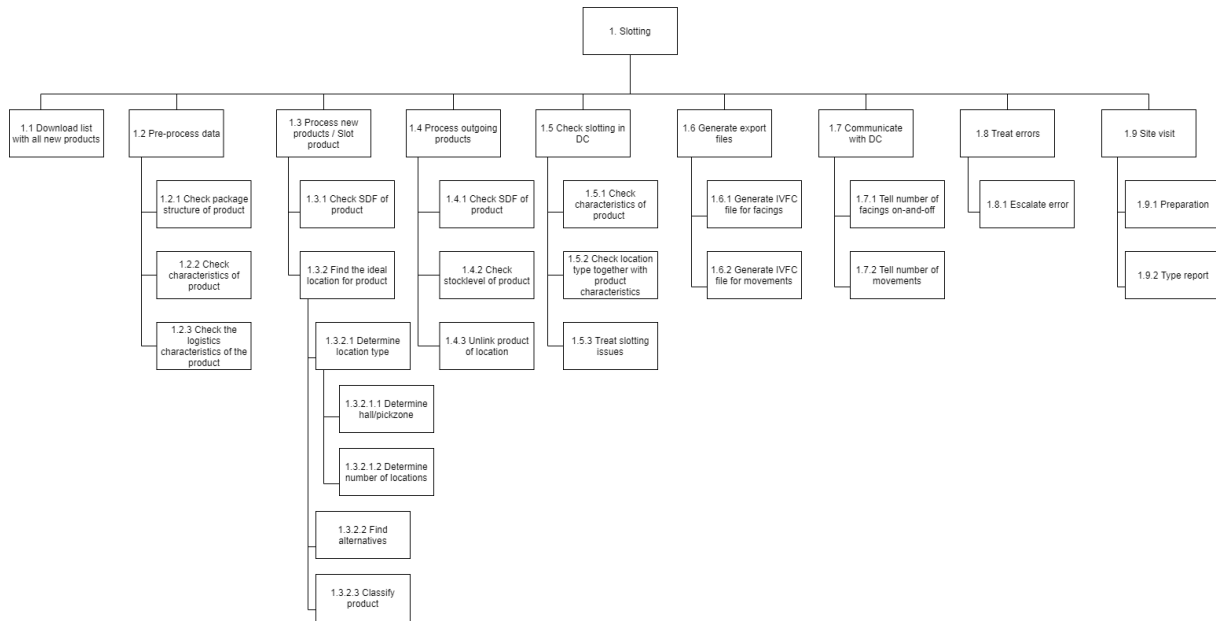


Figure 23: Task decomposition of the slotting process

After assigning the top tasks by HTA, each subtask is classified to be done by the human, the computer, or both, see Table 12.

		Assignment		
Level	Task	Human	Machine	Category
1	Slotting	X	X	
1.1	Download list with all new products		X	
1.2	Pre-process data	X	X	Rule-based
1.2.1	Check package structure of the product		X	Rule-based
1.2.2	Check the characteristics of the product		X	Rule-based
1.2.3	Check the logistics characteristics of the product		X	Rule-based
1.3	Process new product / slot product	X	X	
1.3.1	Check SDF of product		X	Rule-based
1.3.2	Find the ideal location for product		X	Expertise
1.3.2.1	Determine location type		X	Rule-based
1.3.2.1.1	Determine hall/pick zone		X	Rule-based
1.3.2.1.2	Determine the number of locations		X	Rule-based
1.3.2.2	Find alternatives		X	Expertise
1.3.2.3	Classify product	X		Knowledge-based
1.4	Process outgoing products	X	X	
1.4.1	Check SDF of product		X	Rule-based
1.4.2	Check the stock level of product	X	X	Rule-based
1.4.3	Unlink product of location	X	X	Skill-based
1.5	Check slotting in DC	X	X	
1.5.1	Check the characteristics of the product		X	Skill-based
1.5.2	Check location type with respect to slotted product		X	Rule-based
1.5.3	Treat slotting issues	X	X	Expertise
1.6	Generate export files and upload to wms		X	
1.6.1	Generate IVFC file for facing		X	Skill-based
1.6.2	Generate IVFC file for movements		X	Skill-based
1.7	Communicate with DC	X	X	

1.7.1	Tell number of facings on-and-off	X	X	Rule-based
1.7.2	Tell number of movements	X	X	Rule-based
1.8	Treat errors	X	X	
1.8.1	Escalate errors	X	X	Skill-based
1.9	Site visit	X		
1.9.1	Site visit preparation	X		Expertise
1.9.2	Type report after a visit	X		Expertise

Table 12: Assigning decomposed tasks to the human and the AI

After the assigning of decomposed tasks to the computer or human. The tasks are assigned to one of the next categories: Skill-based, Rule-based, Knowledge-based, and Expertise (Cummings, 2014).

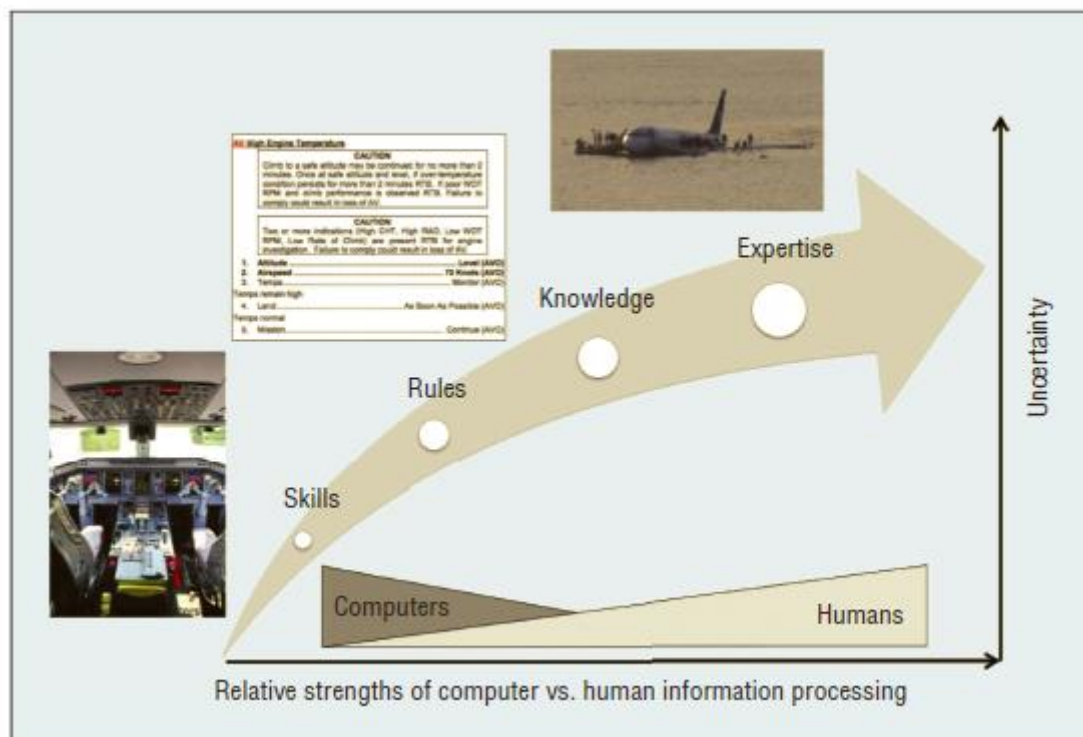


Figure 24: Role allocation for information processing behaviors (skill, rule, knowledge, and expertise) and the relationship to uncertainty

Figure 24 shows the difference between those categories. The x-axis shows the strengths of the computer versus the human in information processing and the y-axis shows the uncertainty.

The skill-based task can be accomplished by sensory-motor actions that require little or no concessions control to perform once the intention is formed (cummings, 2014). Automation is superior in skill-based tasks because such tasks have a clear feedback loop to identify the difference between the desired outcome and the actual results.

The rule-based tasks are highly rehearsed by rules, routines, or procedures to select a course of action (Cummings, 2014). Intelligent agents with optimization algorithms work primarily at the rule-based level. However, when faced with uncertainties, automation may not store the relevant information or doesn't include variables that impact the final solution.

The humans' power of induction is critical in the knowledge-based and expertise tasks. Humans' judgment and intuition are essential to deal with the situations where the goal is ambiguous, uncertainty is high and mathematically optimal solutions are unavailable. The induction of humans is

difficult for computer programming to replicate, especially true expertise. Considering efficiency, humans can make use of the intelligent agents' advantages in speed, calculation accuracy, memory and information processing capacity to complete the knowledge-based tasks. Therefore, tasks can be done by automation, human alone or their collaboration. The category assignment can also be seen in Table 12. There is a possibility that some tasks will fall into more than one category.

Rule-based assignment of task

The number of rule-based tasks in this process 'slotting' is very high. As already mentioned, a rule-based task requires a higher level of cognition since interpretation must occur to determine that, given some stimulus, which set of rules must be declared to attain the desired goal state. By the very nature of their if-then-else structures, rule-based behaviors are also potentially good candidates for automation, but the uncertainty of the rules is the key (Cummings 2014). In this case, humans' high-level cognition is required to decide the criteria and weight of an optimal solution. Hereby, the rule-based tasks need the collaboration of humans and automation to create a better solution or answer on their goals. To incorporate more into the aspects of intelligence amplification, the machine will be used to compute fast and complex situations in the solution. While the human will focus on the rules and the goals. To tackle the problem of automation of rule-based behaviors is to use the human to select the right rule or procedure, or even better to create the right rule or procedure for the given set of stimuli. As what Cummings mentioned, rule-based tasks are the candidate for automation, if the rule set is well established.

Data cleaning; the first step to automation

As already mentioned in the theoretical framework, data cleaning, a labor-intensive and complex process, is estimated that it will account for 30%-80% of the development time in a data warehouse project (Shilakes, 1998). This highlights the need for data-cleaning tools to automatically detect and effectively remove inconsistencies and errors in the data. For plenty of different processes, all the data needs to be accurate and free of errors. For example, in the slotting process, all the data needs to be available to create the most optimum slotting in the different dcs. When one or more product data is not conform the standard, the optimum slotting cannot be created, which will lead to more movements and indirect more costs. Most of the time this task of the subprocess is available in the main process and used as the pre-processing phase. Employing pre-processing steps can greatly improve results (Sorrentino et al., 2010). Which pre-processing task creates the best results depends on each situation and is open for further research because different steps will lead to a different manner of data cleansing (Zhao & Ma, 2017). Furthermore, the need for documenting all the steps of the slotting process is also one step towards an automated process. Automation is putting the process into code. A bullet list in a process document is a code if it is treated that way. This is a form of Intelligence Amplification like the human is writing down exactly step by step what a computer will do and will automate the steps in his head. From there, all the processes can be easier been automated (Limoncelli, 2018).

4.2. Goals

The future goals, identified along the process, must be met when implementing the tool and new process. Goals of the validation tool:

- Data validation of different imported list used in the slotting process.
- Callback which products have errors as well as products without errors.

When the tool is fully implemented, the chance that errors can slip in between is very low. The main goal is to validate every single new product before it will even be seen by the slotter itself. If there are any simple errors, it can quickly be bounced back to the department of replenishment. Other errors, which need more investigation, are already been seen and the action can be made.

4.3. Platform

The platform of the slotting process is multiple applications combined, which are responsible for the whole process. This platform contains the validation tool to validate all the product data before this reached the other application on the platform. When the data is validated and corrected, it will be saved to the database of the platform and it can be used by the other applications. One of those applications is the slotting tool where slotters can allocate products in the distribution or remove it from a location when the revenue is empty. There is a possibility to see different reports over the data with help of Microsoft Azure. Figure 25 shows a simple overview of the platform.

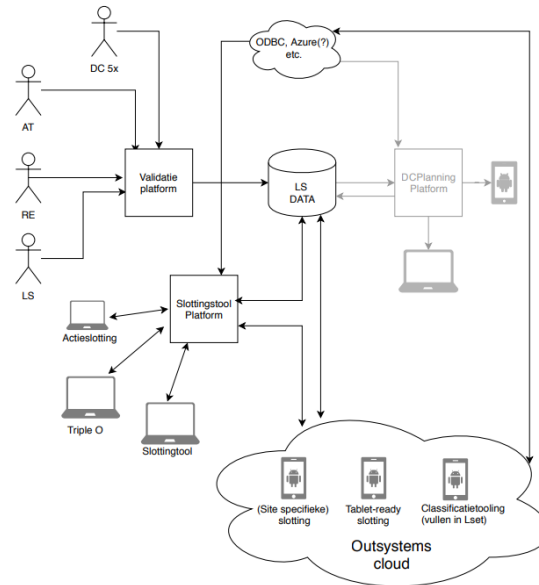


Figure 25: Slotting platform

Different departments of the AH upload documents to the platform, and before the data will be added to the database of the platform, it will be validated. The validation tool is an important tool to ensure that all data is valid and can be used by the employees of the department. One of the goals of the validation tool is to reject data what is not complete and not correct. It needs to be bounced back to the person who upload the document. Furthermore, this platform can be seen as the platform for data exchange from and to the department of Logistics Support and no documents need to be sent anymore via any other media, like the email or sharing it via internet. All documents which passed the validation are useable for the department.

5. Design and Implementation

5.1. Introduction

To enhance the worker's performance and work quality in a different process, the need for a supportive tool is needed. This tool needs to focus on the different rule-based tasks or processes in the main process. The different rule-based tasks or processes can be rewritten to one algorithm which supports the user to achieve its goal. An algorithm is created with help of different rules and is mainly created by people who have knowledge about programming and programming languages (Freeman & Skapura, 1991). It is a sequence of different if-else statements which summarize into a result, see Figure 26.

Usually, people which are not software programmers did not have the knowledge to create a software algorithm as well as creating a tool to create this algorithm with (Freeman & Skapura, 1991). The algorithm is a big black box where 'something' happened to achieve the desired results and goals, for people who did not create this algorithm (Ashby, 1956; Feshbach, 1979; Freeman & Skapura, 1991). The people shall know the behavior of it but will not know how this is reached. The process is described like, human or computer upload the input into a process, function or tool, 'something' happened, and there will be output (see Figure 27).

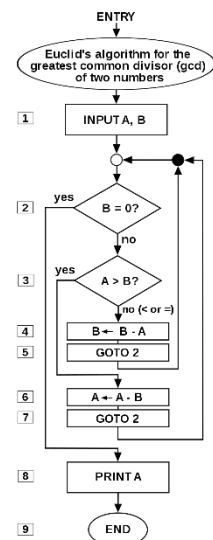


Figure 26: Example of an algorithm

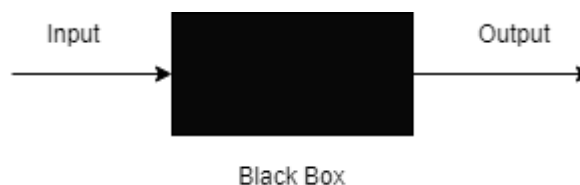


Figure 27: The Black box model

"A computer is a robot that performs any task that can be described as a sequence of instructions"

The focus of this tool provides insight into the different functions and behaviors what happened in the black box as well as to create the rules or algorithm by the user itself. The black box and the unknown working of it will disappear because the user will create this 'black box' by itself, so the black box will become a white box and the internals can be viewed. It can be said that it is a sort of white box 'plus' because the internals can also be changed by the user itself. Rules can be created, read, updated, or deleted, and how those rules are being validated to achieve the goal of the creator. The black box model in figure X will be upgraded to white box plus model as in Figure 28. The architecture in the black box is described in the next sections of this chapter.



Figure 28: The White box 'plus' model

5.2. Reference Architecture

5.2.1. Introduction

In this section, a reference architecture is defined for performing parts of the slotting process using Intelligence Amplification. The purpose of a reference architecture is to provide guidance for future development (Cloutier et al., 2010), provide standardization of concrete architectures and facilitation of the design of concrete architectures (Angelov, Grefen, & Greefhorst, 2012). This is defined as an abstract level. A concrete architecture is an architecture specifically designed for software application (Angelov et al., 2012).

The reference architecture consists of two different phases; goal creation and validation, and the slotting process as defined in chapter 3. After describing the reference architecture, the tool and the implemented architecture will be discussed. This tool will not fully cover the two different phases but gives a good suggestion that the concept works. The development of a suitable solution with an extensive list of stakeholders and their wishes will involve many risks, especially when it will be very complex. Building the right solution will be unattainable, but due to the limit of time and scope of the project, the final solution will be built in two iterations, with a feedback loop and a connection phase for the research findings between every iteration. The number of iterations is set up with the methodology of Scrum in mind which is described in this chapter.

5.2.2. Process

5.2.2.1. Rule creation process

The rule creation process starts with the goal definition process, which is explained in the next section. After the goal is clear, the rules can be created to which the goal can be validated. After testing the rule, the rule can be recreated or updated if it did not meet the specified outcome. Finally, the goal is used to validate all the data which need to apply to this rule.

5.2.2.2. Goal definition process

The goal of the processes needs to be clear before it can be automated. This goal must be known by every party involved in the automated process otherwise things will terribly wrong be interpreted and the goal will not be reached. The goal can be broad, but in this context, the goal will focus on different rules which can be set by the user. For example, the goal is to be sure that every product will be under twenty-three kilograms. The motivation to achieve this goal is the most important aspect of defining a goal. The mindset of the user needs to change before the goal can be reached, the mindset needs to be more towards an automated system which will check your goals by rules instead of doing it manually.

In Figure 29, a simple relationship can be found between the goal, the outcome, the driver, and the assessment. The goal of an improvement of a process will come from a driver which is found by the assessment what can be improved. The outcome of the process will state if the goal is met or not. For example, the assessment is that the workload of employees is too high. The driver of this will be the workload of employees and the goal is to reduce the workload of employees by another process. This process can that some parts of the process will be automated. The outcome of that process will give the answer to it is reduced.

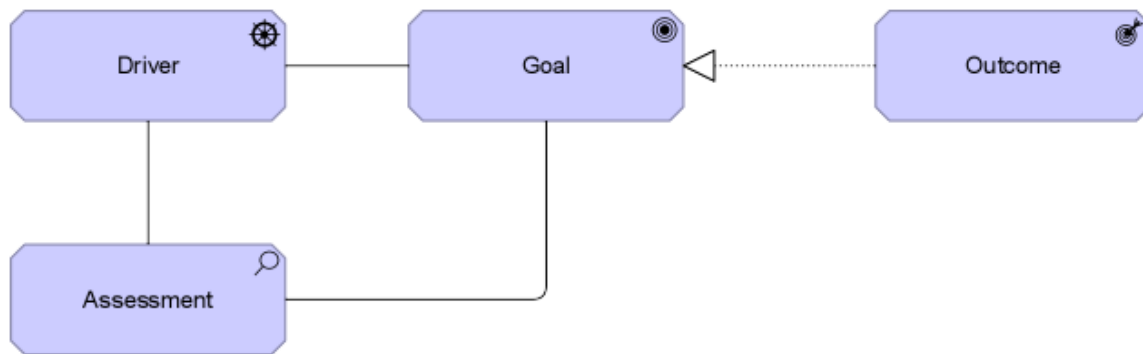


Figure 29: Relationship between the goal and the driver, the assessment and the outcome

The goal definition process consists of four subprocesses where the goal gets defined. Different questions can be asked to specify the goal better. At the end of this process, the goal is defined and the process for automation can start.

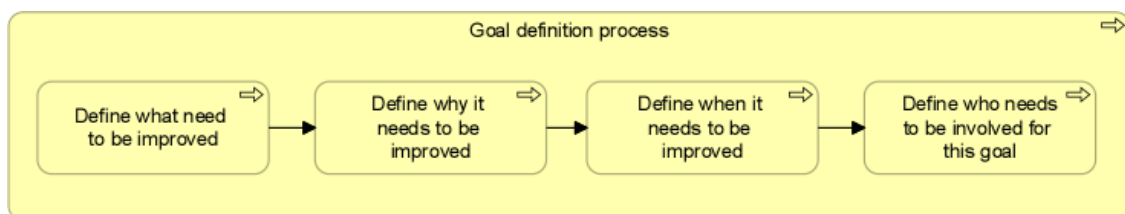


Figure 30: The goal definition process

5.2.2.3. Goal validation process

The goal validation process starts with the request for validation of the goals by triggering the machine. The machine will validate all the rules defined by the user and give back the results which conform the given rule set, and which do not conform this. The user will inspect the results and if some goals are not met, the rules will be updated or improved for the next run to be sure the goals are reached. This will be a process until all goals are reached. The user will have the lead about the rules which need to be checked to reach the goal. The machine is used for his computational supercomputer power to check the rules.

5.2.3. Architecture

The reference architecture is described in the next section. It starts with the business layer, followed by the application architecture, and finished with the technology layer. In the end, all the layers are shown in one figure.

5.2.3.1. Business Layer

The business layer represents the to-be scenario, where goals can be validated by defining them in an application. As explained in the process, the user needs to define a goal, create the rules which can validate the goal and test and refine it to be sure it will achieve the goal, Figure 31.

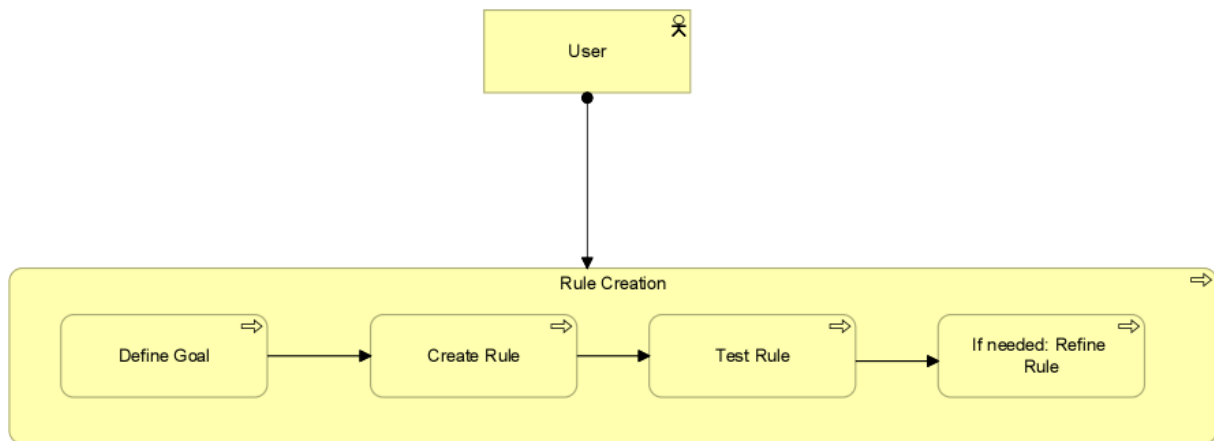


Figure 31: Business layer of the reference architecture, the rule creation process

After creation, removing, updating or deleting the goal, the user can validate its goal by triggering the application which will validate the goal by the defined rules by the user. To be sure, the goal is achieved the user will check the results. When it achieved the goal, the ruleset is well defined and need no further improvements. Otherwise, the user will change the rule-set and validate again to achieve the goal, Figure 32.

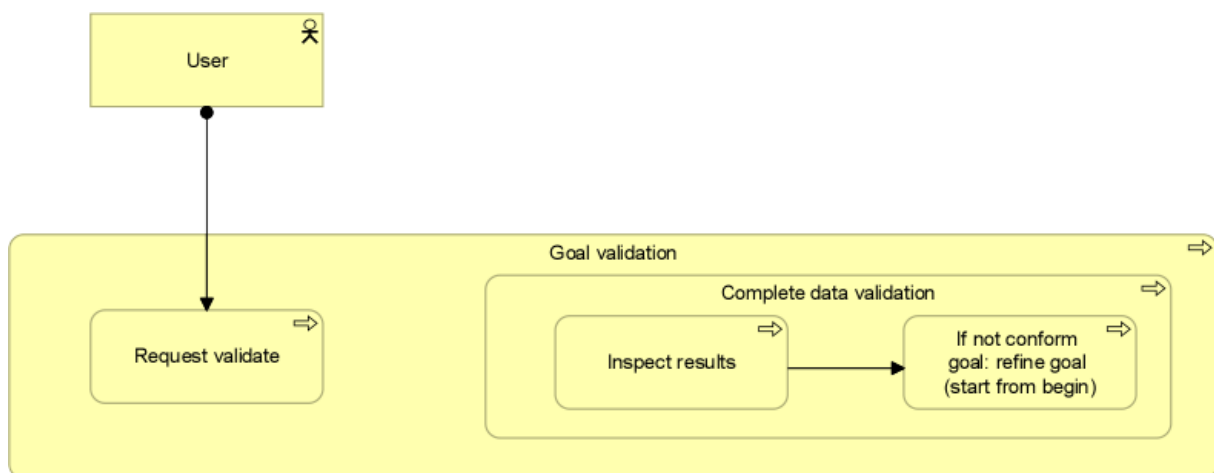


Figure 32: Business layer of the reference architecture, the goal validation process

5.2.3.2. Application Layer

The application layer is the heart of this architecture. The application layer describes the working of the application, the software components, information systems, services, and the data object is used. The application component encapsulates its behavior and data, exposes services, and makes them available through interfaces. One strong part of this application is that is independently reusable for multiple purposes. The user interface is the main access point in which the user can access the application. There are two different functions available in this application, the crud action of a rule, and the validation of the rule. The validation function contains a process in which the rule will be validated among by the user-defined thresholds and types. If the to-be-validated object is below, above, or between the defined thresholds it will give a result back to the user. The data object rule is created in the first function and be used in the second function. In Figure 33 the application layer can be found.

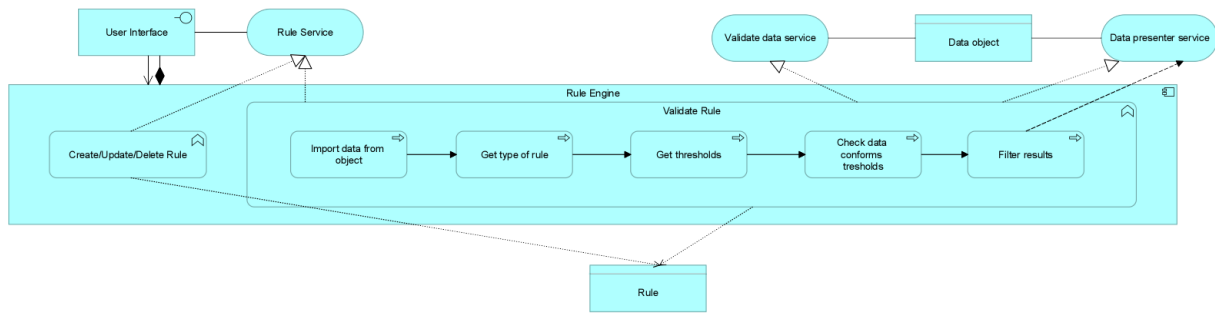


Figure 33: Application layer of the reference architecture

5.2.3.3. Technology Layer

The technology layer is relatively simple and consists of a cloud software system to make that the application is accessible at every time and moment of the day. This cloud software system runs in a cloud environment which has a database what is used by the cloud software system. In this database, all the created rules will be saved. The cloud software system has three different services to process and saves the data. And a presentation service where the data can be seen.

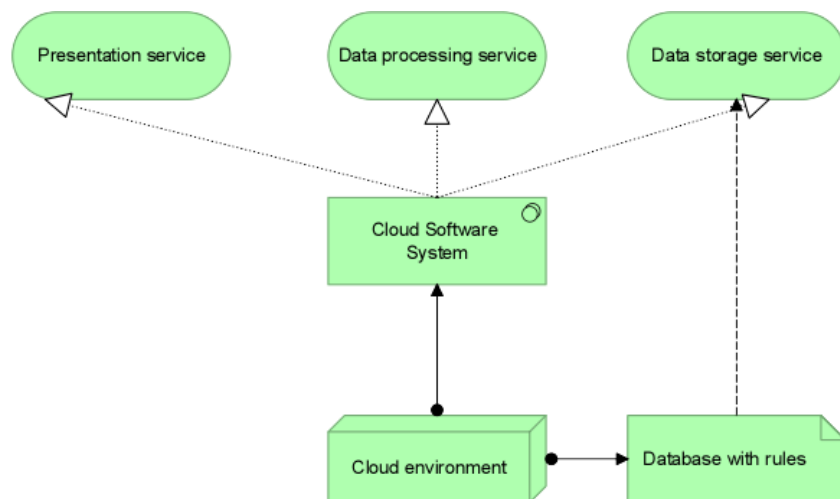


Figure 34: Technology layer of the reference architecture

5.2.3.4. Overview

To draw a comprehensive overview of the reference architecture, the three layers (business, application, and infrastructure) are synthesized in a single architectural overview. The reference architecture is shown in Figure 35. Most of the processes and functions are in the so-called 'black box'. The rules in this black box will be created by the end-user.

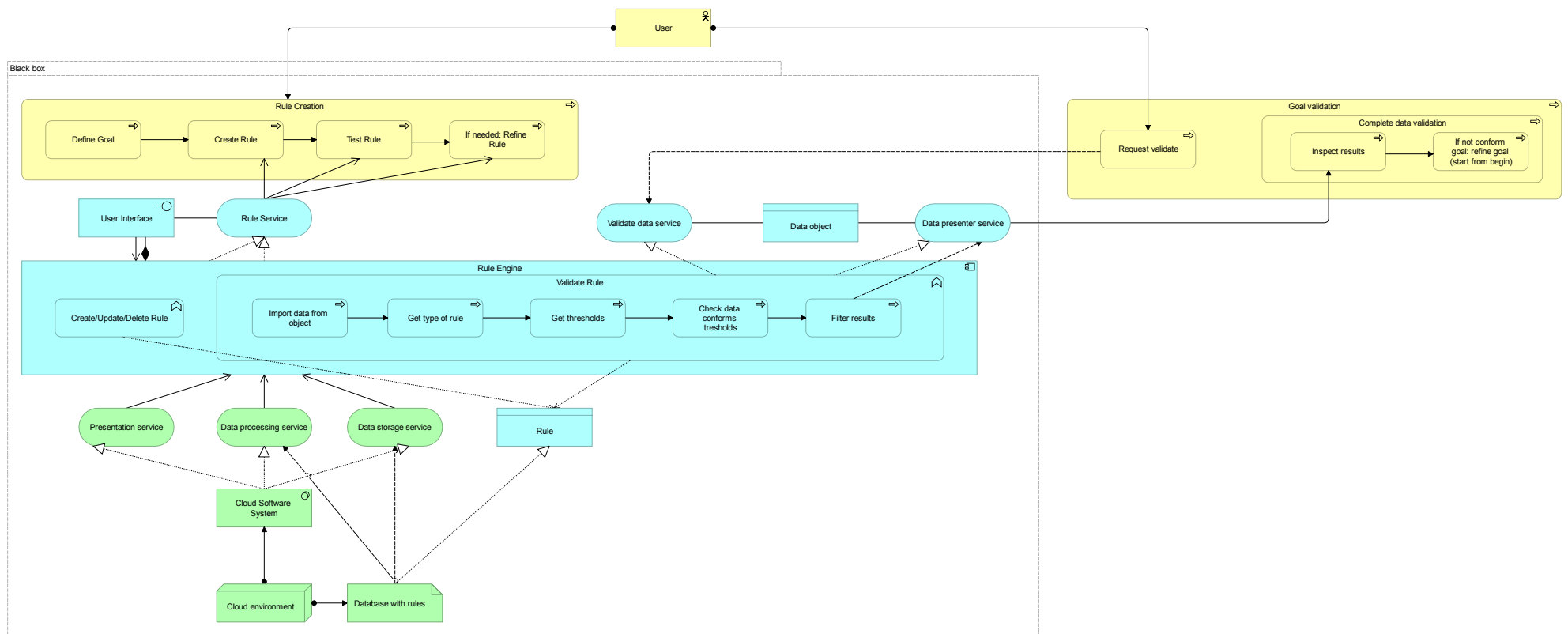


Figure 35: Total view of the reference architecture

5.2.4. Entity relationship model

The entity-relationship model is shown in Figure 36. The entity-relationship model consists of different entities which will form the system. The first is the parent rule. This rule can consist of multiple rules which can be encountered as one rule to check. It can be seen as a definition for example, that this might be true or this might be true to validate the object. Like, the order number must consist of a length of 15 characters or 16 characters. Different types of rules can be used in this system as well defined in the parent rule. The rules in a parent rule can also contain multiple sets of rules which the object needs to be validated, like an and variant. For example, the order number must contain the letter 'A' and also it needs to be 15 characters long. The rule type association is to save the different rule types on which the rule can consists, like decimal, text, date, another field in the dataset, or even a whole SQL Query for the database can be inserted as a rule. The different types of that specific type have a filter in it, which contains the checking threshold in which the rule can be validated, like 'in between', 'greater or equals', or 'smaller than' etc.

The 'ObjectToValidate' is, what the name explained, the object which needs to be validated by all the rules in the system. Some rules can be deactivated if they are not needed for special cases. This object can be every data object which the human will specify.

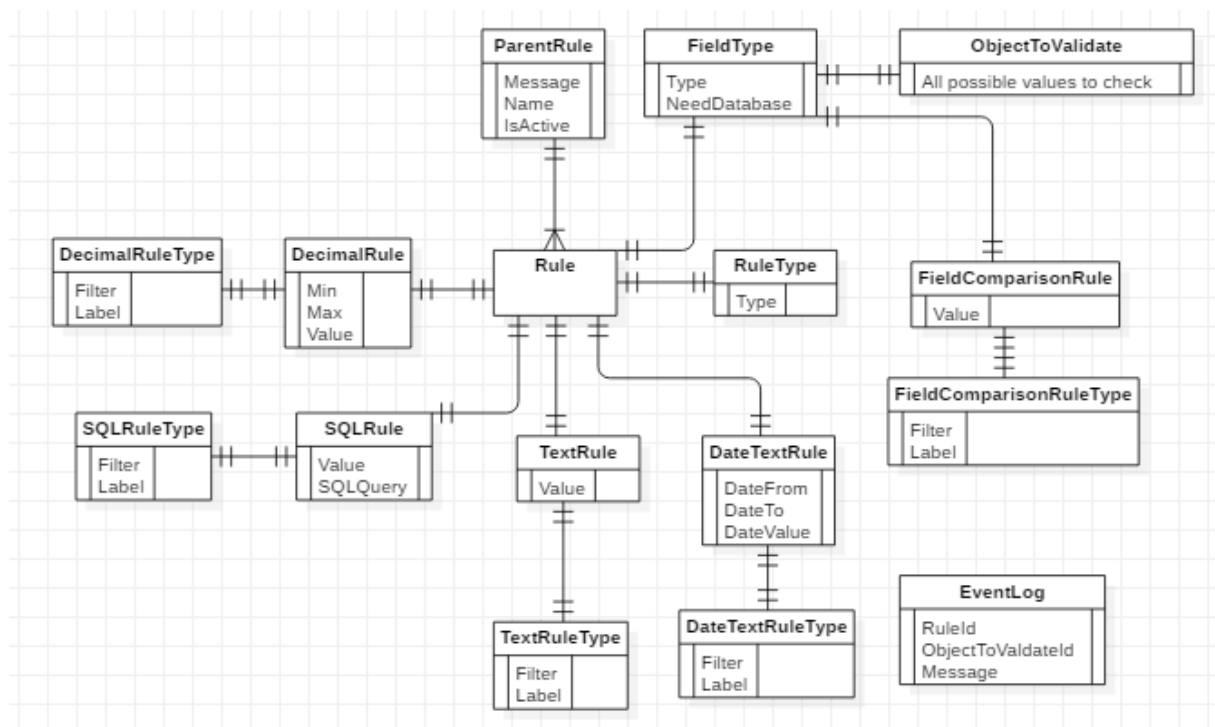


Figure 36: Entity relationship model of the application

5.3. Design and development of application

In this section, the design and development of the application for data cleaning and validation are described. First, start with a small introduction to the tool. Furthermore, the goal of the tool will be described. The approach of the design and development is told. And last, the different implementation phases or sprints are described with their product of the iteration.

5.3.1. Introduction

A hybrid automated situation (IA) needs clean data and so the first focus point will be a cleaning or validation tool (Rahm & Do, 2000; Raman & Hellerstein, 2001; Yoon, 2016). The starting point of the application should be the subtasks of preprocessing data (1.4). However, before starting with that, all the different errors and checks need to be known before it can be rebuilt in a tool. During the project, it appears that not all the errors will occur every time and will not reflect the current situation or interpreted differently within the different checks. To overcome these limitations there are meetings planned with different employees of the department of Logistics Support. Together we have discussed the most important errors and how the tool should respond. After the implementation of the validation tool, the next part of the slotting process will be more automated. This is the step of the allocation of the product. The first part of this process will be automated, the determination of the type of the location in the distribution center (1.3.2.1) as well as the location type for the already slotted product in the dc (1.5.2). At last, all the outgoing products will be checked if they have store demand forecast (SDF) or stock level (1.2).

5.3.2. Goal

The goal of the validation application is clear. With help of the tool, the errors and checks which are fired first and the employee or slotter knows direct what is wrong with the product which needs to be slotted. When the most common mistakes are no longer possible, the slotter can do their work faster. The extra benefit what will come into place is that the system will learn to detect errors and can do better the automatic slotting in the future. Another benefit is that users of the validation tool, in this case, the slotter, can make their own ruleset to check products. There are rules for validation for different categories but there are also general rules applicable to all products. The knowledge of programming or modeling in Outsystems or other tools is not needed.

5.3.3. The approach of the application development

This study is going to use the Scrum framework to build the prototype because it resembles the design cycle of the design science research methodology. DSRM has also the iterative nature of design development. It starts with different stakeholders and their goals and can be iteratively added to the backlog of the project. DSRM aimed at the research and scrum aimed to build the product.

5.3.3.1. Agile

Agile software development is the general term for a group of software development methods like Rational Unified Process (RUP), Extreme Programming (XP), and Scrum and many more. There are about 40 different Agile methods. Agile methods have the characteristics that they are iterative and step-by-step development. Those are different from the traditional waterfall development methods, where the requirements and wishes been identified before the development start. With the help of Agile those wishes, and requirements will be defined while developing the software. The benefits of Agile is flexibility to response fast on change on wishes, priorities, and requirements.

5.3.3.2. Scrum

Scrum is a framework within which people can address complex adaptive problems, while productively and creatively delivering products of the highest possible value (Schwaber & Sutherland, n.d.). Scrum

is a flexible and simple Agile method. It can be used for every product or service (software and non-software). Simple, flexible, and easy communication is the most important characteristics of Scrum. There are three different roles in Scrum:

- Product Owner

The product owner is the role of the internal or external customer. This role has the most authority and responsibility.

- Scrum Master

The Scrum Master is the facilitator of the project. He will not manage a team like a traditional project manager. The most important task is to motivate the team and will try to remove all the impediments of the team members.

- Development team member

The third role of Scrum is that of a team member. The team will define what the team will make the arriving sprint and commit to that. They created for every sprint a new goal what needs to be done at the end of the sprint. Sprints will consist of one to four weeks.

Another important aspect is the use of user stories. User stories describe the requirements for what needs to be made. Those user stories will be saved in the backlog, this is a prioritized list. The business can add more user stories to the backlog but can also remove some stories. The user stories are mostly formulated according to a predefined format like: *As a <type of user>, I want <some goal> so that <some reason>.* The user stories can be split up in tasks to define them precisely enough to estimate the time needed for completion.

Another important to use the methodology or principles of scrum is that the team, where I am part of, use Scrum to develop other tools or projects. The new slottings-platform is one of those projects and my part will be part of that slottings-platform team. Every day will there be a standup to measure the progress and there will be three questions asked to every participant of this standup:

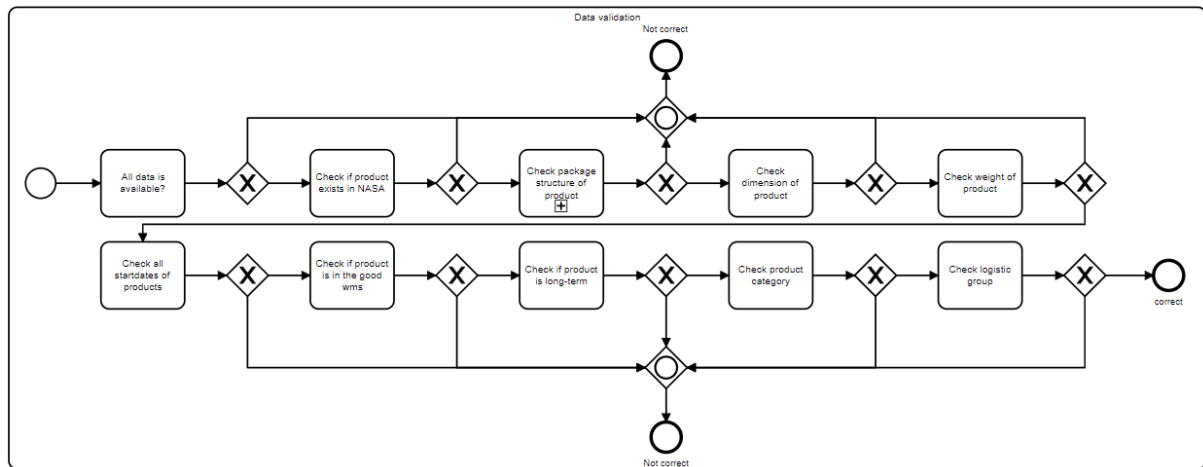
- What did I do yesterday that help the development team meet the sprint goal?
- What will I do today to help the development team meet the sprint goal?
- Do I see any impediment that prevents me or the development from meeting the sprint goal?

The backlog of this tool is created on a real-life brown paper. There is also use made of an online board namely 'Trello'. In this online environment, there can easily swim lanes made to progress the work. One swim lane contains all the user stories which need to be made. The other swim lanes are for the progress. In Figure 37 an illustration is shown for the online environment Trello.



Figure 37: Trello

As mentioned, the first part to improve in the slotting process will be the data validation. The data of the products, which will come into different distribution centers, need to be validated. When the data of a product is correct the slotter can go to the next step of the slotting process, slot the product in the distribution center. The data validation process is as follows:



All those process tasks are summarized in the rules in the next section. With help of those rules, the different tasks can be processed and validated automatically instead of manual.

5.3.4.1. Rules

Based on the process of data validation several rules are created to validate the different data from different source systems. The systems which are used are the Bakker-list, SDRS-list, the AGF-list, and the plus-min-list. Those files are read every Monday and Wednesday in the tool because then the files are prepared by other departments at the AH. The other days the documents can be uploaded manually, so that every day the different files can be checked if needed, by the system. The rules which are created in the system for the first validation part are:

- Product must be known in NASA
- Weight must be below 23 kilograms;
- Length needs to be longer than the width;
- Number of packages of a product needs to be more than zero;
- Length must be smaller than 1,20 meters
- Width must be smaller than 1,06 meters
- Height must be smaller than 2,15 meters
- Wrong logistic group (101 for LDC, 100 RDC)
- Different start date logistic group versus start date of pallet
- No start date in the distribution center
- No classification data like glass, flammable, pressable etc.

5.3.5. The architecture of the validation tool

The architecture of the validation tool is described in the next section. It starts with the business layer, followed by the application architecture, and finished with the technology layer. In the end, all the layers are shown in one figure.

5.3.5.1. Business Layer

The business layer represents the process 'data validation'. The user has defined the goal 'clean data' which have different rules. Those rules will be created one-for-one in the tool by the user. After

creating one rule the user will test this rule if it conforms the expected result. So yes, the rule is accepted, if not the user will refine this rule.

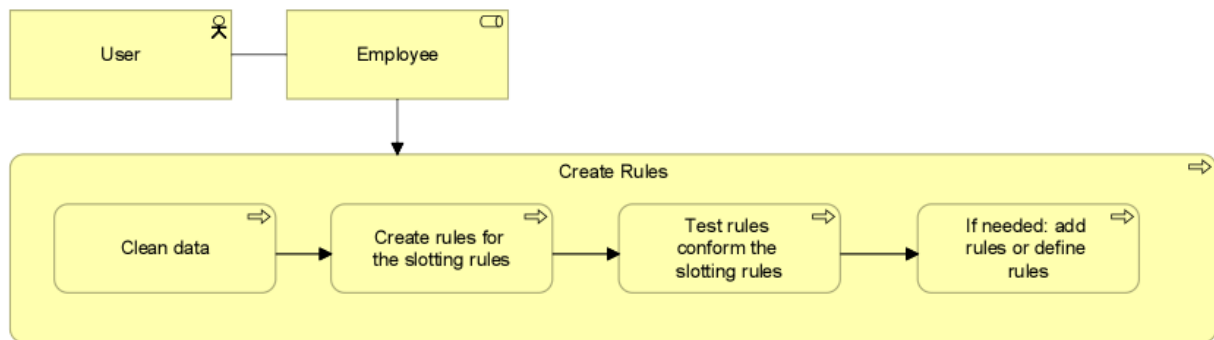


Figure 39: Business layer of the validation tool, the rule creation process

After creating the rules, the user will use the tool for validation of these rules. The user can manually upload the data he wants to be validated or he can create a schedule in which the data automatically will be validated. The different data sources which are used to create the slotting list are now automated and will be read every Monday and Wednesday. After the tool has validated the data, the user will inspect the results and will use this data in the next process step if it is conforming the goal. After the data validation is done, the product will be added to the slotting list and the slotting process can be continued.

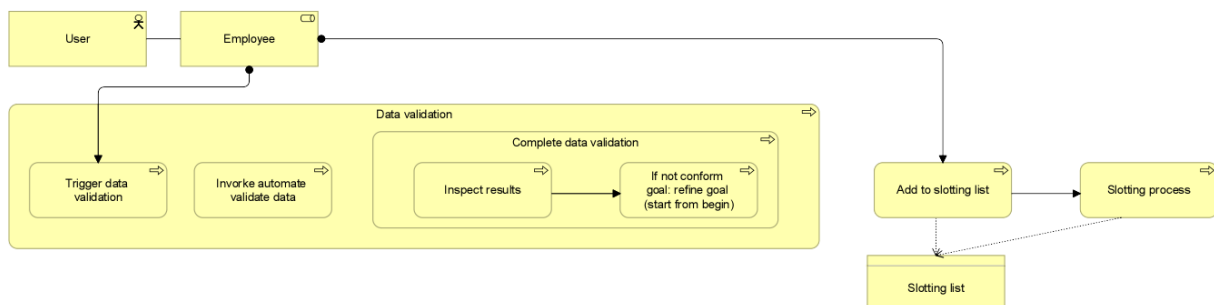


Figure 40: Business layer of the validation tool, the goal validation process

5.2.3.5. Application Layer

The application layer is the heart of this architecture. The application layer describes the working of the application, the software components, information systems, services, and the data object it used. The application is made in Outsystems. The user interface can be accessed by using a browser and surf to the URL of the validation tool. This will be the main access point of the application. There are two different functions available in this application, the crud action of a rule, and the validation of the rule. The action of the rule creation will be used to implement all the defined rules above in this validation tool. The validation function contains a process in which the rule will be validated among by the user-defined thresholds and types. If the to-be-validated object is below, above, or between the defined thresholds it will give a result back to the user. The data object rule is created in the first function and be used in the second function. When the tool is done with the validation the user will get all the results and can inspect the results given by the data presenter service which is accessible by the browser. In Figure 41 the application layer can be found.

Outsystems is a platform as a service (PaaS) intended for developing and delivering enterprise web and mobile applications, which run in the cloud, on-premises or in hybrid environments.

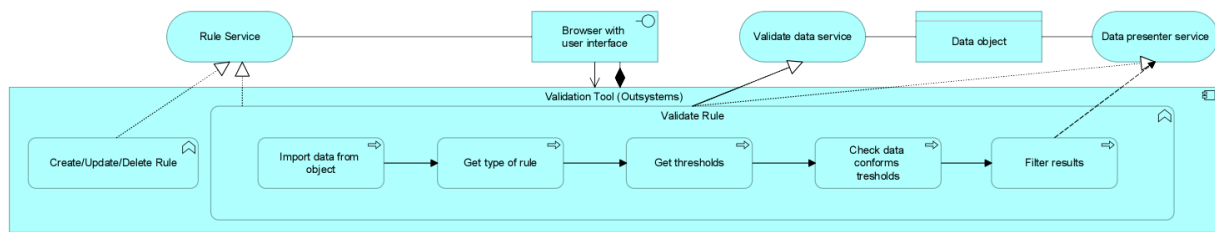


Figure 41: Application layer of the reference architecture, the goal validation process

5.2.3.6. Technology Layer

The technology layer is relatively simple and consists of the Outsystems development platform to make that the application is accessible at every time and moment of the day. This cloud software system runs in a private cloud environment on the Outsystems online cloud environment which has two different databases, one with rules and one with product data which are used by the application. In this database, all the created rules by the user are saved. The cloud software system has three different services to process and saves the data. And a presentation service where the data can be seen.

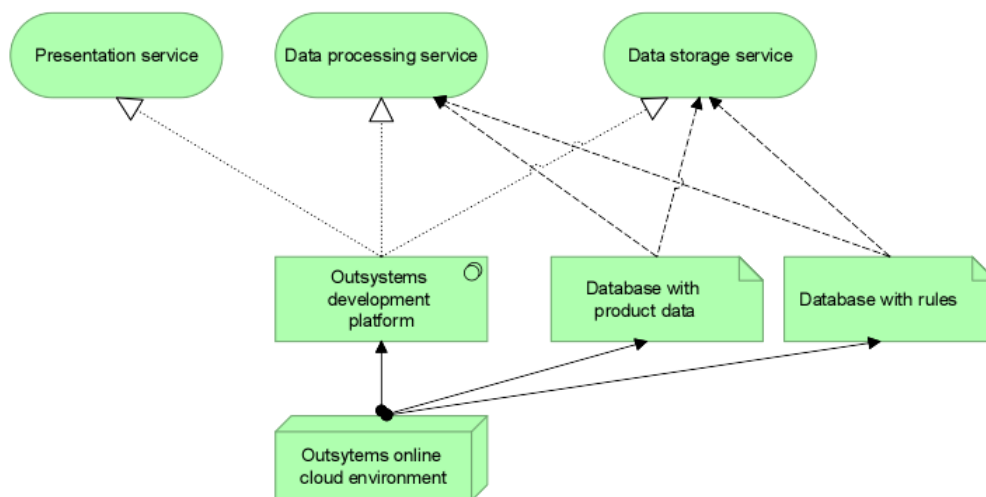


Figure 42: Technology layer of the validation tool

5.2.3.7. Total overview

To draw a comprehensive overview of the reference architecture, the three layers (business, application, and infrastructure) are synthesized in a single architectural overview. The reference architecture is shown in Figure 43. Most of the processes and functions are in the so-called 'black box'. The rules in this black box will be created by the end-user and saved to the database of the cloud environment. Those rules are used to validate the data.

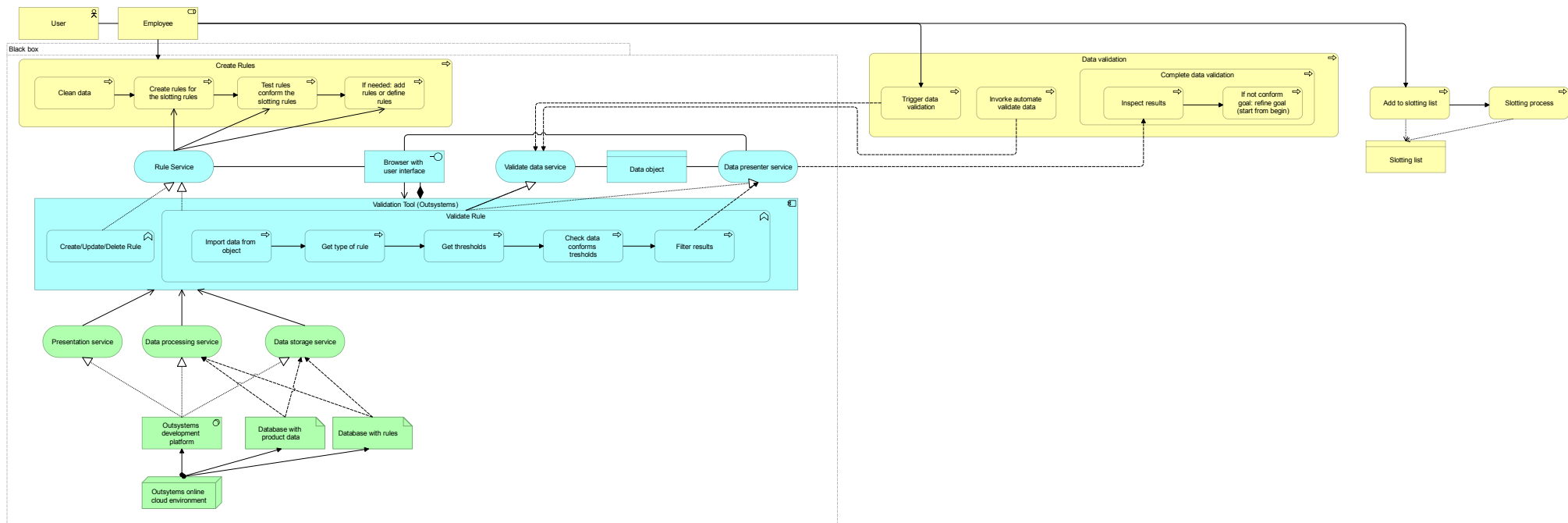


Figure 43: Total overview of all the layers of the validation tool

5.2.6. Sprints

5.2.6.1. Sprint 1 – Iteration (1)

Organization

The scope of this project is not a full-scale development process, but this part has not the resources of an ideal scrum team. The benefit is to join the other scrum team which is currently working at the slotting platform. This team creates a new slotting tool and my part the 'validator' will join this platform. Thereby you can say that there is an ideal scrum team, but the project will be done on my own. However, the initial meeting to determine the stories and the condensed time frame made this first iteration (sprint 1) an Agile way of working.

The initial session was held with several disciplines out of the supply chain of Albert Heijn: employees of the logistic preparation department, their team lead, team members of the scrum team and the manager of the logistics support department. Some stakeholders of the project, like Replenishment in this first sprint.

The results of the meeting are digitalized and put in a shared backlog with all the stakeholders of the initial meeting. During the first sprint, they were all informed of the process to ensure the development heading in the right direction.

Sprint planning

The following user stories will be done in this sprint:

Number	User Story
1	As a Logistics Support employee, I want a validation tool to reduce errors in data
2	As a Logistics Support employee, I want an overview of all product categories so that I can see all the product categories which there are created.
2a	As a Logistics Support employee, I want to create a category to link validation rules to it.
2b	As a Logistics Support employee, I want to remove a category if this not exists anymore or if I make mistakes in it.
2c	As a Logistics Support employee, I want to be able to view a category so that I can see which validation rules are linked.
2d	As a Logistics Support employee, I want to be able to change a category so that the characteristics can be changed.
3	As a Logistics Support employee, I want an overview for all validation rules so that I can see all the validation rules which there are created
3a	As a Logistics Support employee, I want to be able to create a validation rule so that this specific error can be caught and be fixed earlier.
3b	As a Logistics Support employee, I want to be able to delete a validation rule so that when it no longer applies it can be removed
3c	As a Logistics Support employee, I want to be able to view a validation rule so that I can see which categories are linked to it
3d	As a Logistics Support employee, I want to be able to change a validation rule so that the characteristics can be changed.
4	As a Logistics Support employee, I want an overview for all validation rules which applies to all categories so that I can see which validation rules apply for all categories
4a	As a Logistics Support employee, I want to be able to create a validation rule which applies for all categories so that I don't need to create them separately.
4b	As a Logistics Support employee, I want to be able to delete a validation which applies for all categories so that when it no longer applies it can be removed
4c	As a Logistics Support employee, I want to be able to view a validation rule which applies for all categories so that I can see which categories are linked to it

4d	As a Logistics Support employee, I want to be able to change a validation rule which applies for all categories so that the characteristics can be changed.
5	As a Logistics Support employee, I want to be able to link a validation rule to a category so that errors can be captured
5a	As a Logistics Support employee, I want to be able to link a validation rule to a WMS or NASA field so that this can be checked in the concerned database.
5b	As a Logistics Support employee, I want to be able to link a validation rule to flow so that only errors in that flow can be captured
5c	As a Logistics Support employee, I want to be able to link a validation rule to a specific dc so that only errors in that dc can be captured
5d	As a Logistics Support employee, I want to be able to test a validation rule so that I will be sure that the validation rule works

Table 13: Sprint 1 user stories

Technical explanation

As already mentioned, the tool is created in Outsystems. Outsystems is a low-code platform. This platform is chosen because of his ease of use. Furthermore, the tool can be platform independent, you can build an app for mobile as well as web or desktop. The platform is pretty time- and cost-efficient and can be used in no time. Due to the time limit of the thesis this platform is chosen and be used to validate the reference architecture. The only hiccup what is encountered so far is to build more complex user interfaces, but that is not needed in this prototype.

Another big pro for this platform is that everyone in the department can use the tool and program the tool. The intuitive editor helps the employee to build powerful app-components and will lead that employees work faster and efficient.

The data used and imported in the database of this tool was extensively discussed with several data experts of the department. The data must be easily available, and the responsibility must remain to the responsible part of the organization. This is already guaranteed because the validation tool is a part of the bigger slotting platform, which is in control of the department of Logistics Support.

Results and evaluation

In this period of three weeks, the first iteration of the tool was built and accessible through the development server of the platform of Outsystems. The tool contains different screens to create the rules and categories. The rules are very basic and different filters can be applied to be validated. The validation rules cannot be assigned yet to different fields in the database of NASA.

Unless the many missing parts, the tool is functioning well. Based on this result, the same members as part of the first meeting where asked to give their opinion again. Unfortunately, it was not possible to arrange a meeting with all the employees at ones. But different meetings where conducted to meet all the employees and asked their opinion. Different aspects of the tool were discussed and the first look and feel of the tool were positive. The potential experience was good of the tool, but the missing parts were essential to be sure the tool works like expected. Those missing parts will be developed in the next iteration and are described in the next sprint planning.

5.2.6.2. Sprint 2 – Iteration (2)

Sprint planning

The following user stories will be done in this sprint:

Number	User Story
6	As a developer, I want to be able to create a connection between the application and the different systems which the Albert Heijn use
6a	As a developer, I want to be able to create a connection between the application and WMS
6b	As a developer, I want to be able to create a connection between the application and NASA
7	As a developer, I want to retrieve all fields of the different database systems so that validation rules can be linked to that
7a	As a developer, I want to retrieve all fields of WMS so that validation rules can be linked to it
7b	As a developer, I want to retrieve all fields of NASA so that validation rules can be linked to it
8	As a Logistics Support employee, I want to be able to enter products in the systems so that it can be checked for errors.
8a	As a developer, I want to read a list of new products so that this can be checked on errors
8b	As a developer, I want to be able to upload a document to check if this contains errors
8c	As a developer, I want to be able to link fields to the imported document so that it can be checked on errors
8d	As a developer, I want to automatic read the selected documents so that it can be checked on errors.
8e	As a developer, I want to combine different lists so that one list needs to be checked
9	As Logistics support employee I want to have an overview of errors which occur after validation so that I can focus on those errors
9a	As Logistics support employee I want to rerun the validation tool so that I can check if an error occurs again
10	As an intern, I want to have meetings with other departments to tell them the story of validation so that they can prepare for major changes in this process
11	As a developer, I want a connection with the slotting tool so that there a to-do list can be created so that slotters know what they need to do.
11a	As a Logistics Support employee, I want a filter on the overview of the errors so that I can filter the only necessary things
11b	As a Logistics Support employee, I want to prioritize the errors so that I know which error is the most important one.
12	As a Logistics Support employee, I want to add more rules in one rule to combine different rules which each other.
13	As a Logistics Support employee, I want to add a criteria filter to the rules to check them in a pattern made by myself.
14	As a Logistics Support employee, I want to upload every possible CSV file to validate by this system
15	As a Logistics Support employee, I want to create my own template which can be validated by the system
16	As a Logistics Support employee, I want the automation of different files that it is validated in different time slots.
17	As a developer, I want to create different asynchronous processes that the system validates the objects fast.

18	As a developer, I want that the tool is in the slotting environment so that there is a connection between the slotting tool and this validation tool
19	As a Logistics Support employee, I want different rule categories so that I can label the different rules by a category.
20	As a Logistics Support employee, I want to create different templates, so I can validate my own data.
20a	As a Logistics Support employee, I want to link a template to rules to validate this in the validation part.

Table 14: Sprint 2 user stories

The user stories will be executed in three weeks. In this period the second and last iteration of the tool was done. The tool can create own templates, where the user can validate their own data among this template. In addition, different rules can be combined. Another opportunity is to create validation filters in the form of (expression 1 and expression 2) or (expression 1 and expression 3). This to validate multiple rules together and have more possibilities to validate data. The rules can be categorized on the functionality of the rule. Those categories can be made by the user itself.

User interface

The tool has different screens where data can be viewed, created, updated, or deleted. The tool uses the Dutch language as labels and text. All the user interface screenshots can be found in Appendix E.

Results and evaluation

The results and evaluation of the tool made in the last iteration are discussed in the next chapter. This because the tool is validated and discussed if it is production ready and the experience of the tool is described.

6. Validation

In this section, the prototype is evaluated and validated. This starts with the definition of the evaluation and validation criteria. The results are described next and the improvement suggestions are discussed in the last part of this chapter.

6.1. Validation criteria

Software development requires a measurement mechanism for feedback and evaluation (Basili, Caldeira, & Rombach, 2006). The IA driven approach proposed in this thesis, however, stresses the benefit of including the user in the process and therefore it focusses on the easiness of using the tool.

Both are required since measuring productivity cannot be viewed in isolation without any accompanying assessment of product quality (Fenton, Pfleeger, & Glass, 1994). To construct the criteria used to evaluate the tool the goal, question, metric (GQM) method is used (Basili et al., 2006). In our case, we define two goals, namely to evaluate workers performance and work quality. This leads to the following two main questions:

- What is the enhancement of workers performance by this prototype?
- What is the enhancement of work quality by this prototype?

6.1.1. Validation approach

To validate the prototype described in the previous steps of the research, experts have been consulted. Different expert sessions are planned to validate if the tool will enhance the workers' performance in the slotting process, as well as the work quality. Furthermore, in those sessions, different interview will be conducted to validate the tool and the working of it. To further validate the prototype, there is a survey conducted after each expert session. This survey will contain questions focused on the improvement of workers performance and work quality topic of the tool in combination with the process.

The tool will be validated in the daily routine. Every day the validation will be run to validate all the new products which will arrive at the different DCs of the AH. Different performance indices will be used to validate this. The first one is the time what is used to detect an error and when the error is detected. The second is the time to know if the error is fixed. This is the time in the mail conversation to the different departments. The last one is the number of errors which are detected.

6.1.1.1. Expert sessions and interviews

Conducting interviews is the most commonly used method of qualitative research, and especially useful for exploring complicated phenomena that require elaborate descriptions of meaning (Palinkas, 2014). To be sure that exploring those phenomena is done properly with respect to completeness and to be sure there is no bias, the different interviews are created. This can be found in Appendix D.

The first session is organized for the explanation of the working of the tool and how to create, read, update, and delete different rules for validation. After this session, a survey is conducted to validate the working of the tool. After the first session, the next sessions are used to execute different use cases for the validation that the goal is reached, and the data is validated. The goal of all those expert sessions is the validation of the two main questions as stated in the previous section. The enhancement of the worker's performance and work quality is the most important part to validate.

6.2. Results

In this section, the results of the different interviews, surveys, and expert sessions are discussed. Several experts and employees are interviewed, and they have filled in different surveys about the

validation tool. The feedback of the interviews, surveys, and expert sessions are concluded in Appendix D.

6.2.1. Reference architecture - prototype

The reference architecture can be validated among the use of the prototype. It can be concluded that in theory, an implementation of the reference architecture can be built. But there are some limitations to this prototype, causing us to emphasize that it, in theory, the implementation can work, but there is a need for better testing, with a bigger prototype setup.

The Outsystems web application was a success. One of the documents the slotters use is automated and will be read every Monday and Wednesday and will be automatically be validated. Users of the tool can create their own template to validate more documents and can create specific rules on those templates. This gives a big opportunity to validate documents and objects which are not specified in the tool itself. The big disadvantage of this tool is that the slotting process needs to be changed and on beforehand the validation needs to be done by one employee which is stated at the helpdesk on that day. The other implication is that the tool must start living in the department and that rules and other templates are created by the employees. In other words, the power of this prototype will be nothing when there are no rules created by the user because then errors should still exist and not be caught. The potential values this tool can give needs must be fully utilized to reduce the costs and time of the validation. This can be reached when the employees of the department will use the tool and its power.

One limitation of the tool is, what is already mentioned, that only the plus-min list is validated at this moment. To full automate the validation part of the slotting process more rules and files need to be added to the tool. This can be done to create different templates and validate over those templates and create rules based on this template.

The other limitation of the prototype is the performance, there is no realistic case could be tested, this due to the first limitation. The test case did not have many rules but with more rules, the performance and the validation time will increase linearly. There is a need for a solution for this, could be stamped as a part of future research. Furthermore, working with more people at the same time on the tool is not tested. There cannot be said that this implies the performance of the tool. The validation tool runs online on the cloud environment and the speed of the tool can be slow if other applications run simultaneously.

The proof that the validation tool based on the reference architecture can be built. As mentioned, the performance is one of the major issues. When this is not on the level it must be the tool will lose its potential value. The other one is that the rules and templates need to be created to be sure every product is validated. With the technology readiness level can also be checked if the tool is at a certain stage to prove that it could work in a production environment.

6.2.2. Results of the reference architecture in practice

First, the results of the survey show that the tool has the potential value to reach the goal explained in chapter 2. This is said by different employees which are interviewed. The potential value is to take away the boring part of the work as well as the improvement of the detection of different errors by the tool. The speed of the recognition of the data errors is much faster than how it is now. Furthermore, a conclusion what can be made is that the work initiative and straightforward. The tool is production ready but there are several nice to have new features which will only upgrade the tool and its potential value. Those features are implementations for future research or even better the next sprint to implement.

Second, the results of the interviews with the employees concluded that the tool is intuitive in use as well as straightforward. This is the same conclusion as in the survey, but the question is double asked to check if there is no bias. Based on the interviews the conclusion is that the tool has potential value to reach the goal of the department. There are missing features but when those are fixed, the tool will be better.

Third, the results of the analyze of the validation process show that the chosen solution can work. Errors which are caught by the rules, which are created by the employees, are been seen earlier. This also because the plus min list is checked earlier relative to the before situation. One big selling point of the validation tool is that never data of a product can be forgotten because all the data will be validated among all the rules. Based on timings in week 7 in 2019, the speed of error recognition is higher than before the tool is introduced. Because the validation process runs on the background, the employee does not have to wait for the results. When the validation is done, the employee can see the different errors. Because this process runs on the back, it takes less time to check every product one by one.

One thing all the interviewed employees agrees is that there must be confidence that the tool does what it promises. To ensure this will happen the tool should be used more to give the confidence. The opinion that is shared based on a week of testing is that the tool will add value to the department because the process itself no longer needs to be done. The ability to add and remove rules gives you the freedom and grip on the validation as it was already done. Some even say that this tool triggers themselves rethink their process and come to a new way of working. Furthermore, the design of the tool becomes a lively process where everyone contributes to, to be sure the goal of 100% validation is reached.

At last, the department of Logistics Support is seeking for the right documents to validate. In addition, they are searching for a manageable list, which is the base for example the slotting. With help of this tool, the initial list or first phase of this manageable list is realized. In the short interviews, the employees point out that responsibility of those list is not actually in the hands of us, but through this tool it becomes more insightful.

6.2.3. Technology readiness level

Another method to validate the prototype is with the use of the technology readiness level as developed by John C. Mankins (1995). The development of new systems capabilities typically depends upon the prior success of advanced technology research and development efforts. The challenge for different managers is to be able to make clear, well-documented assessment of technology readiness and risks. To measure this, different technology readiness levels are introduced by Mankins (1995) Those levels are a systematic measurement system that supports different assessments of the maturity of technology and consistent comparison of maturity between different types of technology (Mankins, 1995). This will evaluate the tool when this is ready to be used in the daily operation.

TRL 1	Basic principles observed and reported
TRL 2	Technology concept and/or application formulated
TRL 3	Analytical and experimental critical function and/or characteristic proof-of-concept
TRL 4	Component and/or breadboard validation in laboratory environment
TRL 5	Component and/or breadboard validation in relevant environment
TRL 6	System/subsystem model or prototype demonstration in a relevant environment (ground or space)
TRL 7	System prototype demonstration in a space environment
TRL 8	Actual system completed and “flight qualified” through test and demonstration (ground or space)
TRL 9	Actual system “flight proven” through successful mission operations

Figure 44: Technology Readiness Levels summarize

Based on the interviews, the technology readiness level is decided, this on an average of all the employees who conducted the interview. This level will be a 7. The tool is ready for testing in a ‘space’ environment which can be translated to a test environment in the department of Logistics Support, the spotting platform. To ensure the tool can be enrolled in a production environment in the department it need to be test heavily and evaluated. It can be used as the main validation tool of data in the data validation part of slotting when most of the wishes are implemented and the “teething” pains are solved. Besides that, it can be used for more rule-based tasks or processes to use this tool to give advice for different kind of goals or questions. Another possibility is to enroll this application to multiple departments in the company. To reach level 8, different changes need to be done by the developer of the prototype. First, all the rules of the validation must be implemented. All the possible and different fields and rules type must be implemented. The prototype must have all the technologies implemented and everything like the user interface must be done.

6.3. Conclusion

In this chapter the validation of the reference architecture is described. The results show that the reference architecture is a solid base to build applications from. The validation tool is one of the first implementations of the reference architecture which contributes to the goal of a better validation of the data. The first implementation phases show that the tool has the potential value to replace the validation task, which is done now manually. The detection of errors is faster and more accurate because nothing is skipped. When the tool will be fully tested and implemented and prove the potential value it can be enrolled in more processes which are built on the principle of rule-based tasks. Furthermore, the technology readiness level method is used to prove the readiness of this product, this can be set on level 7, which implies that the systems are ready for demonstration in a space environment, in this case, the testing environment of the department of Logistics Support, the slotting platform. One of the key conclusions is that the human is still needed to ensure that the rules are correct, and all the goals can be reached.

What is the enhancement of workers performance by this prototype?

With help of the validation tool, the employees can work faster and accurate. The validation tool validates product data to ensure there are no mistakes. When there are mistakes the employee can fast communicate the error to other departments. The boring task of validation is now done by the machine so that there is more time for the other more creative tasks of the slotting process. The

workers performance will be enhanced with help of intelligence amplification because the machine does the repetitive task of validation. The human in case will update or create more rules to make the machine smarter to detect more errors in the future. As said in an interview the speed of recognition of error is much faster than before this tool. He said: "It costs me at least 3 hours less work per week". That will imply that when more features are added, the employee will be in principle spent less time for this task the rest of the week.

What is the enhancement of work quality of this prototype?

The validation tool upgrades the quality of the data as well as the quality of work. Because the tool validates a product every time according to a certain set of rules, you can say that the quality is improved because no product is skipped. The set of rules, created by the employees, is the equivalent of the measurement of the quality of this data set. When all rules are applied on the data set, one can say that the quality of the data of this product is good, and the product can be slotted in the next steps of the process. The quality of work will improve because you can ensure that there is always good data at your disposal.

7. Conclusion

In this chapter, the conclusion of this research will be given. First the answer to the main research question is given. After that, the answers to each research question is described. These answers will lead to different conclusions. Furthermore, limitations of this research are described. Finally, the contribution of this research is discussed and recommendation for future directions in research and practice is given.

7.1. Introduction

The main research question of this thesis is:

“How to enhance the workers’ performance and quality of work in the slotting process through the adoption of Intelligence Amplification?”

To answer the main research question, this research has resulted in three deliverables. The first product is this report which contains an extensive description of the state of literature, the current situation and the possible future situation of the environment where the prototype can live. As well as a description of iterative development of the prototype.

Eight research questions are stated to answer the main questions. The questions are divided in four groups, which were answered in previous chapters. First, the answers of all the research questions are given, to come to the final conclusion of this research.

All those answers can lead to a new iteration cycle of the design science research methodology to improve the reference architecture. The reference architecture now is based on the rule-based tasks of the slotting process. This architecture can be improved by validating this to other processes. For example, the slotting process or even in other processes in the company or other companies. Furthermore, the new cycle can lead to new insights in how rule-based tasks can be executed with the adoption of Intelligence Amplification. In addition, the next iteration cycle can look to other tasks than only the rule-based tasks. The output can be modification to the reference architecture or a new reference architecture for that types of tasks.

7.2. Literature

The first set of research questions were answered with the literature study. To get more familiar with the concept of Intelligence Amplification and the tasks which must deal with the allocation of products in the warehouse. This lead through the following questions:

SQ1a: What is the current state of the literature of Intelligence Amplification and slotting?

SQ1b: What are the applications of IA in enhancing worker’s performance or reducing workload?

SQ2: What are the applications of Intelligence Amplification in enhancing slotting or the slotting process?

The background of the research problem is explored using the literature review. Two concepts are described in chapter 2, namely slotting and IA. Slotting, the process of allocation of products in the warehouse, can impact the warehouse operations significantly. These operations contribute to the overall cost of supply chain management. Warehouses become more valuable to the company and the supply chain when they contribute to reduced costs and improved service, flexibility, and responsiveness. As described in the first section of chapter 2, different policies are described on how the warehouse should be arranged but there is no uniform way of doing the slotting in it. There is no uniform process of assigning a location to a product in the warehouse.

The conclusion made is that the research of IA is in its infancy. In the current state of the art, there are many studies taking advantage of the human-machine cooperation to complete tasks but those are focused on enhancing the machine instead of the whole system. They don't further realize the concept of Intelligence Amplification. The aim of the researches to involve humans is to help overcome the limitation of automation. It focuses more on 'how can we make the system more intelligent', but not on the part to enhance the intelligence and capabilities of the human. Concerning the concept of IA, the human did not only to involve in the system but also it is important to highlight the human central role in solving complex problems. In summary, all the improvements are improvements which are most of the time based on the improvement of the whole system's performance instead of completing tasks, without studying the amplification effects of humans' intelligence. Humans play a critical role in the Intelligence Amplification system, so there is more research needed which will focus more on amplifying human intelligence instead of the system.

7.3. Current Situation

An analysis of the current situation of the slotting process at the Albert Heijn is made. This to find the current challenges, and as noted in chapter 3, the current challenges became clear by empirical research with the help of interviews, surveys, and site visits. The questions which related to the current process of slotting at the Albert Heijn is stated as follows:

SQ3: How do the current process and complementary processes of slotting work in the warehouse of Albert Heijn?

The high-over view of the slotting process is seen in Figure 45. An extensive description of the whole process of slotting and the complementary processes can be found chapter 3. In addition, the current overview of the architecture can be read in section 3.4.

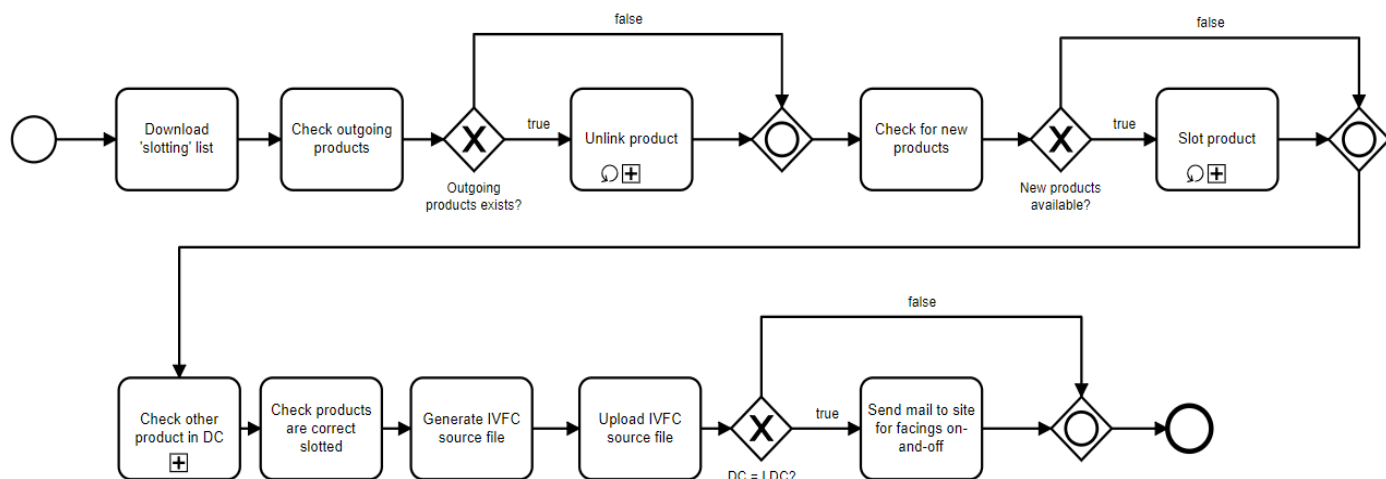


Figure 45: Slotting process

Because of the big scope of the process of slotting, there is an analyze made, which subprocess or tasks is the less optimized or most time-consuming process.

SQ4: What is the less optimized or most time-consuming process in the process of slotting?

Several conclusions can be made with the help of the process analyze of the current situation, surveys, and interviews. All these main conclusions can be found in chapter 3, but the two most important conclusions are presented below:

- The most time-consuming part is to check if every product in the warehouse has the right location.
- The step of data validation or errors will cost mentally the most effort. In addition, with a solution time of around 10 hours, this is one of the most time-consuming tasks of slotting.

The main research question is to enhance the workers performance and quality of work by enhancing the intelligence amplification. To find the part that benefits the most of being amended by the utilization of intelligence slotting, the followed question is designed:

SQ5: Which (sub-) process of slotting benefits most of being amended by the utilization of intelligence amplification?

To give more direction to the research, it was decided to focus on validating the product data and improving the product data quality. This because the validation process is an important task within the slotting process. When different products did not have the right data, the product cannot be allocated to a location in the distribution center.

With help of the IA framework designed by Dobrkovic (2016), the slotting process is split it up in different tasks. All those tasks are assigned to the human, machine, or both. All those tasks are categorized by different categories: Skill-based, Rule-based, Knowledge-based, and Expertise (Cummings, 2014). By this categorizing all these tasks, the conclusion can be made that the number of rule-based tasks is high so the need for a reference framework which can deal with those tasks. In chapter 4, the categorizing, as well as the assigning of the tasks can be read.

Data cleaning is the first step for automation. Data cleaning is a labor-intensive and complex process, is estimated that it will accounts for 30%-80% of the development time in a data warehouse project (Shilakes, 1998). The need for an application which validates and clean the data is present and will be the first goal.

7.4. Reference Architecture

To enhance the workers' performance and work quality in a different process, there is a need for a solution. The following question is stated to find the solution for improving the most time-consuming process improved by IA.

SQ6: How can we design a solution for the most time-consuming process of slotting improved by intelligence amplification?

From the literature study, there is no application found which can help the slotting process to improve the quality of work and the workers performance. Still to improve this part of the process, there is chosen to create a solution which focus on rule-based tasks. As already mentioned, those tasks are the most used tasks in the slotting process and need a solution to automate these tasks more. To find the possible solution to tackle the problem of validation, the need for creating a reference architecture is high. This reference architecture consists of two different phases; rule creation and validation, and the slotting process as defined in chapter 3.

As proven in chapter 6, the reference architecture can be used for every task which is rule-based. The only requirement is that should be a group of users, who can create, update, or even delete the rule. The users should be responsible for reaching the goal they have created in the first phase of the process. Otherwise, the reference architecture and the implementation of those will not work.

7.5. Prototype and validation

The solution what is chosen is a supportive tool. This tool needs to focus on the different rule-based tasks or processes in the main process. The different rule-based tasks or processes can be rewritten to one algorithm which supports the user to achieve its goal. To validate the prototype the following question is stated:

SQ7: How well does to chosen solution meet the original requirements?

The validation requirements focus on two questions:

- What is the enhancement of workers performance by this prototype?
- What is the enhancement of work quality by this prototype?

What is the enhancement of workers performance by this prototype?

With help of the validation tool, the employees can work faster and accurate. The validation tool validates product data to ensure there are no mistakes. When there are mistakes the employee can fast communicate the error to other departments. The workers performance will be enhanced with help of intelligence amplification because the machine does the repetitive task of validation. The human in case will update or create more rules to make the machine smarter to detect more errors in the future.

What is the enhancement of work quality of this prototype?

The tool validates a product every time according to a certain set of rules, you can say that the quality is improved because no product is skipped. The other important improvement in the quality of work is that the data is error free. The set of rules, created by the employees, is the equivalent of the measurement of the quality of this data set. When all rules are applied on the data set, one can say that the quality of the data of this product is good, and the product can be slotted in the next steps of the process. The quality of work will improve because you can ensure that there is always good data at your disposal.

7.6. Final conclusion

With the answers to all the research questions, it is now possible to answer the main research question:

“How to enhance the workers’ performance and quality of work in the slotting process through the adoption of Intelligence Amplification?”

Using our knowledge about the concept of Intelligence Amplification, we designed a reference architecture to handle rule-based tasks in the slotting process. In other words, a solution is created which focusses on a close collaboration between the human and machine to empower the human in the process of rule-based tasks, in this case, the slotting process. This reference architecture is the base for the prototype developed to prove that the reference architecture works. This prototype is the validation tool as described in this research.

Albert Heijn, especially the department of Logistics Support now uses the tool to support their daily tasks and this is an evident result of this research. However, the numerous meetings, conversation, interviews, and questionnaires held to understand the purpose of the need for validation and the need for a reference architecture for rule-based tasks. This not only to identify the features that can be added to the tool in further development but also the importance for triggering the employees of Albert Heijn to rethink their process and inspire them to improve their work processes. The results of the testing of the prototype give enough hints and insights to use it and implement and develop the

full version of the tool to enhance the workers' performance as well as the quality of work in the slotting process.

7.7. Contribution

The value of this research is created by the journey that is made by the various parties of this project. This journey results in a solution for the application of the concept of intelligence amplification on the slotting process. First, the current situation of the slotting process is analyzed, and a critical review is made to find the most time-consuming process of this process. The solution of the problem can be found in a reference architecture which do not focus only on the most time-consuming part but a solution for more tasks which contains the problem. This reference architecture is tested and validated with help of a prototype which focus on the validation part of the slotting process. All the steps of the research have been evaluated and examined, to find a generic solution for the problem.

This research contributes to the field of research by providing a reference architecture for rule-based tasks within the slotting process with help of the user. This reference architecture could support future research in the application of intelligence amplification in different sectors. Without the help of humans, machines will never be smarter. To ensure that control remains with the human, the machine remains as smart as we make it ourselves.

The contribution to practice lies in the explanation of the concept of IA and the prototype of the validation tool based on the reference architecture. In principle, many rule-based tasks could be automatized with help of the reference architecture. With the prototype we contributed to the knowledge of AH, by demonstrating how data validation can be done. This in combination with people who are not programmers by nature, but who can create rules that deliver results or advice.

7.8. Limitations

The reference architecture focusses only on automatization of the rule-based tasks of a process. As Cummings (2014) describes, there are more categories to focus on.

The developed prototype which was subsequently developed has several limitations. It only validates the plus-min list at this moment. To full automate the validation part of the slotting process more rules and files need to be added to the tool. This can be done to create different templates and validate over those templates and create rules based on this template.

As mentioned, the performance is one of the issues. When this is not on the level where it must be, the tool will lose its potential value. This can be mitigated to run the validation by night, so that the validation results are available in the next morning. The performance of the tool can become a problem when the employee wants to validate other lists than the standard and the results are not immediately shown but need to wait a certain amount of time on it. The other one is that the rules and templates need to be created to be sure every product is validated.

This research is done specifically within the context of Albert Heijn, especially the department of Logistics Support. First, no other departments of the Albert Heijn were asked to implement and validate the tool. Furthermore, no competitors of the retail market were asked to test the prototype, and this can be considered as a limitation to this research. When gathering the results of the evaluation it was difficult to get the responses of all the experts involved. Therefore, not all the evaluation results are held with the same number of experts and this makes the results more difficult to interpret.

7.9. Recommendations

The result of this research is a reference architecture supplemented with a prototype. The focus on a reference architecture based on rule-based tasks, and the limitations to the prototype provide recommendations for future research.

7.9.1. General

Based on the limitations we make several suggestions for future work. The prototype should be upgraded to the next level to be sure the potential value of this tool can be fully utilized. Lots of suggestions can be made and are made to improve the quality of the product. It can be small things like a button here or larger adjustments to make the look and feel better. All those suggestions need to be written down to upgrade the prototype and maybe the reference architecture.

Our reference architecture provides a starting set of guidelines for developing applications based on automation of rule-based tasks in the slotting process. When it is used for new applications this should result in feedback to update the document. What already is mentioned, the reference architecture can be upgraded with more categories which are described by Cummings (2014). These vary from simple automatic tasks to complex tasks.

Based on the expert sessions, the potential of this framework can be more utilized. Different other tasks and parts of rule-based processes can be used to validate the potential value of this framework. This research in this context, is not sufficient enough to conclude that it will work in other contexts, other than the current context at the department of Logistics in the Albert Heijn. The implementation and validation of the prototype should be tested in other environments, for instance at the competitors in the retail market.

Furthermore, the cooperation and collaboration between the human and machine can be more researched. For example, the slotting process, and maybe other processes, can be disassembled to small tasks, which can be automated on the manner of Limoncelli (2018). To automate and get rid of the repetitive and easy tasks of processes the human should be involved to document the steps. After that, pseudo code needs to be made to be sure every part of the process can be automated and is clear for automation. The last step is to create rules and validate all these rules to be sure the goal is reached. In conclusion, the human must focus on the more creative tasks of processes. This will be able as soon as more repetitive tasks will be done by the machine.

7.9.2. Ahold-Delhaize

For Ahold-Delhaize, especially the department of Logistics Support the recommendation is to create, update, even delete rules to be sure that everything is validated. Discuss and talk about which rules are needed. Furthermore, update the prototype with all the wishes which are needed to reach above goal. For example, the communication between the departments can be further automated to send automatic email to replenishment for standard problems. Or even better, bounce all the data which will not come to the validation tool. In addition, find a way to further automate the slotting process, it can be by building a new application on the reference architecture.

References

- Agrawal, R., Mannila, H., Srikant, R., Toivonen, H., & Verkamo, a I. (1996a). Fast discovery of association rules. *Advances in Knowledge Discovery and Data Mining*, 12, 307–328. Retrieved from <http://www.cs.helsinki.fi/hannu.toivonen/pubs/advances.pdf>
- Agrawal, R., Mannila, H., Srikant, R., Toivonen, H., & Verkamo, a I. (1996b). Fast discovery of association rules. *Advances in Knowledge Discovery and Data Mining*, 12, 307–328.
- Ahold-Delhaize. (2017). Annual Report. Retrieved February 20, 2019, from https://www.aholddelhaize.com/media/6530/2017_aholddelhaize-annual-report_interactive.pdf
- Ahold-Delhaize. (2019a). Ahold-Delhaize; Brands. Retrieved February 20, 2019, from <https://www.aholddelhaize.com/en/brands/netherlands/our-brands-in-the-netherlands/>
- Ahold-Delhaize. (2019b). Ahold-Delhaize; Where we operate. Retrieved February 20, 2019, from <https://www.aholddelhaize.com/en/about-us/company-overview/where-we-operate/>
- Ahold-Delhaize. (2019c). Ahold Delhaize; Who we are. Retrieved February 20, 2019, from <https://www.aholddelhaize.com/en/about-us/company-overview/who-we-are/>
- Angelov, S., Grefen, P., & Greefhorst, D. (2012). A framework for analysis and design of software reference architectures. *Information and Software Technology*, 54(4), 417–431.
- Annett, J. (2003). Hierarchical task analysis. In *Handbook of cognitive task design* (pp. 41–60). CRC Press.
- Ashby, W. R. (1956). *An introduction to cybernetics*. Chapman & Hail Ltd., London.
- Barca, J. C., & Li, R. K. (2006). Augmenting the human entity through man/machine collaboration. In *2006 IEEE International Conference on Computational Cybernetics, ICC*. <https://doi.org/10.1109/ICCCYB.2006.305689>
- Barzi, F., & Woodward, M. (2004). Imputations of missing values in practice: results from imputations of serum cholesterol in 28 cohort studies. *American Journal of Epidemiology*, 160 1, 34–45.
- Basili, V. R., Caldeira, G., & Rombach, H. D. (2006). The Goal Question Metric Approach. 5th ACM-IEEE International Symposium on Empirical Software Engineering (ISESE'06). *Rio de Janeiro*.
- Baskar, S., Subbaraj, P., & Rao, M. V. C. (2003). Hybrid real coded genetic algorithm solution to economic dispatch problem. *Computers and Electrical Engineering*, 29(3), 407–419. [https://doi.org/10.1016/S0045-7906\(01\)00039-8](https://doi.org/10.1016/S0045-7906(01)00039-8)
- Bechar, A., & Edan, Y. (2003). Human-robot collaboration for improved target recognition of agricultural robots. *Industrial Robot*, 30(5), 432–436. <https://doi.org/10.1108/01439910310492194>
- Blanchard, D. (2007). *Supply chain management : best practices*. *Supply Chain Management* (Vol. 45). John Wiley & Sons.
- Bley, H., Reinhart, G., Seliger, G., Bernardi, M., & Korne, T. (2004). Appropriate human involvement in assembly and disassembly. *CIRP Annals - Manufacturing Technology*, 53(2), 487–509. [https://doi.org/10.1016/S0007-8506\(07\)60026-2](https://doi.org/10.1016/S0007-8506(07)60026-2)
- Bradshaw, J. M., Feltoovich, P. J., & Johnson, M. (2011). *Human-agent interaction*. *The Handbook of Human-Machine Interaction: A Human-Centered Design Approach*. Retrieved from <https://www.scopus.com/inward/record.uri?eid=2-s2.0-84900924030&partnerID=40&md5=f422de5f67d1ba27c62bce751d837710>
- Burstein, M. H., Muivehill, A. M., & Deutsch, S. (1999). An approach to mixed-initiative management of heterogeneous software agent teams. In *Proceedings of the 32nd Annual Hawaii International Conference on Systems Sciences*. 1999. HICSS-32. Abstracts and CD-ROM of Full Papers (Vol. Track8, p. 10 pp.-). <https://doi.org/10.1109/HICSS.1999.773087>

- Cámara, J., Moreno, G., & Garlan, D. (2015). Reasoning about Human Participation in Self-Adaptive Systems. In *Proceedings - 10th International Symposium on Software Engineering for Adaptive and Self-Managing Systems, SEAMS 2015* (pp. 146–156). <https://doi.org/10.1109/SEAMS.2015.14>
- Casini, E., Suri, N., & Bradshaw, J. M. (2015). Leveraging human oversight and intervention in large-scale parallel processing of open-source data. In *Proceedings of SPIE - The International Society for Optical Engineering* (Vol. 9499). <https://doi.org/10.1117/12.2177264>
- Castle, N. (2018). What is Semi-Supervised Learning? *Oracle DataScience.Com*. Retrieved from <https://www.datascience.com/blog/what-is-semi-supervised-learning>
- Ciervo, B. (2018, April). How to Create an Effective Warehouse Slotting Process. *Conveyco*. Retrieved from <https://www.conveyco.com/create-effective-warehouse-slottting-process/>
- Cloutier, R., Muller, G., Verma, D., Nilchiani, R., Hole, E., & Bone, M. (2010). The concept of reference architectures. *Systems Engineering*, 13(1), 14–27.
- Coetzer, J., Swanepoel, J., & Sabourin, R. (2012). Efficient cost-sensitive human-machine collaboration for offline signature verification. In *Proceedings of SPIE - The International Society for Optical Engineering* (Vol. 8297, p. 82970J). <https://doi.org/10.1117/12.910460>
- Cormen, T. H., Leiserson, C. E., Rivest, R. L., & Stein, C. (2009). *Introduction to algorithms*. MIT press.
- Cummings, M. M. (2014). Man versus machine or man + machine? *IEEE Intelligent Systems*, 29(5), 62–69. <https://doi.org/10.1109/MIS.2014.87>
- de Greef, T., van Dongen, K., Grootjen, M., & Lindenberg, J. (2007). Augmenting cognition: Reviewing the symbiotic relation between man and machine. *Augmented Cognition*, 4565 LNAI, 439–448. https://doi.org/10.1007/978-3-540-73216-7_51
- de Groot, A. D. (1969). *Methodologie* (Vol. 6). Hague: Mouton.
- DeAngelis, S. (2018). Is Technology Killing Supply Chain Management? *Business at the Base of the Pyramid*. Retrieved from <https://www.enterrasolutions.com/blog/is-technology-killing-supply-chain-management/>
- den Berg, J. P. va., & Zijm, W. H. M. (1999). Models for warehouse management: Classification and examples. *International Journal of Production Economics*, 59(1), 519–528. [https://doi.org/https://doi.org/10.1016/S0925-5273\(98\)00114-5](https://doi.org/https://doi.org/10.1016/S0925-5273(98)00114-5)
- DiBona, P., Shilliday, A., & Barry, K. (2016). Proactive human-computer collaboration for information discovery. In *Proceedings of SPIE - The International Society for Optical Engineering* (Vol. 9851). <https://doi.org/10.1117/12.2222805>
- Dobrkovic, A., Liu, L., Iacob, M. E., & Van Hillegersberg, J. (2016). Intelligence amplification framework for enhancing scheduling processes. *Lecture Notes in Computer Science (Including Subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)*, 10022 LNAI, 89–100. https://doi.org/10.1007/978-3-319-47955-2_8
- Engelbart, D. C. (1995). Toward augmenting the human intellect and boosting our collective IQ. *Communications of the ACM*, 38(8), 30–32. <https://doi.org/10.1145/208344.208352>
- Ertel, W. (2017). *Introduction to Artificial Intelligence*. Springer. <https://doi.org/10.1007/978-3-319-58487-4>
- Farooq, U., & Grudin, J. (2016). Human-computer integration. *Interactions*, 23(6), 26–32.
- Fenton, N., Pfleeger, S. L., & Glass, R. L. (1994). Science and substance: A challenge to software engineers. *IEEE Software*, 11(4), 86–95.
- Ferreira, P., Doltsinis, S., & Lohse, N. (2014). Symbiotic assembly systems - A new paradigm. In *Procedia CIRP* (Vol. 17, pp. 26–31). <https://doi.org/10.1016/j.procir.2014.01.066>
- Feshbach, D. (1979). What's inside the black box: A case study of allocative politics in the Hill-Burton Program. *International Journal of Health Services*, 9(2), 313–339.

- Frazelle, E. (2001). *World-Class Warehousing and Material*. New York, NY, (1st ed), New York: McGraw-Hill.
- Freeman, J. A., & Skapura, D. M. (1991). Algorithms, applications, and programming techniques. In *Neural Networks*. Citeseer.
- Fumi, A., Scarabotti, L., & Schiraldi, M. M. (2013). Minimizing warehouse space with a dedicated storage policy. *International Journal of Engineering Business Management*, 5(1), 1–8. <https://doi.org/10.5772/56756>
- Goldberg, D. E. (1989). Genetic Algorithm in Search, Optimization, and Machine Learning. *Journal of Applied Sciences*, 16(12), 412 pp.
- Griffith, D., & Greitzer, F. L. (2007). *Neo-Symbiosis. Novel Approaches in Cognitive Informatics and Natural Intelligence* (Vol. 1). <https://doi.org/10.4018/978-1-60566-170-4.ch007>
- Gu, J., Goetschalckx, M., & McGinnis, L. F. (2007). Research on warehouse operation: A comprehensive review. *European Journal of Operational Research*, 177(1), 1–21. <https://doi.org/https://doi.org/10.1016/j.ejor.2006.02.025>
- Harmatuck, D. J. (1976). A comparison of two approaches to stock location. *The Logistics and Transportation Review*, 12(4), 282–284.
- Hausman, W. H., Schwarz, L. B., & Graves, S. C. (1976). Optimal storage assignment in automatic warehousing systems. *Management Science*, 22(6), 629–638.
- Hu, Y., Zhang, S., Jia, X., & Chen, J. (2017). Slotting optimization of warehousing system based on the Hungarian method. *Advances in Intelligent Systems and Computing*, 502, 143–155. https://doi.org/10.1007/978-981-10-1837-4_13
- Iacob, M. E., Jonkers, D., Quartel, H., Franken, H., van den Berg, H., & van den Berg, H. (2012). *Delivering Enterprise Architecture with TOGAF® and ARCHIMATE®*. BIZZdesign.
- Irby, B. T., Sue, C. A., Smith, K. M., & Alwan, M. (2016). Maximizing the Shadowing Experience: A Guidance Document. *Hospital Pharmacy*, 51(1), 54–59. <https://doi.org/10.1310/hpj5101-54>
- Iyer, A. M., & Jayal, R. (2016). *Intelligent slotting for the warehouse. Artificial Intelligence: Concepts, Methodologies, Tools, and Applications* (Vol. 4). <https://doi.org/10.4018/978-1-5225-1759-7.ch112>
- Iyer, N. P., & Jayal, R. (2016). *Intelligent Slotting for the Warehouse. Artificial Intelligence*. <https://doi.org/10.4018/978-1-5225-1759-7.ch112>
- Jacucci, G., Spagnoli, A., Freeman, J., & Gamberini, L. (2014). Symbiotic interaction: A critical definition and comparison to other human-computer paradigms. *Lecture Notes in Computer Science (Including Subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)*, 8820, 3–20. https://doi.org/10.1007/978-3-319-13500-7_1
- Jane, C. C. Storage location assignment in a distribution center, 30 *International Journal of Physical Distribution and Logistics Management* 55–71 (2000). <https://doi.org/10.1108/09600030010307984>
- Jones, E., & Battieste, T. (2004). Warehouse picking productivity increases by slotting inventory in the golden zone. In *IIE Annual Conference and Exhibition 2004* (p. 1541). Retrieved from <https://www.scopus.com/inward/record.uri?eid=2-s2.0-30044444185&partnerID=40&md5=36dadf28d9761996fe36e622bb75f851>
- Kallina, C., & Lynn, J. (1976). Application of the Cube-Per-Order Index Rule for Stock Location in a Distribution Warehouse. *Interfaces*, 7(1), 37–46. <https://doi.org/10.1287/inte.7.1.37>
- Kärkkäinen, M., Holmström, J., Främling, K., & Artto, K. (2003). Intelligent products—a step towards a more effective project delivery chain. *Computers in Industry*, 50(2), 141–151. [https://doi.org/10.1016/S0166-3615\(02\)00116-1](https://doi.org/10.1016/S0166-3615(02)00116-1)
- Keller, S. B., Keller, B. C., & Council of Supply Chain Management Professionals. (2014). *The Definitive Guide to Warehousing: Managing the Storage and Handling of Materials and Products in the Supply Chain*. Pearson Education.

- Kitchenham, B. A., & Brereton, P. O. (2013). A systematic review of systematic review process research in software engineering. *Information and Software Technology*, 55(12), 2049–2075. <https://doi.org/https://doi.org/10.1016/j.infsof.2013.07.010>
- Lange, D. S., Gutzwiller, R. S., Verbancsics, P., & Sin, T. (2014). Task models for human-computer collaboration in supervisory control of teams of autonomous systems. In *2014 IEEE International Inter-Disciplinary Conference on Cognitive Methods in Situation Awareness and Decision Support, CogSIMA 2014* (pp. 97–102). <https://doi.org/10.1109/CogSIMA.2014.6816547>
- Li, M., Chen, X., & Liu, C. (2008). Pareto and niche genetic algorithm for storage location assignment optimization problem. In *3rd International Conference on Innovative Computing Information and Control, ICICIC'08*. <https://doi.org/10.1109/ICICIC.2008.655>
- Licklider, J. C. R. (1960). Man-Computer Symbiosis. *IRE Transactions on Human Factors in Electronics*, HFE-1(1), 4–11. <https://doi.org/10.1109/THFE2.1960.4503259>
- Limoncelli, T. A. (2018). Documentation is Automation. *Commun. ACM*, 61(6), 48–53. <https://doi.org/10.1145/3190572>
- Lyll, A., Mercier, P., & Gstettner, S. (2018). The Death of Supply Chain Management. Retrieved from <https://hbr.org/2018/06/the-death-of-supply-chain-management>
- Mankins, J. C. (1995). Technology readiness levels. *White Paper*, April, 6, 1995.
- McKinsey. (2015a). An Executive Guide to Machine Learning. *McKinsey & Company*. Retrieved from <https://www.mckinsey.com/industries/high-tech/our-insights/an-executives-guide-to-machine-learning>
- McKinsey. (2015b). An Executive Guide to Machine Learning. *McKinsey & Company*.
- Microsoft. (2009). End of life Windows 7. Retrieved from <https://support.microsoft.com/nl-nl/help/4057281/windows-7-support-will-end-on-january-14-2020>
- Muppani (Muppant), V. R., & Adil, G. K. (2008). A branch and bound algorithm for class based storage location assignment. *European Journal of Operational Research*, 189(2), 492–507. <https://doi.org/10.1016/j.ejor.2007.05.050>
- Nakagawa, E. Y., Feitosa, D., & Felizardo, K. R. (2010). Using systematic mapping to explore software architecture knowledge. In *Proceedings of the 2010 ICSE Workshop on Sharing and Reusing Architectural Knowledge - SHARK '10* (pp. 29–36). New York, NY, USA: ACM. <https://doi.org/10.1145/1833335.1833340>
- Napolitano, M. (1998). *Using modeling to solve warehousing problems: a collection of decision-making tools for warehouse planning and design*. Warehousing Education and Research Council.
- Padoy, N., & Hager, G. D. (2011). Human-machine collaborative surgery using learned models. In *Proceedings - IEEE International Conference on Robotics and Automation* (pp. 5285–5292). <https://doi.org/10.1109/ICRA.2011.5980250>
- Palinkas, L. A. (2014). Qualitative and mixed methods in mental health services and implementation research. *Journal of Clinical Child and Adolescent Psychology : The Official Journal for the Society of Clinical Child and Adolescent Psychology, American Psychological Association, Division 53*, 43(6), 851–861. <https://doi.org/10.1080/15374416.2014.910791>
- Pang, K.-W., & Chan, H.-L. (2017). Data mining-based algorithm for storage location assignment in a randomised warehouse. *International Journal of Production Research*, 55(14), 4035–4052. <https://doi.org/10.1080/00207543.2016.1244615>
- Pazour, J. A., & Carlo, H. J. (2015). Warehouse reshuffling: Insights and optimization. *Transportation Research Part E: Logistics and Transportation Review*, 73, 207–226. <https://doi.org/10.1016/j.tre.2014.11.002>
- Peffer, K., Tuunanen, T., Rothenberger, M. A., & Chatterjee, S. (2007). A Design Science Research Methodology for Information Systems Research. *Journal of Management Information Systems*, 24(3), 45–77. <https://doi.org/10.2753/MIS0742-1222240302>

- Petersen, C. G., Siu, C., & Heiser, D. R. (2005). Improving order picking performance utilizing slotting and golden zone storage. *International Journal of Operations and Production Management*, 25(10), 997–1012. <https://doi.org/10.1108/01443570510619491>
- Petersen, C. G., & Thompson, K. R. (2003). A comparison of warehouse slotting measures. In *Proceedings - Annual Meeting of the Decision Sciences Institute* (pp. 2183–2188). Retrieved from <https://www.scopus.com/inward/record.uri?eid=2-s2.0-0442325453&partnerID=40&md5=601e3720c19c256ea930930fbc32f1f9>
- Pettey, C. (2018). Gartner Top 8 Supply Chain Technology Trends for 2018. *Gartner*. Retrieved from <https://www.gartner.com/smarterwithgartner/gartner-top-8-supply-chain-technology-trends-for-2018/>
- Plano, S., & Blasch, E. P. (2003). User performance improvement via multimodal interface fusion augmentation. *Fusion03*.
- Pujawan, I. N., Vanany, I., & Dewa, P. K. (2017). Human errors in warehouse operations: an improvement model. *International Journal of Logistics Systems and Management*, 27(3), 298. <https://doi.org/10.1504/IJLSM.2017.10005117>
- Rahm, E., & Do, H. H. (2000). Data cleaning: Problems and current approaches. *IEEE Data Eng. Bull.*, 23(4), 3–13.
- Raman, V., & Hellerstein, J. M. (2001). Potter's wheel: An interactive data cleaning system (Conference). In *The 27th International Conference on Very Large Data Bases (11 - 14th Sept, 2001)* (Vol. 1, pp. 381–390).
- Ray, S. (2017). Essentials of Machine Learning Algorithms. Retrieved from <https://www.analyticsvidhya.com/blog/2017/09/common-machine-learning-algorithms/>
- Ren, X. (2016). Rethinking the Relationship between Humans and Computers. *Computer*, 49(8), 104–108. <https://doi.org/10.1109/MC.2016.253>
- Richards, G. (2011). Warehouse management. *The Chartered Institute of Logistics and Transport (UK) and Kogan Page, London*.
- Roy, D. (2004). 10x - Human-Machine Symbiosis. *BT Technology Journal*, 22(4), 121–124. <https://doi.org/10.1023/B:BTTJ.0000047590.43185.1b>
- Sankar, S. (2012). The rise of human-computer cooperation. Retrieved September 9, 2018, from http://www.ted.com/talks/shyam_sankar_the_rise_of_human_computer_cooperation
- Saunders, M., Lewis, P., & Thornhill, A. (2011). Research Methods for Business Students, London, Financial Times Prentice Hall. In *RICS Construction and Property Conference* (p. 1191).
- Schwaber, K., & Sutherland, J. (n.d.). Scrum Guide. Retrieved February 21, 2019, from <https://www.scrum.org/resources/scrum-guide>
- Sheridan, T. B., & Verplank, W. L. (1978). *Human and Computer Control of Undersea Teleoperators. ManMachine Systems Lab Department of Mechanical Engineering MIT Grant N0001477C0256*. <https://doi.org/10.1080/02724634.1993.10011505>
- Shilakes, C. C. (1998). Enterprise information portals. *Merrill Lynch In-Depth Report*.
- Smith, L. C. (1983). Machine intelligence VS. Machine-aided intelligence in information retrieval: A historical perspective. *Lecture Notes in Computer Science (Including Subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)*, 146 LNCS, 263–274. <https://doi.org/10.1007/BFb0036351>
- Sorrentino, S., Bergamaschi, S., Gawinecki, M., & Po, L. (2010). Schema label normalization for improving schema matching. *Data & Knowledge Engineering*, 69(12), 1254–1273. <https://doi.org/https://doi.org/10.1016/j.datak.2010.10.004>
- Tamm, T., Seddon, P. B., Shanks, G., & Reynolds, P. (2011). How does enterprises architecture add value to organisations? *Communications of the Association for Information Systems*, 28(1), 141–168. <https://doi.org/10.1287/isre.3.1.60>

- Tan, J. T. C., Duan, F., Zhang, Y., Kato, R., & Arai, T. (2009). Task modeling approach to enhance man-machine collaboration in cell production. In *Proceedings - IEEE International Conference on Robotics and Automation* (pp. 152–157). <https://doi.org/10.1109/ROBOT.2009.5152658>
- te Pas, H. (2018). IRI: Jumbo wint, AH verliest marktaandeel. Retrieved February 20, 2019, from <https://www.distrifood.nl/formules/nieuws/2018/01/iri-jumbo-wint-ah-verliest-marktaandeel-101115146>
- The-open-group. (2016). ARCHIMATE® 3.0.1. Retrieved February 21, 2019, from <https://publications.opengroup.org/standards/archimate/c179>
- Tompkins, J. A. (1998). *The Challenge of Warehousing*.
- Tsamis, N., Giannikas, V., McFarlane, D., Lu, W., & Strachan, J. (2015). Adaptive storage location assignment for warehouses using intelligent products. *Studies in Computational Intelligence*, 594, 271–279. https://doi.org/10.1007/978-3-319-15159-5_25
- Vagia, M., Transeth, A. A., & Fjerdings, S. A. (2016). A literature review on the levels of automation during the years. What are the different taxonomies that have been proposed? *Applied Ergonomics*, 53, 190–202. <https://doi.org/10.1016/j.apergo.2015.09.013>
- van den Broecke, P. (2018). Rol groothandel verandert. Retrieved October 7, 2018, from <http://magazines.evofenedex.nl/data-en-digitalisering/08-rol-groothandel-verandert/>
- Wang, Z., Lin, J., Liu, C., & others. (1996). Order picking optimization for a multicarousel-single-server system solved by improved annealing procedure. *Controls and Decisions*, 11(7), 182–187.
- Webster, J., & Watson, R. (2002). ANALYZING THE PAST TO PREPARE FOR THE FUTURE: WRITING A LITERATURE REVIEW Analyzing the past to prepare for the. *Management Information Systems Quarterly*, 26(2), 3. <https://doi.org/10.1.1.104.6570>
- Wigness, M., & Gutstein, S. (2018). On the impacts of noise from group-based label collection for visual classification. In *Proceedings - 16th IEEE International Conference on Machine Learning and Applications, ICMLA 2017* (Vol. 2018–Janua, pp. 84–91). <https://doi.org/10.1109/ICMLA.2017.0-174>
- Wilson, B., Bounds, M., McFadden, D., Regenbrecht, J., Ohenhen, L., Tavakkoli, A., & Loffredo, D. (2018). VETO: An Immersive Virtual Environment for Tele-Operation. *Robotics*, 7(2), 26. <https://doi.org/10.3390/robotics7020026>
- Wolfswinkel, J. F., Furtmueller, E., & Wilderom, C. P. M. (2013). Using grounded theory as a method for rigorously reviewing literature. *European Journal of Information Systems*, 22(1), 45–55. <https://doi.org/10.1057/ejis.2011.51>
- Wong, C. Y., McFarlane, D., Ahmad Zaharudin, A., & Agarwal, V. (2002). The intelligent product driven supply chain. In *IEEE International Conference on Systems, Man and Cybernetics* (Vol. vol.4, p. 6). <https://doi.org/10.1109/ICSMC.2002.1173319>
- Wooldridge, M. (2002). *An Introduction to MultiAgent Systems* (1st ed.). John Wiley & Sons. Retrieved from <http://www.amazon.com/exec/obidos/redirect?tag=citeulike07-20&path=ASIN/047149691X>
- Xia, C., & Maes, P. (2013). The design of artifacts for augmenting intellect. In *Proceedings of the 4th Augmented Human International Conference on - AH '13* (pp. 154–161). <https://doi.org/10.1145/2459236.2459263>
- Xu, J., Lim, A., Shen, C., & Li, H. (2008). A heuristic method for online warehouse storage assignment problem. In *Proceedings of 2008 IEEE International Conference on Service Operations and Logistics, and Informatics, IEEE/SOLI 2008* (Vol. 2, pp. 1897–1902). <https://doi.org/10.1109/SOLI.2008.4682840>
- Yamaba, H., Matsumoto, S., & Tomita, S. (2008). An attempt to obtain scheduling rules of Network-Based Support System for Decentralized Scheduling of Distributed Production Systems. In *IEEE International Conference on Industrial Informatics (INDIN)* (pp. 506–511). <https://doi.org/10.1109/INDIN.2008.4618153>
- Yoon, A. H. (2016). Focus feature: Artificial intelligence, big data, and the future of law. *University of Toronto Law Journal*, 66(4), 456–471. <https://doi.org/10.3138/UTLJ.4007>

- Zhang, J., Yin, Z., & Wang, R. (2015). Recognition of mental workload levels under complex human-machine collaboration by using physiological features and adaptive support vector machines. *IEEE Transactions on Human-Machine Systems*, 45(2), 200–214. <https://doi.org/10.1109/THMS.2014.2366914>
- Zhao, Z., & Ma, Z. (2017). A Methodology for Measuring Structure Similarity of Fuzzy XML Documents. *Computing*, 99(5), 493–506. <https://doi.org/10.1007/s00607-017-0553-x>
- Zitzler, E., & Thiele, L. (1999). Multiobjective evolutionary algorithms: a comparative case study and the strength Pareto approach. *IEEE Transactions on Evolutionary Computation*, 3(4), 257–271. <https://doi.org/10.1109/4235.797969>
- Zou, S., & Zou, X. (2017). Understanding: How to resolve ambiguity. *IFIP Advances in Information and Communication Technology*, 510, 333–343. https://doi.org/10.1007/978-3-319-68121-4_36
- Zu, Q., Cao, M., Guo, F., & Mu, Y. (2011). Slotting optimization of warehouse based on hybrid genetic algorithm. In *Proceedings - 2011 6th International Conference on Pervasive Computing and Applications, ICPCA 2011* (pp. 19–21). <https://doi.org/10.1109/ICPCA.2011.6106472>

Appendices

A – Taxonomy proposed by Sheridan and Verplank (1978)

Level of autonomy	Description	Explanation
Level 1	Fully manual control	The computer offers no assistance.
Level 2	The computer offers a complete set of decision/action alternatives	Several options are provided to the human who decides.
Level 3	The computer narrows the selection down to a few	Human still has to decide.
Level 4	The computer suggests one alternative	Human decides amongst suggestions.
Level 5	The computer executes that suggestion if the human approves	Human approval needed for execution.
Level 6	The computer allows the human a restricted time to veto before automatic execution	Limited time for veto given to the human.
Level 7	The computer executes automatically, then necessarily informs the human	No human interference, just information at the end.
Level 8	The computer informs the human only if asked	Human gets information only if asks.
Level 9	The computer informs the human only if it decides to	Computer decides whether to give information.
Level 10	Fully autonomous Control	The computer decides everything and acts autonomously, ignoring the human.

Table 15: Taxonomy proposed by Sheridan and Verplank (1978)

B - Survey 'Slotting'

Slotting

Deze vragenlijst is om een indicatie te krijgen hoe lang taken duren en hoe v...

Welke dc(s)? *

Geef aan welke dc(s) je hebt geïnst...

- ☐ DCO
- ☐ DCP
- ☐ DCT
- ☐ DCZ
- ☐ LDC

Welke stroom? *

Welke stroom heb je geïnst...

- ☐ Displays
- ☐ Houdbaar
- ☐ Vers

Tijdsindicatie *

Graag in onderstaand rooster aangeven met welke taak je het langst bezig bent geweest. Dit geef je aan met een 1 en het kortst met een 12. Wanneer niet alle taken zijn uitgevoerd, graag vanaf onderaan prioriteit nummer 12 naar boven werken per taak die niet uitgevoerd is.

	Actie opdrachten	Actie werkzaamheden	Acties afsluiten	Beoordelen korting -	Berekenen kortingstijl	Fouten	IPC	Maken niet Bereken...	Niet aangekondigde artikelen	Nieuwe artikelen	Stapelen = met voorb...	Sorteren Items	N.V.T
1	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐
2	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐
3	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐
4	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐
5	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐
6	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐
7	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐
8	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐
9	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐
10	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐
11	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐
12	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐

Processtappen

In welke volgorde is het proces verlopen?

	Communicatie	Combineren gegevens	Nieuwe producten	Producten afsluiten	Verpakken	Vervoeren
1	☐	☐	☐	☐	☐	☐
2	☐	☐	☐	☐	☐	☐
3	☐	☐	☐	☐	☐	☐
4	☐	☐	☐	☐	☐	☐
5	☐	☐	☐	☐	☐	☐
6	☐	☐	☐	☐	☐	☐

Waarschuwing: Deze vraag is niet verplicht, als je geen antwoord geeft op alle rijen tellt het antwoord niet mee

Extra werkzaamheden

Wanneer er extra werkzaamheden zijn uitgevoerd hieronder graag invullen welke.

Opbouwfouten

Hoeveral opbouwfouten zijn er voorgekomen?

Welke opbouwfouten zijn opgetreden?

Vink hieronder aan welke het waren

- ☐ Afmetingen
- ☐ Gewicht
- ☐ Uitpakindicate
- ☐ Verkeerde logistieke groep
- ☐ Verkeerde product categorie
- ☐ Verkeerde temperatuurzone

Niet aangekondigde artikelen

Hoeveral niet aangekondigde artikelen zijn er voorgekomen?

Artikelen zonder volledige data

Hoeveral artikelen zonder volledige data zijn er voorgekomen?

Product niet bekend in NASA?

Hoeveral producten waren niet bekend in NASA?

Product met SDF van 0 met locatie

Hoeveral producten met een SDF van 0 hadden wel een locatie?

Verstuur

Figure 46: Survey – Slotting

C - Explanation of each slotting task

Judge slotting – general

The judgment of the slotting is a task where every product in the distribution is checked if it has the right number of location as well as the right location. For example, when a product has more than 50% fat than the product needs a location with a sprinkler.

New articles

This task must deal with the assignment of new products to a location in the dc.

Keep track of the slotting list

The employee with this task is responsible for the status of the slotting list. The employee updates the list with new products and products which will be removed in the dc. In addition, in this task the errors in the data are recognized.

Slotting issues

This task must deal with all the issues which arise in the dc. Those issues can be sent by different employees in the dc and via different mediums like PrepChat, google drive, or phone.

Errors

The task to find a solution for every different error which occurs in the slotting process as well in the data which is needed to slot the dc.

Site visit + preparation

The physical visit of a distribution center to check if there are some errors in this dc. Furthermore, check if there are problems in the process of distribution to the shops.

Mail to replenishment

The time-consuming task to communicate all the errors and other problems to the department of replenishment.

Care for promotion

This is the slotting process for the promotions in the dc. Every dc has 'promotion aisles' where the slotting for those products can happen.

IVFC (interface to wms)

In this task all the communication and uploading of different documents to the wms through IVFC.

Not announced articles

The employee who is responsible for this task must deal with the not announced products. Those products are new products which are not announced (duh) and need to have a location in the dc. Those allocations are ad-hoc and not inline with the plan what is created in the start of the week.

Disconnect articles

All products with no SDF and no revenue can be disconnected by the slotter. When this is done, the location is free for another product.

Other activities

Everything else what a slotter needs to do to complete the slotting process.

D - Feedback

Employee 1:

Ik heb vanmiddag de validatie tool bekeken en heb handmatig wat regels kunnen toevoegen. Onder andere de regel aangemaakt "indien artikel = Breder dan Rolco" dat opbouw niet juist / artikel te breed is. Het aanmaken van de regels werkt erg prettig en is niet heel ingewikkeld om te doen. Ik zou als tip wel willen meegeven om hier en daar wat extra uitleg toe te voegen, dit alles om dingen duidelijker te maken.

Mooi om te zien dat we nu op voorhand fouten kunnen ontdekken en dit kunnen terugkoppelen aan de verantwoordelijke partijen.

Dit scheelt gedurende de week enorm veel mail en telefoonverkeer. Als voorbeeld kunnen we zeggen dat wanneer een opbouw niet compleet is dat het aanzienlijk veel tijd gaat besparen omdat deze fouten in principe rechtstreeks teruggekoppeld kunnen worden naar Replenishment. Dit scheelt ons behoorlijk veel uitzoek tijd als het niet alle tijd is. Dus ik zou het willen gokken op 7 uur, want niet alle fouten zullen direct terecht zijn en mail blijft toch nog nodig. Dit ook ter verduidelijking van de meldingen die wij eventueel automatisch gaan doorsturen.

Wanneer de tool in de nabije toekomst nog verder wordt uitgebreid zullen we hier zeker steeds meer plezier aan beleven.

Employee 2:

Voor mij wat punten, wellicht al wel of niet een keer genoemd:

- Algemene navigatie (schermen-mogelijkheden per scherm) is goed.
- Ik mis optie om alle goedgekeurde artikelen in een keer door te zetten. Wellicht nog met optie om dit per HB/vers/LDC te doen.
- Weergave van HB/vers/LDC in de artikel lijsten.
- Als ik naar een artikel zoek die niet in de lijst staat krijg ik een leeg scherm. Dat scherm blijft vervolgens leeg. Alle artikelen zijn nu uit de lijst verwijderd en ik krijg ze niet meer terug.
- Vanuit Rule categories zit nog geen knop terug naar startscherm&vorig scherm.
- Vanuit Locatie overzicht, Opbouw overzicht en Categorie zit alleen functie terug naar hoofdscherm, niet terug.
- Mogelijkheid tot toevoegen Categorie en gebruik op Templates is echt top.
- Verschillende regeltypen en daarbij komende mogelijkheden logisch en simpel.
- Toevoegen van templates en regels moet ik echt nog wat meer feeling bij krijgen door te mee te oefenen. Ga ik op ander moment nog oppakken.

Voor zover het nog niet doorgeschemerd is, ik ben zeker positief over het gebruiksgemak en de mogelijkheden. Zijn mijn inziens nog wel wat verbeteringen in mogelijk, neemt niet weg dat het er goed uit ziet.

De validatie tool heeft de potentie om verschillende doelen te bereiken. Top gedaan! Ik heb hem vandaag als eerste mogen gebruiken om de snelheid van dit proces nog eens een keer goed te doorlichten, maar het scheelt mij behoorlijk veel tijd aangezien ik niet elk product hoeft te bekijken. Ondanks dat deze lijst nu slechts 150 producten bevat kan ik wel zeggen dat het mij zeker een minuutje per product gaat schelen, maar ik denk wel dat de validatie 's

nachts moet gebeuren als er steeds meer regels aanwezig gaan zijn. Misschien kunnen we dit dan gelijk doorvoeren op de andere tools ;)

Wat ik mis is een export per categorie om gelijk meerdere foutieve producten terug te kunnen sturen naar Replenishment. Maar voor nu zie ik zeker de positieve en toegevoegde waarde om deze tool in de toekomst te gebruiken.

Employee 3:

Allereerst bedankt voor jouw bijdrage aan onze afdeling, ik heb je met veel plezier leren kennen. In Eindhoven staat het bier altijd koud! (geen Grolsch overigens)

Dank voor je uitleg vandaag, klinkt goed! Hierbij mijn bevindingen/input:

- ik denk dat we (de slotters) elkaar moeten dwingen je tool te gebruiken en beter te maken. Wij lopen anders het risico een mooie kans te missen.
- zoals regelmatig besproken, vind ik het belangrijk dat werk gedaan wordt door degene die er verantwoordelijk voor is. De echte winst bij data-kwaliteit zit volgens mij in het inbouwen van checks *voordat* het doorgestuurd wordt naar de volgende in de keten.
- Controle op transport type (conditioned/unconditioned) mis ik nog
- Controle op temperatuurzone mis ik nog
- Controle op repack-artiklen
- Op het hoofdscherm zou ik Plus Min lijst eerste/tweede keer ingelezen anders formuleren. Iets in de trant van versie x-2 en versie x-1.

Nogmaals, ik denk serieus dat we dit moeten gaan gebruiken om één van onze vervelende taken te kunnen automatiseren. Zoals je al hebt uitgelegd wordt de computer slimmer gemaakt door ons, maar moeten er goed over na denken welke controles we er in moeten gaan stoppen om dit werk daadwerkelijk niet meer te hoeven doen.

Alvast veel succes met afstuderen!

E - User Interface

Update or create rule

Wijziggen regel

Startpagina

Naam:

Gewicht kleiner of gelijk aan 23 kg

Melding:

Gewicht is groter dan 23 kg

Categorie:

Alle categorieën

Regel toevoegen

Regel Type: *

Getal

Gebruik template?

☐

Veld type:

-

Veld:

-

Filter:

-

Waarde:

0

Opslaan en nieuwe regel

Regel type	Veld	Filter	Waarde
Getal	WEIGHT	Groter of gelijk aan	23

Opslaan

Annuleren

[Wijzigen](#) [Verwijderen](#)

Figure 47: Screenshot 'create/update' rule

Regel overzicht

Startpagina

ZoekResetNieuwe Regel

Name	Melding	Is actief?
Logistieke groep	Logistieke groep is niet correct	✓ Actief Wijzigen Verwijderen
Pallet mag niet breder zijn dan 1.06 meter	Pallet is breder dan 1.06 meter	✓ Actief Wijzigen Verwijderen
Collonummer mag niet groter zijn dan 1 miljoen	Collonummer is groter dan 1 miljoen	✓ Actief Wijzigen Verwijderen
Pallet mag niet langer zijn dan 1.20 meter	Pallet is langer dan 1.20 meter	✓ Actief Wijzigen Verwijderen
Gewicht kleiner of gelijk aan 23 kg	Gewicht is groter dan 23 kg	✓ Actief Wijzigen Verwijderen
Pallet mag niet hoger zijn dan 2.15 meter	Pallet is hoger 2.15 meter	✓ Actief Wijzigen Verwijderen

6 records

Figure 48: Screenshot 'overview rules'

F - Subprocesses of the slotting process

Subprocess: 'unlink product'

Looking at the subprocess 'unlink product' (figure X), it starts with the check of the specific product, which will be out of assortment, has indeed no stored demand forecast (sdf). When it has also no stock anymore for that product the product will be unlinked from that location. When the location is unlinked it will be noted in the slotting list as unlinked. This is a repetitive task until all the products are unlinked.

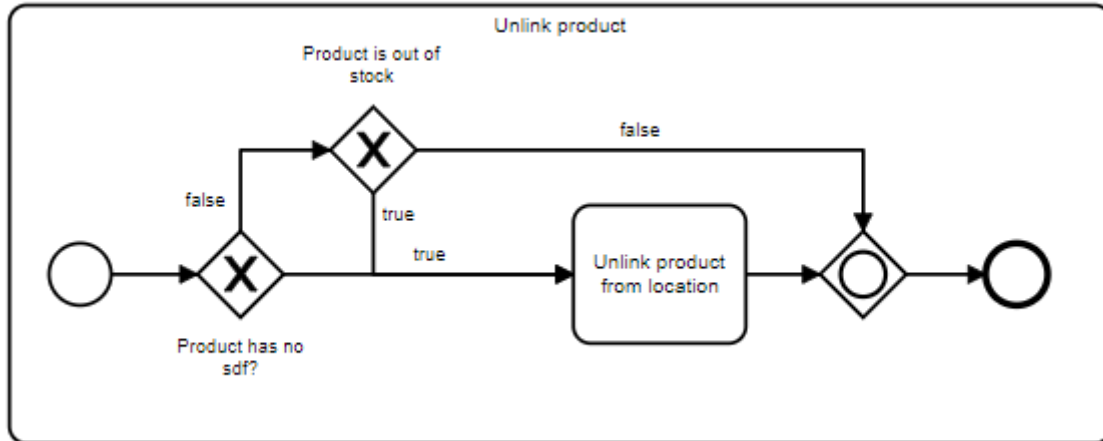


Figure 49: Subprocess: 'Unlink product'

FIGURE X: SUBPROCESS: 'UNLINK PRODUCT'

Subprocess 'Mail to Replenishment for the error'

When the 'slotter' indicates an error, the determination of the impact will be defined. When the impact is high, the slotter will have contact with another person in the AH. If the error contains something that is specific for that flow and not a general error, the flow manager will be contacted, and they will escalate or fix the issue themselves. Otherwise, all the other error will first pass by the front office of the department replenishment and they will escalate it further. The process for escalation is explained in the diagram 'swim lane model for the process'.

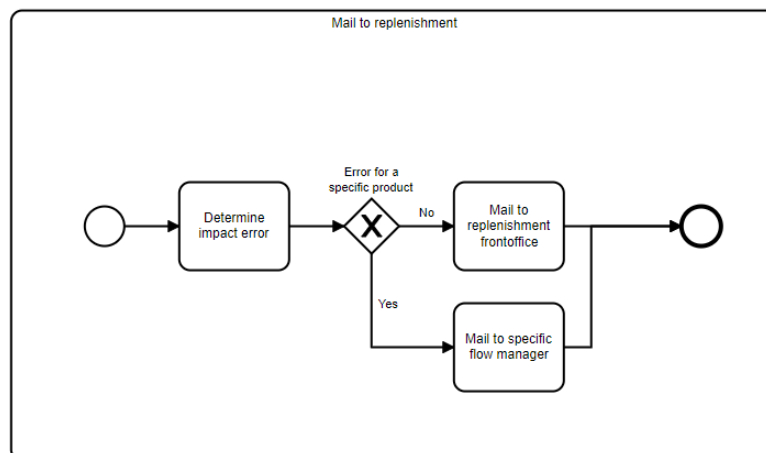


Figure 50: Subprocess: 'Mail to Replenishment for the error/ if the error'

Subprocess: 'check other product in dc'

In this subprocess, the products in the dc are checked. This will be done in the same manner as slotting a new product but most of the time the product is already correct slotted and can be ignored.

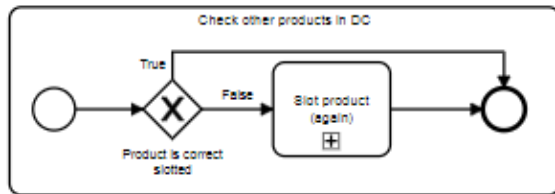


Figure 51: Subprocess: 'Check other product in DC'

G – List of errors occurring the slotting process

This list is prioritized based on the importance of the error:

- Product is not known in NASA
- SDF of zero, but there is a location, no slotting is needed
- Errors in the package structure of the product
- Sizes (height, length, width) are too big
- Weight above the 23 kilograms
- Unknown data (sizes, glas/not glas, flammable/ not flammable etc.)
- Errors with all the RAD display (Dolly → Rolly, Rolly → Dolly)
- Different start dates per product
- Product is incorrectly placed in the wrong system (LDC → RDC, RDC → LDC)
- Long-term products stay in the one-time products and vice versa
- Product category is wrong
- Product is in the wrong logistic group
- Errors in unpacking indication