

CYBERPARK Firma Değerlendirme Platform Tasarımı

Bilkent CYBERPARK

group01/grup_foto.jpeg

Proje Ekibi

Serhat Ciritci, Tuğçe Ecem Ergüven
Özkan Kul, Ersan Efe Semerci
Zeynep Selay Sertyel, Selin Ulupınar
Oğulcan Yılmaz

Şirket Danışmanı

Ferhan Z. ŞENKON
Kiralama ve İşletme Direktörü
İmran GÜRAKAN
İş Birliği ve Teknoloji Transferi Birimi
Müdürü

Akademik Danışman

Yrd. Doç. Dr. Firdevs Ulus
Endüstri Mühendisliği Bölümü

ÖZET

Mevcut durumda, CYBERPARK standardize edilmemiş bir yeni firma proje başvuru değerlendirme sistemi kullanmaktadır. Kararlar, hakem değerlendirme raporlarıyla birlikte karar vericilerin öngörü ve deneyimlerine dayanarak verilmektedir. Sistemin standart olan bir önceliklendirme ve skaler bir temele oturtulmuş bir değerlendirme yöntemi yoktur. Bu projenin ana hedefi matematiksel bir temele dayanan ve kullanıcı dostu bir çok amaçlı karar destek sistemi oluşturmaktır. Bu sistemin çıktıları; karar vericilerin tercihlerine göre oluşturulmuş başvuru skorları, bu skorlara göre sıralandırılmış bir bekleme listesi ve bekleme listesinin istatistiksel bilgileridir.

Anahtar Kelimeler: Çok Kriterli Karar Verme, AHP, DEMATEL, Karar Destek Sistemi, Başvuru Değerlendirme

Design of Company Evaluation Platform for CYBERPARK

1 Company Information

Founded in 2002 by Bilkent Holding and Bilkent University, CYBERPARK is one of the largest techno-parks in Turkey. The primary mission of CYBERPARK is to stimulate the establishment of new companies that are promising the ability to develop advanced technology. The company contains 240 R&D companies, 5 research centers, 1 micro nano-chip factory, approximately 4000 employees, and 115.000 m² indoor space. Bilkent CYBERPARK's aim is to help new and small enterprises to deal with difficulties in their early stages and to foster its mature companies' export capabilities. Together with high quality office space and facilities, providing value-added services for this aim is the powerful side of Bilkent CYBERPARK. Those services include info days, seminars, trade delegations, pre-incubation and incubation services, clustering activities, accelerators, hackathons and online entrepreneurship programs. (CYBERPARK, 2019).

2 System Analysis

Currently, there is a high level of demand from many companies for office rentals in CYBERPARK. Every time CYBERPARK receives an application, it should be immediately evaluated. The assessment process of an application starts by an initial online application form which is filled by the candidate firms. This form includes general information, firm's field of work (software, design or R&D) and firm's project proposal. The process continues as follows:

- Applications that satisfy the first stage conditions are proceeded to a judge or an appointed board of judges.
- Applications which pass the second stage successfully are invited to present their proposed project to CYBERPARK's Supervisory Board of Rental Department.
- Finally, if an application also fulfills the requirements of CYBERPARK's Supervisory Board of Rental Department and the company's demand for space and available amount of space in CYBERPARK are matched, it is directly taken under the organization of CYBERPARK. If not, it is put on a waiting list.

3 Problem Definition

The current assessment process that CYBERPARK uses for the acceptance of new applications for an office within the company is not standardized. The decisions on acceptance or rejection for companies which apply to cooperate with Bilkent CYBERPARK and prioritization of companies from the waiting list are conducted manually by the Rental Department. Furthermore, this process lacks a standardized prioritization of the criteria: the scoring, ranking and acceptance

of applications are not built on a scalar and mathematical basis. Instead, the directors' and judges' evaluations are based on their former experience and prudence. Since there are many criteria to be evaluated, the current system leads to complexity. It also creates unnecessary utilization of time and workforce. Certain parameters can be extracted from the current application forms, however, they are not prioritized in this current context of assessment.

4 Literature Review

There is a need for a system which embodies the priorities of decision makers during the assessment of company evaluations. To understand how the decision makers actually go along with the process and to convert their steps to a scalarized system, a literature review on multi-criteria decision making methodologies is conducted. In the decision support system design, the team has decided to use DEMATEL, AHP and hybrid methods.

Note that AHP is frequently used for decision-making when the field is analyzed. AHP method assigns weights to the criteria and critical inputs from pairwise comparisons and valuations made by the decision-makers. In this method, the pairwise valuations are made on a scale from one to nine, where one stands for equal importance and nine stands for absolute superiority (Pachemska et al., 2014). The even numbers stand for the compromise values between the priorities. After these pairwise comparisons, a matrix is obtained by the ranked pairs (Varajão and Cruz-Cunha, 2013; Saaty and Hall, 1999). This matrix demonstrates the direct and indirect relationships between different criteria and critical inputs. Column totals are calculated for all criteria. Then, each pairings' ranking is divided into the column totals in order to normalize the matrix. The averages of the normalized values in each row are taken in order to obtain the eigenvalues. Ömürbek and Tunca stated that the eigenvalues represent the weights of the critical inputs or criteria (Ömürbek and Tunca, 2013).

The criteria and critical inputs that have more interaction with the others are expected to admit higher weights according to DEMATEL methodology (Falatoonitoosi et al., 2013). Also, DEMATEL is a method for achieving unpredictable relations among criteria or critical inputs. It means that the relations which could not be realized or stated by decision-makers could be detected with the DEMATEL method after calculating the total relation matrix. The initial step for utilizing the DEMATEL method is to ask decision-makers to determine values from zero to four for each dual comparison among criteria, as well as the critical inputs within each criterion. These comparisons represent the causal relationships among them, where zero corresponds to no effect and four corresponds to very strong effect.

According to Saaty, in order to achieve confidential results, the number of criteria and critical inputs should be less than nine. In addition to this, for each criterion consistency ratio should be calculated. Consistency ratio shows whether the de-

cision makers are consistent on the decisions. While finding the consistency ratio, consistency index and consistency ratio should be calculated for each reciprocal comparison matrix. After obtaining the eigenvalues, the non-normalized matrix is multiplied by the eigenvector. The resulting values are divided by the corresponding eigenvalues. The average of the end products of these divisions gives λ_{max} . Corresponding CR formulas are provided below:

$$\text{Consistency Index} = \frac{(\lambda_{max}-n)}{(n-1)}$$

$$\text{Consistency Ratio} = \frac{\text{Consistency Index}}{\text{Random Index Value}}$$

The random index values for n number of criteria are available in Appendix A. If the CR is much higher than 0.1, the decisions may be considered unreliable because they are too close to randomness, and the judgements are either useless or need to be conducted again (Saaty, 1980).

In this project, there are certain affecting-affected relationships between different critical inputs and/or criteria. There are also precedence differences between different critical inputs or criteria. Therefore, an integrated system of different methodologies is used. Particularly, DEMATEL comes in handy since there are certain critical inputs/criteria which have an effect on others such as company capital having an effect on capital per share holder. AHP, on the other hand, comes in handy since there are also certain critical inputs/criteria which have significant importance superiority over others. For instance, Judge Committee Report Decisions have significant importance compared to the other criteria.

We designed a decision support system which has two components: "Setup" and "Operations". We built a software that is able to determine the scalar weights of criteria, which is the Setup component, and obtain standardized scores of applications with respect to criteria weights, which is the Operations component.

5 System Design

For the implementation and execution of the Decision Support System, Excel VBA is utilized in order to create a user-friendly interface. Since Microsoft Excel is highly used in CYBERPARK, the selected decision tool is easily integrated. Utilization of the decision support system is achieved through an interface that is built with clear and simple commands for the convenience of the decision-makers.

5.1 Setup Component

While constructing the setup component, the first step was analyzing the forms that the decision makers use during the evaluation process. After careful consideration with the decision makers, it is concluded that not all information in the forms are considered as critical information. In order to minimize the number of comparisons and consider the system in a categorical manner, the critical information is clustered into 6 criteria. After this point critical information is referred to as critical inputs. They can be accessed in Appendix B. In the Setup

component, scalarized weights of criteria and critical inputs are found with AHP, DEMATEL and Hybrid methods, which combines the AHP & DEMATEL scores into a single score. This way, the users can utilize the system to their preference. Additionally, decision makers may want to change the criteria, critical inputs or dual comparison values in the future. Hence, for the sake of making a fully dynamic and compatible system, the software is prepared so that the users can add, remove and change the valuations of criteria and critical inputs and initialize the platform accordingly whenever they want to. Initially, the decision makers were asked to make the dual comparisons directly on a matrix. However, due to high consistency ratios and feedback from industrial advisors, it is concluded that it was not user-friendly. Hence, the pairwise comparisons are designed in a survey based interface in the software. Furthermore, in order to minimize the possible consistency errors in the future setups, an algorithm that checks the user's answers along the survey for transitivity and disables the upcoming options that would conflict with the already given answers is implemented.

5.2 Operations

In this component; users can calculate a new application score, add or remove an application to the waiting list, delete the score of an application from the system and view the waiting list. When calculating a new application score, critical input information is extracted with respect to the application ID and applicant company name. Each critical input information is scored between zero and five for the sake of compatibility with the judge committee report criterion's critical input which are already on a point basis from zero to five, see Appendix C. Valuation of as much critical input information as possible is automated by assigning appropriate scores to possible information. The ones that require human judgement are left for the users to score. After all critical inputs are scored, the application's overall score is calculated. Each critical input information score is multiplied by its weight. Multiplication of the sum of these weighted critical inputs by the corresponding criterion's weight gives the criterion score. Finally the overall application score is obtained by summing all criteria scores. The final application score is shown to the users categorically, an "Application Summary" is provided. This summary includes some trivial information apart from the critical inputs. These are the information which the industrial advisors stated that even though an application is acceptable based on its critical inputs, these trivial information have to be checked and they have to calculate particular ratios like project budget divided by project duration, by hand to make sure whether the applicant project is realistic and follows all guidelines. Since CYBERPARK stated that automating even one of these calculations would save them significant time and energy, all information that they check is detected, automated and visualized with respect to whether they are positive or negative, see Appendix C. Whenever the waiting list is opened, all applications in the waiting list are sorted from the highest application score to the lowest one with respect to the Hybrid approach by default, see Appendix C. With the "Change List Settings" button the

sorting option and displayed score methodology can be changed. Moreover, the list can be filtered by values such as required office space, project beginning date or different types of scores so that the users can utilize the list for their current specific needs as accurately as possible. Additionally, to get a better grasp of the current structure of the overall waiting list while making decisions, maximum, minimum and average values of all types of scores as well as the required office space in the list are displayed with the “Show List Statistics” button. Finally, with the “See Application Point Distribution” button, pie charts which show how well an application performs on different criteria according to all methodologies can be seen. See, Appendix C.

6 Verification and Validation

Under the regulations of KVKK (Law On The Protection Of Personal Data), CYBERPARK stated that any kind of historical data can not be shared with the group. In order to understand whether these methodologies are working properly or not, first, a verification process is conducted. For this purpose, sample data constructed by the group members is used. Primary aim of the verification process via using the sample data was to ensure that the algorithms were working as intended as well as the quality of software application, design, and architecture. In addition to this, it is checked whether the methodologies are consistent with the designed system.

Note that the criteria and critical inputs with higher importance have to admit higher weights according to AHP methodology. Similarly, the criteria and critical inputs that have more interaction with the others are expected to admit higher weights according to DEMATEL methodology. For example, according to comparisons done by industrial advisors, throughout the project construction, it is stated that the only definitive criteria that stands out for their decision making process is the Judge Committee Decisions and that other criteria’s priorities would take form accordingly, in terms of validating criteria weights. The Judge Committee Decisions criteria takes the highest value (nine) when it is compared to different criteria during the setup surveys. Consequently, it got the highest weight (0,58) according to AHP methodology, second highest weight (0,22) in DEMATEL methodology and again the highest score (0,40) in Hybrid methodology. The default scores are displayed as the hybrid scores. According to the current valuations done by decision makers the ranked criteria weights list is as follows:

On the other hand, for validation purposes, testing and demonstrating are conducted in order to see whether the designed software meets the needs and expectations of the decision makers or not. After the demonstration of the system, the company gave feedback on the delivery of expectations.

- Instead of taking the application number as the identifying number ID for the application the Industrial Advisors stated that they normally identified

the applications also using the company name. However, since the same company may have more than one application in, after deliberating with the Operations Support Services Manager (OSSM) we agreed on generating an application code with the application ID and applicant company name.

- The OSSM also stated that it would be better if the users could select the application from a drop down menu rather than entering its code as they are usually not sure about which application to evaluate without seeing all alternatives.
- In addition to previously automated critical inputs, industrial advisors wanted to automatize the evaluation of the funding source of the project, the nationality of the foreigner company shareholder and the university the academician shareholder works in.
- It is requested that the units of the calculations which are made for checking the trivial information of the application are shown.
- The OSSM also gave feedback on scoring of the university the academician shareholder works in. Bilkent University can be referred as Bilkent University or İhsan Doğramacı Bilkent University. Therefore, the automatization of this critical input is changed in order to be adjustable to different refutation to Bilkent University.
- Finally, the OSSM gave feedback on the size allowed for open ended questions and mentioned the length for the answers of this can vary up to 3 pages. Therefore, the answers for these questions are displayed as scroll through boxes.
- Instead of initializing the system by clearing previously calculated weights and application scores when the setup is changed, industrial advisors wanted to store the results with different setups separately. Therefore, we suggested that they use different copies of the Excel file with the software for different setups. For example, the users will be able to have files with setup of 2019-2021 and 2021-2022 if a change occurs in the criteria or critical inputs.

Industrial advisors are satisfied with the overall system. They mentioned that they wanted to use it with their current evaluation processes as soon as possible. They stated that they found it very convenient and well designed.

7 User Interface

The Setup component covers the services needed for the setup of the support system which are:

- Changing the valuations of criteria
- Changing the valuations of critical inputs
- Adding/removing criteria

- Adding/removing critical inputs

If new criteria and/or critical inputs will be added to the system in the future, scoring of their corresponding data will be done via open-ended answers. However, the system is coded with detailed explanations on how and why the algorithms are constructed the way they are. Due to this the system can be coded in another platform by CYBERPARK's IT department in the future. Similarly, their valuation types can also be changed via coding by professionals.

The Operations component covers the operations needed for the actual utilization of the support system, which are:

- Calculating the score of a new application
- Adding a new application to the waiting list
- Displaying the waiting list
- Displaying the waiting list statistics

Below, the main page of the interface can be seen. The rest of the sections from the interface can be seen in Appendix C.

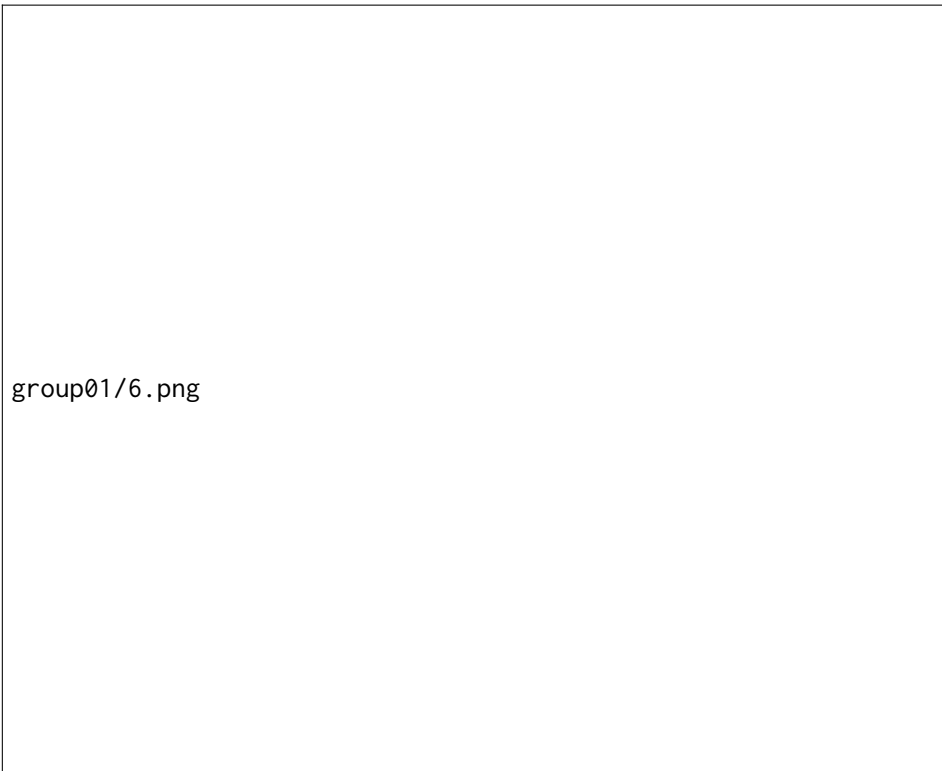


Figure 1: Main Page

8 Project Contributions

As mentioned in the problem definition, the main problems that CYBERPARK faces are the lack of standardization of the whole process. CYBERPARK aims to be consistent in its decision-making process. Within this scope, a transitivity algorithm which prevents the inconsistent setup of the system at the first place is constructed. Therefore, with the help of this project the decisions intended to be non contradictory. Existence of the transitivity algorithm is important because when the first comparisons of the AHP method are examined, there were a significant number of inconsistent and contradictory decisions which resulted in high consistency ratios. With the help of this implementation, the reliability of the decisions and therefore the decision support system is increased.

Furthermore, one of the main goals of the project is to assist the decision makers to make their decisions on a more standardized basis. With the integration of our system, an approach based on well-known multiple criteria decision-making methodologies is utilized. The users of the Company Evaluation Platform are allowed to change and alter criteria and critical inputs which builds the foundation of the decision making system. Therefore, the decision support system is responsive against changing conditions. Hence it is dynamic and sustainable.

With the DSS, the users will be able to view the applications in the waiting list; the average, the maximum and minimum of the scores. They will also be able to view score distributions among the criteria of a chosen application. Therefore, the users are provided with a scored, sorted and categorized model instead of a large unprocessed database when they make a decision.

For a decision-making support system project, quantitative benefits are hard to detect and define due to the immeasurability of the objectives. Reduced time can be taken as the most measurable benefit. In the existing system, it may take multiple days to make hiring decisions and it requires a lot of work force, instead the whole process will be decreased to approximately 15 minutes. After the demonstration of the system, the decision makers were impressed by the system and what it offers during the decision making process as well as its design. They stated that they wanted to use it for the current existing system immediately.

References

- CYBERPARK. Bilkent CYBERPARK. *Bilkent CYBERPARK*, 2019. URL <https://www.cyberpark.com.tr>.
- E. Falatoonitoosi, Z. Leman, S. Sorooshian, and M. Salimi. Decision-Making Trial and Evaluation Laboratory. *Research Journal of Applied Sciences, Engineering and Technology*, 5:3476–3480, Apr. 2013. doi: 10.19026/rjaset.5.4475.
- M. Pachemska, R. Timovski, Martin, and Riste. (PDF) ANALYTICAL HIERARCHICAL PROCESS (AHP) METHOD APPLICATION IN THE PROCESS OF SELECTION AND EVALUATION. *ResearchGate*, 2014.

URL https://www.researchgate.net/publication/276985609_ANALYTICAL_HIERARCHICAL_PROCESS_AHP_METHOD_APPLICATION_IN_THE_PROCESS_OF_SELECTION_AND_EVALUATION.

T. L. Saaty. International hellenic university 2 paraskevopoulos konstantinos. pages 1–5, 1980.

T. L. Saaty and M. Hall. FUNDAMENTALS OF THE ANALYTIC NETWORK PROCESS. 1999, page 14, Aug. 1999.

J. Varajão and M. M. Cruz-Cunha. Using AHP and the IPMA Competence Baseline in the project managers selection process. *International Journal of Production Research*, 51(11):3342–3354, June 2013. ISSN 0020-7543, 1366-588X. doi: 10.1080/00207543.2013.774473. URL <http://www.tandfonline.com/doi/abs/10.1080/00207543.2013.774473>.

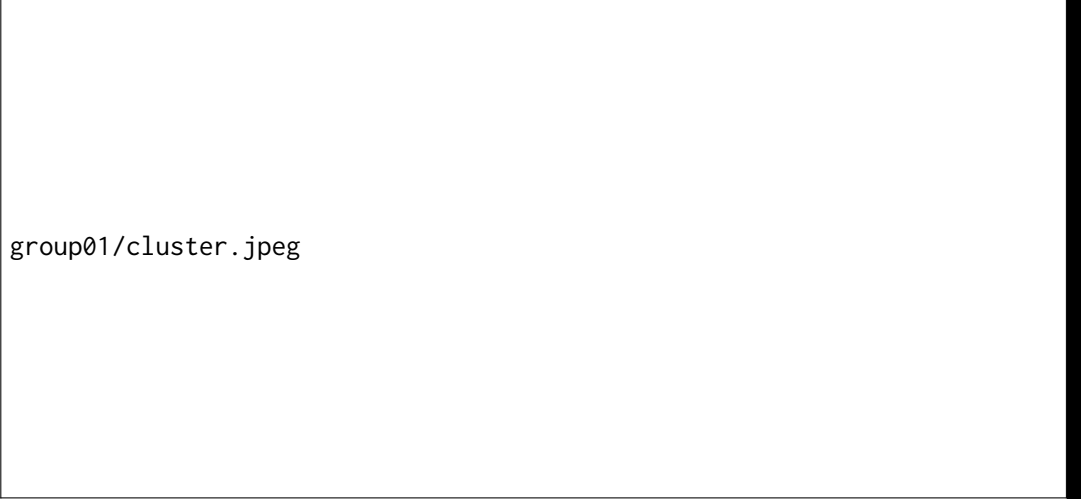
N. Ömürbek and Z. Tunca. A CASE STUDY ON GROUP DECISION MAKING STAGE IN ANALYTIC HIERARCHY PROCESS AND ANALYTIC NETWORK PROCESS METHODS. 2013.

Appendix A Random Index Values

Table 1: Random Index Values

n	1	2	3	4	5	6	7	8	9	10
RI	0.00	0.00	0.58	0.90	1.12	1.24	1.32	1.41	1.45	1.49

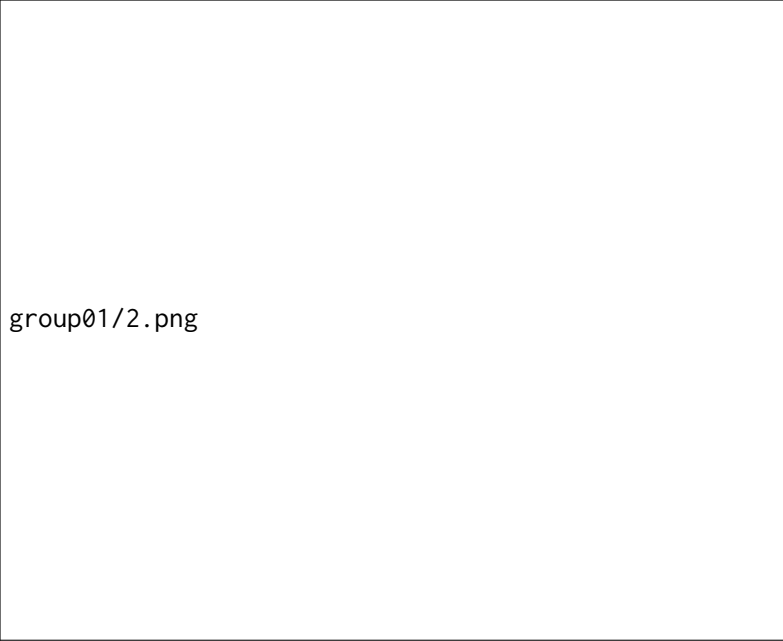
Appendix B Clustered Criteria



group01/cluster.jpeg

Figure 2: Clustered Criteria

Appendix C Interface



group01/2.png

Figure 3: DSS Setup

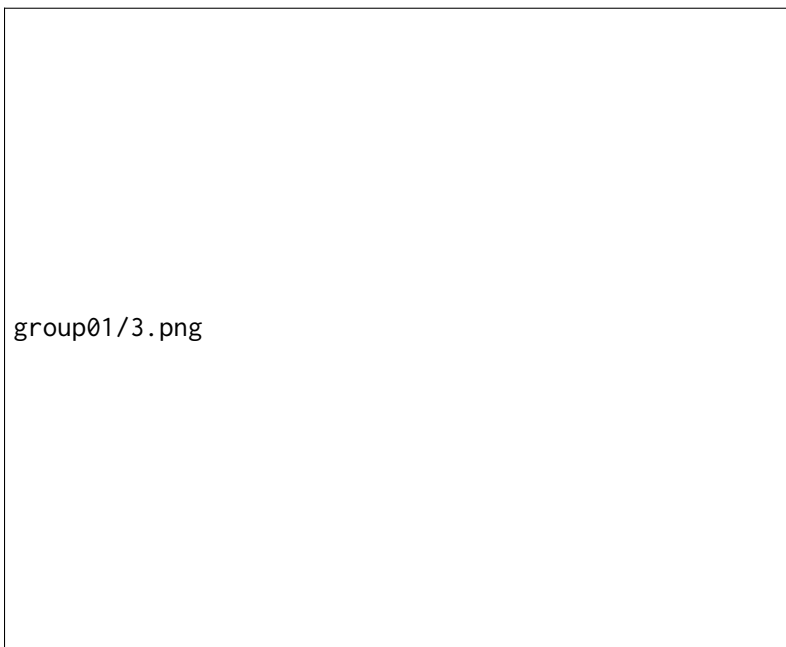


Figure 4: Add or Deduct Critical Input/ Criteria

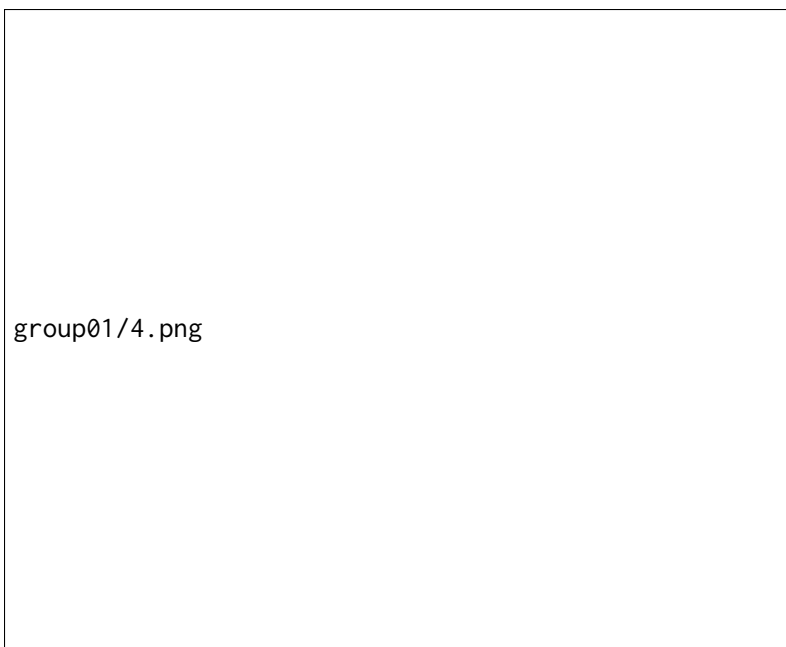


Figure 5: Statistics of the Applications



Figure 6: Valuation of The Criteria/ Critical Input AHP



Figure 7: Scoring Interface

Yarı Mamul Üretim Hatlarında Ara Stok Seviyelerinin Düşürülmesi için Parti Büyüklüğü ve Üretim Planı Belirleyen Karar Destek Sistemi Arçelik Kompresör İşletmesi

group02/group_photo.png

Proje Ekibi

Kürşad Ali Akdoğan, Pelin Erdil, Derya Erkan,
Beste Fırat, Ecenur Oğuz, Göksu Ece Okur, Hakan Dorukhan Yeşilli

Şirket Danışmanı

Kadir Esmeroğlu
Endüstri Mühendisliği Takımı

Akademik Danışman

Prof. Dr. Mehmet Selim Aktürk
Endüstri Mühendisliği Bölümü

ÖZET

Arçelik Kompresör işletmesinde ara stok seviyelerini yönetmek için standart bir sisteme ihtiyaç duyulmaktadır. Projenin amacı Arçelik kompresör işletmesindeki ara stok seviyelerini düşürmektir. İşletmedeki istasyonlar ara stok düşürme metodlarına göre iki gruba ayrılmıştır. Bu gruplar için üretim çizelgeleme-parti büyüklüğü, ve üretim dengeleme olmak üzere iki farklı algortima tasarlanmıştır. Bu algoritmalar kullanıcı ara yüzüne aktarılarak işletmenin kullanımına sunulmuştur. Sistem, işletmenin geçmiş aylardaki talep ve ara stok verilerini kullanarak doğrulanmıştır. Tasarladığımız sistem kullanıldığında aylık talebe bağlı olarak muhafaza ve biyel hatları için ara stok seviyelerinde sırasıyla ortalama %53 ve %57 düşüş öngörülmektedir. Karar destek sistemi firma tarafından onaylanmış ve kullanılmak üzere işletmeye gönderilmiştir. İşletme, karar destek sistemini denedikten sonra sistemin yıllık 140,000 TL finansal getirisi olacağını öngörmüştür.

Anahtar Kelimeler: Karar Destek Sistemi, Parti Büyüklüğü, Üretim Çizelgeleme, Üretim Dengeleme, Ara Stok Yönetimi

Developing a Decision Support System for Integrated Lot Sizing and Scheduling to Decrease WIP Levels in Semi-Finished Production Lines

1 Company Information

Arçelik is the leader of the Refrigeration Appliances Market in Turkey with a 36% retail volume share between the years 2011-2020 (Euromonitor Passport, 2021). It is stated by the company that the Arçelik Compressor Facility supplies 80% of the compressor demand of Arçelik Refrigerator Facility. Arçelik Compressor Facility also supplies spare parts for broken down compressors as an after-sales service. 20 varieties of products are produced with 50-60 SKU's, the products can be compartmentalized into three types of compressors namely; KIK (Compact Inverter Compressor), MINIKIK (Mini Compact Inverter Compressor) and MIDI. The demand of the company consists of the demand that is given by the Arçelik Refrigerator Factory, the international partners, and service requests.

2 System Analysis and Problem Definition

The aim of the project is to decrease the WIP levels of the semi-finished production lines called Machining line and Mechanical line. The demand data of semi-finished production lines is obtained from the production plan of the assembly line. There is seasonality in demand, and the demand of summer months is higher than winter months. A system analysis is conducted and it revealed several reasons for keeping high WIP levels in the facility. The lack of a particular scheduling method in production and a standardized production leveling procedure causes an increase in the WIP levels. It is found that the different lots of the same products are combined in order to minimize the number of changeovers, however this procedure creates more WIP. Also, the lots are not divided efficiently and thus this causes the subsequent stations to stay idle. As a result, idle stations cause high WIP levels. In certain semi-finished production line stations, the capacity of the machines are not enough to satisfy the demand of the assembly line on time. Therefore, the WIP levels are kept high in order to eliminate idleness. These reasons of keeping high WIP levels in the facility are identified as a result of the system analysis. The analysis suggests that different stations need different methods to reduce WIP levels.

Two different cases are identified in the semi-finished product lines. In the first case, the stations feed multiple succeeding stations and thus having certain level of WIP is required in order to avoid idleness in succeeding stations. This idleness will be referred as starvation throughout the report. Housing line located at the mechanical line is chosen as pilot stations for this case. In the housing line, there are two types of WIP. The WIP stored after presses will be referred as housing WIP and the WIP stored after welding stations will be referred as the housing assembled WIP in this report. Flow chart of housing line can be seen in Figure 1.

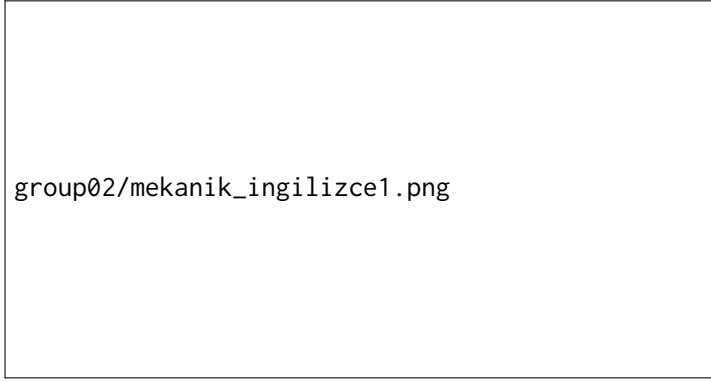


Figure 1: Flow Chart of Housing Line

In the second case, the chosen stations do not feed multiple machines but their capacities are not enough to satisfy the daily demand from a day before without keeping any WIP. Connecting rod line located at Machining line is chosen as the pilot station for this case. Single type of model is processed in each station in the connecting rod line. The flow chart of the connecting rod line can be seen in Figure 2. Each processing station feeds only one honing station, and honing stations cannot fulfill the demand on time for approximately one third of the days in a month. The insufficient capacity of the honing stations cause high WIP levels in the connecting rod line.



Figure 2: Flow Chart of Connecting Rod Line

These cases led to the development of two different algorithms to decrease the WIP levels while satisfying the demand on time.

3 Decision Support System

A Decision Support System (DSS) is developed to schedule production in housing line and the connecting rod line while minimizing WIP. The DSS has several components to provide efficient production planning and WIP management system. The DSS is designed to be a dynamic system since the demand of the facility is

not stable and has seasonality. In order to obtain a dynamic approach, the system updates itself when new information is received regarding demand, workdays, number of shifts etc.. In this section, the components of the DSS will be explained in detail.

3.1 Master Table & Material Master Table

Firstly, the necessary inputs for the DSS are investigated. Currently, two main inputs are used in the production planning of the semi finished product lines. The two main inputs are the monthly production amounts in the assembly and the main order list for the finished products. The company does not have a single main input which combines the two data. However, acquiring the data from a single input would make the production planning more efficient in semi-finished production lines because the combination of two data includes detailed information of semi-finished production lines to create an accurate schedule. Therefore, the DSS provides a main input. This input is created dynamically each month. It is composed of the monthly customer orders list and monthly planned assembly production amounts. The main input is referred as the Master Table. It is used to create the production plan which the DSS provides for the company. Furthermore, the Master Table is expected to be useful not only for the proposed DSS but also for planning the operations in the entire facility. The Master Table is differentiated to create another table, which is referred as Material Master Table. It is created by infusing the bill of material information with the Master Table. This table is used for acquiring information on the materials required for the production of semi-finished products. The materials required for each planned assembly production are listed by the DSS for convenience. Screenshots from Master Table and Material Master Table can be found in Appendix A.

3.2 Safety Stock

In Arçelik Compressor facility, it is stated that unexpected demand may occur during a month. Therefore, in order to minimize the undesirable effects of the unexpected demand, safety stock must be kept. The safety stocks are calculated for effective WIP management. The formula, as discussed in (Nahmias and Olsen, 2015), used for the safety stock calculation is given below. λ denotes the demand rate in unit time, σ_λ denotes the standard deviation of demand in unit time, μ denotes the expected lead time and Z is the z value at the desired service level.

$$Z \times \sigma_\lambda \times \sqrt{\mu} \quad (1)$$

Lead times are calculated according to the batch sizes and material handling. Variations in lead times are near to zero. Therefore, the lead times are fixed and taken as parameters. Service rate is determined as 95% with the approval of the company. Safety stocks are computed dynamically by the DSS for each compressor model. At the beginning of each month, the monthly demand information is received and the system computes the standard deviation of the demand for each

model. Therefore, dynamic safety stock algorithm ensures that the system could satisfy the demand throughout a year.

3.3 Housing Line Lot Sizing & Scheduling Algorithm

In this section, the proposed system for the housing line will be explained. In order to give a better understanding of the system, firstly the buffer stock algorithm will be explained. Buffer stock algorithm is developed to be used when there is a risk of starvation. Starvation can occur if a machine feeds multiple stations. It is observed that in housing line, buffer stocks are needed after presses since they feed multiple welding stations. Therefore, the housing line is selected as a pilot line. In the proposed system, both presses and welding stations produce the demand of the following day. Lot sizing and scheduling are performed in order to reduce the WIP levels between presses and welding stations. The aim of the approach is to find the excess time that would occur if there was zero WIP. Through this approach, the maximum excess time is determined. The algorithm calculates the minimum buffer stock to eliminate the maximum excess time. This excess time is calculated for each day of a month by utilizing a mathematical model. This mathematical model takes the processing and setup times as parameters and finds the lot sizes which will minimize the excess time in a fixed schedule. Minimizing the excess time also minimizes the WIP levels. The algorithm ensures that the production can continue without stopping the line while keeping minimum WIP. The proposed lot sizing and scheduling algorithm is as follows;

- **Step 1:** The assembly production schedule is acquired as a demand input. For possible demand configurations, mathematical solvers were initially developed. The suitable solver is found according to the demand configuration. The Excel solver is utilized to find the excess time at the welding stations assuming there is zero WIP. These processes are repeated for each day according to the given monthly assembly demand data.
- **Step 2:** The maximum excess time for each compressor model for the given month is calculated. The minimum buffer stock amount that is necessary to eliminate the maximum excess time is calculated for each compressor model. The minimum buffer stock level ensures that production can continue without stopping the line.
- **Step 3:** Production amounts are computed for each station using the calculated buffer and safety stocks.
- **Step 4:** It was requested by the company to have the production sequences fixed according to the demand configurations. Therefore, the production sequences were determined beforehand by a detailed what-if analysis. Excel Solver is again utilized to obtain lot sizes according to the given sequence. The model takes production amount as an input instead of daily demand, and computes the lot sizes. The schedule and lot sizes are the final output of the DSS.

The mathematical model is developed and is solved by Excel Solver. Objective of the mathematical model is to determine the lot sizes such that the excess time would be minimized while ensuring that welding stations are used efficiently. There are eight mathematical models for eight different demand configurations. Mathematical model for housing line can be found Appendix B. This is the only mathematical model provided in this report out of eight models due to space limitations. Other mathematical models which contain two product groups and/or single lots are extensions of the model provided.

3.4 Connecting Rod Line Production Leveling Algorithm

In this section, the proposed system for the connecting rod line of Machining line will be explained. There is a difference in the approach to reduce WIP levels compared to the approach applied to the housing line. The main reason for keeping high WIP levels in connecting rod line is identified as the insufficient capacity of the honing stations. The capacity of the honing stations cannot fulfill the assembly demand without any WIP on certain days of the month. Therefore, the production must be carried out ahead of schedule for the days in which the demand cannot be fulfilled. Therefore, production leveling is required to keep constant WIP that satisfies the demand without any delay. The DSS utilizes a mathematical model. In the mathematical model, production amount and the required amount of WIP are identified as the decision variables. The mathematical model finds the production amounts for each day while minimizing the WIP levels and ensuring the demand is met. The details of the mathematical model for connecting rod can be found in Appendix C.

3.5 User Interface

The User Interface (UI) is designed to provide outputs of DSS while considering the company's requests to provide a beneficial product. The monthly production plan of the assembly is taken as an input to the proposed DSS. After the input is entered, the DSS will be run at the beginning of the month in order to compute the corresponding monthly schedule. The production planning department will be responsible of creating the production schedule. Monthly production plan, safety stocks and buffer stocks will be computed. The UI of DSS will be used by the operators through the computers provided in work stations. Workers will be able to see the production schedule of each machine. The first page of the UI is the home page where operators can select the line that they are working. For both Mechanical and Machining lines, a separate UI page is created. The time period and the machine name can be selected from their respective combo boxes. After an operator provides a time period and a machine name, the program lists the production schedule of the selected machine for the corresponding time period. The resulting list contains six columns. The columns show the information that the company requested to see. These provided values are the dates, the production line names, the machine names, the SKU numbers, the SKU names, and finally their production amounts in a particular ordering to fit the schedule that is found

by DSS. The operators have to enter the production amounts they produced at the end of the each day in order to keep track of the WIP levels. The DSS consists a "WIP Monitor" page which the users are able to see the WIP levels of each product type and each day. The company is not able to keep track of the WIP level daily with their current system. The DSS offers this feature for an effective WIP management. The UI has the feature to update the production amount. For instance, if the actual produced amount does not match the amount given in the schedule found by the DSS, then the actual produced amounts should be logged into the system. The system will compensate this change in the production of the following shifts. Therefore, an update page is constructed. This feature offers flexibility to the system. It is stated by the company that demand data may change during a month due to new customer or service orders. Therefore, the DSS is able to reflect to these updates. When a change occurs in demand during a month, then the "update assembly demand" button can be selected in order to update the data. Screenshots from UI can be seen in Appendix C. Please note that information regarding the products are blurred due to the confidentiality issues.

4 Sensitivity Analysis

A sensitivity analysis is conducted to see how the proposed system works under different scenarios. The scenarios in which the average monthly demand is much higher or lower than the expected demand are considered. Also, scenarios are generated while considering the seasonality in demand. The monthly demand is lower in winter and it is increasing in summer months. The detailed information of the scenarios will be given in the next section.

4.1 Sensitivity Analysis for Housing Line

Four different demand data are generated and used as four scenarios for the sensitivity analysis of the housing line. The demand of high and low scenarios are compared with the expected average monthly demand of summer and winter seasons. Scenario 1 represents a very low average monthly demand whereas scenario 4 represents a very high average monthly demand. When past annual demand data is analyzed, it is observed that these two scenarios are highly unlikely to occur. Scenarios 2 and 3 represent the average expected monthly demand for winter and summer respectively. The scenarios are generated based on the demand data of the assembly line for January, 2021. The demand is multiplied by various constants to generate the demands of scenarios 1, 3 and 4. The ratio between the demand of each model is preferred to be kept constant in order to create a realistic data. Table 1 is created for sensitivity analysis for housing parts' WIP levels. Table 2 represents the results of the sensitivity analysis in housing assembled parts' WIP levels.

Table 1: Results of the Sensitivity Analysis for Housing Parts

Scenario	Demand	WIP Level Before	WIP Level After	Improvement
1	3790	7580	1400	82%
2	6317	12634	5490	54%
3	8717	17434	14119	19%
4	11219	22438	18709	17%

Table 2: Results of the Sensitivity Analysis for Housing Assembled Parts

Scenario	Demand	WIP Level Before	WIP Level After	Improvement
1	3790	4738	2200	53%
2	6317	9000	4293	52%
3	8717	10896	10177	6%
4	11219	14024	13617	3%

The WIP Level Before column represents the WIP level if the company were to operate under the current system with the given demand data. In the current system, the housing WIP level is double the amount, and the housing assembled WIP level is kept 1.25 times of the daily demand of assembly. The WIP Level After represents the required WIP levels found by the proposed system. It can be observed from Table 1 and Table 2 that as the demand level increases, the percentage of improvement decreases. As a result of the sensitivity analysis, the system was able to give an improvement of 17% and 3% under high demand scenarios for housing and housing assembled parts' WIP levels respectively. Through this result, the proposed system is verified and expected to provide an improvement to the facility even under different scenarios.

4.2 Sensitivity Analysis for Connecting Rod Line

Three different demand data are used as three scenarios similar to the sensitivity analysis conducted in housing line. Scenarios 1 and 3 are not expected to occur and thus enabled the verification of the system under unlikely circumstances. Scenario 1 represents a lower demand case compared to the expected demand data. In scenario 2, demand of January, 2021 is used. Scenario 3 represents the high demand case compared to scenario 2. Demand of Scenario 1 and 3 are generated by multiplying the demand with a constant. Table 3 represents the results of the sensitivity analysis in connecting rod line. The WIP level in the current system is stated as 13500 units and kept constant by the company. It is provided in the column called "WIP Level Before". The expected WIP levels and improvement percentages can be seen in the other columns. When demand increases, improvement percentages decrease. However, 35% decrease in WIP levels can still be obtained even under an extreme scenario.

Table 3: Results of the Sensitivity Analysis for Connecting Rod Line

Scenario	Demand	WIP Level Before	WIP Level After	Improvement
1	3158	13500	4488	66.75%
2	6317	13500	5841	56.73%
3	10632	13500	8703	35.53%

Due to safety reasons, the company did not allow the entrance for any guests to the factory. However, the company was able to compare the results of the system to the current practice of the company through April 1st to 3rd, remotely. The produced amount in these days and production amounts of the DSS are compared. The production amounts of the system are found to be able to satisfy the demand and thus verified the applicability of the system. After comparing the results of the systems, the company confirmed that they produce more than they need and thus the resulting stock levels of DSS are found to be less than the current system. The comparison tables can be found in Appendix E.

5 Validation

The outcome of the DSS and the current practice of the company is compared in order to validate the DSS. The validation is performed by using the demand data of January, 2021 for three compressor models. Details of validation can be seen in following sections.

5.1 Housing Line

When the proposed system is applied, it is observed that the housing WIP level will decrease from 15000 to 6894 units in a week and will stay stable around that level. It is also observed that the housing assembled WIP will decrease from 9000 to 4293 units.

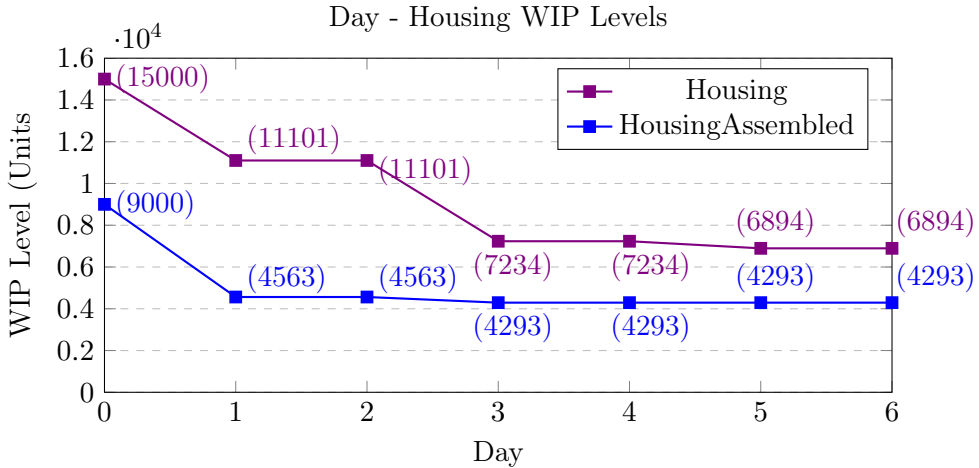


Figure 3: Change in WIP Levels in Housing Line with January Data

Figure 3 shows the transition to the suggested housing WIP levels and the suggested housing assembled WIP levels when the DSS is applied. After the implementation of the DSS, a 54.04% decrease is predicted in the housing WIP levels and a 52.3% decrease is predicted in the housing assembled WIP levels in January, 2021. The DSS is approved and decided to be implemented by the company after the validation results are shared.

5.2 Connecting Rod Line

The proposed system for connecting rod line is validated and explained in this section. Two different types of WIP's are analyzed which are processing WIPs, and honing WIPs. When January WIP levels in the current system and proposed system are compared, 62.22% decrease in honing WIPs and 52.32% decrease in processing WIPs are observed. The effect of production leveling on WIP levels are shown in the graph below. The graph is constructed with the demand data of January. Processing times of processing stations are lower than honing stations. Therefore, the processing stations are able to feed their subsequent stations continuously without carrying any buffer stock. When an increase is seen in the graph, it means that production levelling is performed. The demand of subsequent days are decided to be produced early due to capacity issues. If the DSS was implemented in January, then the WIP levels would gradually decrease to the safety stock levels.

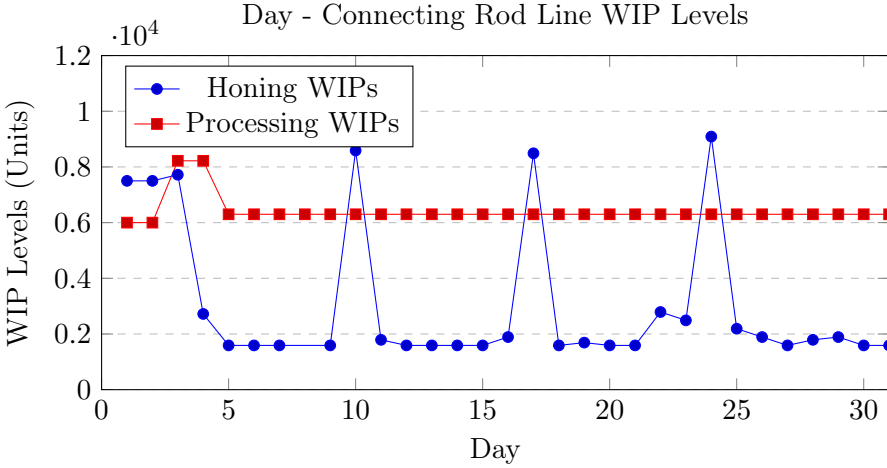


Figure 4: Change in WIP Levels in Connecting Rod Line with January Data

6 Project Contributions and Implementation Details

The main motivation of the project is to decrease WIP levels in semi-finished production lines. For this purpose, a DSS is developed. It creates a production schedule while minimizing WIP levels. The DSS is developed in a way that is unique to the company. It is personalized according to the needs of the Arçelik

compressor facility. It is observed that the company was in need of a dynamic system due to the fluctuations and uncertainty in demand. Therefore, a dynamic DSS is developed to respond to the changing conditions in the production smoothly. In the facility, production scheduling is performed manually and company is in need of a standardized system. That is why, the DSS provides a formal production plan. The company cannot keep track of the WIP levels easily within a month. Therefore, having a formal production planning system also helps the company to keep track of the WIP during a month. The DSS also computes the necessary safety stocks to ensure an efficient WIP management and to eliminate the impact of the demand uncertainty. Also, a Master Table and a Material Master Table are created by the DSS according to the requests of the company. The other major traits of the DSS are that it is easy to use and modify, extendable and it has no implementation cost. A template explaining the details of the algorithm for the other stations is prepared and sent to the company as a guidance for extension. User-friendliness of the DSS is given a high priority. Therefore, as the company utilizes Excel for the majority of their operations, Excel is used in all aspects of the project as well so that the system can be easily operated. The excel solver is used to solve the mathematical models of the system. Also, the User Interface is made simple and clear. It gives instructions and warnings to guide users. It is worth to be mentioned that the DSS is able to give an output in approximately 23 seconds. The time measurement was performed on a laptop with 64-bit Operating System, x64-based Intel(R) Core(TM) i5-8265U processor, CPU @1.60GHz to 1.80 GHz, and with installed memory (RAM) equal to 8.00 GB. It is important for the execution time to be reasonable since it is expected that the company will use the DSS daily. As it is explained in Section 5, the system is expected to decrease the housing WIP by %54.04, while the housing assembled WIP is expected to decrease by %52.3. For the connecting rod line, the total honing WIP is expected to decrease by %52.32 and the total processing WIP is expected to decrease by %62.25.

Implementation of the DSS and the performance of the algorithm are approved by the company. The DSS is delivered to the company during a meeting on 24.03.2021. A user manual is prepared and delivered to the company to explain how to operate the proposed system in detail. The company has decided to examine the system in the facility between the dates April, 3 and April, 14. In April, 14, a meeting was conducted with the company in order to receive feedback regarding their examinations. In the following days, necessary adjustments on the DSS are done based on feedback and the updated DSS is sent to the company in April, 19. Some physical adjustments in the housing line are recommended to the company to increase the efficiency of the DSS. These recommendations are approved by the company and it was stated as an easy and applicable operation by the production engineers. The company stated and verified that the number of days in stock is expected to decrease from 4.2 to 2.6 days approximately when the DSS is used. This decrease in the stock days is a measure of the benefits of the

DSS to Arçelik and it corresponds to approximately 140,000 TL annual financial profit. Also, Arçelik Compressor facility has plans to use SAP for their operations in the upcoming years. Therefore, they plan to integrate the proposed system to SAP and thus make further use of the system.

References

Euromonitor Passport. Refrigerator appliances: Company share of arçelik in turkey, 2021. URL <https://www.portal.euromonitor.com/portal/statisticsevolution/index#>.

S. Nahmias and T. L. Olsen. *Production and operations analysis*. Waveland Pr, 2015.

Appendix A Screenshots from Master&Material Master Table

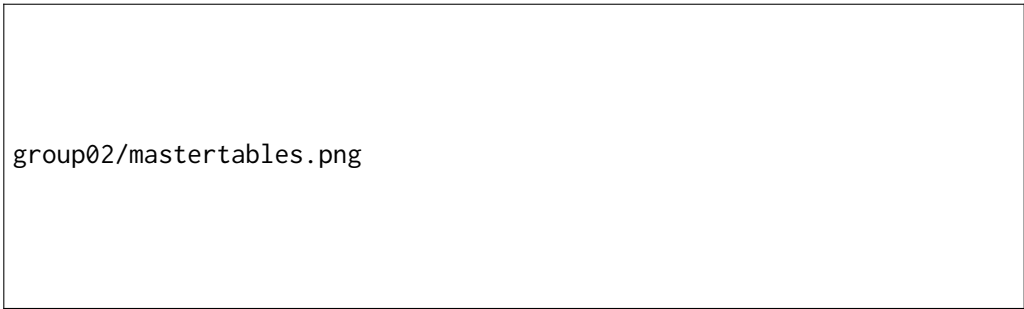


Figure 5: Screenshots from Master&Material Master Tables

Appendix B Mathematical Model for Housing Line

Sets:

I : Lots in Press 1 $I = \{1, 2, 3, 4, 5, 6, 7, 8\}$
 L : Set of Products $K = \{1, 2, 3, 4\}$ where 1 = KIK, 2 = MINIKIK lower
3 = MINIKIK upper, 4 = MIDI lower
 Z : Set of Products $Z = \{1, 2, 3\}$ where 1 = KIK & MINIKIK lower
3 = MINIKIK upper, 4 = MIDI lower
 N : Set of welding groups $N = \{1, 2, 3\}$ where 1 = Welding stations 1 and 4,
2 = Welding station 2, 3 = Welding station 3

S : Defines a set for lot pairs that belong to the same part where (i,j) correspond to $(i,j) \in \{(1,8), (2,6), (3,5), (4,7)\}$

F : Defines a set for triplet where $z \subseteq Z$, $n \subseteq N$, and $i \subseteq I$

$$(z, n, i) = \{(1, 4, 8), (2, 2, 6), (3, 2, 5)\}$$

B : Defines a set for triplet where $z \subseteq Z$, $n \subseteq N$, and $i \subseteq I$ where

$$(z, n, i) \in \{(1, 1, 1), (1, 2, 4), (1, 3, 4), (1, 4, 8), (2, 1, 2), (2, 2, 6), (3, 1, 3), (3, 2, 5)\}$$

M : Defines a set for triplet where $z \subseteq Z$, $n \subseteq N$, and $i \subseteq I$ where

$$(z, n, i) \in \{(1, 2, 1), (1, 3, 4), (1, 4, 7), (2, 2, 2), (3, 2, 3)\}$$

Decision Variables:

$$x_i : \text{Lot size of lot } i \text{ in Press 1} \quad \forall i \in I$$

$$s_i : \text{Start time of product of lot } i \text{ in Press 1} \quad \forall i \in I$$

$$m_{zn} : \text{Start time of production } z \text{ in welding group } n \quad \forall k \in K$$

$$\theta_z : \text{Excess time of production in welding station } z$$

Parameters:

$$D_l : \text{Demand of product } l \quad \forall l \in L$$

$$p_0 : \text{Processing time of Press 1 (minutes)}$$

$$p_z : \text{Processing time of welding group } z \text{ (minutes)}$$

$$c : \text{Changeover time for Press 1 (minutes)}$$

$$l : \text{Minimum lot size}$$

$$a : \text{Allocated time per day (minutes)}$$

$$t : \text{Time needed for products to arrive to welding stations from Press 1 (minutes)}$$

Linear Model:

$$\min \sum_{z \in Z} \theta_z \quad (1)$$

$$\text{subject to: } s_1 = 0 \quad (2)$$

$$s_i = s_{i-1} + p_0 * x_{i-1} + c \quad \forall i \in I - \{1\} \quad (3)$$

$$x_i \geq l \quad \forall i \in I \quad (4)$$

$$\theta_z \geq m_{zn} + p_z * x_i - a \quad \forall (z, n, i) \in F \quad (5)$$

$$\sum_{i \neq j} x_i + x_j \geq D_l \quad \forall (i, j) \in S \quad (6)$$

$$m_{zn} \geq s_i + t \quad \forall (z, n, i) \in B \quad (7)$$

$$m_{zn} \geq m_{z(n-1)} + p_z * x_i \quad \forall (z, n, i) \in M \quad (8)$$

$$x_i \geq 0 \quad \forall i \in I \quad (9)$$

$$s_i \geq 0 \quad \forall i \in I \quad (10)$$

$$m_{zn} \geq 0 \quad \forall z \in Z \quad \forall n \in N \quad (11)$$

$$\theta_z \geq 0 \quad \forall z \in Z \quad (12)$$

Appendix C Mathematical Model for Connecting Rod Line

Sets:

I : Each day of a month $I = \{1, 2, \dots, 31\}$

Decision Variables:

w_i : WIP level during day $\forall i \in I$

p_i : The production amount for day i $\forall i \in I$

Parameters:

D_i : Demand of subsequent station in day i $\forall i \in I$

c : Daily production capacity

s : Safety stock

$k_i = \begin{cases} 1 & \text{if day } i \text{ is a workday} \\ 0 & \text{otherwise} \end{cases} \quad \forall i \in I$

Linear Model:

$$\min \sum_{i \in I} w_i \quad (1)$$

$$\text{subject to } p_i \leq c \times k_i \quad \forall i \in I \quad (2)$$

$$w_i \geq s \quad \forall i \in I \quad (3)$$

$$D_i \leq w_i + p_i \quad \forall i \in I \quad (4)$$

$$w_i = w_{i-1} + p_{i-1} - d_{i-1} \quad \forall i \in I \quad (5)$$

$$p_i \geq 0 \quad \forall i \in I \quad (6)$$

$$w_i \geq 0 \quad \forall i \in I \quad (7)$$

Appendix D User Interface



Figure 6: Screenshots from UI

Appendix E Comparison of the Results with the Current System



Tedarik Zinciri Kompleksite Analiz Aracı Geliştirme Projesi

DHL Supply Chain

group03/group_photoDHLSC.jpeg

Proje Ekibi

Sıla Aksoy, Fatih Selim Aktaş, Orçun Asıl,
Berçem Dicle Canlı, Göktuğ Doğan, Selin Uryan, Ecem Uzun

Şirket Danışmanı

Niyazi Ömer Usta
Veri ve İş Zekası Analisti
Can Yörür
Değer Yaratma Bölge Müdürü

Akademik Danışman

Prof. Dr. Bahar Yetiş Kara
Endüstri Mühendisliği Bölümü

ÖZET

DHL LLP (Lider Lojistik Ortağı), iyileştirme ve maliyet tasarrufu için müşterilerinin tedarik zincirinde kompleksite yaratan unsurları belirlemeyi amaçlamaktadır. Ancak, kompleksiteye neden olan unsurların belirlenmesi şirket için oldukça zaman alan bir süreçtir. Bu projenin temel amacı, DHL LLP'nin tedarik zincirlerinde kompleksite yaratan unsurları ve iyileştirme noktalarını belirlemek için harcadığı zamanı azaltmak amacıyla tedarik zincirinin olası kompleksite etmenlerini analiz eden bir sistem geliştirmektir.

Anahtar Kelimeler: Tedarik Zinciri Kompleksitesi, Entropi, Tahminleme, Ağ Analizi, Kompleksite Etmenleri

Development of Supply Chain Complexity Analysis Tool

1 Company Information

DHL Supply Chain, which is part of the DPDHL (Deutsche Post) Group, is one of the most prominent logistics service providers in the world. It offers specific, proven expertise in the Engineering and Manufacturing, Life Sciences and Healthcare, Retail and Technology sectors. DHL Supply Chain increases efficiency and quality while also creating a competitive advantage by combining value-added and management services with traditional order management and distribution (DHL Supply Chain, 2020).

2 Brief System Description

DHL utilizes a framework called Lead Logistics Partner (LLP) which gives advanced digital-integrated supply chain management solutions. LLP has 4 fundamental modules which are design, manage, operate, and continuous improvement. *Design* module contains several sub-modules that provide strategic approaches about warehouse and transportation network development. These create market strategies and reduces costs by 15-20% on average for contracted lifetime. *Manage* module is utilized for sourcing of warehouses, transportation, managing third-party logistics, procurement and quality. *Operate* module controls all operational units including control towers, vendors, warehouses, ocean and air freight (DHL LLP, 2020). *Continuous improvement* is a data-driven module that enables DHL LLP to communicate all users with one database to standardize their performance reporting, total logistics cost management and inventory analytics. The module assists the global partner network to manage process reviews that eliminate inefficiencies and identify continuous improvement opportunities.

3 Problem Definition

According to Isik (2009), supply chain complexity can be defined as all operational uncertainties and structural varieties caused by the information and material flows along the supply chain partners. Currently, DHL LLP has a few methods to identify the problems in the supply chain of the customer, however, it is still a difficult and time-consuming task. Thus, DHL LLP needs to pinpoint the complexity drivers for each supply chain separately. According to our industrial advisor, it takes approximately 90 days for each project to determine these drivers. Hence, it is a very time-consuming task to identify these drivers without a quantifiable measurement tool. Furthermore, there is no standard outlook to evaluate the supply chain complexity which makes it time-consuming to adapt to different systems. Since each system may have a different complexity driver that could increase the cost, there should be a decision support system that will guide engineers to focus on which complexity driver to implement the related improvement lever that can decrease the supply chain's complexity.

4 Proposed System Approach

The Complexity Measurement Tool (CMT) that will be delivered to DHL LLP may be referred as *system* in the following sections in the report. CMT is composed of three modules which are *Forecasting*, *Network* and *Transportation* modules. These modules are decided based on the work scope of DHL, which is decomposed into focused part of supply chain. Also, the data that DHL LLP has provided may be referred as *shipment data* or *transportation data*. The input of the system is the shipment data. The outputs of the system are the complexity levels for each module.

4.1 Forecasting Module

We apply general time series analysis in forecasting and hence we used the past values up to certain look-back windows in our methods. Four subpackages are suggested to check whether forecasting a variable is a difficult task; Autoregressive Integrated Moving Average (ARIMA), Neural Networks, Discrete Fourier Transform with extensions and conditional polynomial regression.

In our analysis, we split data into 3 parts; training, validation and test data. Our models are trained on the training data and within each model, best set of hyper-parameters are chosen based on performance on the validation set. Finally, best model is chosen among models that are validated for each method and based on performance on the test data which shows the generalization power of the method to unseen dataset.

After our analysis; if the standard deviation of the noise is larger than a third of mean value of the forecast, we will classify it as complex. This coincides with testing significance of the forecasts at significance level $\alpha = 0.0025$ assuming Gaussian Noise. Furthermore, even if the noise is not Gaussian, we can still use this test because Gaussian distribution has the maximum entropy for a given variance (Cover and Thomas, 2006, p. 411). Hence, if the forecasts are significant assuming Gaussian noise, it will be significant for any other distribution.

In order to do find meaningful results and make comparisons from complexity caused by volatility of variables DHL can identify important variables and calculate complexity values for each variable. After that, a weighted average can be computed based on importance of the variables that will be determined by DHL.

Autoregressive Integrated Moving Average (ARIMA)

Autoregressive Integrated Moving Average model is a well known statistical model used in time series analysis (Hyndman and Athanasopoulos, 2018, p. 223). General ARIMA model with order (p,d,q) has the following form;

$$y_t^{(d)} = \sum_{i=1}^p \phi_i y_{t-i}^{(d)} + \sum_{j=1}^q \theta_j \varepsilon_{t-j} + \varepsilon_t \quad (1)$$

where y_t denotes the value of variable of interest, ε_t denotes the error term at time t and $y^{(d)}$ denotes d^{th} order differencing. However, ARIMA may give very limited results if there are hidden variables. Hence, natural extension of ARIMA is seasonal ARIMA with exogenous variables (SARIMAX). Thus, SARIMAX models may be deployed as well.

Neural Networks

If ARIMA fails to explain the data, it may be caused by the nonlinear relationships with the past values. These nonlinear relationships can be identified by complex forecasting methods like neural networks (Hyndman and Athanasopoulos, 2018, p. 335) which has theoretical approximation guarantees (Cybenko, 1989; Wang et al., 2020). Our system will be using Long-Short Term Memory (LSTM) architecture because we are doing time series analysis which LSTM is best suited for.

Nevertheless, neural networks are known to overfit the data which could make predictions meaningless (Hastie et al., 2001, p. 398). To avoid it, we added regularization terms using elastic-net. The final loss function of the neural network is as follows;

$$\min_{\theta} L(\theta) = E(\theta) + \beta_1 \|\theta\|_1 + \beta_2 \|\theta\|_2 \quad (2)$$

where θ denotes the network parameters, $L(\theta)$ is the full loss function with the regularization, $E(\theta)$ is the least squares error and β 's are the regularization parameters.

Discrete Fourier Transform

Fourier Transform gives the frequency spectrum of a signal which means that it identifies the signals underlying observations (McClellan et al., 2003, p. 308). If computed values have spikes, it can be interpreted as variable behaving in a seasonal fashion. However, if the frequency spectrum of forecast errors or response value does not have spikes, it can be interpreted as modeling of noise. This noise is inherent to the data creating process of response value. If the data creating process is noisy, it means that forecasting this particular value is difficult and adds complexity to the system.

Similar to ARIMA and SARIMAX, one natural generalization of Fourier Transform is Laplace transform. Our system performs Laplace transform for a unique value of exponential parameter found from least squares fit of exponential term. Moreover, after identifying the exponent and sinusoids, a linear regression is done for better approximation. In addition to extension of Fourier analysis with exponential term, a linear term multiplied by a sinusoid is also considered.

For example; assuming that we identified only one sinusoid in the discrete fourier transform and it has period of 5 days. Then the linear regression problem will be in the form;

$$\min_{\beta_{c5}, \beta_{s5}} \|y - \beta_{c5} \cos(\frac{2\pi}{5}t) + \beta_{s5} \sin(\frac{2\pi}{5}t)\|_2^2 \quad (3)$$

where $y, t \in \mathbb{R}^n$, y denotes the forecast value and t is the time steps.

Conditional Polynomial Regression

In forecasting analysis, sometimes only a subset of observations is of interest. For example, if we only want to forecast number of transports between two cities, we may want to use only the relevant data for it, instead of adding exogenous variables because it would require so many parameters such as LSTM and SARIMAX, to embed that information. Therefore, the proposed forecasting methods get very complex, which makes the model overfit and therefore their performance gets poorer. Therefore model is implemented where our system only uses data points conditioned on the given information. Nevertheless, since conditioned data will be generally very small, we proposed that a small model like polynomial regression should be used.

Polynomial term will be the time and since the model should be relatively small and being inspired by cubic splines and their smoothness property according to James et al. (2013), we propose a cubic polynomial to fit the data. Thus, we will be fitting the following model;

$$\min_{\beta} ||y - \beta_0 - \beta_1 t - \beta_2 t^2 - \beta_3 t^3||_2^2 \quad (4)$$

where $\beta = [\beta_0 \ \beta_1 \ \beta_2 \ \beta_3]^T$, $y, t \in \mathbb{R}^n$, y denotes the forecast value and t is the time steps as before. Finally, depending on the data size; larger polynomial, and previously proposed methods may be used and validated.

4.2 Network Module

Density and Connectivity

To our knowledge, there is no previously proposed Laplacian matrix based measures. We propose two measures, density and connectivity, based on Laplacian matrix of the supply chain network to quantify complexity value of it, that can be calculated in the order of $O(n^3)$ and $O(n^2)$ respectively.

We advocate the use of density and connectivity of the undirected graph of a supply chain to measure complexity of network of a supply chain. To measure the density of a graph, we calculate the number of spanning trees in the graph using Laplacian Matrix. Similarly, in order to measure connectivity of the network; we use the algebraic connectivity of its graph which is the second smallest eigenvalue of the Laplacian Matrix (Fiedler, 1973). These values can be used to understand how close the system is to the "perfect" and "chaotic" network. A perfect network is the most dense network possible. Hence a perfect network's undirected graph is fully connected. Moreover, by Cayley's formula; a complete graph with n vertices K_n has n^{n-2} spanning tree (Cayley, 1889, p. 76) and its Laplacian's second smallest eigen value is n . Nevertheless, this definition assumes there is transportation between every center which may not be true. Instead, we could only have transportation from sources to targets and no transport within

the source and target nodes. Therefore we also define "semi-perfect" network. Similar to Cayley's formula, a complete bipartite graph has $n^{m-1}m^{n-1}$ spanning trees and second smallest eigenvalue $\min(m, n)$, where n and m are number of vertices in parts of bipartite graph $K_{m,n}$ or $K_{n,m}$ which is given by Scoins' formula (Scoins, 1962, p. 16) and using Kirchoff Matrix Tree theorem (Kirchoff, 1847). Similarly a network is chaotic if its undirected graph has only one spanning tree. Hence, we define density and connectivity of a graph $G = (V, E)$ as follows;

$$d(G) = \frac{\kappa(G)}{\kappa(P)} \quad , \quad c(G) = \frac{\lambda_{|V|-1}(L_G)}{\lambda_{|V|-1}(L_P)} \quad (5)$$

where $\kappa(G)$ is the number of spanning trees of graph G , P is either a complete graph $K_{|V|}$ or a complete bipartite graph $K_{a,b}$, $a + b = |V|$, λ_i is the i^{th} largest eigenvalue and L_G is the Laplacian matrix of graph G .

We assume dense graphs are better from network perspective because then most of the transportation is done directly which is simpler to manage. Using multiple hubs could make it harder to manage the transportation of a product. Therefore, if the number of spanning trees turns out to be large, it could mean that there are many connections between nodes and most of the transportation is done directly. These information can be used to compare complexity of two different supply chains relatively. List of measures used to classify a network is given in Appendix A.1.

In addition, our system will also calculate which edges are the most important ones based on gradient of the measures defined above;

$$D = \frac{dd(G)}{dL} = L^\dagger \quad , \quad C = \frac{dc(G)}{dL} = v_{|V|-1} v_{|V|-1}^T \quad (6)$$

where D is the gradient with respect to density, C is the gradient with respect to connectivity, $L^\dagger$ denotes the pseudoinverse of the Laplacian matrix, v_i is the i^{th} eigenvector. Derivations can be found in Appendix A.2.

Entropy

Information theory which states that entropy is the expected total amount of information about the system conditions that can measure uncertainty and probability is used by several authors to measure supply chain complexity in the literature (Isik, 2009). Shannon's entropy function is as follows: $H = -\sum_i p_i \log_2 p_i$, where p_i is the probability of observing state i . Cheng et al. (2013) apply Shannon's information theory to analyze the directed network structure of the supply chain and calculate a complexity index.

Complexity index calculation algorithm suggested by Cheng et al. (2013) is used to calculate the complexity level for static network structures of DHL's customers' supply chains. Start locations, intermediate locations and end locations of each shipment in the shipment data of two customers provided by DHL are used to

form directed supply chain network structures in the matrix form. Structural uncertainty of supply chain networks of these customers are calculated with the following formula:

$$C_{st} = R(I, O)TST + H(type)n \quad (7)$$

TST represents the total system throughput and n is the total number of nodes in the network. $R(I, O)$ gives the degree of disorder in a supply chain network and as the value of it increases the system uncertainty will be higher, i.e. the supply chain will be more complex. $H(type)$ gives information about how diverse are supply chain members. As the value of $H(type)$ increases, more information is required to describe the supply chain network structure.

The algorithm takes directed supply chain network structure matrix, which is obtained from the shipment data of customer, as an input and a complexity index value of the network structure is calculated as an output. The complexity index value is not labelled as complex or not complex based on a threshold value, since defining this kind of threshold may not be valid for all customers. Therefore, the complexity value calculated by the algorithm cannot be evaluated by itself. However, the complexity index value can be used to compare the network structures with each other because it has a relative meaning. So, a network structure with higher complexity index value than other is more complex in terms of network structure. In their supply chain analysis, DHL can use this complexity index to compare complexity levels of different customers in terms of network structures.

4.3 Transportation Module

A transportation maturity matrix is developed which provides a clearer view of how well the system is managed and how complex it is, by revealing the advanced and open-to-improvement aspects of the system (Proenca and Borbinha, 2016). The maturity matrix method requires a system to be divided into different categories and different maturity levels are defined for each category. These categories are created based on the system's specific features such as different processes and constraints. The user selects the appropriate maturity level for each category based on its state on the specified category.

To create the maturity matrix, a supply chain transportation process is divided into categories based on the constraints, processes and requirements with the help of DHL experts. Then different maturity levels are defined for each category. These categories are 3PL Procurement & Consultancy, Order & Transport Management, Freight Bill & Audit, Custom Clearance, Complaint Management, Return & Repair Process, Quality management, Labor Management. 6 different levels are defined to quantify each category. Different aspects of each category, such as automation, documentation, standardization and IT system integration are defined in different competencies at each level. Level 0 means the category is not in the scope of DHL's work packages. From level 1 to level 5 competence in defined aspects increases, so the complexity level for the category decreases.

Weights of the categories are calculated according to Analytical Hierarchy Process (AHP) which is a pairwise comparison method. It ensures that importance of the subject is ranked according to an hierarchical order (Karaoglan and Sahin, 2016). To get these results, a form is sent to DHL experts and weights are calculated based on the form results. VBA is used to calculate categorical and total complexity value, which are defined in percentage value, of a customer. The complexity value is calculated with following formula:

$$C = \sum_{j=1}^8 \sum_{i=0}^5 r_{ij} w_j / (1 - \sum_{j=1}^8 A_{0j} w_j) \quad (8)$$

where r_{ij} is the rate of the level i which is under category j , w_j is the weight of the category j and A_{0j} is a binary variable which returns 1 if the level 0 is chosen for category j , 0 otherwise. Total value of weights adds up to 1. Rates of the categories starts from 1 for level 1, which is the worst case for a category, and decreases by 0.2 as the level increases, i.e. the complexity of the category approaches to the optimal case. The weights of the categories that are not included in the business model, i.e not in scope level, are redistributed over other categories to ensure that the weights of included categories are summed up to 1.

Total complexity, in percentage, is printed after the calculations are done. Additionally, the contribution of each category to the total complexity can be observed with a categorical spider diagram and pie chart, which aims to give a ranking of categories based on the individual complexity values. DHL is suggested to prioritize the categories that are in the top positions when they are analyzing a customer's supply chain. Based on the expert suggestions, entropy based complexity values are also provided in this module. Python algorithm is developed that takes variables from the data provided by DHL. These variables have different options and how often and differently these options are used in shipments determines the complexity level of transportation of a customer. Entropy values for selected variables are calculated separately and results are provided in descending order of variables. The entropy values are between 0 and 1. As the result gets closer to 1 that means the variable selected has too much uncertainty and contribution to total complexity.

5 Verification and Validation

Verification of the algorithms are done with the shipment data sets of 3 customers of DHL which will be referred as customer A, B and C. Each method for each module is tested and ensured that everything works properly.

Forecast

For forecasting, number of daily transportation count and daily transportation in loading meters are given as valuable information that DHL may want to know their values for their daily works by the DHL experts and hence we predicted them in time series fashion. Based on the complexity levels of each different

method used in forecast module, the best method to forecast stated variables is found. In addition, the library we developed in Python for Forecast module automatically generates plots of forecasts for the best method of each variable for visualization. Figure 1 in Appendix B.1 shows the forecast plot of best method found by the algorithm for daily transport count. Default hyperparameters used in the algorithms may change to obtain better results, at the cost of computational time.

To verify conditional forecasting methods, transport count and loading meters for a route between Thayngen and São José Dos Campos, which has a significant number of shipment, are forecasted in a time series fashion although observations are not equally spaced. Similar analysis is done for the forecasting part. Validation is not applicable for forecast module because it is sufficient to compare the predicted values with the actual values and observe the performance of the model. Additionally, DHL has no prior forecasting practice, therefore it is not possible to compare the performance of the methods with DHL’s forecasting methods.

Network

For verification of the network module, stage data of Customer A is used because it was the first data that DHL sent. Stage data contains starting and ending points, the intermediate locations such as hubs for each shipment.

Density and connectivity of network structures are analyzed and the results are stated between 0 and 1. Threshold values were not defined to label a network structure as complex based on the result that we obtained, because such thresholds may not be valid for every customer. Additionally, DHL experts wanted to see the results without any labels as complex and not complex, they wanted to interpret the results as they are. Customer A appears to not have a dense and connective network structure which states that the network structure may be defined as complex. The effects of the hubs for Customer A is also measured. The value obtained from the analysis was negative, where we generally expect these measures to be small positive numbers. In general, hubs make the graph more sparse, not more dense unless the hubs are highly connected with each other which increases the density of the graph. Since the number of hubs in Customer A’s network is high, and the transportation is done by using hubs, this could be the reason of the increase in the complexity of supply chain network.

Network analysis for customers A and B, which are the ones that DHL provided complete data for network analysis, are done and results are shared with DHL experts. The results of each measure is interpreted as the closeness to boundary values; 0 and 1. A result close to 1 for divergence from dense graph, divergence from bipartite graph and dis-connectivity from dense graph measures are interpreted as complex for all customers. DHL experts validated that the results of these two measures are as they are expected, since the two customers have complex network structures. For the measures that include hub effect we obtained

negative values for both customers which means connectivity of hubs negatively affected network structures for these customers. DHL experts validated this result by confirming the complicating effect of hubs on network structures for both customers. Since entropy value has a relative meaning, the entropy values of two customers can be directly used for comparison of two network structures. DHL experts stated that since the entropy value is calculated based on the information theory, they cannot interpret the result with their own experience and intuition but they can use the value in their analysis.

Transportation

DHL experts filled the questionnaire for Customer A, B and C based on the characteristics of customers and then the complexity levels are printed in a percentage value. To provide insights for the end-users, spider diagrams and pie charts are created. These charts demonstrate the contribution of each complexity driver to the overall complexity. Figures 2 and 3 in Appendix B.2 are the real-time outputs of the complexity matrix, after all the inputs are entered. In Appendix B.2, the distribution of the complexity among categories can be observed in Figure 3, where the Return & Repair Process is the leading complexity driver, followed by Custom Clearance.

For validation, we provided a ranking of categories based on the contribution to the total complexity level. Since the weight of each category is calculated using the AHP method and the complexity matrix is filled by DHL experts, the results are very much in line with their expectations. Spider and pie charts provided for each customer are also found explanatory and effective to illustrate the results and effect of each category. For entropy calculations, since the entropy value is calculated based on the information theory which has no prior practice in DHL, experts could not interpret the result with their own experience and intuition. However, the results are easy to understand with simplistic visualizations and are compatible with DHL's expectations.

6 User Interface and Implementation

A user-interface is developed in Microsoft Excel VBA for the ease of use. Network, Forecast and Transportation modules are all combined in this user-interface. Since the algorithms for forecast and network modules were written in Python environment, they needed to be called by Excel VBA. Moreover, the output of the algorithms were printed on a HTML file which is automatically directed by the VBA. In the home-page of interface, the path of the python.exe file has to be given as compulsory input to run the forecast, network and entropy analysis. Moreover, the home-page is also a navigation panel where the user can be directed to 3 different modules that are transportation, network and forecast, respectively and also to settings and the user-manual. The home screen of interface can be found in Figure 4 at Appendix C.

In the transportation module, a user can fill out the maturity matrix that has

8 different categories and each category has 6 states where every category and level is explained in detail. User can choose one level for each category based on the maturity of customer's supply chain. After filling out the maturity matrix, the user will receive a complexity level percentage as an output which can be used to compare different supply chains. Furthermore, the user may print out the respective categorical spider and pie charts to pinpoint the more complex parts in a supply chain. A part of transportation complexity matrix can be found in Figure 5 at Appendix C.

In the network module, there are 2 different complexity analysis which are entropy-based and network. For the entropy-based complexity analysis, the user first needs to input the path of the python script of entropy analysis and a customer's shipment data. Then, the retrieve button shows the variables that can be inputted for entropy analysis. After that, the user may choose the variables that he/she wants to calculate the level of entropy for. Similarly, in the network part, the user needs to input the path of python script of network analysis and a customer's stage data. The results of both of these analysis are printed on an HTML file. Network screen of the interface can be found in Figure 6 at Appendix C.

In the forecast module, there are 2 different complexity analysis which are forecast and conditional forecast. For both analysis, the path of customer's shipment data and python scripts of forecast and conditional forecast are given as input. Then, the retrieve/publish button brings the variables that can be predicted and given as features. After choosing the variables to be predicted and to be given as feature, the respective categories and values are displayed depending on which analysis the user wants to make. Lastly, user has the opportunity to specify the speed for running algorithms which determines the comprehensiveness of the analysis. After all inputs are given, VBA runs the python code and the output is displayed on an HTML file which is automatically directed. The output both illustrates the results graphically and numerically. Forecast screen with example inputs can be found in Figure 7 at Appendix C.

7 Benefits to the Company

The decision support system provides insights about the complexity of three different modules in less than 30 minutes which is significantly lower than the regular analysis process which may take up to 90 days. Niyazi Usta from DHL states that the CMT will enable them to analyze supply chain in a shorter time and detect the complexity drivers. Besides saving time for analysis of the supply chain of DHL's customers, the CMT provides an indirect positive cash flow for DHL. By reducing the time spent for analysis part, the required labor may decrease depending on the complexity of the customer's supply chain. The CMT will be guiding the experts to the most complex parts of the customer's supply chain.

8 Conclusion

The expectation of the company was a tool that can pinpoint complexity drivers by examining different parts of the supply chain. The main aim of the CMT is to shorten the time spent on the analysis of the customers' supply chains. The CMT provides a combination of quantitative and qualitative methods to analyze supply chain complexity. The CMT consists of three modules. Each module aims to calculate complexity independently. Since the uncertainty is an important factor in literature for supply chain complexity, the forecast module measures the complexity by calculating the standard error of the forecast variable while the network module measures the complexity by calculating the density, connectivity and entropy of the network structure. Apart from the other modules, the transportation module analyses the complexity by a complexity matrix which is designed for DHL's business model. The matrix is filled for each customer separately and returns a complexity level which can be used to compare different customer's supply chains. In the end, CMT will provide significant insights to DHL LLP with the outputs of network, forecast and transportation modules.

8.1 Further Improvements

Since there has not been a complexity measurement practice before, CMT can be considered as a prototype. The possible improvements are as follows. For the forecasting module, additional periods such as weekly, monthly and yearly can be added. Moreover, new forecasting methods can be added to improve the quality of the output. For the network module, new measures can be added to broaden the perspective in network analysis. For the transportation module, new categories can be added to the maturity matrix, therefore the output can reflect a deeper analysis. Also, scope of the categories can be broaden to increase the accuracy of the output. The output of the data and comparison of different outputs can be shown in a matrix analysis format to improve the visualisation of the data and enable end users to compare the complexity levels of different customers.

References

- A. Cayley. A theorem on trees. *The Quarterly Journal of Pure and Applied Mathematics*, 23:376–378, 1889.
- C. Cheng, T. Chen, and Y. Chen. An analysis of the structural complexity of supply chain networks. *Applied Mathematical Modelling*, pages 2328–2344, 2013. doi: 10.1016/j.apm.2013.10.016.
- T. M. Cover and J. A. Thomas. *Elements of Information Theory (Wiley Series in Telecommunications and Signal Processing)*. Wiley-Interscience, USA, 2006. ISBN 0471241954.
- G. Cybenko. Approximation by superpositions of a sigmoidal function. *Mathematics of Control, Signals, and Systems*, 2:303–314, 1989. doi: <https://doi.org/10.1007/BF02551274>.

- DHL LLP. Dhl supply chain lead logistics partner portfolio march 2020, 2020.
- DHL Supply Chain. About dhl supply chain, 2020. URL <https://www.dhl.com/tr-tr/home/bolumlerimiz/tedarik-zinciri/dhl-tedarik-zinciri-hakkinda.html>.
- M. Fiedler. Algebraic connectivity of graphs. *Czechoslovak Mathematical Journal*, 23(2):298–305, 1973. URL <http://eudml.org/doc/12723>.
- T. Hastie, R. Tibshirani, and J. Friedman. *The Elements of Statistical Learning*. Springer Series in Statistics. Springer New York Inc., New York, NY, USA, 2001.
- R. J. Hyndman and G. Athanasopoulos. *Forecasting: principles and practice*. OTexts, Melbourne, Australia, 2nd edition, 2018.
- F. Isik. An entropy-based approach for measuring complexity in supply chains. *International Journal of Production Research*, 48(12):3681–3696, June 2009. doi: 10.1080/002075409028105938.
- G. James, D. Witten, T. Hastie, and R. Tibshirani. *An Introduction to Statistical Learning: with Applications in R*. Springer, 2013. URL <https://faculty.marshall.usc.edu/gareth-james/ISL/>.
- S. Karaoglan and S. Sahin. Dematel ve ahp yöntemleri İle İşletmelerin satın alma problemlerine bütünlük bir yaklaşım, dslr kamera Örneği. *Journal of Business Research Turk*, 2016. doi: 10.20491/isarder.2016.183.
- G. Kirchhoff. Über die auflösung der gleichungen, auf welche man bei der untersuchung der linearen verteilung galvanischer ströme geführt wird. *Annalen Der Physik und Chemie*, 72:497–508, 1847. doi: <https://doi.org/10.1002/andp.18471481202>.
- J. McClellan, R. Schafer, and M. Yoder. *Signal Processing First*. Number 1. c. in Signal Processing First. Pearson/Prentice Hall, 2003. ISBN 9780130909992. URL <https://books.google.com.br/books?id=j81pQgAACAAJ>.
- D. Proenca and J. Borbinha. Maturity models for information systems - a state of the art. *Procedia Computer Science*, 100(12):1042–1049, 2016.
- H. I. Scoins. The number of trees with nodes of alternate parity. *Mathematical Proceedings of the Cambridge Philosophical Society*, 58:12–16, 1962. doi: <https://doi.org/10.1017/S0305004100036173>.
- Z. Wang, A. Albarghouthi, and S. Jha. Abstract universal approximation for neural networks, 2020.

Appendix

Appendix A Network Appendices

A.1 Complexity Measures for Network

$$m_1(G) = \frac{\log(\kappa(G))}{(|V| - 2) \times \log(|V|)} \quad (9)$$

$$m_2(G) = \frac{\log(\kappa(G))}{(m - 1) \times \log(n) + (n - 1) \times \log(m)}, m + n = |V| \quad (10)$$

$$m_3(G) = \frac{\lambda_{|V|-1}(L_G)}{|V|} \quad (11)$$

$$m_4(G) = m_1(G) - m_1(\hat{G}) \quad (12)$$

$$m_5(G) = m_2(G) - m_2(\hat{G}) \quad (13)$$

$$m_6(G) = m_3(G) - m_3(\hat{G}) \quad (14)$$

$G = (V, E)$ is the graph of supply chain network, $\hat{G} = (\hat{V}, \hat{E})$ is the graph of supply chain network ignoring hubs, i.e every transport is assumed to be direct. $\hat{G}$ is constructed by transportations' initial and final destination only.

We compute relative closeness of density and connectivity of networks to complete and complete bipartite graphs. However, for numerical concerns some of the computations are done in logarithmic scale. If there are multiple subgraphs that are disjoint, i.e it is not possible to find a spanning tree for the original problem, the previously proposed methods are applied to every disjoint subgraph and weighted averages based on number of nodes of each graph are taken. In addition, this process is repeated for the modified network, where hubs are removed and transports are assumed to be done directly. Finally, difference between the original and modified networks' density and connectivity measures, effect of hubs on the supply chain network is computed.

A.2 Gradient Calculation

Let $\kappa(G)$ denote number of spanning trees of graph G . We can ignore the denominator in the formula of measure m_1 , see Appendix A.1, since it is constant. By Kirchoff Matrix Tree theorem (Kirchoff, 1847);

$$\kappa(G) = \frac{1}{n} \prod_{i=1}^{n-1} \lambda_i \quad (15)$$

where λ_i is the i^{th} dominant eigenvalue of Laplacian matrix of graph G and $n = |V|$. Let L denote Laplacian matrix of graph G and D denote gradient of this function. Taking the derivative of this function with respect to matrix entries;

$$D = \frac{d\kappa(G)}{dL} = \frac{1}{n} \times \sum_{i=1}^{n-1} \frac{d\lambda_i}{dL} \left[\prod_{j \neq i}^{n-1} \lambda_j \right] = \frac{1}{n} \times \sum_{i=1}^{n-1} v_i v_i^T \left[\prod_{j \neq i}^{n-1} \lambda_j \right] \quad (16)$$

where v_i is the eigenvector corresponding to i^{th} dominant eigenvalue. Moreover we can modify and simplify the formula because we are not interested in actual magnitude of gradient, relative importance and sign of the gradient is sufficient. So multiplying the equation by $\frac{1}{\kappa(G)}$ gives;

$$\hat{D} = \frac{\widehat{d\kappa(G)}}{dL} = \sum_{i=1}^{n-1} \frac{1}{\lambda_i} v_i v_i^T = L^\dagger \quad (17)$$

where $L^\dagger$ denotes the pseudoinverse of the Laplacian matrix. Similarly for the gradient of the connectivity measuring function after ignoring constant;

$$C = \frac{d\lambda_{|V|-1}(L_G)}{dL} = v_{|V|-1} v_{|V|-1}^T \quad (18)$$

Appendix B Verification Results

B.1 Forecast Results on Customer A's Shipment Dataset



group03/jj_forecast_count_mk2-eps-converted-to.pdf

Figure 1: Forecasts for Daily Transport Count on Unseen Data

B.2 Transportation Results on Customer A’s Shipment Dataset

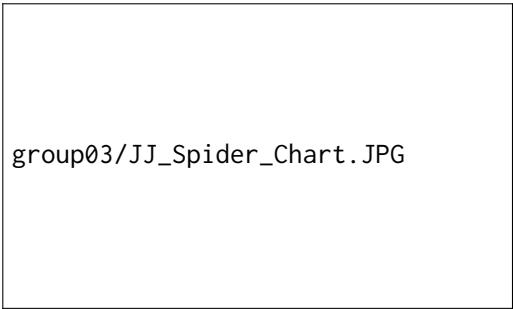


Figure 2: Categorical Spider Diagram

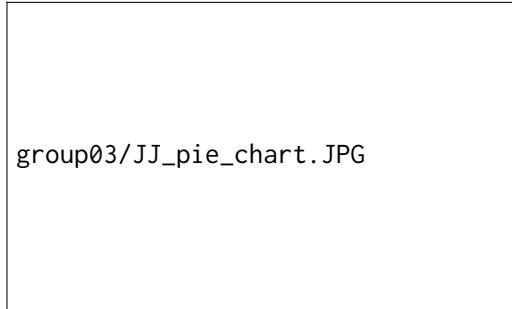


Figure 3: Categorical Pie Chart

Appendix C User Interface

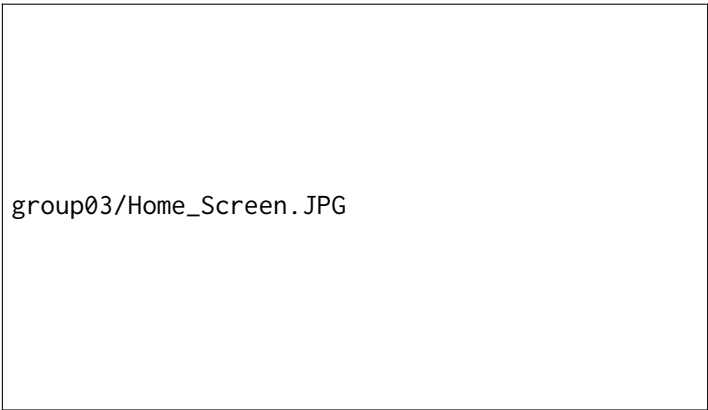
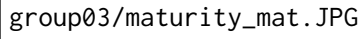
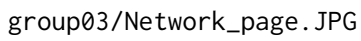


Figure 4: Home Screen

A screenshot of a software interface titled "group03/maturity_mat.JPG". The interface is mostly blank, suggesting a data visualization or analysis tool that is not fully rendered in this view.

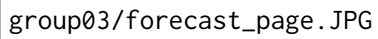
group03/maturity_mat.JPG

Figure 5: Transportation Analysis Screen

A screenshot of a software interface titled "group03/Network_page.JPG". The interface is mostly blank, suggesting a data visualization or analysis tool that is not fully rendered in this view.

group03/Network_page.JPG

Figure 6: Network Analysis Screen

A screenshot of a software interface titled "group03/forecast_page.JPG". The interface is mostly blank, suggesting a data visualization or analysis tool that is not fully rendered in this view.

group03/forecast_page.JPG

Figure 7: Forecast Analysis Screen

Elektriksel Güvenlik Test Sonucu Tahminleme Sistemi Tasarımı

Arçelik A.Ş. Kompresör Fabrikası

group04/group_photo.png

Proje Ekibi

Zeynep Akdağcık, İlayda Aksoy, Haya Alkolak, Muammer Akın Aydoğan, Ecem Budak, Mahmut Sefa Ergündüz, Selin Özen

Şirket Danışmanı

Kal. Kont. Müh. Mücahit Olmaz

Endüstri Mühendisliği Takımı

Akademik Danışman

Prof. Dr. Mustafa Çelebi Pınar

Endüstri Mühendisliği Bölümü

ÖZET

Arçelik Kompresör Fabrikası'nda üretilen kompresörlerin stator kısmını oluşturan parametrelerin, kalite kontrol departmanı tarafından yapılan Hi-Pot testlerini nasıl etkilediği bilinmemektedir. Şirket, stator parametrelerini iyileştirerek ürünlerin sahada arıza oranını düşürmek ve yeni stator tasarımlarında Hi-pot test sonuçlarını sahaya inilmeden öngörebilmek istiyor. Bu proje, statorun 27 tasarım parametresi ile Hi-Pot testinin sonuçları arasındaki ilişkiyi gösteren matematiksel bir modelin yanı sıra, şirketin gelecekte yeni veriler için kullanacağı bir tahmin sistemi oluşturmayı amaçlamaktadır. Bu raporda mevcut sistem ve örnek veriler incelenir, çözüm yaklaşımları sunulur, uygulanan metodoloji açıklanır, sonuçlar ve elde edilen modeller analiz edilir ve bir proje planı sunulur.

Anahtar Kelimeler: Matematiksel Modelleme, Veri Analizi, Regresyon Analizi, Elektriksel Güvenlik, Hi-pot Testi

Predictive System Design for The Electrical Safety Test Results

1 Company Information

Arçelik A.Ş is a global company which is founded in 1955. It is one of the leading companies in both Turkey and foreign countries. It maintains its manufacturing services in the white goods industry, producing a wide range of appliances like refrigerators, washing machines, dishwashers, compressors, etc. The company has proved its globality with its wide range of export and import network among 145 countries. Considering its 50% rate of domestic market share in the white goods industry, Arçelik is a major company in Turkey Arcelik (2020). It has both international competitors, like Haier, Samsung, Electrolux, Bosch, Gorenje, and LG, and local competitors like Vestel. In Arçelik Eskişehir Compressor Plant, the total number of workers is 400, 75 of them belong to white-collar workers. The annual production capacity of compressors in Arçelik Eskişehir compressor plant is nearly 3 million Arcelik (2020).

2 System Analysis & Problem Definition

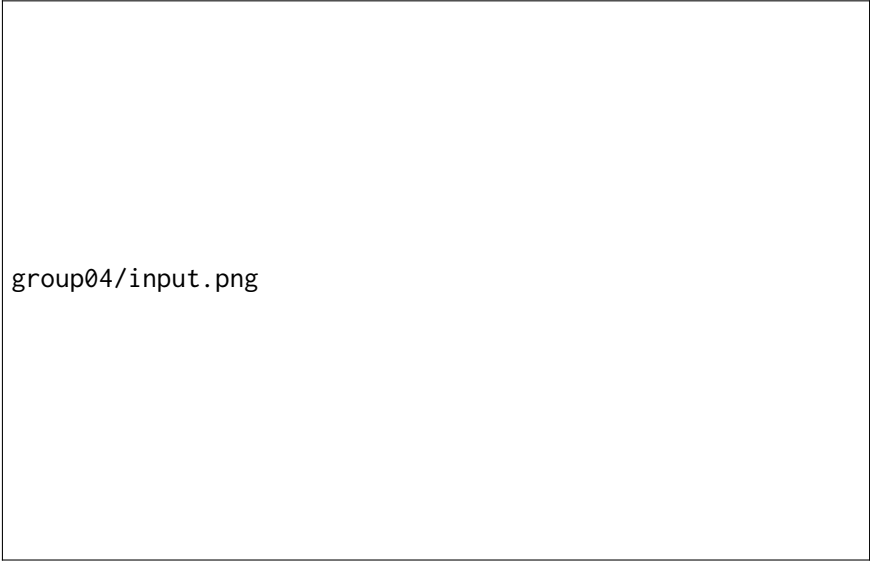
2.1 Analysis of the Current System

In the current system, the Entrance Quality Control Department is responsible for conducting quality tests and measurements on the produced compressors. The electrical safety of the compressors produced by Arçelik is defined according to the results of the Hi-Pot test, which was mentioned in the previous chapter. The leakage current is measured by applying a high AC voltage between the windings and the core. Thus, the electrical short circuits formed between the winding and the core are tested. The conditions of the test should be as follows; the applied voltage is 2000V, the duration of the test is one second and the maximum current is 1 mA . The leakage current measured against the applied voltage should be a maximum of 1 mA EDC (2020). In the end, a verifying procedure is made during the test to make sure the leakage current ($I_{HV\ res. [mA]}$) results are lower than the maximum set value ($I_{HV\ res. \ max. \ [mA]}$). In case the leakage current is higher than the maximum acceptable set value, the test will fail. If a failed stator is detected at the end of Hi-Pot test, 2 different procedures are applied. If there is an easily accessible wire waste or a replaceable insulating paper error in defective stators, it is repaired and put to the Hi-Pot test again. In other cases, all defective stators are scrapped. For this test, some parameters of the stator are used, 27 of them being the main context of the project. These parameters are the stator core height; plastic wrap holder pin diameters; thickness, length, and width of isolation paper; stator solidity ratio; thickness of wire enamel; the steepness of stator; plane angle; and the upper, lower, middle diameters of the stator. All of the parameters' quality controls are made during the production lines. By looking at the results of the Hi-pot Tests, the stators' performance is evaluated.

2.2 Problem Definition

In Arçelik, there are mainly two types of problems related to the project. The first problem is arising from the traceability issue. Some materials are both outsourced from the suppliers and produced inside the factory. The ones that are outsourced from the suppliers and controlled by the Hi-pot Test in the Outcoming Quality Assurance Department have no problems. On the contrary, some of the products produced inside have been found defective in the Hi-pot Tests that are conducted by the Incoming Quality Department. When the products are produced inside and then sent to the assembly line, some parameters become difficult to control. The parameters that need to be examined in the project are Core Height; Upper, lower, middle diameters; Inner Diameter Steepness; two Upper Holders (A&B); two Lower Holders (A&B); Thickness; Width. The company desires the tests to be done 100 percent on all materials and be fast. Thus, the Quality Assurance Department cannot trace the products' quality properly. Furthermore, during production (of the stator), there exists some waste of materials, scrap, time, and labor costs that may hinder the improvement of the process. Due to this trace-ability problem, the firm faces another issue concerning the failure rates of the production field. In order to avoid this problem, defective products should be detected before going out of the plant. Even if the results from Hi-pot test are acceptable, there may still be refrigerators returning from the production area.

The problem scope of the project is on the quality optimization of compressors' electrical safety. It is assumed that the parameters of the stator affect the result of the Hi-pot test. However, it was not clear which parameter/s are more significant in affecting the electrical safety test results. This situation created problems for the company to analyze the data of parameters. Although the problems listed above cannot only be solved by mathematical modeling of the Hi-pot test, the project aims to increase the company's control over the process by finding the relation between the parameters. In other words, to optimize their electrical motor design and minimize their material costs, Arçelik desires to know how each of these specified 27 parameters affect the result of the Hi-pot test, and how the parameters are correlated with each other. Therefore, this project intends to emphasize the problem of generating a predictive model that reveals how the result of the Hi-pot test changes depending on the selected parameters, fitting the data obtained from Arçelik and finding correlations between these parameters if there are any. The flowchart of the project that demonstrates the input, the output and the methodology can be seen below:



group04/input.png

Flowchart of the project that demonstrates the input, the output and the solution methods.

3 Proposed Solution Approach

The main suggested approach on the project is to build the mathematical model using different regression techniques. Regression analysis is a predictive modeling technique that investigates the relationship between a dependent (target) and independent variable (predictor) Mirkin (2011). In this case, the target and the dependent variable is the Hi-pot test results data while predictors are the design parameters' data. Therefore, the regression analysis seems suitable for the problem.

There are various kinds of regression techniques available to make predictions based on the number of independent variables, type of dependent variables, and shape of the regression line. In this case, the variables are continuous and there are several independent variables. It is also not for sure that each parameter has an effect on the Hi-pot test results. Hence, it is useful to build the model using Multi-Linear Regression (MLR), Ridge Regression, Least Absolute Shrinkage and Selection Operator (LASSO) Regression, and Elastic-Net Regression. MLR is a statistical technique that uses two or more independent variables to predict the outcome of a dependent variable. Ridge Regression is a technique that shrinks coefficients to non-zero values to prevent over-fitting, but keeps all variables. LASSO is a type of linear regression that uses shrinkage, where data values are shrunk towards a central point, like the mean Oleszak (2019). Lastly, Elastic-Net Regression uses the penalties from both the lasso and ridge techniques to regularize regression models. Ridge Regression, Lasso Regression and Elastic-Net Regression have been used particularly to avoid the multicollinearity problem that arises from the high correlation between some of the parameters Montgomery (2007).

The correlation plot can be found in Appendix A.

3.1 Model Building

Within the scope of the project, only the stator model 20Z1 and the design parameters associated with this stator model are analyzed, upon the request of Arçelik. The dataset provided by the company regarding the stator(20Z1) and the mathematical model consisted of 87 observations, with one dependent variable, which is the Hi-pot Test Result, and 27 independent variables, which are the parameters. Different regression techniques using R were applied to build the model and understand the effect of the parameters on the Hi-pot tests. The Root Mean Square Errors (RMSE) of each model have been compared and are as follows for each model:

- Multi-linear Regression: 0.008774423
- Ridge Regression: 0.009403651
- LASSO Regression: 0.009459675
- Elastic-Net Regression: 0.009403651

The RMSE's of the models are close, all of them are approximately 0.009. However, Linear Model has the least RMSE value. It can be tentatively concluded that the Linear Model gives the most statistically significant results and consequently, becomes the most suitable model for the problem. The parameters have been chosen based on step-wise regression (step-wise selection) consists of iteratively adding and removing predictors, in the predictive model, in order to find the subset of variables in the data set resulting in the best performing model, that is a model that lowers prediction error. In the final model, the thickness parameters turned out to be the most statistically significant parameters affecting the Hi-pot test results. Middle and upper diameters followed the thickness parameters in terms of importance. The importance plot of the parameters can be found in Appendix B. The results have been shared with the industrial advisor and the advisor approved that the results were parallel to the expected results. The advisor agreed that the thickness is known as an important factor affecting the electrical insulation, therefore it is expected that predictors related to thickness turned out to be more significant in the response variable which are Hi-pot test results.

3.2 Model Improvement

Three methods have been considered in order to improve the Linear Regression model: Log-Transformation, Examination of the Diagnostic Plots, and Exploring Interactions. To understand whether the enhancement techniques are useful, Akaike's Information Criteria (AIC), Bayesian Information Criteria (BIC) and Adjusted R-square have been used as metrics to measure the goodness of the model. The basic idea of AIC is to penalize the inclusion of additional variables to a model. It adds a penalty that increases the error when including additional

terms. The lower the AIC, the better the model. BIC is a variant of AIC with a stronger penalty for including additional variables to the model. The lower the BIC, the better the model. The adjusted R-squared is a modified version of R-squared that has been adjusted for the number of predictors in the model. The adjusted R-squared increases only if the new term improves the model more than would be expected by chance so higher the Adjusted R-square, the better the model.

Log-transformation is a data transformation technique that is used to reduce the variability of the data and to make the non-normal data closer to a normal distribution, which is expected to increase the accuracy of the statistical analyses Changyong et al. (2014). The model inputs have been log-transformed in the beginning. Since the AIC and BIC of the model has increased and the Adjusted R-square has decreased, we decided to proceed with the original model Keene (1995).

Exploring interactions involves adding interaction terms to a regression model and can greatly expand understanding of the relationships among the variables in the model and allows more hypotheses to be tested. The resulting multi-linear regression model has been extended by including interaction terms. As a result, most of the combinations turned out to be statistically insignificant. Adding interaction terms to the model turned out to be ineffective since the AIC and BIC of the model has increased while Adjusted R-square has decreased James et al. (2013).

The diagnostic plots have been checked to make sure that the main linear regression assumptions are met in the model. These assumptions are that the model should be linear; the residual errors should be distributed normally; the variance of the residual errors should be constant across all models (homoscedasticity); errors of residuals should be independent. The model satisfies all of the assumptions therefore does not need further improvement. Further details can be found in Appendices C,D,E and F.

3.3 Model Validation and Verification

To validate the model, Cross Validation has been used in the model selection process. Cross validation is a technique that tests how a statistical analysis will result in an independent data set. Its main use is to predict with what accuracy a predictive system will work in practice. In a prediction problem, the model is usually trained with a "known dataset" ("training set") and tested with an "unknown dataset" ("validation set" or "test set"). The purpose of this test is to measure the model's ability to generalize to new data and to identify problems of overfitting or selection bias James et al. (2013).

Besides Cross Validation, the model performance has been tested by a test data of size 10 provided by the company. The test data included the measured parameter values for all predictors of 10 stators and the real Hi-pot test results corresponding

to each of these stators. The test data has been uploaded to the model and the prediction function has been used in order to obtain the Hi-pot test results predicted by the model. The predicted Hi-pot test results are compared with the real Hi-pot test results for each stator and the differences between the real and the predicted test results have been checked. Root Mean Square Error has been used as a measure of accuracy using the below formula:

$$RMSE = \sqrt{\sum_{i=1}^n \frac{(\widehat{y}_i - y_i)^2}{n}}$$

where $\widehat{y}_i$ is the forecast value, y_i is the actual value and n is the sample size. According to the results, it is observed that the model is able to predict the Hi-pot test results with a root mean square error (RMSE) of 0.008938686. Normalized RMSE has been computed by dividing the RMSE value by median of the Hi-pot values. The normalized RMSE turned out to be 0.01895396. This confirms that the model's performance is satisfactory and the model predicts with high accuracy.

The mean absolute percent error (MAPE) has also been used as another metric which expresses accuracy as a percentage of the error by using the formula below:

$$MAPE = \frac{1}{n} \sum_{t=1}^n \left| \frac{A_t - F_t}{A_t} \right|$$

where A_t is the actual value, F_t is the forecast value and n is the sample size. As a result, MAPE turned out to be 0.01551805 which is approximately 1 percent. Smaller values indicate a better fit so we can conclude that on the average, the forecasts are off by 1 percent.

In addition, the differences between the real and the predicted Hi-pot test results have also been analyzed using two sample t-test, as requested by the industrial advisor. The two-sample t-test is a method used to test whether the unknown population means of two groups are equal or not. In other words, the null hypothesis is that the means of the two groups are equal and the alternative hypothesis is that they are not Montgomery et al. (1982). The test statistic has been computed by using the formula below:

$$t = \frac{\bar{x}_1 - \bar{x}_2}{\sqrt{s^2 \left(\frac{1}{n_1} + \frac{1}{n_2} \right)}}$$

where $\bar{x}_1$ is the sample mean of the first dataset, $\bar{x}_2$ is the sample mean of the second dataset, s is the standard deviation, n_1 is the sample size of the first dataset and n_2 is the sample size of the second dataset. The resulting test statistic is -0.01 and the resulting p-value of the test turned out to be approximately 98 percent,

therefore we fail to reject the null hypothesis on a 5 percent significance level and conclude that there is no significant difference between the real Hi-pot test results and the results predicted by the model.

The validation results have been shared with the industrial advisor and as a result of the statistical analysis conducted by the industrial advisor, it is concluded the model is valid.

In order to verify the model, two random values have been generated for each parameter considering their statistical distributions as well as their lower and upper tolerances. Considering different combinations of the 2 values of the 8 parameters which are included in the model, a dataset has been created containing 256 values. Hi-pot values have been found for each data using the original predictive function. A new linear regression model has been found using the new dataset. The original predictive model and the obtained model have been compared and the results turned out to be exactly the same. This way, it has been verified that the code for prediction is correctly working, hence, the model is correct.

4 User Interface and Implementation Plan

The primary subject of the implementation plan is the predictive system user interface that has been designed. To implement the proposed model, a user interface is designed using R (see Appendix G). As the first step of the implementation plan, the interface is shown to the relevant company employees as a demo in a meeting done at the end of the March. The company did not suggest further improvements and the system is confirmed by the industrial advisor. In the company, it is suggested that the final and confirmed system will first be tested by the quality assurance employees and it is suggested that the company launches explanatory demos/tutorials in order to teach and explain the system's working principle to the employees. Before starting to use the system for the real production line, it is strongly suggested that a sample of stators with measured predictors should be produced in order to test the system result and to compare the system's hypothetical Hi-pot test output with the real Hi-pot test results of the stators. If the test results are reliable enough to implement the system for the real production line, the interface can be used by the company for predictive purposes.

The user interface contains a manual parameter input part as well as a file installment part to upload data. It is planned that when Research and Development department proposes a stator design with new design parameters, the employees of the quality assurance department will first select the stator model from the "Model" menu, and they will enter the hypothetical parameter values to the system as input by entering the measurements for each design parameter. When user clicks to the " See Hi-pot test results button", the system will provide the expected Hi-pot test response as well as the error percentages of the four regression models, and the best model among them as the output. The system will also allow for a new dataset installment in the form of an Excel file, in order to

update the models and data when necessary. In the future, if Arçelik desires to collect new/actual data from the production line in order to update the models, the new file will be uploaded using the "Upload file" button. The system will read the new dataset and run the regression models again to find a new mathematical model which is more suitable to new data. As a result, the system will be more sustainable and will remain up to date.

5 Conclusion and Project Contributions

Considering the problem scope of the project, the company will mainly benefit from the project deliverables in terms of customer/employee safety. For Arçelik, enhancement of customer safety and satisfaction is the most important part of the project. Thus, since the project focuses on understanding and predicting the results of an electrical safety test (Hi-pot), cost optimization is not the primary benefit expected by the company, and Arçelik did not provide any monetary detailed data. When the company understands how the design parameters affect the electrical safety and which of them are significant, they will be able to optimize their stator design in a more reliable way, which will be less likely to result in a current leakage and safer for the customers.

In addition, using the predictive system designed by the group will enable company workers to foresee the expected Hi-pot results based on a given parameter data. The company employees will enter the design parameter measurements and the system will provide the predicted Hi-pot test result based on the input. This means that the company will have a larger perspective and control over the electrical safety test results. In addition to safety regulations, this predictive system design will provide benefits in terms of time and performance usage as well. In the current system, whether the stators designed by the Research and Development department will fail the Hi-pot test or not is understood after the stators are actually produced and tested. A test sample is first designed, produced and sent to the production line and at the end of the line, the stators are tested for electrical safety via Hi-pot test. This process takes around 2-3 months for the company to see the results of the Hi-pot test. After the predictive system is delivered, the company will be able to see the expected results by just entering hypothetical values for the parameters, meaning that they will save from time and performance that was needed for the production and test procedure. The elimination of this production stage will help the company to save from the associated costs as well. Along with that, it is expected that the waste of material and scrap will be reduced by providing a better understanding of parameters affecting the Hi-pot test results, because the number of stators that fail the test and become scrapped will be reduced. This means that, although the cost enhancement is not the primary benefit of the project, it can be said that the implementation of the project will indirectly result in some monetary benefits as well.

References

- Arcelik. Arçelik investor presentation, 2020.
- F. Changyong, W. Hongyue, L. Naiji, C. Tian, H. Hua, L. Ying, et al. Log-transformation and its implications for data analysis. *Shanghai archives of psychiatry*, 26(2):105, 2014.
- EDC. Ac dielectric strength test, 2020.
- G. James, D. Witten, T. Hastie, and R. Tibshirani. *An introduction to statistical learning*, volume 112. Springer, 2013.
- O. N. Keene. The log transformation is special. *Statistics in medicine*, 14(8): 811–819, 1995.
- B. Mirkin. *Core concepts in data analysis: summarization, correlation and visualization*. Springer Science & Business Media, 2011.
- D. Montgomery, E. Peck, and G. Vining. Introduction to linear regression analysis. new york: John wiley and sons.", 1982.
- D. C. Montgomery. *Introduction to statistical quality control*. John Wiley & Sons, 2007.
- M. Oleszak. Regularization: Ridge, lasso and elastic net (2018). URL <https://www.datacamp.com/community/tutorials/tutorial-ridge-lasso-elastic-net>, 2019.

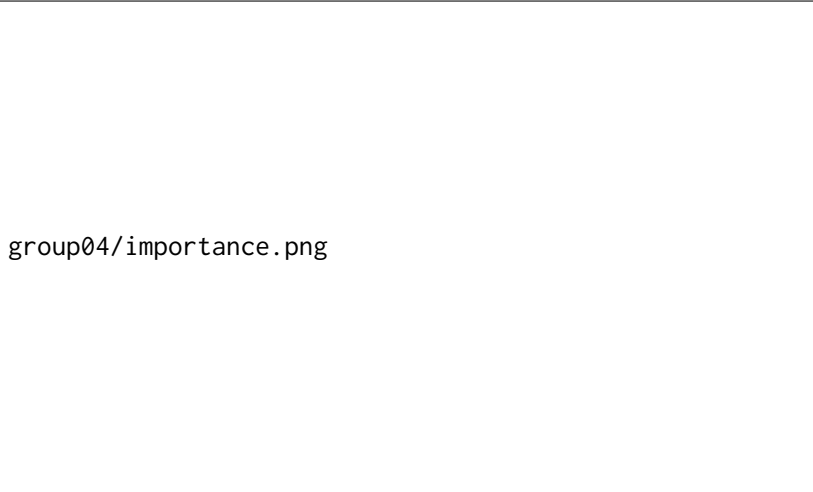
Appendix A Correlation Plot



group04/corplot.png

The correlation plot between 27 parameters in which the dark blue color indicates the high correlation between parameters. When the correlation between the parameters are high, this means that some parameters can be explained by the other parameters and it increases the variance of the model.

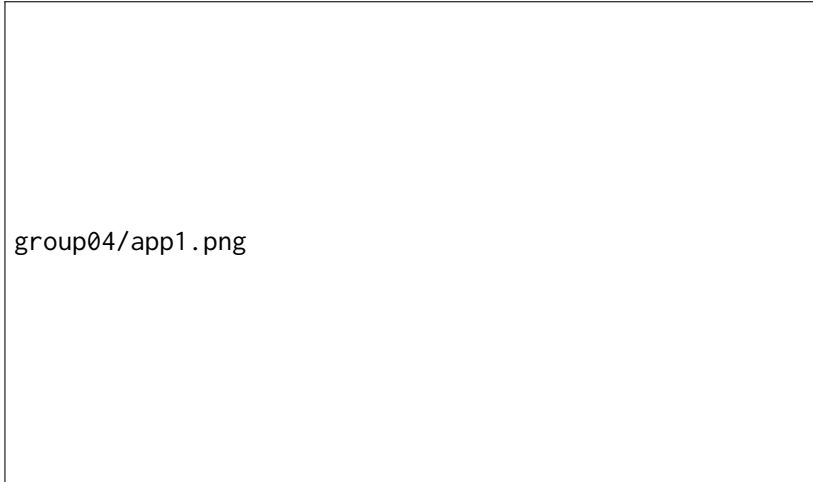
Appendix B Importance Plot



group04/importance.png

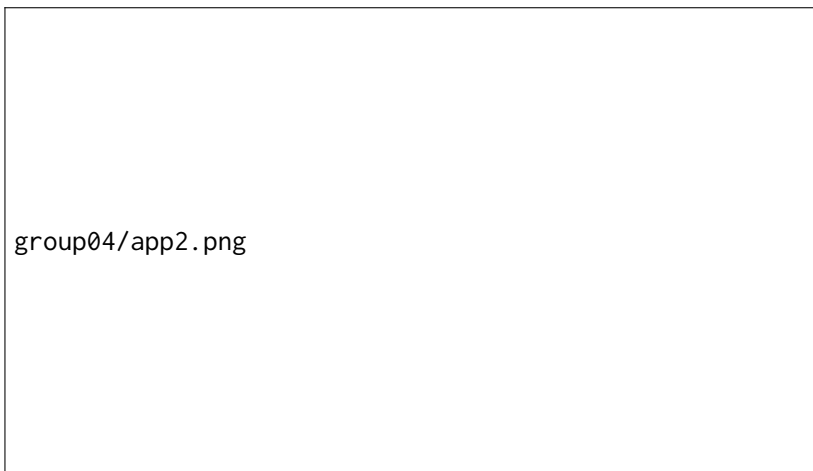
The “varImp” function in R tracks the changes in model statistics, such as the Generalized Cross Validation, for each predictor and accumulates the reduction in the statistic when each predictor’s feature is added to the model. This total reduction is used as the variable importance measure.

Appendix C Residuals Plot



For the linearity of the model, the Residuals vs. Fitted plot is checked if the line is horizontal and straight at 0 value. Since least squares is used to fit the regression model, it is expected the residuals and fitted values to be uncorrelated. In other words, any linear combination of predictor and predictors are not expected. Here, the line follows the straight dashed line, so it satisfies this assumption.

Appendix D Normal Q-Q Plot



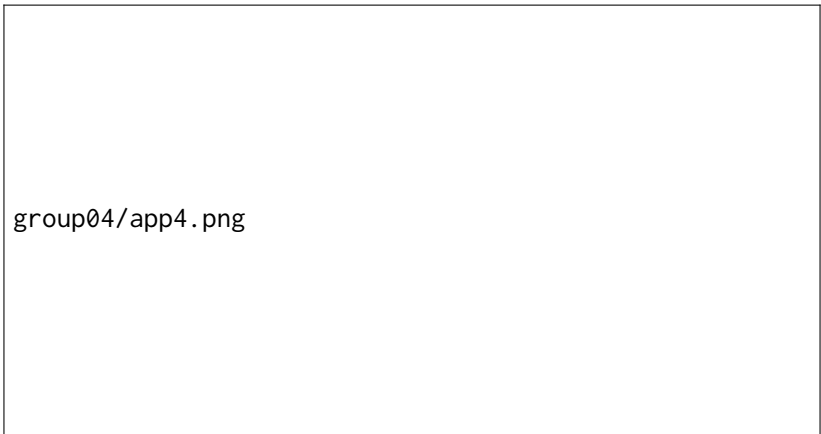
For the normality assumption, Normal Q-Q plot is checked to see whether the residuals are normally distributed. The smoother line again follows the dashed line, so it does satisfy the normality assumption.

Appendix E Scale-Location Plot



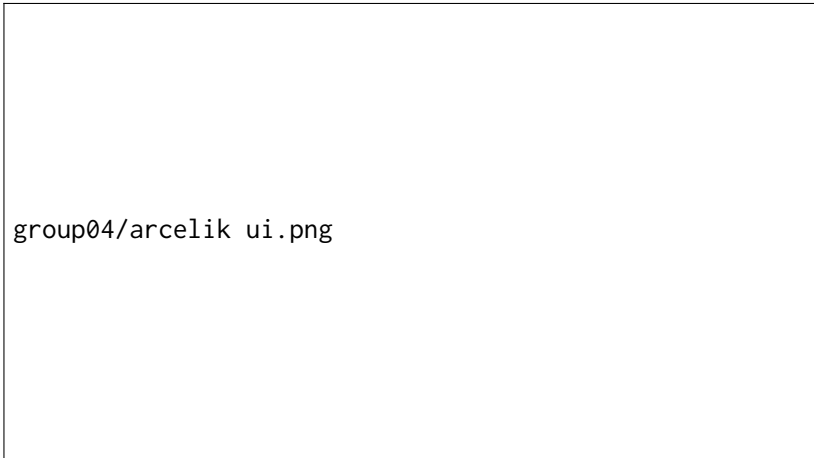
For homoscedasticity, Scale-Location plot is checked to see a horizontal line with homogenic scattered residuals points. The square root of the standardized residuals represents the standard deviation, and it is expected to look stable. The line is almost horizontal so it can be concluded that the variance is constant.

Appendix F Residuals vs Leverage Plot



Residuals vs Leverage plot is checked to examine influential values among outliers/high-leverage points using Cook's distance. There seems to be no influential points.

Appendix G User Interface



User interface where quality assurance engineers in Arçelik can enter parameter values and foresee the hi-pot test results.

Servis Arıza Oranı Tahminlemesi

Arçelik A.Ş. Kompresör İşletmesi

group05/group photo.png

Proje Ekibi

Emir Ataman, Yağmur Censur, İrem Elçin Sayın,
Osman Kağan Yayla, Aylin Yıldırım, Ali Burak Yüksel, Mehmet Mert Yüksel

Şirket Danışmanı

İlker Yılmaz
Endüstri Mühendisliği Takımı

Akademik Danışman

Prof. Dr. Savaş Dayanık
Endüstri Mühendisliği Bölümü

ÖZET

Arçelik Kompresör Fabrikası'nın Kalite Kontrol Departmanı Servis Arıza Oranı adı verilen bir ölçüm ile servis arızalarının sayısını tahmin etmektedir. Bu tahminin doğruluğu şirketin kendi sistemleri üzerindeki yetkinliğinin artması ve üretim hatlarında önleyici ve engelleyici önlem alınabilmesi için önemlidir. Projenin amacı mevcut tahmin yöntemini değerlendirmek, farklı tahmin yöntemlerini mevcut yöntemle karşılaştırmak ve en başarılı istatistiksel sonucu veren yöntemi seçmektir. Bu amaç doğrultusunda sistem analizi ve literatür taraması yapılmıştır. Birçok tahmin yöntemi değerlendirilmiş ve en başarılı yöntem grup tarafından geliştirilmiştir. Belirlenen iş paketleri üzerinden çalışma 2020-2021 yılı içerisinde yürütülmüş ve sistem girdileri ile çıktıları raporda sunulmuştur.

Anahtar Kelimeler: Optimizasyon, Sezonsallık, Tahmin yöntemleri, Trend, Zaman serisi analizi

Service Call Rate Forecasting

1 Identification of the Problem

1.1 System Description and Analysis

Arçelik Compressor Plant started its operations in 1977 as an internal subsidiary of Arçelik A.Ş. Today with its 336 blue-collar and 65 white-collar workers, Arçelik Compressor Plant is able to produce 4 million compressors in a year.

The facility has mainly two different SKUs: Inverter and conventional compressors. The main difference between the two SKUs is inverter compressors can work at different rotational speeds. Therefore, inverter compressors are electronically controlled, working with smart algorithms, and by low voltage consumption, it is a much more energy-efficient system than conventional compressor systems. In detail, there are eight main models and 29 sub-models that are considered as Arçelik Compressor Plant's product.

Three types of datasets are received regarding the production, installation, and complaint. The datasets have been analyzed both individually and intersectionally. It is analyzed that the production amount for each year is greater for the conventional type of compressor. Also, looking at the origin of compressors, it is seen that there are more locally produced compressors than imported ones. General data analysis can be seen in Appendix A.

1.2 Problem Definition

In Arçelik Compressor Plant, one of the most crucial metrics measured is Service Error Rate (SER). SER is defined as the ratio of defective compressors to the compressors in use and under warranty. A compressor's warranty starts when the refrigerator is placed in, is installed to the customer. Moreover, its warranty term is three years. Hence, SER is calculated in this three-year period. An increase in the SER leads the quality and manufacturing team to change their current processes and take preventive actions to reduce those defect rates accordingly. The method used by the company for predicting the number of service calls is called the drift method (Hyndman and Athanasopoulos, 2014). The company implements the drift method as follows:

Indices

i : observation month after the production date

j : production batch

Parameters

$X_{i,j}$: Total cumulative number of compressor complaints observed until month i from production batch j

$Y_{i,j}$: Total cumulative number of refrigerators sold until month i from production batch j

$$SER_{i,j}^{\text{observed}} = \frac{X_{i,j}}{Y_{i,j}}$$

1. Calculating the total number of installed (sold) refrigerators and constructing Table 2 in Appendix B. This step can be easily done by pivoting the Production Batch, Installation Month and Count of Installed/Sold Refrigerator.
2. Calculating the total number of compressor complaints and constructing Table 3. This step can be easily done by pivoting the Production Batch, Complaint Month and Count of Complaints.
3. By using Tables 2 and 3 and $SER_{i,j}^{\text{observed}}$, Table 4 is constructed which provides data points for the prediction model that are currently using.
4. In the Table 5, SER predictions are calculated as follows,

$$\begin{aligned} SER_{i,j}^{\text{prediction}} = & SER_{i-1,j}^{\text{observed}} \\ & + \frac{(SER_{i-1,j}^{\text{observed}} - SER_{i-2,j}^{\text{observed}})}{3} \\ & + \frac{(SER_{i-2,j}^{\text{observed}} - SER_{i-3,j}^{\text{observed}})}{3} \\ & + \frac{(SER_{i-3,j}^{\text{observed}} - SER_{i-4,j}^{\text{observed}})}{3}. \end{aligned}$$

By making some simplifications we get the following formula for $SER_{i,j}^{\text{predicted}}$

$$SER_{i,j}^{\text{prediction}} = SER_{i-1,j}^{\text{observed}} + \frac{(SER_{i-1,j}^{\text{observed}} - SER_{i-4,j}^{\text{observed}})}{3}.$$

This formula is used for the time period between months 4 and 15. In month 16, a slight change is applied to the formula in order to increase the precision of the model. This is recommended by the Industrial Advisor (IA).

$$SER_{i,j}^{\text{prediction}} = SER_{i-1,j}^{\text{observed}} + \frac{(SER_{i-1,j}^{\text{observed}} - SER_{i-4,j}^{\text{observed}})}{6}$$

As it can be seen from the $SER_{i,j}^{\text{predicted}}$, the prediction model can be started four months after the production month. Therefore, the currently used model cannot predict the first four-month. Missing data in Table 4 are filled with the predicted data. This is used to create Table 5.

5. In this step $SER_{i,j}^{\text{predicted}}$ of each column is multiplied with 0.9 as a partial adjustment. This adjustment is done because the complaints data are not clean (Recommended by the IA).
6. In the final step, $SER_{i,j}^{\text{predicted}}$ is multiplied with initial production amount for the batch j to find the annual complaint number for that batch. It can be found in Table 6.

According to the IA, the company only looks at the final SER value of each year without looking at the final SER values of other months. In other words, their aim is to learn the final SER value in the 16th, 28th, and 40th month because those are the last month of each year (Since they do not make any estimation for first four months, it is $(12 + 4) = 16$, $(24 + 4) = 28$ and $(36 + 4) = 40$). An example of this calculation can be observed in Appendix B.

SER calculation is important for the company because it allows them to analyze whether the production has any concerning problems that they need to address. In this way, the company can take preventive precautions before time, money, and human power are invested into the production. Moreover, it also allows them to see whether a new product line (e.g., a new segment or a model) has any technical or production-related problems, so it is a crucial metric for Arçelik. Nonetheless, the current system is inadequate for a number of reasons that will be discussed below.

The current method does not account for seasonality. However, refrigerators' workspace is affected by the temperature of the environment; hence more complaints may be observed in the summer. No estimations are made for the first four months. Because of this, SER calculation is starting after four months have passed. This situation can cause the company to miss specific patterns and/or effects that can be observed in the first four months.

Finally, the IA stated that they don't track the performance of the current forecasting method. Therefore, they don't know about the accuracy of their predictions. This makes it harder for them to decide whether they should change their current forecast method or not. Also, it is almost impossible to compare different methods with each other without a KPI for forecast methods such as forecast errors. The problem with the company's forecasting system is that it is not effectively keeping track of forecast errors via RMSE, MAPE, and MASE.

2 Proposed System

R, Python, Excel, and KNIME are used to draw conclusions from data and used for implementation and validation of the proposed system. To overcome the issues stated above, we have decided to use a hybrid system that tackles all of the matters. To solve the seasonality problem, we used forecasting methods that enabled the usage of seasonality like Holt-Winters and SARIMA. We have used all the months, contrary to Arçelik's method that starts at the fourth month, to

incorporate all the observations.

2.1 Solution Approach

In order to find the best forecasting method for Arçelik and their use case, we tested different forecasting methods. Since there are a lot of different methods in forecasting demand, it is useful to have different sources to compare alternatives. Comparing current methods such as single exponential smoothing, Holt's method, and SARIMA, is helpful in finding the optimal result. However, since our data are only for three years, it may not be sufficient enough to see the compared results. According to research, for univariate time series data, Holt's method is the most appropriate method to forecast one-variable demand. (Armstrong and Collopy, 1992) Nonetheless, since Holt's method cannot take seasonal factors into account, we can use the Holt-Winters Additive Seasonal Method to capture the seasonality. (Wang, 2019). Also, since our data are not plenty, it is useful to have an evaluation of different techniques for comparing errors across time series.

We have come up with a Heuristic Approach that tackles all the problems mentioned above. Since this approach is a mix of different forecasting methods, it is necessary to give background information on the forecasting methods used in Heuristic Approach.

Background Information on Used Forecasting Methods

Double Exponential Smoothing/Holt's Method: This method is used to capture the trend and/or seasonality. In general, the series comprises of two parts that are a systematic part, including level, trend, and seasonality, and a non-systematic part that is simply the noise term (Brownlee, 2020). Apart from this, we also have two options that are additive and multiplicative when modeling. The multiplicative model is useful when the trend variation increases over time (ABoS, 2017). However, since no increase in trend variation has been observed, an additive model has been used.

ARIMA: This method combines AR (Autoregressive) and MA (Moving Average) models. In the AR part, the variable of interest is forecasted using a linear combination of past values of that variable. While in the MA part, uses past forecast errors in a regression-like model rather than using past values of the forecast variable in a regression. Moreover, ARIMA models are used to describe the auto-correlations within the data (Hyndman and Athanasopoulos, 2014).

SARIMA: SARIMA models are mostly represented with both its non-seasonal and seasonal parts (Yang et al., 2017). It yields accurate results for forecasting the short and medium term.

Triple Exponential Smoothing/Holt-Winters' Method: The idea behind triple exponential smoothing is to apply exponential smoothing to the seasonal

components in addition to level and trend. Holt-Winters' method can be used for forecasting a series with trend and seasonality.

2.2 Implementation of the Heuristic Approach

Since the number of data points for each production month varies from 1 to 36 at any given time, it is impossible to apply the same time series method to all the production months without sacrificing forecasting accuracy. The reason is that at least 25 data points are necessary for methods that utilize a seasonal component like SARIMA or Holt-Winters. Therefore it is necessary to treat production months differently or break them into segments according to the data points that they have. We have taken the segmentation approach and divided the production months into three parts:

1. The first segment consists of production months with more than 25 observations. SARIMA or Holt-Winters method can be used for this segment since they have enough observations to fit into a seasonal method.
2. The second segment consists of production months with at least 13 observations and no more than 25 observations. ARIMA and Holt's Method (Double Exponential) is considered for this segment. However, since the data are seasonal in its essence, and these methods cannot handle seasonality, it is required to feed the methods with percentage based observations, SER.
3. The third segment consists of production months with less than 13 observations. In order to apply seasonal methods, previous years' production months and production months in this segment were concatenated to get a new time series that have more than 25 observations. For example, for the production month "2020-01", there are only 12 observations available, however if the first 12 months of data from production months "2018-01" and "2019-01" are combined together with "2020-01", we obtain a new series that has 25 observations. This allows the usage of the same models for the first segment in this segment.

In order to get as much information from the data as possible, this is necessary. With implementation in mind, we have come up with a Python code that can segment the production months and conduct the forecasts when a single Excel file is inputted into the code. Following steps are taken in the code:

1. Read the data and split the production months into segments as described above.
 - (a) Input Excel file can be found in Appendix C.
 - (b) Data are split into test and train sets. Last four observations are used in the test set and the remaining observation for each production month is in the train set.

- (c) Then production months are distributed into segments according to the number of observations after the train set is created. Names of the months are stored in three lists that correspond to segments.
 - (d) Date information is removed from the data since each production month and the time series corresponding to it start from a different month. This will allow us to create data frames with ease.
2. A loop for the list that corresponds to the first segment and third segment begins. Note that the third segment has been turned into time series with more than 25 observations before.
 - (a) Using the items of the list and the separated date information a data frame is created.
 - (b) Then a model is fitted. For this segment SARIMAX and Holt-Winters Exponential Smoothing from the statsmodel library are used. More information can be found in the following sections.
 - (c) Output is stored in an Excel file called “trainresults”.
 3. A loop for the list that corresponds to the second segment begins.
 - (a) Using the items of the list and the separated date information a data frame is created.
 - (b) Then a model is fitted. For this segment, ARIMA and Holt Method from the statsmodel library are used. More information can be found in the following sections.
 - (c) Output is stored in an Excel file called “trainresults”.

Technical Aspect of the Heuristic Approach

Double Exponential Smoothing: Holt function from statsmodel is called with the following arguments: the data frame, damped trend, smoothing level as 0.75 and smoothing trend as 0.1. Forecast is initialized with “legacy-heuristics” from statsmodels. Results are added to the “testresults” Excel file. It is important to note that results can be further improved through optimizing the smoothing constants.

Triple Exponential Smoothing: Exponential Smoothing from statsmodels, a Python library is imported to the notebook. Then as described above, a loop for each production month begins and it is matched with correct dates and a new data frame is created. It is very important to note that this way of creating data frames have been used throughout the approach. Then four different parameter sets are used. They are as follows:

1. Additive trend, additive seasonality.

2. Additive trend, multiplicative seasonality.
3. Additive trend, additive seasonality with damped trend.
4. Additive trend, multiplicative seasonality with damped trend.

Note that initialization to the forecast is automatically estimated by the 'ExponentialSmoothing' function, and Box-Cox is used to serialize the series in all methods. After the methods have been fit, we plot them and add the results to the "trainresults" Excel file. In our calculations with a four month test set, the lowest average MAPE score came from the third method which is the one with additive trend, additive seasonality with damped trend. It is possible to implement some code that would automatically choose the best method for different production months but without a green light on the proposed approach this was seen as unnecessary.

ARIMA: Again utilizing the excellent library that is statsmodel, just by calling its ARIMA function called `tsa.arima`. ARIMA with a data frame and order parameters as arguments is enough. Then the results are added to the "testresults" Excel file. Note that order parameters are optimized through Akaike Information Criterion (AIC). This is done through an iteration over (0,0,0) through (1,1,1). Results are added to the "testresults" Excel file.

SARIMA: Utilizing statsmodel library, `tsa.statespace.SARIMAX` function is called with arguments as the corresponding data frame, order parameters and seasonal parameters. Parameters are optimized through their AIC score. This is done through an iteration over (0,0,0)x(0,0,0,12) through (1,1,1)x(1,1,1,12). Results are added to the "testresults" Excel file.

3 Validation of the Proposed System

We validated our methods with cross-validation technique to make sure that the methods we use work correctly, consistently, and reliably. It is a technique that examines the results of statistical analysis in an independent data set. Data sets are created by selecting samples in different amounts and the sets are divided into training and test sets. The accuracy of the proposed system is calculated by taking the average of the test sets.

In cross-validation, we used both one-step and multi-step errors. In multi-step errors, the horizon was between 2 to 12 months according to the availability of the data points.

For 2018, the values up to and including the 25th month are taken as the training sets, and the remaining months are used as the test sets with both one-step and multi-step ahead forecasts. Then we repeated the same procedure, but this time the values up to and including the 26th month are the train sets, and the remaining were the test sets. Finally, we repeated the same procedure by taking

the values up to and including the 27th month. This way, we took three different samples for all production months while completing the validation step so that we evaluated the success of both the Arçelik’s Method and Heuristic Method more deeply.

For 2019, the values up to and including the 13th month are taken as the training sets, and the remaining months are used as the test sets in the first sample. Finally, for 2020, the values up to and including the 4th month are taken as the training sets, and the remaining months are used as test sets in the first sample.

We followed this approach because we wanted our data to reach a level of maturity so that the efficiency of the method would be analyzed more successfully. Hence, for example, in 2018, we used the first 25th months as the train sets within the first sample.

In the end, we have come up with the RMSE, MAPE, and MASE values of the current Arçelik’s Model and the Heuristic Approach. Arçelik’s Model gives 168.105, 35.183%, and 11.106% respectively in RMSE, MAPE, and MASE. The Heuristic Approach yields 43.165, 17.248%, and 3.628% respectively in RMSE, MAPE, and MASE. Hence, it is possible to conclude that we have improved the system by more than 50% in all performance measures. The validation results can also be seen in Appendix D.

4 Implementation of the Proposed System

We used the KNIME platform in the implementation stage because the company wanted us to use either Excel or KNIME in this stage. We are providing a user interface and a user handbook. Moreover, within the KNIME platform, several notes enable the user to always keep in mind the essential details without feeling the need to open the user handbook.

We have two main steps in the implementation part. In the first step, we are arranging the data, and in the second step, we are using the prediction code. The workflow of both of these steps can be found in Appendix E.

To arrange the data, firstly, Arçelik opens the input files in KNIME. KNIME transforms these files into a form that is implementable to the prediction code. If the user has already concatenated the data, then s/he can directly import the file into KNIME without any concentration. All possible input areas are marked in the user interface. In the Data Arrangement step, final python codes make the last changes in the data sets. Two output data sets are joined, and in the end, the user receives the final version of the data-set combined with complaints and production data.

To complete the second step, where we have our prediction code, we insert arranged data set into the final python script, which includes the heuristic approach code. In the end, the KNIME platform gives an output table in Excel, and a shortened example of this output table can be found in Appendix F. Finally, we are

writing this output into an excel file which can be saved where the user wants. To sum up, it is possible to read the Excel data and give predictions to a new Excel file that the user specifies. In addition, users can use this output table in the KNIME platform for further operations without any outputting.

5 Project Benefits

The main benefit that the project provides to Arçelik is the correct calculation of monthly service call rates. The main objective of the project, as stated by the company, is to evaluate the accuracy of the current model and compare its accuracy to the models the project team will decide on.

As discussed above, we managed to improve the system by more than 50% in all performance measures. Hence, it is possible to state that the Heuristic Approach will benefit Arçelik in various ways.

The accuracy of the model will help the company to improve their production lines according to service rates since they will be able to know where the problem occurs, for example, in which compressor type or refrigerator model. Also, they now have the tools to predict future SER values accurately. Hence, they will be in more control about the breakouts of the compressors.

Even though the company did not provide a monetary measure about the SER forecasting, this benefit will return as the monetization of forecasting errors. More specifically, the more accurate the company can predict the number of breakouts, the more chance it has to repair the breakouts. Hence, increasing the lifetime of the refrigerators. When lifetime increases, the money that Arçelik spends for breakouts from the guarantee decreases. In this way, the company will have a cost reduction.

6 Conclusion

Throughout the project, we tried to satisfy the company's expectations as much as we can. Our approach takes seasonality into account; hence it gives more reliable and accurate forecasting results. Moreover, we managed to compare several different statistical methods, which can be listed as Logistic Regression, Moving Average, SARIMA, ARIMA, Holts Winter, Hybrid Linear Regression & ARIMA, Hybrid ETS & ANN and Regression Tree & Machine Learning Model with the current Arçelik method as requested by the company. By trying these methods, we managed to find good learners such as SARIMA, ARIMA, Holts, and Holts Winter, so that we combined the strengths of these methods with a segmentation approach and come up with the Heuristic Method. With this method, we managed to get more than a 50% increase in the forecasting accuracy according to the cross-validated performance measures that can be found in Appendix D.

6.1 Future Work

Future work can deepen the current analysis. Our model does not consider the differences in refrigerator type, compressor origin, or the place where the refrigerator is installed. Furthermore, the developed model cannot take regionality into account; hence, further analysis could include the regionality aspect.

References

- A. B. o. S. ABoS. Australian Bureau of Statistics Web Site, 2017. URL <https://www.abs.gov.au/websitedbs/d3310114.nsf/home/time+series+analysis:+the+basics>.
- J. Armstrong and F. Collopy. Error measures for generalizing about forecasting methods: Empirical comparisons. *International Journal of Forecasting*, 8(1):69 – 80, 1992. ISSN 0169-2070. doi: [https://doi.org/10.1016/0169-2070\(92\)90008-W](https://doi.org/10.1016/0169-2070(92)90008-W). URL <http://www.sciencedirect.com/science/article/pii/016920709290008W>.
- J. Brownlee. A gentle introduction to exponential smoothing for time series forecasting in python, 2020. URL <https://machinelearningmastery.com/exponential-smoothing-for-time-series-forecasting-in-python/>.
- R. J. Hyndman and G. Athanasopoulos. *Forecasting: Principles and Practice*. OTexts.com, 2014. ISBN 9780987507105.
- X. Wang. The Short-Term Passenger Flow Forecasting of Urban Rail Transit Based on Holt-Winters’. pages 265–268, 2019. doi: 10.1109/ICECTT.2019.00067.
- Y. Yang, C. Liu, and F. Guo. Forecasting Method of Aero-Material Consumption Rate Based on Seasonal ARIMA Model. *3rd IEEE International Conference on Computer and Communications (ICCC)*, pages 2899 – 2903, 2017. ISSN 978-1-5090-6352-9. URL <http://search.ebscohost.com/login.aspx?direct=true&db=edsee&AN=edsee.8323062&site=eds-live>.

Appendix A: Analysis and Interpretation of the Data

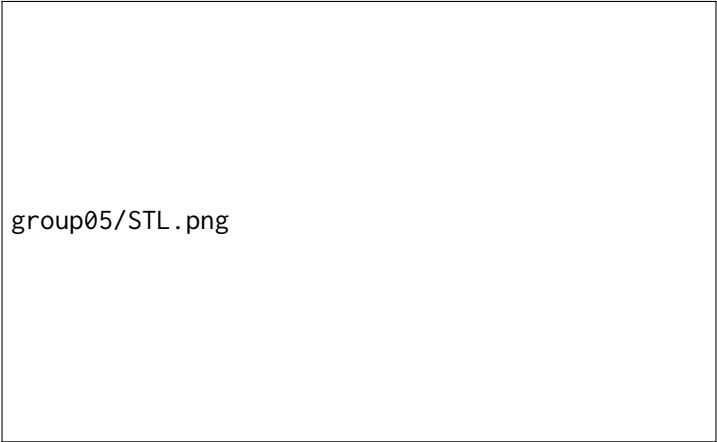


Figure 1: STL Decomposition

Appendix B: Tables For SER Calculations

Table 1	Production	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16
2019-09	104.090	8.486	30.350	41.222	48.232	52.781	57.136	61.542	65.119	69.469	74.632	80.187	83.077	84.644	85.561	86.098	86.343
2019-10	79.348	7.685	20.165	27.804	32.419	36.459	41.088	44.644	48.753	53.341	58.371	61.306	62.639	63.416	63.960	64.257	
2019-11	87.155	5.652	19.207	27.587	34.117	40.936	45.745	51.521	58.831	66.176	69.771	71.590	72.730	73.441	73.747		
2019-12	38.862	4.291	8.684	11.807	15.205	17.385	19.930	23.229	26.648	28.352	29.087	29.507	29.708	29.814			
2020-01	77.170	3.839	13.896	24.181	30.792	38.085	46.133	56.797	62.704	65.005	66.328	67.064	67.385				
2020-02	85.282	5.523	17.009	23.997	31.615	39.887	52.853	59.590	62.096	63.349	64.020	64.388					
2020-03	145.020	5.634	19.833	39.390	59.291	86.386	99.361	104.125	106.490	107.791	108.262						
2020-04	122.475	5.597	26.382	47.628	74.409	85.913	90.097	92.072	93.086	93.484							
2020-05	99.850	9.475	37.897	64.699	74.549	77.811	79.190	79.979	80.330								
2020-06	162.795	16.515	85.216	113.411	123.390	128.204	130.954	132.095									
2020-07	217.353	46.155	119.509	146.433	159.594	165.906	168.424										
2020-08	160.143	25.821	72.840	91.226	99.824	103.263											
2020-09	131.368	21.547	53.945	66.872	72.134												

Table 2: 16-Month Number of Installations

Table 2	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16
2019-09	3	24	66	94	126	149	178	215	244	292	332	381	434	470	502	518
2019-10	-	13	29	41	54	68	75	84	117	142	171	199	227	238	245	
2019-11	1	20	48	65	85	94	114	128	149	167	183	210	231	238		
2019-12	3	7	10	23	27	34	49	62	73	86	91	99	103			
2020-01	-	6	21	32	49	72	100	132	167	183	207	218				
2020-02	1	12	32	51	86	122	169	205	231	248	255					
2020-03	5	18	42	109	165	207	264	300	333	344						
2020-04	-	10	48	88	161	209	249	271	288							
2020-05	2	46	116	281	394	449	484	498								
2020-06	10	79	194	282	324	368	397									
2020-07	14	171	384	512	590	628										
2020-08	6	104	195	257	291											
2020-09	4	53	116	141												

Table 3: 16-Month Number of Compressor Complaints

Table 3	Production	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16
2019-09	104.090	0.035%	0.079%	0.160%	0.195%	0.239%	0.261%	0.289%	0.330%	0.351%	0.391%	0.414%	0.459%	0.513%	0.549%	0.583%	0.600%
2019-10	79.348	0.000%	0.064%	0.104%	0.126%	0.148%	0.165%	0.168%	0.172%	0.219%	0.243%	0.279%	0.318%	0.358%	0.372%	0.381%	
2019-11	87.155	0.018%	0.104%	0.174%	0.191%	0.208%	0.205%	0.221%	0.218%	0.225%	0.239%	0.256%	0.289%	0.315%	0.323%		
2019-12	38.862	0.070%	0.081%	0.085%	0.151%	0.155%	0.171%	0.211%	0.233%	0.257%	0.296%	0.308%	0.333%	0.345%			
2020-01	77.170	0.000%	0.043%	0.087%	0.104%	0.129%	0.156%	0.176%	0.211%	0.257%	0.276%	0.309%	0.324%				
2020-02	85.282	0.018%	0.071%	0.133%	0.161%	0.216%	0.231%	0.284%	0.330%	0.365%	0.387%	0.396%					
2020-03	145.020	0.089%	0.091%	0.107%	0.184%	0.191%	0.208%	0.254%	0.282%	0.309%	0.318%						
2020-04	122.475	0.000%	0.038%	0.101%	0.118%	0.187%	0.232%	0.270%	0.291%	0.308%							
2020-05	99.850	0.021%	0.121%	0.179%	0.377%	0.506%	0.567%	0.605%	0.620%								
2020-06	162.795	0.061%	0.093%	0.171%	0.229%	0.253%	0.281%	0.301%									
2020-07	217.353	0.030%	0.143%	0.262%	0.321%	0.356%	0.373%										
2020-08	160.143	0.023%	0.143%	0.214%	0.257%	0.282%											
2020-09	131.368	0.019%	0.098%	0.173%	0.195%												

Table 4: 16-Month Observed SER

Table 4	Production	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16
2019-09	104.090	0.035%	0.079%	0.160%	0.195%	0.239%	0.261%	0.289%	0.330%	0.351%	0.391%	0.414%	0.459%	0.513%	0.549%	0.583%	0.600%
2019-10	79.348	0.000%	0.064%	0.104%	0.126%	0.148%	0.165%	0.168%	0.172%	0.219%	0.243%	0.279%	0.318%	0.358%	0.372%	0.381%	0.392%
2019-11	87.155	0.018%	0.104%	0.174%	0.191%	0.208%	0.205%	0.221%	0.218%	0.225%	0.239%	0.256%	0.289%	0.315%	0.323%	0.345%	0.354%
2019-12	38.862	0.070%	0.081%	0.085%	0.151%	0.155%	0.171%	0.211%	0.233%	0.257%	0.296%	0.308%	0.333%	0.345%	0.362%	0.380%	0.388%
2020-01	77.170	0.000%	0.043%	0.087%	0.104%	0.129%	0.156%	0.176%	0.211%	0.257%	0.276%	0.309%	0.324%	0.346%	0.369%	0.389%	0.400%
2020-02	85.282	0.018%	0.071%	0.133%	0.161%	0.216%	0.231%	0.284%	0.330%	0.365%	0.387%	0.396%	0.418%	0.436%	0.452%	0.471%	0.479%
2020-03	145.020	0.089%	0.091%	0.107%	0.184%	0.191%	0.208%	0.254%	0.282%	0.309%	0.318%	0.339%	0.358%	0.375%	0.394%	0.412%	0.421%
2020-04	122.475	0.000%	0.038%	0.101%	0.118%	0.187%	0.232%	0.270%	0.291%	0.308%	0.333%	0.354%	0.376%	0.398%	0.420%	0.441%	0.452%
2020-05	99.850	0.021%	0.121%	0.179%	0.377%	0.506%	0.567%	0.605%	0.620%	0.658%	0.688%	0.716%	0.748%	0.778%	0.807%	0.838%	0.853%
2020-06	162.795	0.061%	0.093%	0.171%	0.229%	0.253%	0.281%	0.301%	0.325%	0.348%	0.371%	0.394%	0.418%	0.441%	0.464%	0.487%	0.499%
2020-07	217.353	0.030%	0.143%	0.262%	0.321%	0.356%	0.373%	0.410%	0.439%	0.467%	0.499%	0.528%	0.558%	0.588%	0.618%	0.648%	0.663%
2020-08	160.143	0.023%	0.143%	0.214%	0.257%	0.282%	0.328%	0.366%	0.403%	0.443%	0.481%	0.519%	0.558%	0.597%	0.635%	0.674%	0.693%
2020-09	131.368	0.019%	0.098%	0.173%	0.195%	0.254%	0.307%	0.351%	0.403%	0.452%	0.501%	0.550%	0.600%	0.649%	0.698%	0.748%	0.772%

Table 5: 16-Month Predicted SER

Table 5	Production	16-Adjusted	Total Complaint
2019-09	104.090	0,540%	562
2019-10	79.348	0,353%	280
2019-11	87.155	0,319%	278
2019-12	38.862	0,349%	136
2020-01	77.170	0,360%	278
2020-02	85.282	0,431%	368
2020-03	145.020	0,379%	549
2020-04	122.475	0,407%	498
2020-05	99.850	0,768%	767
2020-06	162.795	0,449%	731
2020-07	217.353	0,597%	1298
2020-08	160.143	0,624%	999
2020-09	131.368	0,695%	913

Table 6: 16th Month SER Data Table with Coefficient Correction and Production Weights

Appendix C: Heuristic Method

Row Label	2019-01	2019-02	2019-03	2019-04	2019-05	2019-06	2019-07	2019-08	2019-09	2019-10	2019-11	2019-12	2020-01	2020-02	2020-03	2020-04	2020-05	2020-06	2020-07	2020-08	2020-09	2020-10	2020-11	2020-12
1	1																							
2	1	1																						
3	1	1	1																					
4	1	1	1	1																				
5	1	1	1	1	1																			
6	1	1	1	1	1	1																		
7	1	1	1	1	1	1	1																	
8	1	1	1	1	1	1	1	1																
9	1	1	1	1	1	1	1	1	1															
10	1	1	1	1	1	1	1	1	1	1														
11	1	1	1	1	1	1	1	1	1	1	1													
12	1	1	1	1	1	1	1	1	1	1	1	1												
13	1	1	1	1	1	1	1	1	1	1	1	1	1											
14	1	1	1	1	1	1	1	1	1	1	1	1	1	1										
15	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1									
16	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1								
17	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1							
18	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1						
19	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1					
20	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1				
21	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1			
22	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1		
23	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	
24	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
25	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
26	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
27	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
28	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
29	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
30	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
31	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
32	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
33	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
34	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
35	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
36	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
37	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
38	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
39	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
40	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
41	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
42	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
43	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
44	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
45	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
46	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
47	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
48	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
49	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
50	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
51	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
52	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
53	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
54	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
55	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1			

Appendix D: Cross Validated Results Table

Arçelik Method			Heuristic Method		
RMSE	MAPE	MASE	RMSE	MAPE	MASE
168,105	35,183%	11,106	43,165	17,248%	3,628

Table 8: RMSE, MAPE, and MASE values of the Arçelik’s and Heuristic Method

Appendix E: KNIME Deployment



Figure 2: Workflow Group that will be Set for Deployment

Appendix F: Deliverable of the KNIME Deployment



Figure 3: Shortened Example of the Output Table in Excel

Ham Madde Envanter Yönetim Sistemi

Türkiye Şişe ve Cam Fabrikaları A.Ş.

group06/Grupfotosu.png

Proje Ekibi

Çağla Albayrak, Özlen Ataç, Said Ay, Alperen Nazım Demir,
Sinem Karaömeroğlu, Kerem Neşet Sarıtepe, İrem Ünal

Şirket Danışmanı

Fatih Altuner
Tedarik Zinciri Müdürü
Ekin Baştürk
İşletme Mühendisi
Veysel Köybaşı
İşletme Mühendisi

Şişecam Düzcamlar Kırklareli Fabrikası

Akademik Danışman

Prof. Dr. Nesim Kohen Erkip
Bilkent Üniversitesi Endüstri
Mühendisliği Bölümü

ÖZET

Şişecam Düzcamlar, Avrupa'da lider Dünya'da ise 5. en büyük düzcamlar üreticisidir. Şişecam Düzcamlar Kırklareli fabrikasındaki hammadde sipariş politikalarının ve belirlenen güvenlik stoğu seviyelerinin gelişime açık olduğu tespit edilmiştir. Şirket yetkililerinin düşünceleri ve kendi bulgularımız bu tespiti destekler niteliktedir. Float ve ayna hatlarının üretim planları, stok seviyeleri, lojistik kısıtlamaları ve belirsizlikler dikkate alınarak hammadde envanter maliyetlerini düşürmek ve Şişecam için tutarlı, sistematik bir sipariş verme politikası oluşturmak amacı ile bir proje geliştirilmiştir. Envanter yönetimini etkileyen parametreleri göz önünde bulundurarak her hammadde için envanter maliyetlerini enazlayan matematiksel modellerin geliştirilmesiyle sürdürülebilir ve sistematik bir çözüm sunulmuştur. Çeşitli yazılım araçlarının desteği ve kullanıcı dostu arayüzü entegrasyonu ile dijital bir karar destek sistemi inşa edilmiştir. Oluşturulan sistem ile Soda özelinde yapılan simülasyonlarda %3,65 iyileşme görülmüştür.

Anahtar Kelimeler: Envanter yönetimi, üretim planlama, tedarik zinciri yönetim, sürekli üretim, matematiksel model

Raw Material Inventory Management System

1 Company Information

Şişecam was founded by Türkiye İş Bankası with the request of the Council of Ministers in 1935. In terms of production capacity, Şişecam stands in the leader position in Europe and it is the world's fifth-largest flat glass producer. Currently, Şişecam operates in the flat glass industry with 10 production facilities located in six countries. Furthermore, Şişecam is the world's third and Europe's second-largest producer of glassware. In addition, as Turkey's largest automotive glass producer, and a sector leader, Şişecam delivers world-class products to automobile manufacturers nationwide, while serving as a supplier for Europe's leading auto manufacturers. Şişecam's clients in the automotive industry include pioneers of car manufacturing such as BMW, Renault, Volkswagen, Mercedes-Benz, Audi, Toyota, GM, Hyundai, and others.

2 System Analysis

Şişecam Trakya is the part of Şişecam that includes three production lines: TR1, TR2, and TAB. Two of these are the continuous production lines of the flat glass production, and the other is the production line of the mirror. RZ, YS, Y1, KY, FP, and FT are abbreviations of the flat glass types for both lines. Şişecam uses monthly production plans. We observe these from their S&OP plans which are updated each month for the next 12 months. These plans vary every time they are updated based on the forecasted and arrived customer demands. Supply chain managers decide their ordering policy according to the production plan of the company. Raw materials have different constraints such as maximum order amount per a specified period or the lot sizes depending on the transportation way and its capacity.

In Silica Sand, on the 15th of each month; the next month's sand request are made. If the train is considered to be insufficient according to the amount of sand consumption, the seaway is preferred. In the seaway shipment, the stock amount is tried not to be below a determined amount

Soda's lead time differs according to its suppliers. In highway shipment, while its lead time can be increased up to 4 days, in general, the shipment time is 6 hours in a day. In seaway, the lead time differs in summer and winter due to the weather conditions. In seaway shipment, ship capacity should be considered.

For Feldspar, every week, the next week's needs are notified to the supplier and Feldspar can be supplied from one supplier when it is required. If there is an order, the quantities vary approximately from 300 tons to 900 tons per month.

In Hematite, ordering amounts should be integer multiples of some determined numbers. Since it is a colorant raw material, the amount of broken glass used, its content and the color campaign that is worked on change the amount of Hematite usage.

Sodium Nitrate order is opened 3 months before the production date for both suppliers. It is used in the Fume privacy color campaign, which is carried out once a year. Warehouse capacity depends on the production demand.

Cobalt Oxide is used in the Fume privacy color campaign, which is carried out once a year. The order is opened 3 months before the production date. The warehouse capacity depends on the production demand.

For Silver Nitrate, there are two suppliers with different lead times one being an overseas supplier whereas other being a domestic supplier. On the average, the order amount is 1200 kg if there is an order given in a month.

Silver reducer has two overseas suppliers and agreements are made annually. If any order is given in a month, the shipment amount is varied in between 500 L and 1000 L to balance the chemical. Orders are given with intermediate bulk containers (IBC) and IBCs should be fully filled.

For Ammonia, lead times are short and the company gives orders when necessary. Although there is not any certain limit for orders, the orders vary in between 3000 L and 5000 L. Orders shipped with IBCs and each IBC capacity is 1000 L.

Palladium Solution has two overseas suppliers. The ordered chemicals from supplier 1 and supplier 2 are used separately, they cannot be mixed. To balance the chemicals, the order amounts vary in between 50 L and 180 L.

Passivator 1 has two overseas suppliers. Different amounts are ordered to balance the chemical. While the ordering amounts differ in between 704 L and 2112 L for supplier 1, they differ in between 100 L and 300 L for supplier 2.

3 Problem Analysis

The evaluation of the inventory amounts with consumption amounts and safety stock values displays the fact that there may exist excessive stock of raw materials at the end of each month as also the company states. The company has pre-determined safety stock values from history and they order considering these levels. We expect the average end inventory levels to be close to predetermined safety stock levels in a good-working inventory and ordering system. (Cohen and Halperin (1980).) However, we observe that their end-inventory levels are almost always more than what they decide as in Appendix A. From the data provided by the company, we see that for some raw materials, the difference might be acceptable but others need improvement and this is one of the reasons for high costs resulting from excess inventory. Their ordering policy is not based on a systematic approach and changes from one raw material to other, depending on the properties of suppliers and raw materials. These are main problems which indicates that the company's inventory management is open for improvements. Other aspect of problem is the safety stock levels. These values are determined by the company based on historical knowledge. The analysis we performed based

on literature showed us that for some materials these values needed to be re-estimated and improved.

4 System Design

4.1 System Summary

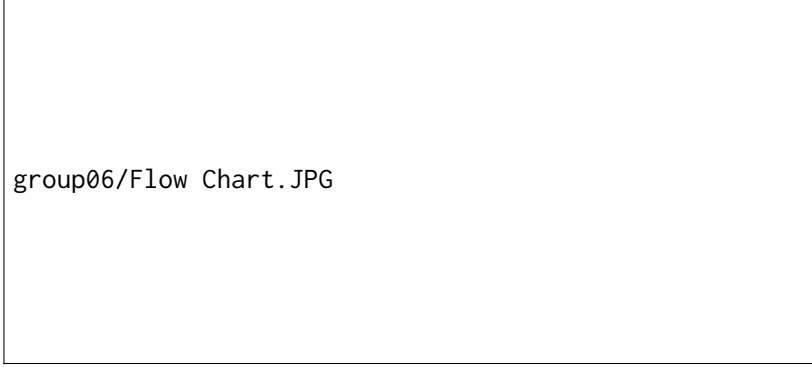


Figure 1: Flow chart for the inputs and outputs of the proposed system

At the center of the system, there is Supplier & Quantity Selection algorithm (SQS) that contains multiple mathematical models for each raw material. Calculated safety stock levels for each raw material, supplier features, relative costs, and facility parameters will be the inputs of SQS. The main output of the system will be the order release schedule with supplier and quantity selection. Total inventory-related cost is the secondary output, however, since our purpose is not to calculate all costs related to inventory, we will use it only to compare and improve our method's results. Proposed system aims to eliminate high inventory levels with a systematic approach by deciding the order release schedule considering relevant constraints to provide logical and safe decisions. By doing so, it also considers inventory holding and ordering costs to prevent monetary losses while decreasing the inventory levels. Therefore, objective functions we use aim to minimize total costs. Our system will provide consistent decision support mechanism for the company.

4.2 Input Operations

S&OP plans come to the system operator for all production lines which include all information about the production such as maintenance, holidays, setups, cleanings, etc. System operator analyzes data and specifies for desired production lines that are TR1, TR2, and TAB. And also there is two type of S&OP plans: one of them is based on weight of the glass, other one is based on the surface area of the glass. Weight based S&OP plans are made for the flat glass production lines, TR1 and TR2, that contain monthly production amounts of flat glass with different types of colors. For every type of colors, the consumption amount of raw

materials is different. Surface area based S&OP plans are made for the mirror production line, (TAB), because mirror production is a covering operation of the flat glass. There is not any consumption difference for the type of glass.

Consumption levels of each raw material for each period are calculated from SOP plans by using consumption amounts of raw materials. Our system only needs S&OP plans to calculate consumption levels of each raw material for each period. By this, the requirement amounts based on S&OP are automatically estimated. If there is a change in consumption amounts of raw materials or S&OP plans, it can easily be integrated into system by the company via Excel template we will provide to them. Our software automatically pulls the raw material requirement data from this file as it captures the current date and uses them as the inputs for its solution.

4.3 Safety Stock

In order to mitigate the effect of lead time and demand fluctuations in our decision model; we calculated the safety stock levels for each raw material. For Şişecam, since the lead time and demand are both stochastic, we utilized the following Equation 1 that are used for stochastic lead time and demand:

$$\begin{array}{ll}
 SS = \text{Safety Stock} & Z = \text{Service Level Factor} \\
 L = \text{Mean Lead Time} & d = \text{Standard Deviation of Lead Time} \\
 s_d = \text{Mean Demand} & s_L = \text{Standard Deviation of Demand}
 \end{array}$$

$$s' = \sqrt{(L + 1)s_d^2 + d^2s_L^2} \quad (1)$$

$$SS = Zs' \quad (2)$$

By using last three years' demand data, we calculated average demand and standard deviation of demand for each raw material. For the lead time information, the company has given the mean lead time and an interval around the mean to represent variation for each raw material, then these lead time data were fitted into triangular distribution.

When there are multiple suppliers for a raw material, mean lead time and standard deviation of lead time are calculated by taking weighted average of different suppliers. Moreover, when there is a seasonality for raw materials, mean demand and standard deviation of demand are calculated by only using the months that the seasonal raw materials are used.

By using the data for the parameters provided by the company, s' values are calculated for each raw material. Then, by dividing the safety stock levels that are determined by the company to s' values, the service levels of the current system are derived.

Table 1: Safety Stock Calculations of the Company

Raw Material	s'	Safety Stock	Z	Non-stockout Probability
Soda (ton)	570.52	1500	2.63	99.57%
Silica Sand (ton)	2032.36	10000	4.92	99.99%
Feldspar (ton)	48.76	300	6.15	99.99%
Hematite (ton)	21.47	60	2.80	99.74%
Cobalt Oxide (ton)	971.76	4000	4.12	99.99%
Sodium Nitrate (ton)	14.95	50	3.34	99.99%
Silver Nitrate (kg)	315.43	1350	4.28	99.99%
Ammonia (L)	351.06	3000	8.55	99.99%
Silver Reducer (L)	1985.67	5000	2.52	99.74%
Palladium Solution (L)	18.13	75	4.14	99.99%
Passivator 1 Solution (L)	223.14	1050	4.71	99.99%

As it can be seen from the Table 1, Z values are different for each raw material. Since all raw materials are required for production, the service levels ought to be same for raw materials of same end product. Therefore, in SQS, the service level factors of mirror raw materials (silver nitrate, ammonia, silver reducer, palladium solution and passivator 1 solution) is equal to 2.52 (the service level determined for silver reducer) which assures 99.41% non-stockout probability. Furthermore, since float production is continuous and stockout issue can be highly costly, in SQS, we determined the service level factors of float raw materials (soda, silica sand, feldspar and hematite) as 2.80 (the service level of hematite) which decreases the stock-out probability less than 0.26%.

4.4 Supplier & Quantity Selection

The SQS aims to decrease the total inventory, unit and fixed costs by deciding order amounts, suppliers and supplier mode selection and release times. In the general version of the SQS, it is assumed that;

- One supplier can have different modes of transportation with different costs and lead times, different transportation methods of a supplier will be considered as different suppliers.
- The model will have weekly periods which will be also used with periodic review.

The sets, parameters, decision variables can be in the Appendix B.1, the general mathematical model is in Appendix B.2 and the explanation for substitute constraints and decision variables is in Appendix B.3.

System will operate in weekly periodic review policy. Weekly periods are used in the SQS. Since one working week's starting and ending dates might be variant in a year, periods are described as "monthly quarters" instead of weeks. The model will be rolled every quarter of a month at least once and its planning horizon will be 24 quarters. This time horizon will be adequate for models to work properly since maximum lead time of raw materials are between 2 and 3 months.

Currently, if we use all constraints of the mentioned general model on a raw material's model; it has 1368 continuous and 236 integer constraints. Solving such model takes approximately 4 seconds in the CBC solver engine which is used in this project. Therefore, the system we are building will be highly applicable in terms of solving speed.

4.5 Model Adjustments and Specified Models for Raw Materials

As observed from the general model in Appendix B.2, the last period index of the left hand side of Equation 9 differs among suppliers. In order to overcome this problem, we introduced substitute constraints. They take place of the Equation 9 of the general model during the solution phase (in the code), but for mathematical notation Equation 9 is sufficient. These constraints do not increase the solution time of the model significantly. This change contributed in terms of applicability and flexibility since Şişecam will be able to insert or remove a new supplier and update the lead time of a supplier without modifying the software code.

In order to integrate the general version of SQS to the each raw material, some constraints need to be re-evaluated in conjunction with corresponding parameters. When the mathematical model in Appendix B is analyzed for each raw material properties, the contract constraint, the Equation 6 will be eliminated from specified model for every raw material. A new decision variable and parameter G_{ip} and b_i are introduced as in Appendix B.1 respectively. The corresponding Equations 25 and 26 are introduced as in Appendix B.2.

In addition, there are some changes in minimum and maximum order amount constraints. For each raw material the model will be updated as in the Table 2. In the table, (+) sign indicates the constraint required for that raw material.

Table 2: Updates in the specified models for raw materials

Raw Material	Minimum Order Quantity Eq. 7-8	Maximum Order Quantity Eq. 4	Integer Constraint Eq. 25-26
Soda	+	+	+
Silica Sand	+	+	+
Feldspar			+
Hematite	+		+
Cobalt Oxide			
Sodium Nitrate	+		+
Silver Nitrate	+	+	+
Ammonia			+
Silver Reducer	+	+	+
Palladium Solution	+	+	+
Passivator 1 Solution	+	+	+

5 Verification and Validation

5.1 Verification

In order to demonstrate that the constructed general model that decides the order amounts with supplier(s) and supplier mode is potentially useful, we need to make sure that every element is represented correctly and the general model represents the trade-offs. Thus, we will analyze the change in the decision variables and objective function by changing the parameters. Initially, we will have two different suppliers with varying contract terms such as fixed cost (f_i), contract constraint (c_i), lead time (l_i), unit cost (u_i), minimum amount to order (s_i) and maximum amount to order (j_i) for each supplier i in conjunction with safety stock, initial raw material inventory, raw material capacity and holding cost.

Table 3: Verification of SQS

Check of	Parameters	Consequence
Contract constraint c_i	- Decrease contract constraint	- More bulk orders.
Minimum amount s_i	- Increase minimum amount	- Increase in ending inventory.
Maximum amount j_i	- Decrease maximum amount	- Increase in objective value due to inventory holding costs.
Fixed cost f_i	- Include fixed cost	- More bulk orders - Increase in objective value due to fixed and holding costs.
Fixed cost f_i and Unit cost u_i	- Trade of between fixed and unit cost	- At higher enough fixed cost, supplier with higher unit cost is chosen.
Safety stock z	- Include safety stock	- Ending inventory level becomes at least z .
Lead time l_i	- Trade-off between lead time and unit cost	- When tight deadlines, supplier with higher unit cost and shorter lead time is chosen to meet the demand.
Holding cost h	- Trade-off between fixed and holding cost	- With low enough holding cost holding inventory is more preferred, but with high enough holding cost having fixed cost is more preferred.

Table 3 indicates a summary for the trade-offs and the behavior of the SQS according to elements of the SQS. As a result, components of the mathematical model are represented correctly and it is potentially useful.

5.2 Validation

For validation, the approach was to understand whether the outputs that our system gave were consistent with the realized ordering policies' results. Our system gives two main output: costs and ending inventory levels. Firstly, we have compared ending inventory levels. We placed the ordered amounts and the consumption levels of 2019 starting from January until May given by the company. We ran our model without optimization. We aimed to see whether our ending inventory constraints were valid without trying to improve by our objective function and other optimization constraints. Hence, we checked if our model was working correctly to calculate ending inventory levels when we entered the consumption levels each period manually. The comparison of model's output of ending inventory levels vs realized ending inventory values are provided in the Table 4. The resulting ending inventory levels were almost the same with the end of inventory records obtained by the company's ERP system.

Table 4: Validation of SQS with raw material soda

Month	Jan.19	Feb.19	Mar.19	Apr.19	May 19
Şişecam Ending Inventory (ton)	1888.5	2180.35	431.65	3175.34	2603.31
SQS Ending Inventory (ton)	1887.86	2179.65	226.84	3174.65	2602.62

Moreover, we repeated the same procedure to validate purchase costs for the same months. We used the unit costs of Şişecam and SQS with the purchases of Şişecam to compare. We obtained average of 0.41% deviation in purchase costs. As a result, the purchase costs and ending inventory levels for each month are close enough to conclude that SQS is valid.

6 User Interface (UI)

To control and implement the system, an UI (named Thrace) is developed Using Python and VBA. CBC solver is integrated into this UI to solve raw material models. There are 11 raw material modules connected to the main page which can be observed from the Appendix C. Users are able to switch between raw material modules using this page. For each raw material module; yearly S&OP, BOM, supplier features, facility parameters, safety stock, current date, and previous run’s outputs will be inputs. Since the previous run inputs are preserved by the UI, users will not need to re-enter these if there are no changes in the current trajectory.

When the UI is delivered to the company, it will contain the parameters of the current operations. However, since all parameters are subject to change in the future, we have integrated them into the UI so that all parameters can be changed easily. If desired; activating, deactivating, updating, adding, and removing suppliers and their parameters are possible for up to six suppliers. All of the remaining parameters are also controlled from the UI. If users do not want to include specific constraints for a supplier due to changes in the supplier contracts, they will be able to deactivate some of the constraints.

After all of the inputs are set, the user clicks the “Solve” button to run the model. The solving phase lasts two to six seconds in average. Then, the user will be able to get the current period’s orders, the upcoming six months’ predicted MRP table, and a detailed report. The outputs of the model are stored in the system for the next period’s execution to use as an input. Moreover, a user manual for each of the modules is available to support the user to utilize all of the functions of the system.

The system has been designed to ensure a convenient tool for facilitating the decision-making process for ordering policy. We have demonstrated the user interface to the company and showed how to change the parameters. This demonstration was not only limited to the description of the system but also the explanation of how the system retrieves the required amount of raw materials from the SOP

automatically. Since, the system was designed compatible with the commonly used software tool Excel VBA, the verbal instructions for utilization was found clear and simple to understand. The company has already asked to use the model prepared for the silver nitrate for order suggestion.

7 Project Contribution

The main contribution of this project is to reduce the total inventory cost of Şişecam Trakya Factory by providing systematic and consistent ordering policies via the models prepared per each raw material with a holistic perspective. The main performance measures are selected as the cost that includes purchase cost as well as cumulative inventory holding cost and end inventory levels. The data from first 5 months of 2019 for Soda, Silica Sand, Feldspar and Hematite which are float raw materials is analyzed by calculating the ending inventory levels, purchases and using consumption levels. The obtained values are compared with the realized values of Şişecam by using the cost performance measure. As a result, we are dedicated an improvements as in Table 5. These results prove that the decision system built gives a cost-effective, solution.

Table 5: Percent improvement in total cost with comparison of SQS and Şişecam

Soda	Silica Sand	Feldspar	Hematite
3.65%	17.92%	17.60%	8.03%

In the current system there are multiple suppliers and different supply methods. However, our models can distinguish different supply methods even though they are from the same supplier which enables the system to choose the best supplier and supply method for a particular period. Also, users are able to put new suppliers with different parameters if necessary, at any time. Parameters can be updated by the company through time easily via the user interface which makes the system more sustainable. Thrace works weekly and each week required parameters are updated. Thus, it considers these changes giving a more accurate ordering plan for next periods. Thrace considers scheduled orders, which prevents any redundant number of new orders. Sometimes, it might be required to give a different decision than the Thrace output due to unexpected external conditions. If the users decide to order a raw material differently than the output of models, they can integrate this as a new scheduled order to the model for further decision periods. Since output is received quickly, Thrace will be able to solve a half-year planning horizon in every period. One of the other advantages the system provides is its potential to be improved for future implementations, since the system uses open solver, which is a free and open software tool. Furthermore, if the system is found to be effective, it can also be implemented to all the other facilities of Şişecam rather than only Trakya Glass Factory.

References

M. A. Cohen and R. Halperin. Optimal inventory order policy for a firm using the lifo inventory costing method. *Journal of Accounting Research*, 18(2):375–389, 1980. ISSN 00218456, 1475679X. URL <http://www.jstor.org/stable/2490584>.

Appendix A Consumption and Total Inventory Levels

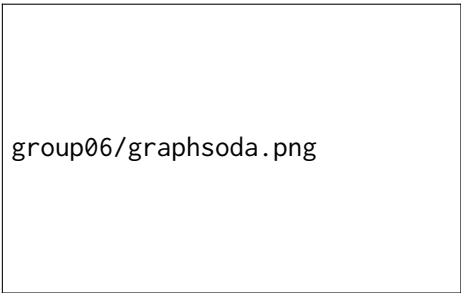


Figure 2: Graph for Soda

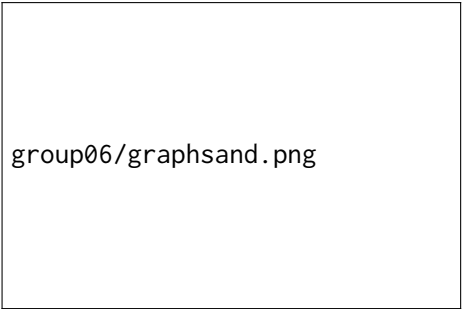


Figure 3: Graph for Silica Sand



Figure 4: Graph for Feldspar

Appendix B Mathematical Model

B.1 Sets, Parameters and Decision Variables

Table 6: Sets, Parameters and Decision Variables of the Model

Sets	
P : Set of production periods, $p \in P = \{1, 2, \dots, m\}$	Y : Set of suppliers, $i \in Y = \{1, 2, \dots, n\}$
Parameters	
c_i : Contract constraint (Use supplier i max c_i times)	q : Warehouse capacity
h_p : Holding cost per unit for period p	w : Holding cost period conversion (0.020833)
r_p : Required amount of raw material at period p	l_i : Lead time for supplier i (in periods)
s_i : Minimum amount that can be ordered from supplier i .	I_0 : Initial inventory ($p=0$)
j_i : The maximum amount we can order from supplier i .	M : Sufficiently large number
t_p : Scheduled (outstanding) orders for period p	z_p : Safety stock for period p
u_{ip} : Unit cost for supplier i for period p	f_i : Fixed ordering cost for supplier i
e_i : Ordering frequency prediction parameter for supplier i .	v : Holding cost multiplier
Additional Parameter	
b_i is an integer parameter (We can order in multiples of b for each supplier i)	
Decision Variables	
A_{ip} : Amount arrived from supplier i in period p	I_p : Inventory level at the end of period p
O_{ip} : Amount ordered from supplier i in period p	X_{ip} : $\begin{cases} 1, & \text{if ordered from supplier } i \text{ at period } p \\ 0, & \text{otherwise} \end{cases}$
Substitute Decision Variables	
K_{ip} : Frozen zone initial controlling variable	D_{ip} : Frozen zone controlling variable
y_1 : $\begin{cases} 1, & \text{if Equation 15 holds} \\ 0, & \text{if Equation 16 holds} \end{cases}$	y_2 : $\begin{cases} 1, & \text{if Equation 17 holds} \\ 0, & \text{if Equation 17 holds} \end{cases}$
y_3 : $\begin{cases} 1, & \text{if Equation 18 holds} \\ 0, & \text{if Equation 6 holds} \end{cases}$	y_4 : $\begin{cases} 1, & \text{if Equation 20 holds} \\ 0, & \text{if Equation 21 holds} \end{cases}$
Additional Decision Variable	
G_{ip} is an integer variable for supplier i in period p	

B.2 General Mathematical Model

$$\min \sum_{p=1}^m \sum_{i=1}^n (X_{ip} f_i) + \sum_{p=1}^m (h I_p) + \sum_{p=1}^m \sum_{i=1}^n (A_{ip} u_i)$$

$$\text{s.t.} \quad I_p = I_{p-1} + \sum_{i=1}^n A_{ip} + t_p - r_p \quad \forall p \in P \quad (1)$$

$$I_p < q \quad \forall p \in P \quad (2)$$

$$A_{ip} \leq j_i \quad \forall i \in Y, \forall p \in P \quad (3)$$

$$I_p \geq z_p \quad \forall p \in P : p \geq \min\{l_1, l_2, \dots, l_n\} \quad (4)$$

$$c_i \geq \sum_{p=1}^m X_{ip} \quad \forall i \in Y \quad (5)$$

$$A_{ip} \leq M X_{ip} \quad \forall i \in Y, \forall p \in P \quad (6)$$

$$A_{ip} \geq s_i - M (1 - X_{ip}) \quad \forall i \in Y, \forall p \in P \quad (7)$$

$$\sum_{p=1}^{l_i} A_{ip} = 0 \quad \forall i \in Y \quad (8)$$

$$h_p = H \vee \sum_{p=1}^m \sum_{i=1}^n (u_{ip} \ e_i) \quad \forall p \in P \quad (9)$$

$$A_{ip} = O_{(i,p-l_i)} \quad \forall i \in Y, \forall p \in P \quad (10)$$

$$X_{ip} \in \{0, 1\} \quad \forall i \in Y, \forall p \in P \quad (11)$$

$$I_p, O_{ip}, A_{ip} \geq 0 \quad \forall i \in Y, \forall p \in P \quad (12)$$

Substitute Constraints

$$K_{i1} = l_i \quad \forall i \in Y \quad (13)$$

$$K_{(i,p-1)} - 1 \leq M y_1 \quad \forall i \in Y, \forall p \in P - \{1\} \quad (14)$$

$$K_{ip} - K_{(i,p-1)} + 1 \leq M(1 - y_1) \quad \forall i \in Y, \forall p \in P - \{1\} \quad (15)$$

$$2 - K_{(i,p-1)} \leq M y_2 \quad \forall i \in Y, \forall p \in P - \{1\} \quad (16)$$

$$K_{1p} \leq M(1 - y_2) \quad \forall p \in P \quad (17)$$

$$K_{1p} \leq M y_3 \quad \forall p \in P \quad (18)$$

$$D_{ip} \leq M(1 - y_3) \quad \forall i \in Y, \forall p \in P \quad (19)$$

$$1 - K_{ip} \leq M y_4 \quad \forall i \in Y, \forall p \in P \quad (20)$$

$$M - D_{ip} \leq M(1 - y_4) \quad \forall i \in Y, \forall p \in P \quad (21)$$

$$A_{ip} \leq D_{ip} \quad \forall i \in Y, \forall p \in P \quad (22)$$

$$D_{ip} \geq 0 \quad \forall i \in Y, \forall p \in P \quad (23)$$

$$y_1, y_2, y_3, y_4 \in \{0, 1\} \quad (24)$$

Additional Constraints

$$G_{ip} = A_{ip}/b_i \quad \forall i \in Y, \forall p \in P \quad (25)$$

$$G_{ip} \in \mathbb{Z}^+ \quad \forall i \in Y, \forall p \in P \quad (26)$$

B.3 Explanation for Substitute Constraints and Decision Variables

As observed from the general model, upper bound of the left hand side of Equation 9 differs among suppliers. In order to overcome this problem, we introduced substitute constraints. They take place of the Equation 9 of the general model during the solution phase (in the code), but for mathematical notation Equation 9 is sufficient. Although these constraints seem complicated and difficult to solve, they do not increase the solution time of the model significantly.

Appendix C User Interface



Figure 5: Main page for Thrace

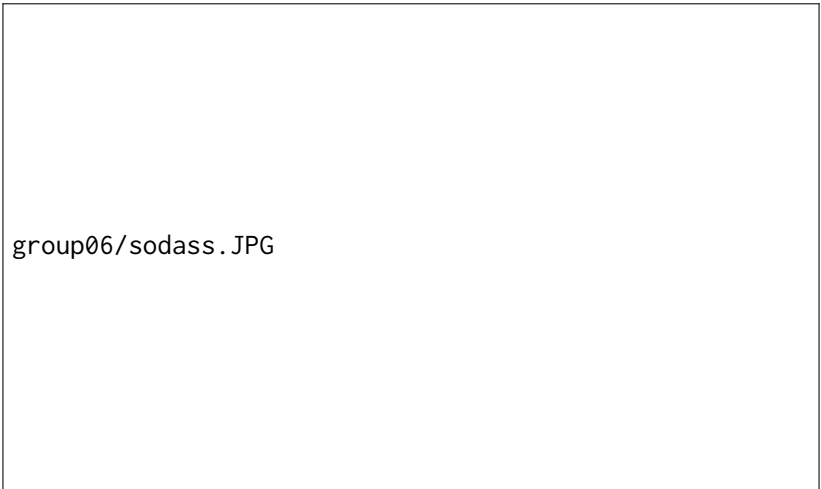


Figure 6: Page of Thrace for Soda

Üretim Hatlarının Performans Takibi ve Yapay Zeka Bazlı Öngörücü Bakım Karar Destek Sistemi Supply Chain Wizard

group07/group_photo.png

Proje Ekibi

Wissam Al-Ali, Mert Ersöz, Ayşe Buse Karabıyık, Nergis Kerimov
Yarkın Özmen, Fahaar Mansoor Pirani, Yunus Emre Yurdagül

Şirket Danışmanı

Evren Özkaya
Ph.D. Kurucu ve CEO
Murat Sağlam
Ph.D. VP Ürün
Yiğiter Çolakoglu
Direktör, Danışmanlık Hizmetleri
Yalçın Arslan
Kıdemli İş Analisti, Danışmanlık
Endüstri Mühendisliği Takımı

Akademik Danışman

Doç. Dr. Arnab Basu
Endüstri Mühendisliği Bölümü

ÖZET

Üretim hatlarında gerçekleşen uzun süreli ve sıklıkla tekrarlayan arızalar ve planlanmamış duruşlar ciddi üretim kaybı ve maliyete neden olur. Projede, SCW'nun Dijital Fabrika platformuyla ilaç üretim hatlarından elde edilen IoT sensör ve operatör verileri kullanılarak hat arıza zamanlaması ve nedenini tahmin eden algoritmaların geliştirilmesi ile hatlara proaktif, hızlı ve yerinde müdahale edilmesi amaçlandı. Makine Öğrenimi (ML) tekniği ile öngörücü bakım yöntemleri konularında yapılan literatür incelemeleri üzerine Destek Vektör Makinelerinin (SVM) uygunluğuna karar verildi. Bu metot kullanılarak bir sonraki arıza tipini tahmin etmek için bir sınıflandırma modeli ve bir sonraki arızaya kadar makinenin ömrünü tahmin etmek için bir regresyon modeli oluşturuldu. Modeller kullanıcı için bir raporlama arayüzüne entegre edildi ve model çıktıları SCW'nun ilaç sektöründeki bir müşterisine sunuldu.

Anahtar Kelimeler: Öngörücü Bakım, Nesnelerin İnterneti, Destek Vektör Makineleri, Makine Öğrenmesi, Veri Görselleştirme

Performance Tracking of Production Lines and AI-Enabled Predictive Maintenance Decision Support System

1 Company Information

Supply Chain Wizard is a management consulting, digital innovation and solutions firm, and a global leader specializing in serialization and traceability, supply chain strategy and operational transformation programs with global headquarters in Princeton, New Jersey, USA and R&D Center in Izmir, Turkey. It has more than 50 employees in four major offices in the USA, Germany, the Netherlands, and Turkey. Despite being a relatively young company, established in 2014, it was named among the top 5000 fastest-growing private companies in the USA in 2018 and 2019 by Inc Magazine and top 500 fastest-growing company in Americas by Financial Times. Also, it was selected as the "Cool Vendor" in Supply Chain Execution Technologies by Gartner in 2019. SCW serves some of the world's largest manufacturers, contract manufacturers and packagers in Life Sciences, Oil & Gas and Consumer Goods industries.

The company has three different business units which are: Consulting, Solutions, and Academy. Under SCW consulting business unit, SCW provides consulting services in the fields of Track & Trace, Supply Chain and, Digital Transformation. As part of SCW Solutions unit, SCW partners with organizations in designing, developing, and implementing digital solutions utilizing state-of-the-art technologies such as Internet-of-Things, Artificial Intelligence (AI), Machine Learning (ML), and Blockchain through its Cloud Platform to enable end-to-end Digital Supply Chain Transformation. Some digital solution platforms include Digital Supply Chain, Digital Factory, and Track & Trace. Lastly, as part of SCW Academy, SCW organizes digital supply chain summits, serialization roundtables, CMO summits and training programs, and runs frequent webinars to support clients with their regulatory and compliance challenges, as well as capability building in the digital age.

2 System Analysis

2.1 Digital Factory - Overall Equipment Effectiveness (OEE) Tracker

OEE Tracker is a production line efficiency tracking solution powered by IoT & advanced analytics to manage and improve line efficiency. It enables data collection from people, IoT sensors, and machines; provides real-time production visibility via mobile app and dashboards; helps identify and eliminate bottlenecks/downtime losses via reports and advanced analytics. OEE Tracker aims to improve OEE performance, reduce costs and increase throughput.

OEE Tracker operates in 4 main steps:

- Collect data from multiple operating units such as panels, tablets, sensors, Programmable Logic Controllers (PLC) machines
- Structure the data by consolidating multiple lines, machines, and user data
- Visualize data via dashboards on the shop floor and mobile devices and applications
- Analyze data and enable decision-making via advanced analytics powered reports

An example of OEE Tracker dashboard can be seen in Appendix A, whereby workers provide input and get updated dynamically regarding the production performance and line status.

3 Problem Definition, Objectives, and Approach

The pharmaceutical companies, working in collaboration with SCW, are facing unexpected and frequent failures on the production lines. These failures lead to a chaotic situation on the shop floor as the operators and the production lines are left idle and the maintenance workers/repairmen are not informed and equipped beforehand to address the failure/stoppage. Amidst the chaos, a considerable amount of time is wasted, which decreases productivity, increases costs and reduces throughput. According to the data collected by the SCW, unexpected machine stoppages and failures cater to between 20-40% of the total downtimes and can cause serious production capacity issues and cost increase.

To proactively eliminate the failures on the production lines, this project aims to create a Predictive Maintenance Decision Support System/Algorithm that could anticipate the timing and type of the next failure based on the sensor and operator data obtained from production lines via OEE Tracker. The output of the algorithm is incorporated via a dashboard, which is designed to be installed on the shop floor to update and notify maintenance, engineering, and production teams regarding the current line status and expected future failure timing and types. By being prepared beforehand for the future failures, this project will help companies achieve an increase in the throughput rate, decrease in the unplanned downtimes, and production & quality costs.

4 System Design

4.1 Data Sets

SCW provided us with four different data sheets collected from a single pharmaceutical production line via OEE Tracker solution. The data sets included the activity history data coming from operator inputs, the sensor data coming from IoT sensors installed on the production line, the shift schedule, and the theoretical product speeds. The activity history data provides an interpretation of core activities occurring on the line: run time, idle time, and downtime along with

the duration and other details. The sensor data displays the number of products coming out of the production line per minute. The shift schedule represents the plan for the allocation of different shifts and workers to available time slots of the line. Theoretical product speeds include the maximum amount of units that can be produced per minute for each product.

4.2 Data Cleaning

According to the data provided by the company, there are eight causes of failures on the production line. We sorted and tabulated the percentage occurrence of each of the failure types, which can be seen in Table 1. However, about 20% of the failures in our data had an unreported failure cause (missing data). The primary reason for the missing failure types in the data was that the operators have to manually change the line status via OEE Tracker and sometimes, during rush hours, operators forget to change the line status.

No.	Failure type	% of failures in the data
1	Blistering	31.7%
2	Cartoner	31.6%
3	Checkweigher	0.4%
4	Foil movement	0.9%
5	Ink Jetter	1.3%
6	Overwrapper	5.4%
7	Popping	8.0%
8	T/E Sealer	0.08%
9	Others (Blanks)	20.6%

Table 1: Causes of Failures on the Production Line

To clean our data, initially, we started with those entries that could be handled manually. Repetitions amongst the rows, across the same period, were removed. Moreover, there was only a single instance corresponding to the failure type T/E Sealer. Therefore, we decided to remove this instance, and we are left with only seven types of failure. Next, using an algorithm that we coded in Python, we randomly populated the blank failure types in proportion to the frequency of the existing seven failure types mentioned in Table 1.

As an example, if a certain failure type has 10 instances in our data out of 100 instances with non-blank entries, the probability of having that type of failure is 0.1 amongst the blank failure types. Therefore 10% of blank failure types would be filled out randomly with that type of failure. This was a crucial exercise to ensure that our data is continuous, consistent, and sufficient in size to build a model.

4.3 Model Building

Predictive Maintenance uses Machine Learning (ML) techniques to learn from historical data and predict the future failure timing and types. It can be formulated using regression and classification models. Each of these models is used for a different purpose. A classification model, in our case, predicts the next failure type in the next cycle. A regression model, on the other hand, predicts the meantime to failure (MTTF).

Model Attributes

Once the data sheets were cleaned, we started extracting attributes to be used as inputs in machine learning models. The data was adjusted according to the attribute table that we had designed. In the attribute table, the rows correspond to cycles; a cycle is defined as the time between two consecutive failures, regardless of the type. The columns are attributes (e.g. shift type, product type, the failure type of the previous five cycles, the time to failure of the previous five cycles, and the time to repair of the previous five cycles), and the two output columns (time to failure and failure type in the current cycle).

The machine learning models conduct both classification and regression analysis. The output of the former is the failure type, which is more likely to occur in the next cycle, whereas the output of the latter is the meantime to failure (MTTF) in the next cycle. The sample output generated is as follows; **There is a 60% chance that the next failure will be a Blistering and with a 66% chance it will occur after 100 ± 60 minutes.**

Machine Learning (ML)

For our project, we primarily used Support Vector Machines (SVM) as the ML tool. SVM is a useful ML technique for data classification and regression analysis. This method is applied to higher-dimensional non-linear classification problems that make it useful for many situations. Support Vector Machines can efficiently conduct a non-linear classification using the kernel trick; it works by mapping the inputs it receives into higher dimensional feature spaces (Bousqaoui et al. (2017)). Support Vector Machines can be implemented using the Scikit-Learn package, which is a Python module that integrates many Machine-Learning algorithms (Cavalcante et al. (2019)). Moreover, different python packages such as Numpy, Matplotlib, and Pandas can be used to conduct data pre-processing, data analysis, and visualization for the dataset provided (Cavalcante et al. (2019)).

The attribute table was constructed in such a way as to ensure that the data is continuous and consistent. However, the size of the data provided is extremely small (about 1000 instances). Moreover, there is a huge class imbalance in the data when it comes to our classification problem, where the classes refer to the failure types. Blistering and Cartoner are the dominant classes which can be seen in Appendix B. The data entries were fed into the model numerically.

Furthermore, to train the model for both classification and regression, we have to make the numerical features to be suitable and consistent within a range for the model. We cannot have data where each attribute or feature has a different range of values. Therefore, we need to scale our data into the same range for each feature with the mean of each feature as 0. Standardization of a dataset is a common requirement for many machine-learning estimators: the majority classes might overwhelm the minority classes. The results can be biased if the individual features are not scaled. Hence, we standardized each feature column, in the training set, according to a standard normal distribution (Hsu et al. (2003)).

Feature Selection Technique

The next thing to consider was the direction of our execution for the model. Once the direction was confirmed, the number of features to be used in each of the two models was selected. Not all the features are used simultaneously as they may not yield the maximum accuracy for the model. To select which direction and which of the features to use, we used a python built-in feature selection technique called "Mutual Information".

Mutual Information (MI) of a feature (attribute) with an output means how much the change in this feature contributes to the change in the output. The higher the mutual information of a feature is, the more significant it is towards the calculation of the output. The result of the feature selection technique was that the model should be executed using classification first and then combining its output to execute the regression model.

The design for executing the model can be seen in Appendix C. For the classification model, the maximum accuracy is attained when the first 19 best features (high MI values) are used out of 73 features. For the regression model, the maximum accuracy is attained when the first 46 best features (high MI values) are used out of 73 features.

The Mutual Information between two random variables X and Y is defined as follows:

$$MI(X;Y) = \sum_{x,y} p(x,y) \log \frac{p(x,y)}{p(x)p(y)}$$

where p is the observed probability of an event from the input table. We created a bar plot of the mutual information given by each of the features for both outputs. In accordance with the execution design in Appendix C, the resulting MI bar plot can be seen in Appendix D. The MI bar plot shows that all the classes (failure types) contribute significantly towards the expected lifetime value.

5 Model Training and Verification

5.1 Model Training

We split the data into two parts; 75% of it was used for training purposes while the remaining 25% was used for the testing purpose. Using cross-validation with

a greater number of folds has drawbacks. Firstly, the size of the data is small (around 1000 entries). Secondly, Blistering and Cartoner are the two failure classes that have a greater proportion of data. Therefore, dividing the data into a greater number of folds/sets may allow the majority classes to overwhelm the minority ones. Therefore, we executed the ML model using a 3-fold cross-validation technique.

Classification Model

We chose the first 19 features/attributes/inputs based on the Mutual Information (MI) graph, in Appendix D, and the fact that they attain the greatest accuracy value. The data points for each of the 19 features were standardized. Again, the output of this model is an array of the likelihood of having a different type of failure in the next cycle. As an alternative, we tried other machine learning models such as Support Vector Classification (SVC) and Random Forests. However, we attained the highest accuracy using the SVC model. Thus, we decided to use the SVC for the classification model.

Regression Model

We chose the first 46 features/attributes/inputs based on the Mutual Information (MI) graph. The data points for each of the 46 features were standardized. Again, the output of this model is a range for the remaining time to the next failure. We tried other machine-learning models such as Support Vector Regression (SVR) and Random Forests models. Upon initial training, the accuracy value for SVR was the highest. Thus we decided to use SVR for the regression model.

5.2 Model Verification

For the classification model, when we ran our SVC model on the testing data, we got an accuracy of 80.6%. For the regression model, when we ran our SVR model on the testing data, we got the accuracy which can be seen in the graph in Appendix E. In the graph, the y -axis corresponds to the percentage accuracy, while the x -axis represents the time allowance or deviation from the predicted time in minutes. For instance, if the time allowance is ± 60 minutes, from the predicted time, the accuracy of our model is 67.7%.

6 User Interface

To implement our model, we built a real-time user-interactive dashboard via Python, which is integrated to our models and dataset. The dashboard aims to summarize and display our model outputs on the shop floor to enable prediction of timing and types of line failures before they occur, and proactively address the issues on the production lines. The dashboard (preliminary version is in Appendix F), consists of the following modules:

- Information cards showing the Shift, Work Order, Product, Activity Status & Type (e.g. Down Time Failure – Blistering) of the selected line

- Last 5 failures with details such as timestamp, failure type, duration and variance
- A pareto chart displaying the breakdown of next failure types with probabilities based on the classification model
- Estimated mean time to failure (MTTF) based on the regression model
- A performance metric, shaped like a speed tachometer and characterized by color, that displays the regression model's confidence in the prediction
- Combined sensor values and activity types (run/down/idle) over time
- A recommendation box stating the necessary course of action to be done on the line as a result of the model outputs (e.g. be prepared for a cartoner failure which has a high likelihood of happening in 60 minutes)

7 Project Contributions

7.1 Benefits to SCW

- The prediction algorithms and visualizations will be utilized by SCW to set the foundation for the new Digital Factory – Maintenance Module. The maintenance module will be integrated into all other Digital Factory modules (e.g. OEE Tracker, Labor Tracker, Asset Tracker, IoT Hub) and will contribute towards enabling end-to-end visibility and integration of manufacturing operations. This will create a competitive advantage against the point solutions focusing on only one manufacturing problem or area.
- The future Maintenance module will increase software revenues of SCW by selling the new module as a standalone tool as well as creating opportunities for cross-selling to the existing Digital Factory customers.

7.2 Benefits to SCW's Clients

The models and the dashboard will provide the following benefits to SCW's clients:

- **Increased throughput/profit:** The predictive maintenance module will decrease unplanned downtimes by proactively addressing the issues on the production lines before the failures occur. This will result in time savings and increased OEE performance. These time savings can be utilized to produce more units and increase profits (especially relevant for factories experiencing capacity issues).
- **Decreased costs:** Time savings can also be utilized to decrease the headcount and direct labor costs if the factories choose to produce the same number of units with fewer labor hours. Better utilization of maintenance workers and timely/fewer large-scale repairs will result in reduced maintenance costs. The life expectancy of assets will increase and will reduce the

Capital Expenditure (CAPEX) investments of the company. The tool will also automate some manual reporting and data collection which saves time and reduces the indirect labor costs.

- **Decreased risk:** The tool will reduce quality events and will improve safety and reliability.
- **Speed-to-market:** The tool will reduce end-to-end cycle times and enhance response times to customer orders.

Among these benefits, increased throughput/profit and decreased costs (direct labor cost) can be quantified based on the existing datasets, model predictions and some assumptions. To illustrate, there are three different cases considering the prediction accuracy of the failure timing and type (classification and regression models), as shown in Table 2. It should be noted that the probability of occurrence for each case is calculated by multiplying the accuracies of the two models (classification and regression models) and the "Failure Time % Changes" are based on SCW's domain expertise.

Case	Description	Probability	Failure Time % Change
1	The timing and type of the failure are accurately predicted by the models	0.55	↓ 40%
2	The timing of the failure is accurately predicted, whereas failure type is NOT predicted accurately	0.13	↓ 20%
3	The timing of the failure is NOT predicted accurately regardless of the correct or incorrect prediction of failure type	0.32	↑ 10%

Table 2: The different cases based on the prediction accuracies of the models and the probabilities/failure time change % per case

Given the three different cases and how they affect the failure time, our models incorporated via dashboard is expected to decrease the Total Failure Time by:

$$(0.55 \times 40\%) + (0.13 \times 20\%) - (0.32 \times 10\%) = 21.4\%$$

Based on our dataset, the total failure time per year is around 94,000 minutes for a single pharmaceutical production line. As the use of dashboard can decrease the total failure time by 21.4%, SCW's clients can save 20,116 [21.4% x 94,000] minutes per line annually. The executives can utilize these time savings a) to reduce headcount which results in decreased labor cost or b) to increase throughput which leads to increased profit. To calculate the financial impact of these two scenarios and to generalize our findings for a single line to an entire mid-size pharmaceutical factory in the US, we gathered the following information from SCW:

Parameter Name	Abbreviation	Value
Number of Lines	NL	10
Number of Operators per line	NO	5
Labor cost (\$/min)	LC	\$0.20
Production per minute (units/min)	PPM	25
Revenue (\$)	R	\$50M
Gross margin (%)	GM	50%
Total annual production (units/year)	TP	20M
Expected annual time savings per line (min/year)	TS	20116

Table 3: Typical parameters of a mid-size pharmaceutical factory in the US

Based on the above parameters, by allocating the time savings obtained from the use of dashboard, a mid-size factory in the US can reduce the direct labor costs by \$201,160 or increase profit by \$6.3 M as shown in the below table:

Scenario	Result	Equation	Amount
Reduce Headcount	Decreased Labor Cost	$TS \times NL \times NO \times LC$	\$201,160
Increase Throughput	Increased Profit	$\frac{TS \times NL \times PPM \times R \times GM}{TP}$	\$6.3M

Table 4: Potential financial impact areas of the Dashboard for a mid-size pharmaceutical company in the US

While the executives can choose to completely reduce the headcount or increase profit, it is also possible to mix and match these two strategies. In that case, total annual time savings can be partially distributed among these two strategies (e.g. 80% of total time savings is allocated for direct labor cost savings and 20% is allocated for profit increase) to satisfy their needs.

7.3 Project Implementation

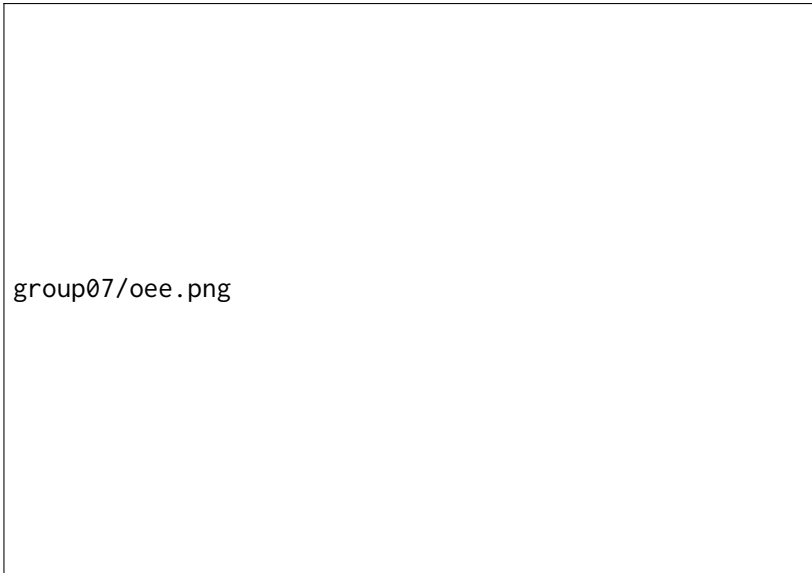
Since SCW does not currently have predictive maintenance system capabilities in place, they will utilize our research, models and dashboards as a starting point to build it in the near future. SCW invited us to a meeting with one of their clients in pharmaceutical industry in Turkey to discuss the co-development opportunities for SCW’s future maintenance module.

We presented our project in the meeting and asked for SCW and its client’s upper-level management feedback on our model and the dashboard. We received positive feedback on our model and dashboard and executives suggested enriching the dataset by including the full maintenance/failure history of assets, leveraging other sensors (vibration, temperature) and incorporating process information to improve the model accuracy.

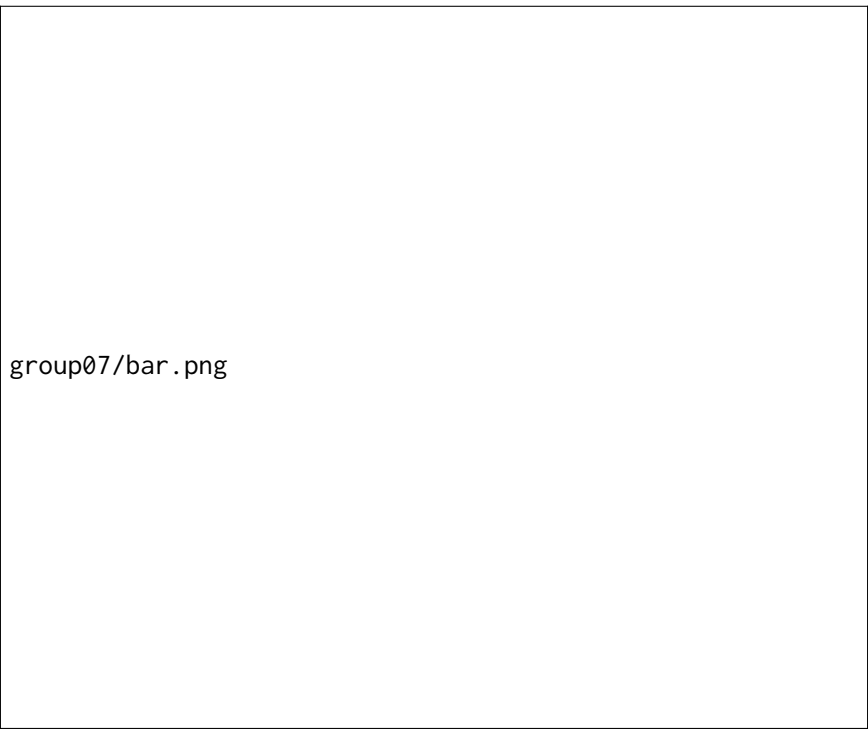
References

- H. Bousqaoui, S. Achchab, and K. Tikito. Machine learning applications in supply chains: An emphasis on neural network applications. In *2017 3rd International Conference of Cloud Computing Technologies and Applications (CloudTech)*, pages 1–7. IEEE, 2017.
- I. M. Cavalcante, E. M. Frazzon, F. A. Forcellini, and D. Ivanov. A supervised machine learning approach to data-driven simulation of resilient supplier selection in digital manufacturing. *International Journal of Information Management*, 49:86–97, 2019.
- C.-W. Hsu, C.-C. Chang, and C.-J. Lin. A practical guide to support vector classification"(<http://www.csie.ntu.edu.tw/~cjlin/papers/guide/guide.pdf>). 2003.

Appendix A OEE Tracker Dashboard



Appendix B Proportion of Classes in the Classification Model



Appendix C Model Execution Flow-Chart



Appendix D Results of Mutual Information

group07/c2rMI.png

Appendix E Regression Model Test Accuracy

group07/regres.png

Appendix F Preliminary Predictive Maintenance Dash- board Design



Konfigürasyon Yönetimi için Karar Destek Sistemi Uygulaması

Nurol Makina ve Sanayi A.Ş.

group08/grpfoto.png

Proje Ekibi

Deniz İrem Çavuş, Zeynep Çulhaoğlu, Begüm Ercan,
Berkan Erdil, Mert Erünsal, Miraç Kuru, Defne Küçükpilakçı

Şirket Danışmanı

Dilek Bilgiç
Konfigürasyon Yönetimi Takım Lideri
Betül Özmen
Konfigürasyon Yönetimi Mühendisi

Akademik Danışman

Doç. Dr. Yiğit Karpat
Endüstri Mühendisliği Bölümü

ÖZET

Nurol Makina, müşteri siparişlerine göre üretim yapan bir savunma sanayi şirkettir. Firmanın mühendislik bazlı veya müşteri ile ilgili çeşitli değişiklik talepleri bulunmakta ve her talebin farklı olmasından dolayı her talebin değerlendirilmesinde çeşitli kriterler kullanılmaktadır. Konfigürasyon Kontrol Kurulu, uluslararası standartlarda tanımlanan belirli kriterler doğrultusunda değişiklik taleplerine karar verir. Bu proje, Konfigürasyon Kontrol Kurulu kararlarını desteklemek için AHP'yi kullanan çok kriterli bir karar destek sistemini içerir.

Anahtar Kelimeler: Karar Destek Sistemi, Konfigürasyon Yönetimi, AHP

Implementation of DSS for Configuration Management

1 General Information About the Company

Nurol Makina works for the defense industry since 1976. The company produces 4x4 tactical wheeled armored vehicles in its facilities in Ankara. Five different types of armored vehicles are produced in these facilities which have a capacity of 800-1,000 pieces per year. During the manufacturing process, Configuration Management (CM) is responsible for parts traceability, change tracking, configuration control and management.

2 System Analysis

2.1 Current System Information

Nurol Makina has five different modifiable products. Engineering change request(ECR) for these products can be created at every stage of production. The first decision point regarding the change request created is to decide whether a publication for ECR exists, and then whether it can be written. The ECR, written according to the existence of the publication, is evaluated according to its applicability in the Research and Development(R&D) board. If R&D makes the applicable decision for this ECR,the next stage is the Configuration Control Board (CCB) meeting.The whole process is shown in Figure 1.

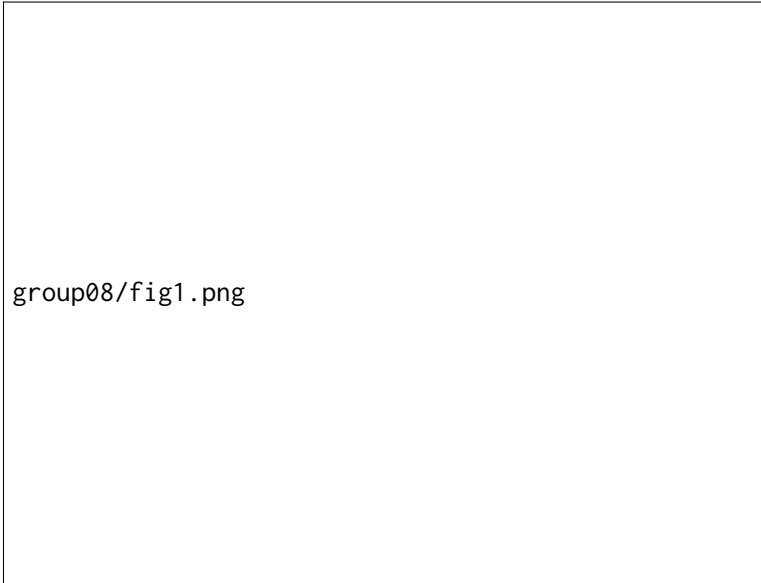


Figure 1: ECR Process

One of the current applications of Nurol Makina regarding ECR is the fast decision cycle. In order to get an ECR into a fast decision cycle, there must be a very serious urgency in production. If vehicle production is slowing down significantly due to a part, that part should be taken into a fast track. According to the

information received from the company, ECRs can be categorized as follows:

1. *Formatting*: This type of change should be made in cases such as very small errors in the drawing, not putting a comma where it should be and typographical error. No department needs to take action when such changes are required. The only thing to do is to inform the CCB about these changes.
2. *Improvement*: These are engineering changes that provide labor or cost savings or increase customer satisfaction. Actions vary according to the decision taken at the CCB.
3. *Repair-Fix*: It is the compulsory engineering to be done due to any malfunction or design error.

2.2 Problem Definition

The weekly generated engineering change requests are high, so the evaluation process takes a long time for both R&D and the Configuration Control Board. In this process, evaluations are carried out manually and based on employee opinion. As a result, human error rate increases and evaluation periods lead to waste of time. At the same time, since the incoming ECRs are unique, the requests should be evaluated according to each criterion and request. In similar situations, re-evaluations lead to a process that does not add value. In addition, in the evaluation of engineering change requests, the decisions that are rejected or deemed unsuitable could not be examined in detail. Long meetings, high number of ECRs and lack of pre-meeting evaluations are the reasons for this. As a result, decisions rejected in terms of the budget, although technically applicable, cannot be evaluated. Different options could not be created with a budget expansion.

2.3 Literature Research

Since it is expected that the priority decision will be determined with R&D and CCB evaluations for weekly ECR, an improved version of the R&D ECR selection model is required for the solution Bin et al. (2015) The purpose of the R&D ECR selection is to select the decision that provides maximum benefit in accordance with the criteria and constraints in the decision-making process where there are too many variables and decision factors. There are multiple methods in the literature to assist in the decision-making process such as TOPSIS (Order of Preference Technique with Similarity to Ideal Solution), VIKOR (ViseKriterijumsa Optimizacija I Kompromisno Resenje), AHP (Analytical Hierarchy Process) and Knapsack-Outranking method. The reason for choosing these methods is that multiple options can be examined in evaluations with too many parameters.

3 Proposed System

3.1 Inputs and Outputs of the System

The overall model consists of two main stages as following, an R&D Elimination Model, and a CCB Decision Support Model. Each model requires many inputs

to run and all inputs are collected each week. The output of the system gives two main outputs for each model. For the first model, output is number of approved ECRs. For the second model, outputs is appropriate decision for each ECR. The interface of these two models also shows the percentage of budget increase for each ECR for being accepted and it gives flexibility to the company.

3.2 R&D Elimination Model

Each week, approximately 40 ECRs are coming to the company. At the first stage, these ECRs are evaluated in RD department to decide which ECRs are approved for evaluating in CCB meeting. Main purpose of this model is properly select appropriate ECRs.

Indices:

N: Number of ECRs that will be evaluated in the RD meeting: $ECR_i, i = 1, 2, \dots, N$.

M: Index of the criteria that the ECR will be evaluated accordingly stock availability, open production order, open purchase order, on vehicle, testing requirements, affected projects, technical applicability, configuration item: $k = 1, 2, \dots, M$.

Parameters:

$b_{k,i}$: Available resources for each criterion k for each ECR i. This limitation represents the resource capacity can be allocated for each criterion for the particular ECR and will be a $M \times N$ matrix.

ov_i : On vehicle state of the each ECR i,

$$ov_i : \begin{cases} 1 & \text{if the changed parts are not assembled yet} \\ 0 & \text{otherwise} \end{cases}$$

ap_i : Affected projects from the each ECR i,

$$ap_i : \begin{cases} 1 & \text{if other projects are affected} \\ 0 & \text{otherwise} \end{cases}$$

ta_i : Technical applicability of the each ECR i,

$$ta_i : \begin{cases} 1 & \text{if the ECR is technically applicable} \\ 0 & \text{otherwise} \end{cases}$$

$v_{k,i}$: Resources consumed by each constraint k, for each ECR i, will be a $M \times N$ matrix. This constraint demonstrates how much will the ECR i will consume from the resources allocated for this criterion for the particular ECR

q_i : Percentage of budget increase for ECR i for being accepted.

$$\sum_{k \in M} v_{k,i} = \sum_{k \in M} b_{k,i} q_i \quad \forall i \in N \quad (1)$$

Decision Variable:

$$y_i = \begin{cases} 1 & \text{if budget constraint is satisfied} \\ 0 & \text{otherwise} \end{cases}$$

Model:

$$\max \sum_{i \in N} y_i \quad (2)$$

$$\text{s.t. } v_{k,i} y_i \leq b_{k,i} \quad \forall i \in N, \forall k \in M \quad (3)$$

$$ta_i \geq y_i \quad \forall i \in N \quad (4)$$

$$d_i \in \{0, 1\} \quad \forall i \in N \quad (5)$$

By using a parameter q_i , the company shows percentage of budget increase for ECR i for being accepted (1). Objective function (2) of this model which is objective function maximizes the number of ECR which satisfies budget constraint. Budget constraint (3) of this model makes sure that consumption of each ECR of given criteria cannot exceed available resource of that criteria. Applicability constraint (4) shows that given ECR should satisfies technical applicability.

3.3 CCB Decision Support Model

CCB Decision Support Model is developed for using in CCB meetings for making appropriate decision about each ECR Tüminçin (2016). ECRs that are evaluated is coming from approved ECRs from the RD Elimination Model. As an input, criteria weights have to be determined to give proper decision. Therefore, AHP method is utilized to calculate criteria weights. AHP is a method that determines each criterion importance between each other according to the decision maker.

1. AHP was filled comparison table by using Table of Relative Scores with the company which is shown in Appendix A.1 and Appendix A.2.
2. Excel was used to calculate the w_k Firstly, column totals obtained shown in Figure A.3 and all the values in the table were divided by the column totals to obtain normalized matrix which is also shown in A.4.
3. The “Priority Vector” which is shown in Appendix A.5 was calculated by taking the average of each row of the normalized matrix.
4. The resulting vector was multiplied by the comparison matrix given initially, and the “All Priorities Vector” was built. This vector was used for computing “Consistency Rate” which shows criteria comparisons consistency. This rate should be less than 0.1 to obtain consistent w_k values.
5. Each element of the matrix of all priorities was divided by the elements of the priority vector to obtain λ values.

6. Taking average of λ values.
7. Consistency index was calculated by using the following formula:

$$CI = \frac{\lambda_{max} - n}{n - 1} \quad (6)$$

8. Random index value was determined by considering Random Index Table (Appendix A.2) which is 1,12.
9. Consistency rate was calculated using this formula and all of these steps between 4 and 9 are shown in Appendix B:

$$CR = \frac{CI}{RI} \quad (7)$$

Indices:

N: Number of decisions can be made for the current ECR retrofitted for factory, call in retrofitted, reworked for factory, call in reworked, rejected: $d_i, i = 1, 2, \dots, N$

M: Index of the criteria that the ECR will be evaluated accordingly in the CCB meeting schedule, production state, cost, safety critic, vehicles on the field: $k = 1, 2, \dots, M$

Parameters:

w_k : The set of M criteria weights that specifies the importance of each criteria in the decision-making process. (will be derived from AHP and normalized.)

$val_{k,i}$: Value assignments for each criterion k in each decision type i, MxN matrix. Each criterion will be ranked from 1 to M by its importance for the considered decision type. (will be normalized.)

$v_{i,k}$: Resources consumed by each constraint k, for each decision i, NxM matrix. This constraint will demonstrate how much will the decision i will consume from the resources allocated for this criterion for the decision.

$b_{i,k}$: Available resources for each criterion k for each decision i for the particular ECR.

$w_k val_{i,k}$: Overall weight of each criteria k for each decision type i. This will be calculated by taking the weighted sum of each criteria's importance by considering the company's usual standards w_k and company's approach to current ECR $val_{i,k}$. While looking for the decision types that fit criteria, it's important to choose the one that fit company's overall and the particular criteria.

$$sc: \begin{cases} 1 & \text{if the current project is safety critic.} \\ 0 & \text{otherwise} \end{cases}$$

q_i :Percentage of budget increase for decision i for being selected.

$$\sum_{k \in N} v_{k,i} = \sum_{k \in N} b_{k,i} q_i \quad \forall i \in N \quad (8)$$

Decision Variable:

$$d_i : \begin{cases} 1 & \text{if decision } i \text{ is chosen.} \\ 0 & \text{otherwise} \end{cases}$$

Model:

$$\max \sum_{i \in N} d_i (w_k val_{k,i}) \quad (9)$$

$$\text{s.t.} \quad (1 - sc) v_{i,k} d_i \leq b_{i,k} \quad \forall i \in N \quad (10)$$

$$sc \leq d_1 + d_2 + d_3 + d_4 \quad (11)$$

$$\sum_{i \in N} d_i = 1 \quad \forall i \in N \quad (12)$$

$$d_i \in \{0, 1\} \quad \forall i \in N \quad (13)$$

By using a parameter q_i the company shows percentage of budget increase for decision i for choosing(8). Main purpose of this model is giving proper decision for each ECR. Objective function of this model (9) selects the decision which maximizes overall weights while it satisfies constraints. Budget constraint (10) shows that the resources consumed for each criteria cannot exceed the available resources for that criteria. In this constraint, if given ECR is safety critic, this constraint is not worked since safety critic is the most important criteria. In safety constraint (11), if this ECR is safety critic, the decision of rejected cannot be chosen. Last constraint (12) provides that this model can be select only one decision.

3.4 Validation and Verification

In this step, Excel Solver is utilized. These two models are coded and solved via Excel Solver. First of all, different scenarios were created by using sample data to observe whether the model outputs is true or not. By using these scenarios, parameters are changed in the Excel Solver and the code is observed whether working correctly or not. After sample data, real data which is provided by the company is used to observe whether the model works correctly or not. Whether the decisions made by the Excel Solver are similar to those made by the company will be evaluated at a meeting to be arranged with the company. To determine the accuracy of criteria weights, consistency rate is used. In this step, Excel is utilized. In the AHP method, consistency rate shows whether these criteria weights

are consistent or not. After several calculations which is shown in Appendix B, if consistency rate is higher than 0.1, these weights will not be considered as consistent. At this point table is filled with the company and our consistency rate is 0.27 which is greater than 0.1. This rate will be decreased after another meeting as our plan.

4 Implementation and Integration to the Company's Systems

4.1 Implementation

For the models to be used practically by the company, we carried the input collection and output display processes over Excel tools. Process starts in the R&D department's regular meeting, by filling input tables on Excel Sheets for the R&D department's elimination stage. After these inputs are fulfilled, the model will work by clicking on just one button as seen on Appendix C. Results of this stage will additionally provide the needed improvement of budget; RD representatives will review the results of the model during their meeting and manually transfer the selected ECRs to the next stage, to the Configuration Control Board. CCB will fulfill the input sheets for each ECR during their weekly meetings, and the decision support system will output the suggested solutions and possible improvements for budget, the interface of this stage can be found in Appendix D. CCB will consider these suggestions when concluding the final decisions. Briefly, both the models will work simultaneously with the meetings to support their decisions with numerical data. For all representatives to be able to use the system, we provided short explanations of both the inputs and the outputs, explaining the range and properties of these parameters and variables. We've also presented the setup calculations that are processed with AHP on the Excel tools in case of need, which is accessible on Appendix E.

4.2 Project's Contribution

Our project turns the fully human-controlled ECR evaluation system into an autonomous system, which will help company representatives to do the elimination and decision making more consistently. We expect to increase the time efficiency of both departments' evaluations, minimizing human errors within the evaluations and, reducing costs arising from erroneous decisions by this new interface. Although it is very soon to show the quantified savings due to calendar of the project, we expect easier and quicker decision-making processes with less errors, which would generate honorable savings. Furthermore, our reports and the implementation interface will be a part of the onboarding process to help new employees to understand the processes of the CC department.

5 Conclusion

As a conclusion, after the examination of the current system we build a decision support system for Nurol Makina. We determined the company's criterias for eval-

uating each engineering change request and developed two models. First model is selecting applicable ECR's for the company and presenting them for further evaluation. Thus, we exclude human interpretation from the selection process and provide required budget enhancement ratios for reassessment of results. Second model evaluates each ECR individually by criterias company used. Through the second model, Nurol Makine will be able to choose the most suitable decision. We decrease total time spent for elimination and evaluation of ECR's and provide more stable results to the company. Also prevent human-driven mistakes during this process. Also our program presents possible budget changes to users. If a better decision or ECR is eliminated because of trivial budget insemination, users can reevaluate the decision for maximizing decision returns. As Nurol Makina requested, we present an adjustable evaluation and selection system independent from ECR's. Because each ECR is unique and can be in any format. So our models only use the data users provide and the criterias our team determine. We build our model in Excel for company. Excel provides fast solutions with a user-friendly interface and enables input data easily. Also data based decision systems provide collection of data and this will give a prospective about decision making in future stages.

References

- A. Bin, A. Azevedo, L. Duarte, S. Salles-Filho, and P. Massaguer. Rd and innovation project selection: Can optimization methods be adequate? *Procedia Computer Science*, 2015.
- F. Tüminçin. Analitik hiyerarşik proses(ahp) ile bir karar destek sistemi oluşturulması: Bir Üretim İşletmesinde uygulama. 2016.

Appendix A AHP Method

A.1 Scale of Relative Importance

Importance Values	Definition
1	Equal Importance
3	Moderate Important
5	Strong Important
7	Very Strong Important
9	Extreme Important
2,4,6,8	Intermediate Values

A.2 Comparison Table

Criteria	Schedule	Prod. State	Cost	Safety Critic	Vehicles on the field
Schedule	1	3	7	1/9	1/3
Prod. State	1/3	1	3	1/9	1/5
Cost	1/7	1/3	1	1/9	1
Safety Critic	9	9	9	1	9
Vehicles on the field	3	5	1	1/9	1

A.3 Column Summation

Criteria	Schedule	Prod. State	Cost	Safety Critic	Vehicles on the field
Schedule	1,00	3,00	7,00	0,11	0,33
Prod. State	0,33	1,00	3,00	0,11	0,20
Cost	0,14	0,33	1,00	0,11	1,00
Safety Critic	9,00	9,00	9,00	1,00	9,00
Vehicles on the field	3,00	5,00	1,00	0,11	1,00
TOTAL	13,48	18,33	21,00	1,44	11,53

A.4 Normalized Matrix

Criteria	Schedule	Prod. State	Cost	Safety Critic	Vehicles on the field
Schedule	0,07	0,16	0,33	0,08	0,03
Prod. State	0,02	0,05	0,14	0,08	0,03
Cost	0,01	0,02	0,05	0,08	0,09
Safety Critic	0,67	0,49	0,43	0,69	0,78
Vehicles on the field	0,22	0,27	0,05	0,08	0,09

A.5 Priority Vectors(Weights of Criteria)

Criteria	Schedule	Prod. State	Cost	Safety Critic	Vehicles on the field	Priority Vector
Schedule	0,07	0,16	0,33	0,08	0,03	0,14
Prod. State	0,02	0,05	0,14	0,08	0,02	0,06
Cost	0,01	0,02	0,05	0,08	0,09	0,05
Safety Critic	0,67	0,49	0,43	0,69	0,78	0,61
Vehicles on the field	0,22	0,27	0,05	0,08	0,09	0,14

Appendix B Calculations of Consistency Rate



Figure 2: Consistency Rate Calculations

Appendix C R&D Model Interface

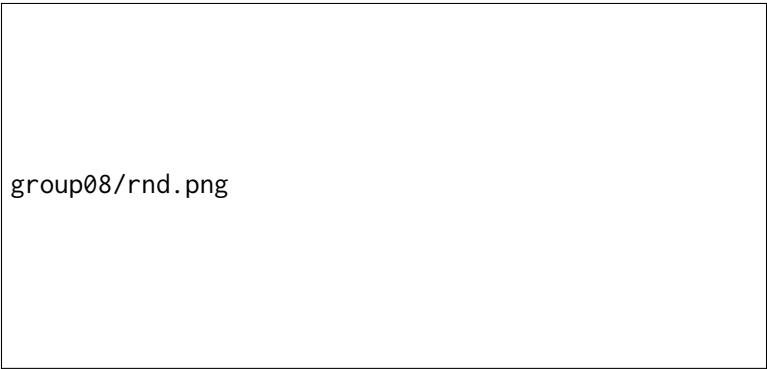


Figure 3: R & D Model

Appendix D CCB Interface

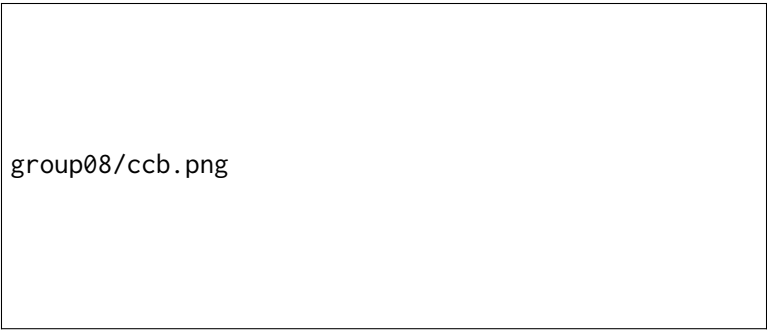


Figure 4: CCB Interface

Appendix E AHP Interface

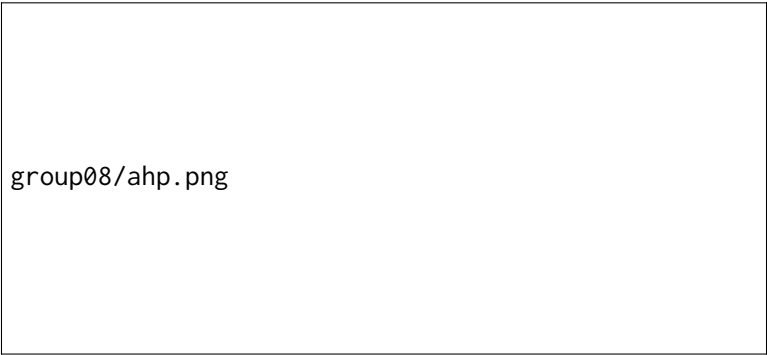


Figure 5: setup calculations that are processed with AHP on the Excel tools

Güvenlik Görevlisi Atama Süreci için Karar Destek Sistemi

Bilkent Sivil Savunma ve Güvenlik Müdürlüğü

group09/groupphoto.png

Proje Ekibi

Berk Berensel, Ata Denizoglu, Selen Erbakan,
Ahmet Umut Gurkan, Cagan Kuru, Omer Oncu, Emir Turkoğlu

Şirket Danışmanı

Gülşah Polatoğlu, Belemir Tamtabak
Bilkent Sivil Savunma ve Güvenlik
Müdürlüğü, Endüstri Mühendisliği
Takımı

Akademik Danışman

Dr. Emre Uzun
Endüstri Mühendisliği Bölümü

ÖZET

Projenin amacı, Bilkent Üniversitesi'nde görevli güvenlik görevlilerinin çalışma saatlerini ve yeterliliklerini gözeterek adil bir karar destek sistemi ile 6 haftalık bir atama çizelgesi oluşturulmasıdır. Mevcut sistemde atama çizelgesi yetkili kişi tarafından sezgisel olarak yapıldığından anlık vardiya değişimlerinde, güvenlik görevlilerinin izin günlerinin, fazla mesai ve çalışma saatlerinin belirlenmesinde sorunlar yaşanmaktadır. Proje kapsamında yöneylem metodları kullanılarak matematiksel model geliştirilmiştir. Matematiksel model, problemin büyüklüğünden ötürü sezgisel yöntem ile çözülmüştür. Sezgisel modelin sonucu çeşitli optimizasyon programları ve geçmiş atama verileri ile karşılaştırılarak doğrulanmıştır. Sezgisel model kolay kullanılabilirliği gözeterek geliştirilen arayüze entegre edilmiş ve gerçek veriler doğrultusunda çeşitli senaryolarda test edilmiştir.

Anahtar Kelimeler: Sezgisel Model, Karar Destek Sistemi, Adillik

Decision Support System for Security Guard Assignment Process

1 General Information

Bilkent University is located in Çankaya/Ankara and the campus area is 50,000 square meters. There are 12,000 registered students and 1,000 faculty and staff at Bilkent University. To enforce rules and provide security to students, staff and faculty. SSGM's main responsibilities are; checking ID cards of people who enter the campus, registering cars, preventing noise pollution inside the buildings, dormitories, and classrooms, providing security. To implement those requirements, SSGM is responsible for 50 posts with 134 employees; Ahmet Özban is the director, 5 chiefs of the unit, 14 administrative staff, and 114 security guards.

2 System Analysis and Problem Definition

In this part, system analysis of SSGM is emphasized; the symptoms and complaints related to the problem, and the problem definition and scope are defined.

2.1 System Analysis

SSGM is responsible for managing 12 security zones consisting of 50 posts which can be seen in Appendix A. SSGM assigns security guards to the posts and patrols. Some security guards can only be assigned to specific posts, some of the security guards can switch among other posts. For example, security guards in the East Campus high schools' regions (10th and 12th zones) or security guards in housing zones where most of the residents are foreign are generally fixed and do not change. The reason behind this is to get to know the students, teachers, parents and be able to recognize any foreign person in the area. These 12 zones are controlled 24/7 in five shifts (08:00-16:00 / 16:00-00:00 / 00:00-8:00) and (08.00-20.00, 20.00-08.00). 8-hour shifts are indexed as 1, 2 and 3, while 12-hour shifts are indexed as 4 and 5. There are 24/7 patrols by motorcycle, on foot, and by vehicle (two for the Main campus and one for the East campus). The weekly standard working time for security guards is between 45 and 60 hours. They have one day off per week and 5 days of administrative leave each year. SSGM considers primarily four criteria before assigning. These are related to security guards, their experience, education, certificates and knowledge of foreign language. SSGM determines past four experiences while hiring, where and how many years security guards have worked, and positions. At initialization, every rookie security guard is assigned to the gates after the orientation program to develop their communication skills and adapt to the system. Security guards who know foreign languages become prominent in regions where foreign faculty members live. Determining the areas where they were successful before and analyzing their expertise in any field are among the criteria that SSGM considers during assignments. These criteria will be taken into consideration while assigning security guards to the posts.

2.2 Problem Definition

There are three critical problems in the current system: fairness, the time spent on the assignments and uncertainty. The main problem is that the assignment process is being handled manually. If unexpected issues occur, such as a security guard gets sick, the assignment schedule needs to be updated manually. Detailed information about the three problems is below.

2.2.1 Assigning Security Guards Manually and Fairness

SSGM follows a policy where assignments to the posts are being done by evaluating the security guards' performance measures. These performance measures are pre-defined by SSGM and considered common for every security guard. Assignments require a system in which a security guard assigned to a post should meet the requirements of that post. Currently, this procedure is being done manually, which means the current system may not provide fair assignments due to the complexity of the problem. Furthermore, assigning security guards to posts manually may create an unbalanced working schedule.

2.2.2 Time Spent on the Assignments

The manual assignment procedure that SSGM does takes a minimum of two days to finish for a two-week cycle. The time spent on the assignments is quite excessive since it is being handled by a single officer.

2.2.3 Unexpected Staff Related Issues

It is not uncommon for security guards to face problems like illness and fail to be present in the assigned post. In such cases, SSGM should replace the security guard with a suitable alternative as soon as possible. However, this process can get quite complicated and needs to be done carefully. The replaced security guards should meet the requirements of the post. If that security guard is already assigned to a post, that security guard should be replaced by another security guard as well. This problem might be the most important and complicated one due to being related to uncertainty.

2.3 Data Analysis

SSGM provided each security guard's experience, requirements of each security post, and several two-week assignment data. Considering the precise assignment data, most of the security guards worked six days a week. After adding the overtime hours, most of them worked close to 60 hours, which is the maximum working hour. The average working hour per week is 54.55 and the standard deviation is 7.9. The working hours of each security guard can be seen in more detail in Appendix B. The pattern is if a security guard is assigned to a post every two weeks, it is likely to be assigned to that post for 5-6 days. For example, if a post requires two security guards throughout the week, the same two or three security guards are assigned to that post. Different security guards are not assigned every day to the same post. We need to adjust the assignment

process to this approach while trying to optimize the process. Most of the security guards work in the same posts and similar shifts during the week. However, there are approximately five security guards who worked in various posts during the week. These security guards have enough capability to work in those posts and have at least five years of experience. We will be referring to them as "versatile security guards". They are security guards who are generally assigned to vehicled patrolling units. Most experienced security guards are assigned to posts requiring a relatively higher level of experience, such as the university's rectorate building. These security guards work all 6 working days of the week at the same post and the remaining one day is covered by the versatile security guards mentioned above. This case is explained to this extent because ongoing procedures like this one make the optimization process not very implementable to the real-world problem Chu and Beasley (1997). These will be discussed to further extent in the sections below.

2.4 Critical Objectives

- Ensuring that every assigned security guard meet the minimum requirements for the given security post.
- In case of an unexpected event or the unavailability of a security guard, creating a candidate list of security guards that can be reassigned within the constraints. Also minimizing the time required to complete this daily reassignment process.
- Defining mathematical fairness for the model.
- Assigning suitable security guards to the posts fairly.
- Delivering a decision support system with a user-friendly interface.

2.5 Critical Assumptions

Assumptions below are decided during the meetings with SSGM or the constraints that SSGM already applies for the current process. These assumptions are fundamentals of the mathematical model. Values are assigned to security guards' attributes such as years of experience, whether knowing a foreign language or not and education level. These values range from one to ten for experience, one to five for education and binary values for foreign language knowledge. Every security post has minimum requirements in terms of the attributes of security guards. These minimum requirements vary from post to post. This does not mean that different security posts have different levels of importance in terms of security. Each security guard hired by the SSGM meets the minimum requirements for security posts that need a minimum level. Gates to the university campus are the security posts that have minimum requirements. A mathematical assumption of "fairness" is defined numerically to evaluate whether the assignments are "fair" among security guards. We defined a variable U , which is the difference between the total working hours and requested working hours by security guards. It is

decided that to have a “fair” assignment schedule, it should be aimed to keep U smaller than the number which is selected. This part will be evaluated extensively in the Model part.

3 Model

The mathematical model formed within the aspect of operational research concepts and forms the basis of our solution Lim et al. (2012). The model is run on CPLEX and provides an optimal solution to benchmark and improve our heuristic outcome. We compare the resulting assignment schedule of heuristic and the output schedule of CPLEX. In the model, parameter T_k is being used as the coefficient matrix and indicates the priority of an individual guard to be assigned to the specified post. Summing the multiplications of parameter “ T_k ” and variable “ X ” with the indices i, j, k , gives the objective function and we are trying to maximize this value. Parameter “ T_k ” is the three-dimensional coefficient matrix, where its values are between 0 and 1. “ X ” on the other hand is a decision variable where it is a 4-dimensional Boolean matrix that decides whether guard “ i ” is assigned to post “ j ”, for shift “ k ”, on the day “ t ” or not, Diaz and Fernández (2001). The model can be seen in Appendix C.

4 Heuristic Approach

4.1 Algorithm of Heuristic

Since CPLEX can not be used by SSGM, because of expensiveness of the license and SSGM uses Excel currently a decision support system, which works with Excel VBA needed to be created. All security guards have 4 qualification which are experience, foreign language, communication skills and certificates, and all security posts have requirements, security guard to be assigned. After deduction of which security guard can be assign to which post, a matrix will be acquired Ang et al. (2019). Total assignment score matrix is defined as a binary matrix and it gives a security guard 1 if s/he can be assigned to the post; otherwise, it gives 0. It sorts the security guards according to their total assignment scores as highlighted in yellow in Appendix D. The aim is to determine how many posts a security guard can be assigned. For example, if a security guard satisfies only post 1’s requirements, s/he should be assigned there since if s/he is not assigned there, s/he will not be able to work in any other post. If the security guard can work at more than one post, the algorithm checks the posts’ assignment scores shown at the bottom of the table which can be seen in Appendix D highlighted in light green to compare which post has fewer scores. post has fewer scores. Therefore, it assigns security guards to the post which had fewer points than the others. For example, while assigning security guard 4, it checks the scores of Posts 1,4 and 5. Then algorithm assigns s/he to Post 4 which has the smallest score among these three. The aim is to assign all the security guards to all security posts while satisfying the posts’ requirements. The number of security guards required for the posts may differ. Besides, while checking these scores, the algorithm also

considers the number of security guards assigned to the posts. It provides the exact number of security guards required.

4.2 Qualities of Heuristic

The algorithm allows to user to assign a security guard to a specific security post. The user just needs to select the security guard and security posts to be restricted in Appendix E. If a security guard resign or be absence for a time the user can remove s/he from the list by using a window can be seen in Appendix F. If SSGM hires a new security guard, the user is just add s/he in the list by using window can be seen in Appendix F. Beside that security guards urgent absence is an another problem and if security guard can not to work in a day, the user select the guard and day that s/he can not come and see all of the available security guards that can be assigned to the post can be seen in Appendix G. We mentioned that there are 2 main shifts however, in the future a new shift can be added or existing shift need to be removed, the user can add or remove a shift by clicking a check box. It can be seen in Appendix H.

5 Deliverables

The project aims to deliver a decision support system, referred to as the "tool" for the upcoming parts, embedded in the Microsoft Excel program. The authorized person is currently spending about two days on the assignment process. The tool is aimed to improve the current system's assigning process while considering major constraints and critical objectives. It is aimed to accomplish the improvement with a heuristic. The deliverable as the final product is an Excel tool/spreadsheet to carry out the heuristic.

6 Benchmarking

To test the heuristic algorithm that we use, we decided to run our model in an optimization software called "CPLEX" where we can find optimal assignments. In the model that we use in CPLEX, the objective function is to maximize the multiplication of the assigned days for security guard, and the preference points (range between 0 and 7) of each security guard which can be accessed and changed at any time through a 3-dimensional matrix (security guard, security post, day). This approach allows us to make sure there are no unrelated security guard/post assignments that are not wanted. Also, using the matrix provides us more flexibility in case of a change in preferences. Before we test the heuristic approach though, we had to make sure both heuristic and CPLEX models were run in the same conditions. These conditions are mostly the constraints that we use, and its provided below;

1. Number of Security Guards
2. Number of Security Posts
3. Number of Guard/Post preferences

4. Number of Working hours allowance
5. Number of Working days allowance
6. Different Shift types used in each Security Post
7. Missing Shifts for each Security Post
8. Missing Days for each Security Post
9. Maximum allowed hours for difference between assigned hours and requested hours

After the above conditions are met for both the heuristic and the CPLEX models, we can compare the results and see the efficiency of the heuristic model. The detailed model of CPLEX can be seen in Appendix C. For the comparison of the CPLEX and the heuristic models, we will use the criteria below;

- The objective function value: After running both models, the CPLEX and Heuristic results are 3282 and 3129 respectively. These results tell us that in terms of efficiency, heuristic model is 96.22 percent close to the optimal solution.
- Total working hours for every security guard: Although the total working hours for some of the security guards differs for the CPLEX and heuristic, the difference in hours are considerably small which is something we were expecting. The working hours data for both CPLEX and heuristic can be seen in Appendix I.
- Security guard locations after assignments: The Security posts locations for security guards after running both models, turned out to be much more similar than we anticipated. There are only a handful of different assignment locations.

We have also tested different variations of the model where the number of Security Guards, number of Security Posts, and the Guard/Post preferences have been changed. This approach further supported the heuristic model and helped to make sure it is running correctly.

7 Validation

In order to tell if the heuristic model results are true, in other words whether the real data and heuristics results are under the same circumstances or not, we need to make sure every condition and constraints that are provided to us by SSGM are satisfied in heuristic results. We will check one by one whether above conditions are satisfied or not to validate the heuristic model.

- Minimum and maximum allowed working hours: The minimum and maximum allowed hours for both heuristic and real data are 45 and 72 hours respectively.

- Maximum number of working days in a week: The maximum working days is 6 for both heuristic and real data.
- Minimum and maximum allowed number of security guards who work at security posts.
- Maximum allowed consecutive working shifts: There are two different scenarios for this condition; a security guard can either work two 8-hours shift consecutively or he/she can work only one 12-hours shift before taking a shift off.

As a result, validation process is done and the all of the given conditions and constraints are satisfied by the heuristic model results that we provided.

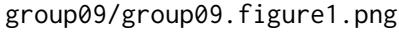
8 Verification and Conclusion

One of the main goals of this project is assigning security guards to appropriate security posts. We compared our heuristic result with data acquired from the company and the results are significantly close to each other. It can be seen in Appendix J and Appendix K. Besides that working hour of security guard another criterion for comparison. We have run the heuristic with different numbers of security guards and posts such as 20 guards and 9 posts, 50 guards and 25 posts and 114 guards and 50 posts which is the same as the current real-life situation, working hours of security guards comparison between heuristic result and real data can be seen in Appendix L. The average working hour per guard per two-week period is 54.55 hours with a standard deviation of 7.9 hours from the actual data between 16-30 November. The average working hour per guard per two-week period is 50.15 with a standard deviation of 8. As a result, the decision support system can run with many employees and posts and cover every requirement asked while considering every constraint of the model mentioned above. It decreases the time required for the assignment process from nearly two days to less than 5 minutes.

References

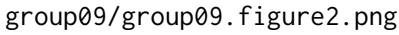
- S. Y. Ang, S. N. A. M. Razali, and S. L. Kek. Optimized preference of security staff scheduling using integer linear programming approach. *COMPUSOFT: An International Journal of Advanced Computer Technology*, 8(4), 2019.
- P. C. Chu and J. E. Beasley. A genetic algorithm for the generalised assignment problem. *Computers & Operations Research*, 24(1):17–23, 1997.
- J. A. Diaz and E. Fernández. A tabu search heuristic for the generalized assignment problem. *European Journal of Operational Research*, 132(1):22–38, 2001.
- G. J. Lim, A. Mobasher, and M. J. Côté. Multi-objective nurse scheduling models with patient workload and nurse preferences. *Management*, 2(5):149–160, 2012.

Appendix A Security Zones



group09/group09.figure1.png

Appendix B Working hours of the Security Guards Between November 16 – November 30



group09/group09.figure2.png

Appendix C Mathematical model

Indices

$i = 1, 2, \dots, 114$: *Security guards*

$j = 1, 2, \dots, 50$: *Security post*

$k = 1, 2, 3, 4, 5$: Shifts
 $t = 1, 2, \dots, 7$: Days in a week
 $s = 1, 2, 3, 4, 5$: Skills

Parameters

T_k : Coefficient matrix

Q_i : Historical data of working hours of security guard "i"

S_{is} : Skill points "s" for the security post "j"

S_{sj} : Requirement points of skill "s" for the security post "j"

$S_{ijs} : S_{is} - S_{sj}$

T_j : Minimum number of security guards that can work at post "j"

Z_j : Maximum number of security guards that can work at post "j"

R_i : Requested working hours by guard "i"

$A_{ij} : 1, \text{if the security guards work 8 hours shift}$

$B_{ij} : 1, \text{if the security guards work 12 hours shift}$

D : Difference between requested and assigned hours for every security guard ■

Decision Variables

$X_{ijkt} : 1, \text{if guard "i" is assigned to post "j", for shift "k", on day "t"}$

Y_i : Working hours of security guard "i"

$$\max \sum_{i \in I} \sum_{j \in J} \sum_{k \in K} T_k * X_{i,j,k,t}$$

subject to: $D \geq U$

$$Y_i - R_i \leq U \quad \forall i \in I$$

$$R_i - Y_i \leq U \quad \forall i \in I$$

$$Y_i = \sum_{j \in J} \sum_{k \in K} \sum_{t \in T} X_{i,j,k,t} * (8A_j + 12B_j) + Q_i \quad \forall i \in I$$

$$Y_i \leq 60 \quad \forall i \in I$$

$$Y_i \geq 40 \quad \forall i \in I$$

$$\sum_{k \in K} X_{i,j,k,t} \leq 2 \quad \forall i \in I, \forall j \in J$$

$$X_{i,j,1,t} + X_{i,j,4,t} \leq 1 \quad \forall i \in I, \forall j \in J$$

$$X_{i,j,1,t} + X_{i,j,5,t} \leq 1 \quad \forall i \in I, \forall j \in J$$

$$X_{i,j,2,t} + X_{i,j,4,t} \leq 1 \quad \forall i \in I, \forall j \in J$$

$$X_{i,j,2,t} + X_{i,j,5,t} \leq 1 \quad \forall i \in I, \forall j \in J$$

$$X_{i,j,3,t} + X_{i,j,4,t} \leq 1 \quad \forall i \in I, \forall j \in J$$

$$X_{i,j,3,t} + X_{i,j,5,t} \leq 1 \quad \forall i \in I, \forall j \in J$$

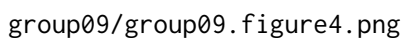
$$X_{i,j,4,t} + X_{i,j,5,t} \leq 1 \quad \forall i \in I, \forall j \in J$$

$$\begin{array}{ll}
\sum_{t \in 6} X_{i,j,k,t} \leq 6 & \forall i \in I, \forall j \in J, \forall k \in K \\
\sum_{i \in I} X_{i,j,k,t} \leq T_j & \forall j \in J, \forall k \in K, \forall t \in T \\
\sum_{i \in I} X_{i,j,k,t} \leq Z_j & \forall j \in J, \forall k \in K, \forall t \in T \\
X_{i,j,k,t} * (S_{i,j,s}) \geq 0 & \forall j \in J, \forall k \in K, \forall t \in T \\
X_{i,j,k,t} \in \{0, 1\} & \forall i \in I, \forall j \in J, \forall k \in K, \forall t \in T \\
A_{i,j} \in \{0, 1\} & \forall i \in I, \forall j \in J \\
B_{i,j} \in \{0, 1\} & \forall i \in I, \forall j \in J
\end{array}$$

Appendix D Total Assignment Score Matrix

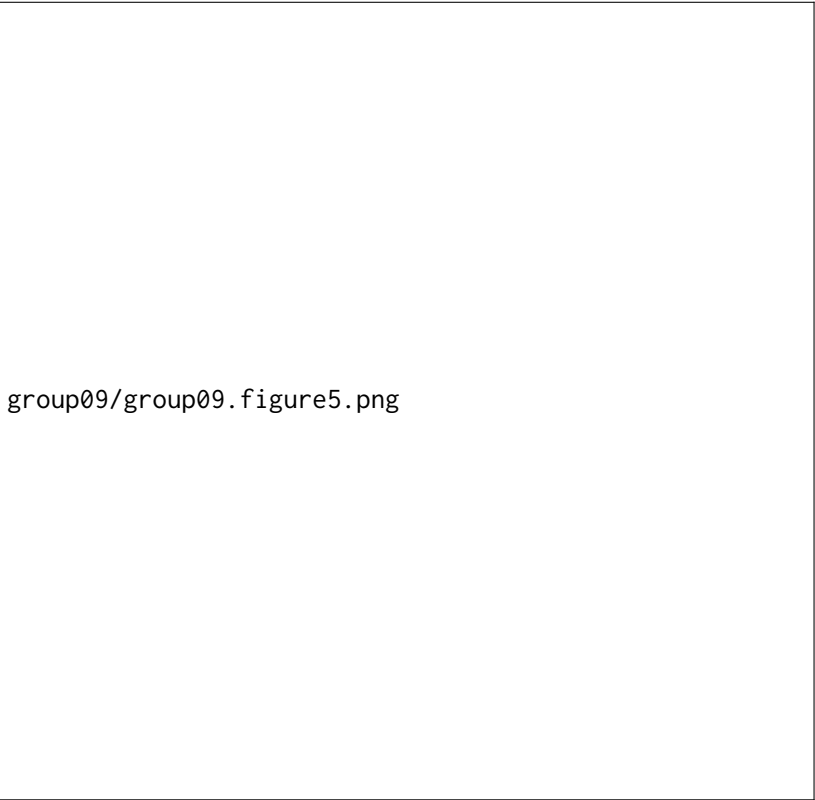
group09/group09.figure3.png

Appendix E Security Guard Selection Interface

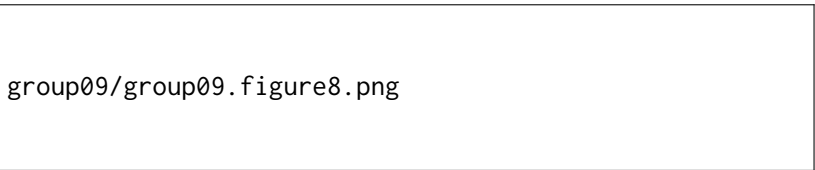
The image area is mostly blank, with a file path label 'group09/group09.figure4.png' located in the lower-left quadrant. This likely represents a screenshot of the Security Guard Selection Interface that failed to load or is a placeholder.

group09/group09.figure4.png

Appendix F The Interface of the Adding/ Removing Guard



Appendix G Example of Adding the New Shifts

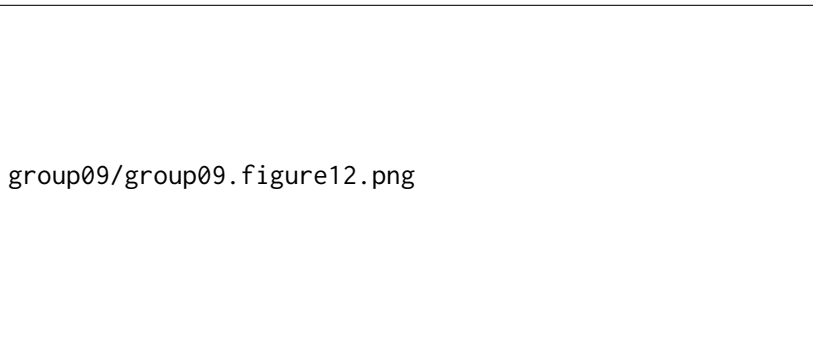


Appendix H The Interface of Selecting Guard Unable to Work



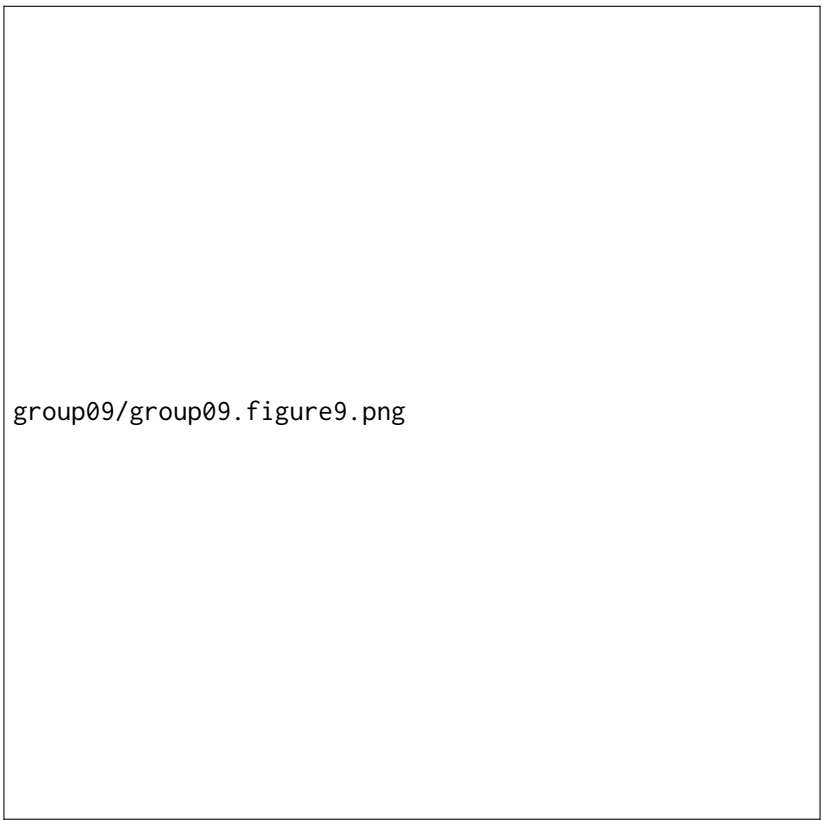
group09/group09.figure6.png

Appendix I Working Hours of Security Guards

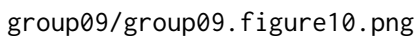


group09/group09.figure12.png

Appendix J Heuristic Algorithm's Assignment Schedule

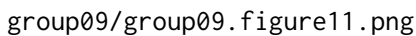


Appendix K SSGM Assignment Schedule



group09/group09.figure10.png

Appendix L Comparison of the Working Hours



group09/group09.figure11.png

Giriş Kalite Kontrol için Örneklem Stratejileri

ARÇELİK A.Ş.

group10/group_photo.png

Proje Ekibi

Eylül Ayhaner, Göksu Çağlayan, Doğukan Danışman, Belen Eyüpoğlu
Sercan Gür, Büşra Sülün, Ömer Faruk Şahin

Şirket Danışmanı

Hakan Haskılıç, Kalite Güvence
Müdürü Hazret Yıldız, Kalite Güvence
Uzmanı Volkan Kaya, SAP Uzmanı
Endüstri Mühendisliği Takımı

Akademik Danışman

Prof. Dr. Savaş Dayanık
Endüstri Mühendisliği Bölümü

ÖZET

Arçelik Bulaşık Makinesi Fabrikası'nda mevcut kalite kontrol sisteminin ve numune alma stratejilerinin verimliliği incelenip iyileştirme yapılabilecek alanlar belirlenmiştir. Bu proje ile kalitenin artırılması, üretim hatlarındaki aksaklıkların azaltılması, hatalı parça sayısı ve müşteri şikayetlerinin azaltılması için karar destek sistemi oluşturulması hedeflenmiştir. Her materyalin her karakteristik özelliğindeki en uygun örneklem büyüklüklerini bulabilen matematiksel model geliştirilmiştir ve karar destek sistemi oluşturulmuştur. Geliştirilen karar destek sistemi, kullanıcı dostu bir arayüz ile şirketin kullanımına sunulmuştur.

Anahtar Kelimeler: Kalite Kontrol, Örneklem Stratejileri, Örneklem Büyüklüğü

Sampling Strategies for Incoming Quality Control

1 Company Information

Arçelik, founded in 1955, is now one of the world's largest white goods companies. The first dishwasher factory in Turkey was established by Arçelik in 1959, and Arçelik Ankara Dishwasher Factory was established in Ankara in 1992, and the factory has a huge capacity to produce 10% of all dishwashers in the world. Arçelik also owns global brands Beko, Blomberg, Grundig, Flavel, Elektra Bregenz, and Altus. Today Arçelik, worldwide, Turkey, Romania, China, South Africa, Thailand, have a total of 18 plants in 9 different countries. Arçelik is a leader in its sector with a figure of over 50% in white goods, built-in, and air conditioners. Arçelik is the third-largest white goods company in Europe and the first in England, Poland, and France in terms of total sales (Arçelik, 2020).

2 Current System Analysis

13,000 dishwashers are produced per day in Arçelik. Therefore, quality control of components supplied from auxiliary industries is a critical step and the strategies of the quality control system are important. The current input quality control strategy is based on the decision that whether to accept and reject samples from each lot according to relevant control characteristics of planned groups for the specific component. Control characteristics imply processes such as packaging, appearance and dimensional checks, and functional test and planned groups specify components with respect to their types. Arçelik decides to accept and reject samples determined by the global tables in TS 2859-1 (TSI, 2012). Operators take samples based on sample sizes shown on the SAP screen. According to our analysis on data that we conducted, it is seen that sample sizes taken from batches have values of 2, 3, 5 and 8, within approximately 97% of samplings. Also, on the screen, the number of defectives is also displayed in the case of acceptance and rejections. After that, SAP waits for the accept/reject decision from the input quality control team. According to these decisions, some of rejected materials are reworked in the factory, accepted ones are directly transferred to the the production line.

3 Problem Definition

The main problem of Arçelik is production line stoppages and reworks that are caused by defective items which are not detected in incoming quality control. According to data analysis that we conducted, it is seen that the materials possibly causing line stoppages are not adequately controlled by the related sample sizes. It means the sampling procedures with regards to sample sizes are not enough and should be altered.

4 Proposed System

4.1 Solution Approach

In this project, our main aim is to reduce the total cost by reducing mainly the production line stoppages and rework. If we tighten input quality control procedures by taking more samples, Arçelik can detect defective items in a batch with higher probability but, if it is done, sampling cost of the factory increases. On the other hand, Arçelik takes smaller samples from incoming batches, which reduces the sampling cost, but defective items may not be detected and the production line stoppages may increase because of defective items. Eventually, it causes higher production line stoppage costs. There is a trade-off between sampling and production line stoppages cost. The aim is to minimize the total cost of the factory. Whether an item will cause the production line to stop or not depends on a conditional probability that changes according to the defective item rate in the batch, the number of samples taken, and the batch size. In our solution, the expected value of this probability is taken in order to calculate possible production line stoppages and their possible costs.

4.2 Mathematical Model

The goal of the project is to minimize the total cost spent on input quality control by maximizing the quality level of each item. Therefore, while minimizing the total cost, the model tries to optimize the sample sizes of each characteristic k of the item i . The model is given in Appendix E.

The total cost consists of inspection cost and expected total stoppages cost as it can be seen in the objective function of the model. The inspection cost of characteristic k of an item i is indicated by I_{ik} and stoppages cost of item i with characteristic k and sample size is indicated by C_{ik} .

In the first constraint, X_{ikj} is the binary variable, and each characteristic of each part must only equal one sample size value because the size of the sample taken for the same characteristic of a product is constant.

The second constraint is associated with the limited total man-hour. Since time needed to inspect characteristic k of an item of part i is shown by t_{ik} for part i with characteristic k , and the total time taken to inspect all the characteristics of all parts should be less than or equal to the total man-hour time available for the inspection.

Moreover, the equations are showing the calculations of coefficients in the objective function. First equation indicates the total stoppage cost of part i with characteristic k in the j^{th} sample size by changing the number of nonconforming items in the batch which is taken for each item i with batch size N_i . First of all, it should be considered the probability of acceptance of the batch having the

size of N_i with consisting m nonconforming items if j^{th} sample size is taken. It is assumed that if a defect found in the sample size, the entire lot is rejected, which means that the number of defective items can be accepted is zero. Batch acceptance probability has probability mass function of the hypergeometric distribution. Number of nonconforming items in the batch with the size of N_i has binomial distribution (Sampaio et al., 2016).

4.3 Flowchart

The major steps of the decision support system and its integration with the current system can be seen in the flowchart in Figure 1. In this flowchart, blue processes

group10/flowchart.PNG

Figure 1: Flowchart

denote the current system in Arçelik, green processes denote the decision support system step that is taken by users, and purple ones denote decision support system steps that are taken by itself. All these steps will complement each other and provide Arçelik with a more effective incoming quality control strategy.

4.4 Decision Support System and User Interface

The mathematical model was coded by using R. Arçelik provides us with 10-months input quality control data and batch sizes, sample sizes of each control characteristics of items, acceptance and rejection decisions can be seen on the data.

Firstly, the code takes this raw data and prepares automatically an input file that includes important parameters for the mathematical model such as defective item rate in a batch, control frequency of control characteristics of items, inspec-

tion time, inspection cost, etc. The input file is given in Appendix A. The model works on this input file.

After this step, users of the decision support system select this file from the computer and enter required data such as monthly worker wage, the number of workers in the factory, the number of shifts in the factory, and cost parameters by using a user interface. The user interface is given in Appendix C.

The model provides an opportunity to select the solver between two options: IBM Cplex and R lpSolver. After entering inputs, if the user clicks on the button, the model begins to run, and it gives the optimal result in 2-3 minutes. Some important KPIs such as total sampling cost, total worker cost, worker utilization can be seen on the screen.

Finally, the user can download the results file that includes the optimal sample sizes for each control characteristic of items. The result file is given in Appendix B. The user has a chance to compare the current sample sizes of each control characteristic of items and the optimal one.

Decision support system can be run with different scenarios. In the first scenario possible sample sizes are 2-3-5-8, in the second scenario possible sample sizes are 0-2-3-5-8. Zero sample size means that Arçelik can skip some items without sampling in incoming quality control. For both scenarios decision support system can be run with different number of workers and the optimal one with minimum total cost can be selected from these results. The result graphs of scenarios can be seen in Appendix C.

5 Verification and Validation

For verification of the model, the possible values of the parameters were tried in the model. Values of parameters such as N_i (lot size of item i), the number of control characteristic of items, t_{ik} (inspection time needed for characteristic k of item i), I_{ik} (inspection cost for characteristic k of item i), d_{ik} (defect rate for characteristic k of item i), and c_{ik} (stoppage cost caused by one item i with nonconforming characteristic k), and T (total labor hours that available for quality control), were written into the model according to insights that obtained from explanatory data analysis. It had been analyzed how the model works in extreme points and frequently encountered points by trying small and large values for each parameter within reasonable limits. For the validation of the model, three items with the different number of control characteristics and different batch sizes were selected from the data given by Arçelik and they were used as inputs of the model. Also, their d_{ik} 's were calculated by using the formula explained in 4.2 Mathematical model of this report. For the other parameters (t_{ik} and I_{ik}), reasonable values were used. The output of the model can be seen in Appendix D.

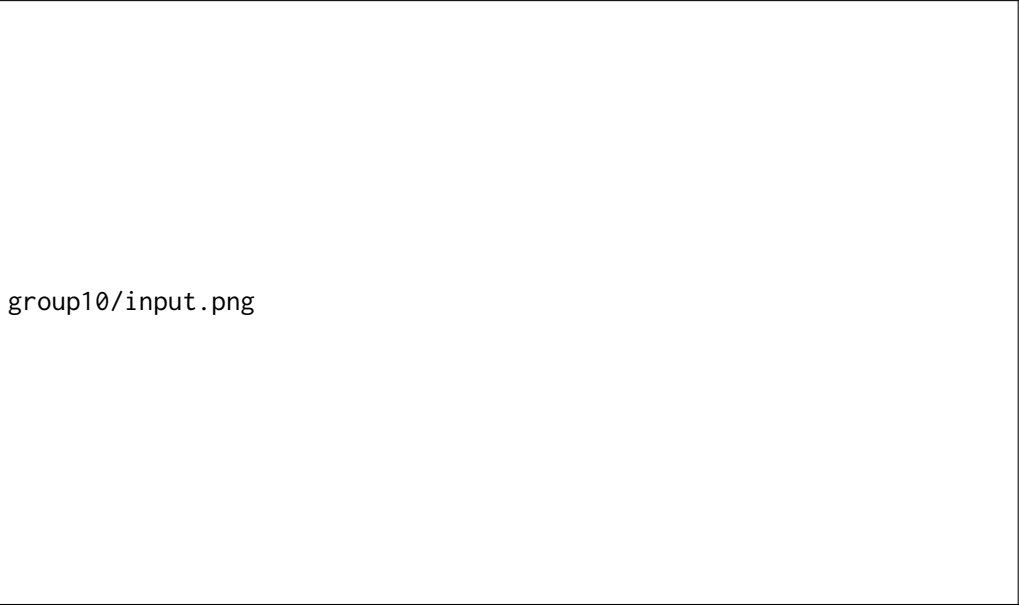
6 Benefits of the Project

The decision support system has been run with real data and differences between the current sample sizes and optimal ones are observed. To compare objective function results, we wrote a code that takes sample size data, inspection costs, stoppage and rework costs, control frequencies as an input, and calculates the objective function values. c_{ik} values are estimated by using real data given by Arçelik. According to importance of quality control characteristics and items, high, medium and low c_{ik} values are used. Sample sizes above 8 are taken as 8 and below 2 are taken as 2 in order for our model to run properly. The code has been run with Arçelik current sample sizes and optimal ones taken from decision support system. Since number of workers, worker salary and c_{ik} parameters in our model is determined as input from user (Arçelik), the cost of these are not included in the benchmark. The important improvement is on sampling costs either way. According to the results, 40.8% decrease can be observed in the total cost and it creates a big improvement in Arçelik Dishwasher plant.

References

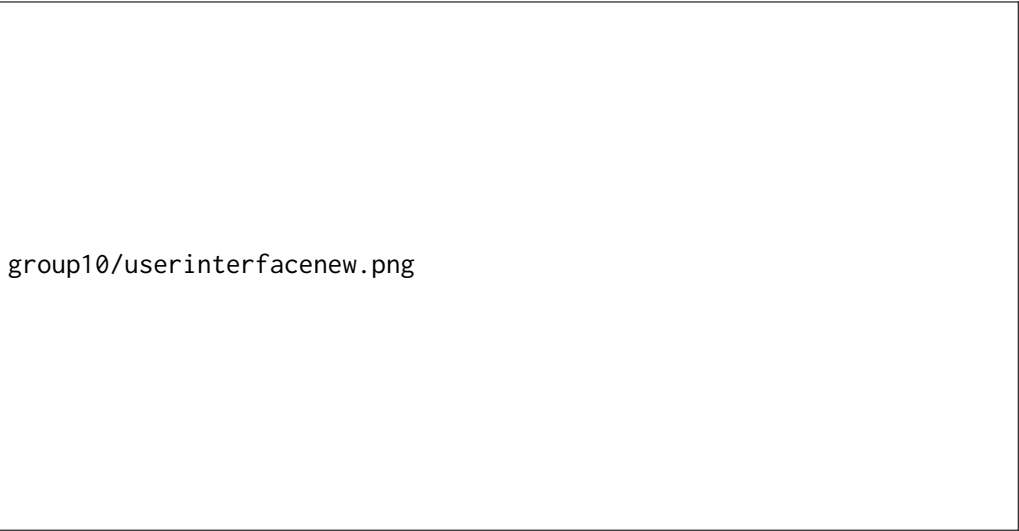
- Arçelik. Arcelik a.s. hakkinda. http://www.arcelikas.com/sayfa/10/ARCELIK_AS_HAKKINDA, 2020. Accessed: 2021-04-16.
- P. Sampaio, M. S. Carvalho, A. C. Fernandes, E. A. Cudney, R. Qin, and Z. Hamzic. Development of an optimization model to determine sampling levels. *International Journal of Quality & Reliability Management*, 2016.
- TSI. Ts 2859-1 sampling procedures for inspection by attributes. Publicly available on Internet, July 2012.

Appendix A Input File



group10/input.png

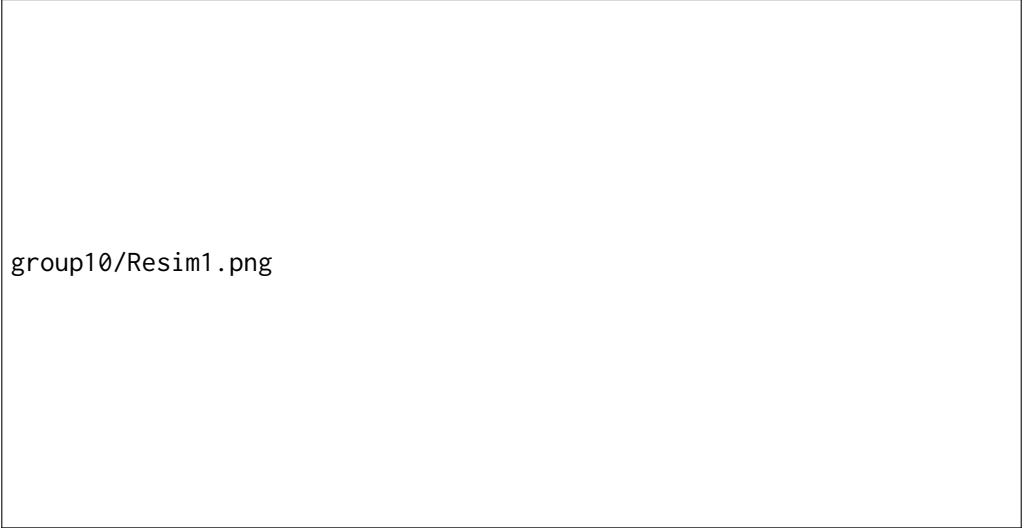
Appendix B User Interface



group10/userinterfacenew.png

Appendix C Objective Function Results of Different Scenarios

C.1 Scenario 1



group10/Resim1.png

In the first scenario, we allow sample sizes 2-3-5-8 and change the number of workers in the incoming quality control department. The figure shows the total cost (sampling cost and worker cost) versus the number of workers. In this scenario, the model gives an infeasible result until 8 workers, because a smaller number of workers than eight are not enough to inspect all characteristics of items even if the model select sample size two for each characteristic of items. After eight workers, the model gives the feasible results and the minimum total cost is achieved with 10 workers.

C.2 Scenario 2 (Skipping)

group10/Resim2.png

In the second scenario, we allow no sampling and the model gives feasible results even if there is only one quality control worker in Arçelik. The minimum total cost is observed with three workers.

Appendix D Results File

item	char	sj	E_{nc}	Cikj	OBJikj
1510150005	DONMA TESTİ	2	2,34	439,11	2679,68
1510150005	ETİKET KONTROLÜ	2	0,81	7,59	167,36
1510150005	FTC TESTİ	2	2,90	218,09	4642,97
1510150005	GRUPLAMA KONTROLÜ	2	1,26	23,76	514,72
1510150005	İKLİMLENDİRME TESTİ	8	7,05	1322,62	4057,85
1510150005	KOMPONENT KONTROLÜ	2	0,81	45,56	1004,18
1510150005	MAKİNE TESTİ	2	2,90	817,85	17411,1
1510150005	PCB KART BOY ÖLÇÜSÜ	2	0,61	5,76	129,02
1510150005	PCB KART EN ÖLÇÜSÜ	2	0,61	5,76	129,02
1510150005	REJENERASYON TESTİ	8	18,52	34,3985	350,39
1510150005	YÜZEY KONTROLÜ	2	0,81	15,18	334,72
1510150010	DONMA TESTİ	8	12,03	2255,84	20572,5
1510150010	ETİKET KONTROLÜ	2	0,12	1,17	347,88
1510150010	FTC TESTİ	2	0,79	59,49	13999,9
1510150010	GRUPLAMA KONTROLÜ	2	0,32	6,03	1518,81

Appendix E Mathematical model

Decision Variable:

$$x_{ikj} = \begin{cases} 1, & \text{if } j\text{th sample size is used for inspecting characteristic } k \text{ of item } i \\ 0, & \text{otherwise} \end{cases}$$

Parameters:

s_j = the number of sample of part in j th sample size, $j=1, \dots, 4$

d_{ik} = probability that an item i has unacceptable characteristic k

N_i = the number of items in the batch for part i

D_{ik} = the number of items in the batch with unacceptable characteristic k

c_{ik} = stoppage cost caused by one item i with nonconforming characteristic k

t_{ik} = time needed to inspect the characteristic k of one item of part i

f_{ik} = Control frequency of characteristic k of item i

T = total man hour available for quality control

I_{ik} = inspection cost of characteristic k of an item of part i

$$\min \sum_{i=1}^M \sum_{k=1}^{k_i} \sum_{j=1}^J (C_{ikj} + I_{ik} \cdot s_j) \cdot f_{ik} \cdot x_{ikj} \quad (1)$$

$$\text{subject to: } \sum_{j=1}^J x_{ikj} = 1 \quad \forall i = 1, \dots, M, \forall k = 1, \dots, k_i \quad (2)$$

$$\sum_{i=1}^M \sum_{k=1}^{k_i} \sum_{j=1}^J (t_{ik} \cdot x_{ikj} \cdot s_j) \leq T \quad (3)$$

$$x_{ikj} \in \{0, 1\} \quad \forall i = 1, \dots, M, \forall k = 1, \dots, k_i, \forall j = 1, \dots, J \quad (4)$$

$$C_{ikj} = c_{ik} \cdot \sum_{m=0}^D m \cdot P(D_{ik} = m) \cdot P(0|N_i, D_{ik} = m, s_j)$$

$$P(D_{ik} = m) = \binom{N_i}{m} \cdot (1 - d_{ik})^{N_i - m} \cdot d_{ik}^m$$

$$P(0|N_i, D_{ik} = m, s_j) = \binom{N_i - m}{s_j} \cdot \binom{N_i}{s_j}^{-1}$$

Geleneksel Kanal Müşterilerinin Segmentasyonu ve Hedef Atama Problemi

ETİ Gıda San. ve Tic. A.Ş.

group11/team11resim2.png

Proje Ekibi

Damla Yakut, Eda Aka, Ege Öndeş,
Irmak Özdemir, İpek Özgüven, Merve Yılmaz, Mustafa Özkök

Şirket Danışmanı

Uğur Akkalkan
Tedarik ve Talep Planlama Müdürü
Endüstri Mühendisliği Takımı

Akademik Danışman

Yrd. Doç. Dr. Selin Kocaman
Endüstri Mühendisliği Bölümü

ÖZET

Müşterilere hedef atarken müşterilerin satış dışı özelliklerine göre değerlendirilmemesi satış temsilcilerine uygun hedeflerin atanmamasına neden olmaktadır. Bu projenin çıktısı, İstanbul'da bakkal grubunda yer alan her bir ürün ve segment müşterisi için aylık satış verilerini, demografik ve lokasyon bazlı özelliklerini analiz edip yeni bir hedef atama sisteminin geliştirilmesidir. Müşteri segmentasyonunda yaygın kümeleme yöntemleri kullanılmıştır ve silüet skoruna göre en makul kümeleme sonuçları benimsenmiştir. Müşteriler için ürün bazında tahmin oluşturulmuş, müşterilerin segmentleri ve büyüme oranları değerlendirilerek ürün bazında hedef verilmiştir. Hedef atama sistemimiz için bir kullanıcı arayüzü oluşturulmuş, bu arayüz ve çıktıları ETİ satış departmanına sunulmuştur.

Anahtar Kelimeler: Müşteri segmentasyonu, hedef atama, satış tahminleme

Traditional Channel Customer Segmentation and Sales Target Assignment Problem

1 Company Description

ETİ was founded in 1961 by Firuz Kanatlı in Eskişehir as a sole proprietorship under the name of ETİ Biscuit Factory. It expanded its product range by an innovative approach and became one of the snack industry leaders. The company offers many different product groups, including biscuits and wafers, crackers, cakes and tarts, chocolates, breakfast products, frozen products, etc. ETİ has 300 products of various kinds and an annual production capacity of 250,000 tons. The production is performed on eight manufacturing facilities located in Eskişehir, Bozüyük, Konya, and Romania. ETİ has over 7.000 employees and net sales of over 5 billion TL. ETİ exports 30% of its products produced in Turkey, and it corresponds to sales of over 147,000 thousand dollars. The remaining 70% of the production is distributed to the whole country by more than 200,000 vendors. The vendors' supply is managed by approximately 150 distribution points that are served by 17 different warehouses. Turkey's Top 500 Industrial Enterprises shows that ETİ is the 33rd biggest Turkey company in 2019 (ISO, 2019).

2 System Analysis and Problem Definition

2.1 Analysis of the Current System

In the current system, there are three types of sales channels: distributors, national customers, and direct customers. The distributors are responsible for reaching customers in traditional channels/school cafeterias, modern channels, and special customers. There are around 150,000 traditional channel/school cafeterias customers. As traditional channels/school cafeterias have the highest number of sales, the project's focus is targeting customer groups in traditional channels. Traditional channels are groceries, kiosks, tobacco shops, delicatessens, and gas stations. To forecast the sales, currently, ETİ is using linear regression and the Holt-Winters method. By using these forecasts ETİ can decide on truck filling strategies for the upcoming period (i.e. month). At present, ETİ has a sales planning system that involves three main stages:

1. Collecting two different sales data, one for SKUs and the other for customers.
2. Assigning the same sales target to 2,100 sales representatives.
3. Analyzing the collected data and forecasting sales amounts for future months. ■

2.2 Problem Definition and Scope

In the current system, the sales director determines a target sales amount of - 120,000 TL worth SKUs in total- for a sales representative. From this practice, it can be inferred that there is no flexibility in terms of the target sales level assigned to each sales representative.

Moreover, the sales director wants to forecast the future sales amounts of each SKU and the cumulative sales amounts of each customer. However, ETİ has a complaint about having the same sales targets for sales representatives. The company is expecting a new system that allows ETİ to make a proper segmentation of customers in traditional channels for each product. The improved system is expected to enable the company to assign the sales targets in a more accurate way considering the customers' characteristics that are affecting the sales potential other than past sales information and the corresponding sales forecasts. In this way, the company can understand the achievable potentials in sales volume, assign the customers sensible sales targets and conduct the production operations relying on this system.

According to the information provided by ETİ, the objective of this project is to provide a better sales target utilizing a customer segmentation that runs dynamically and accurately. This algorithm aims to increase accuracy in the target assignments and therefore provide achievable yet challenging targets. After the analysis of the data we obtained, it was understood that the majority of the customers were in the grocery category. Therefore, further analysis was based only on the grocery customers. Additionally, Istanbul was chosen as the pilot zone because the sales data were accumulated in this city (approximately 80% of the sales).

Furthermore, the number of SKUs was 476 initially, in later stages they were grouped under 69 products. In the end, the graphical user interface is delivered after the segmentation, sales forecasting, and sales target assignment components.

3 Proposed System

The new sales assignment system has three components: segmentation, forecasting and targeting. Segmentation is necessary to gather the customers with similar characteristics in terms of population, development, volume of sales. The future sales predictions of each customer on the product basis are obtained by forecasting step. Then, considering both forecast results and the segments of the customers, the sales targets are assigned on product basis.

3.1 Segmentation

In the segmentation step, sales data of 2018 and 2019 is used. Additional to the sales data provided by the company, demographic and location-based features are included to the segmentation process. Our model is based on district level and we have a development index of districts, the population by age groups and sex, educational level by age groups and sex. The population and education information is gathered by contacting TURKSTAT.

For segmentation of the customers, K-Means and Agglomerative clustering methods are used due to their convenience, and the fact that they are already being

used for customer segmentation in the industry. The data is scaled and clustered using Python. In order to evaluate the goodness of fit of the clusters, we considered several methods such as Elbow Method evaluating within cluster sums of squares (WSS) , Silhouette Method considering average Silhouette Scores (Shahapure and Nicholas (2020)). WSS measures the squared average Euclidean distance of all the points within a cluster to the cluster centroid. The WSS for different numbers of clusters (k) are plotted and called “Elbow graph”. The first k which WSS becomes to diminish is the “elbow” of the graph. For our customer clusters with respect to products, the “elbow” might not be easily distinguishable, therefore, ambiguous (Kodinariya and Makwana (2013)). Therefore, we lastly used average silhouette scores to measure how acceptable our cluster fits are. Those scores are calculated as following:

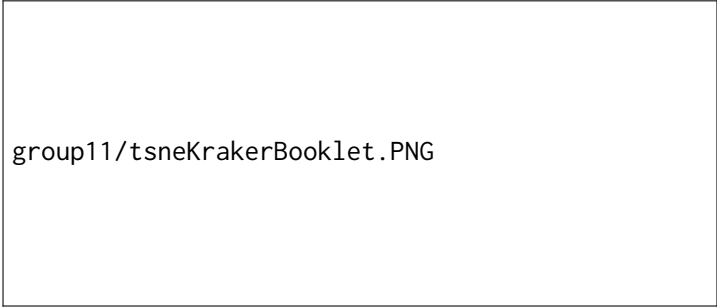
$$s = \frac{b - a}{\max(a, b)}$$

where

- a: The mean distance between a sample and all other points in the same class
which measures the closeness of points in the same cluster
 - b: The mean distance between a sample and all other points in the next nearest
cluster -which measures the distance of the points of different clusters
- Note that silhouette score is between -1 and 1, the interpretation of which is that the higher the better.

For Agglomerative Clustering, we tried the linkage parameters “complete”, “average” and “ward”. The best performance in terms of silhouette score obtained using complete-linkage where distance is measured between the farthest pair of observations in two clusters, therefore, we only kept complete-linkage while implementing user interface.

To represent our customers -and their clusters- while having high-dimensional features, we utilized T-Distributed Stochastic Neighbour Embedding (t-SNE). t-SNE is a tool for visualizing high-dimensional data by embedding it into low dimensions. In this way, customers and their clusters are fitted and demonstrated in two-dimensional space while having many features –more than a hundred in our project. The t-SNE representation of the clustering examples before and after the demographic and location-based information can be observed in Figure 1. Visually, the clusters are much more distinctive.



group11/tsneKrakerBooklet.PNG

Figure 1: Visual representation example of clusters for a specific product customers

Using Spectral Decomposition, the variability in the customer data is analyzed. 24-month sales data explain approximately 20% of the variability. District development index, male population age between 0-24 and elderly female population 60-90+ demonstrate the variability in the customer data most, therefore shaped the clusters. In light of these, one can say that the demographic information enhanced the segmentation process and helped to distinguish the customers in order to understand the potentials of the customers.

3.2 Forecasting

We used four different forecasting methods while predicting future sales with 2 years sales data of groceries in İstanbul. These are Autoregressive Integration Moving Average (ARIMA), naive forecasting, simple exponential smoothing, and Holt-Winters' methods.

In these models we have applied various train-test sets (i.e. 12-12, 18-6) to compare the performances. In the end we compared the performances of these 4 models with 3 error metrics. These are Mean Absolute Error (MAE), Mean Squared Error (MSE), Mean Absolute Percentage Error (MAPE). For the user interface, we used the first 18 months of the sales data for training and the remaining 6 months for testing the performance of each model in order to use as much as the data to train the models. While implementing the user interface, we experimented hyper-parameter combinations for the time-series forecast methods and adopted parameters giving least mean MAPE results for randomly selected products. After having analyzed the performances of each forecasting model, it turned out that the best model for each product differ considering the values of error metrics (MAPE). Still, the company has the opportunity to observe the best forecasting method for the products using the user interface to compare the outputs.

3.3 Sales Target Assignment

Sales targets of a customers is assigned by considering the following:

- The forecast growth rate of the customer for the desired target month
- Forecast average growth rate of other customers in the same cluster for the

desired target month

- Forecast maximum and minimum growth rates in the cluster of the customer

The average of the forecast growth of the customer and the forecast average growth of the cluster in which that customer is located, is given as the growth target. At the same time, the maximum and minimum growth information in the cluster helps to have information about the location of the targeted grocery store in that cluster.

To show that the current targeting system is plausible, one can observe the graphs in Appendix A. In the graphs, we can observe the best and worst groceries in terms of the expected sales growth rate of Crax in January 2020. For the worst customers, we can observe that predicted amounts of sales are less than targeted amounts. Here, by taking the average of forecast sales and the cluster mean, we are providing tolerance to customer sales. The worst customers are encouraged to sell more. Similarly, the best customers are targeted less than the forecast sales amount due to provide a tolerance level. In this way, targets become more achievable for all customers since we are not only considering the personal sales history but also the cluster status.

4 Validation

To validate the segmentation part, we can consider the Silhouette scores. This score is between -1 and 1, the interpretation of which is that the higher the better. The silhouette scores of the product clusters are between 0.56 and 0.91. The mean silhouette score of product types is found satisfactory as almost 50% of the customer clusters have a mean score of above 0.80. In Appendix B, the first ten products with the highest silhouette scores can be observed.

Validation of forecast can be done by examining the error rates of products. We chose a pilot district due to the heavy data load, which is Bağcılar since it has the largest number of customers in Istanbul. There are about 42,500 groceries in Bağcılar, whereas there are a total of 232,000 groceries in Istanbul. Bağcılar has approximately 20% of the grocery stores of Istanbul. We then randomly selected and forecast 12 out of 69 products. The error rates can be observed in the Appendix C.

Based on the results we received from the random sample, we can say that our error rate (MAPE) is between 10% and 40%. We have only 24-month sales data for each customer, showing the total monthly sales. Although we have 24-month sales data, the forecast results appear to be satisfactory. When this system runs with more data, it can give much better results. The target assignment system is approved by the company.

5 Implementation and Integration

We started the implementation in a large product groups basis. After the feedback of the company, we altered the implementation in the product basis discussed in Scope. We implemented a user interface in order to provide company a user-friendly environment to use this targeting system ,therefore, to observe detailed information such as the customers' estimated sales, the status of their clusters, and the assigned target.

The user interface inputs are as follows:

- CSV file containing sales data and demographic data of the past years
- Province
- Customer Type (for future use, we used only the grocery stores)
- Product Type
- Requested Months
- Forecasting Method
- Clustering Method

User interface outputs are as follows:

If a single month is selected,

- Forecast sales of grocery stores
- Average sales forecast for the cluster of grocery stores
- Sales target given to grocery stores
- Maximum and minimum growth percentage of the cluster of grocery stores

If multiple-month is selected,

- Forecast sales volumes for the selected months
- Sales targets for the selected months

The user interface can be observed in Appendix D.

6 Benefits to The Company and Suggestions

To understand the benefits to the company, it is useful to mention the current system and the proposed benefits of the new system. In the system which is currently in use, all sales representatives are given the same amount of sales targets, and traditional sales channels are not classified in the product basis. The improvements we have achieved with the new system are as follows:

- Sales representatives are targeted by considering the unique sales potential of the customers.

- Traditional sales channels are classified by product groups.
- Customers are clustered considering not only their sales but also their demographic characteristics.
- Both the various forecasting results and the average sales of their respective clusters were taken into account when giving sales targets to customers.
- There is an interface that allows accessing various information of all customers for the desired month.

As ETI assigns the same target to all of its customers, a new targeting practice is introduced to the company as mentioned in the Sales Target Assignment section.

We conducted our analysis on the district level as the neighbourhood information is not provided. This study may be extended by approaching the customers on a neighbourhood basis, in other words, all the demographic information may be gathered from TURKSTAT on the neighbourhood basis.

References

- T. Kodinariya and P. Makwana. Review on determining of cluster in k-means clustering. *International Journal of Advance Research in Computer Science and Management Studies*, 1:90–95, 01 2013.
- K. R. Shahapure and C. Nicholas. Cluster quality analysis using silhouette score. In *2020 IEEE 7th International Conference on Data Science and Advanced Analytics (DSAA)*, pages 747–748, 2020. doi: 10.1109/DSAA49011.2020.00096.

Appendix A Target Assignment



Figure 2: Target Assignment of Best Groceries in Growth for January 2020



Figure 3: Target Assignment of Worst Groceries in Growth for January 2020

While assigning the sales targets of the customers, both customer growth rate and cluster growth rate are considered. Because the target assignment relies on forecast values, therefore, deviate from the real future sales, segmentation and targeting method are to calibrate the forecast reflecting the real circumstances.

Appendix B Silhouette Scores

Product	K-Means Silhouette Score	Agglomerative Silhouette Score
TOPKEKTEKIL	0.91	0.86
SUSAMLICUBUK	0.91	0.88
FORM	0.90	0.90
ETICIKOLATA	0.88	0.86
SULTANIBUR.	0.86	0.81
CRAX	0.85	0.83
GONG	0.85	0.81
KARAMCIKOLATA	0.85	0.81
BURCAKKURABI	0.84	0.83
TARTINI	0.84	0.83

Table 1: The First Products with Best Mean Silhouette Scores

Appendix C Error Rates

		ARIMA	Holt-Winters'	Naive	Exp. Smoothing
CICIBEBE	MAE	22.24	22.88	20.58	20.25
	MSE	3007.13	5017.06	4391.01	2523.53
	MAPE	0.32	0.34	0.3	0.3
WANTED		ARIMA	Holt-Winters'	Naive	Exp. Smoothing
	MAE	16.13	13.4	10.81	14.48
	MSE	715.47	748.76	670.58	786.41
	MAPE	0.6	0.45	0.37	0.49
SUSAMLI Ç.		ARIMA	Holt-Winters'	Naive	Exp. Smoothing
	MAE	8.61	8.26	7.28	7.51
	MSE	236.9	251.31	250.63	204.06
	MAPE	0.23	0.23	0.19	0.2
FORM		ARIMA	Holt-Winters'	Naive	Exp. Smoothing
	MAE	15.54	13.43	10.8	12.7
	MSE	1590.58	1032.29	601.48	837.12
	MAPE	0.36	0.33	0.27	0.3
TUTKU		ARIMA	Holt-Winters'	Naive	Exp. Smoothing
	MAE	29.39	31.32	29.21	28.68
	MSE	2085.29	2849.25	2854.14	2081.57
	MAPE	0.34	0.41	0.36	0.34
BALIK		ARIMA	Holt-Winters'	Naive	Exp. Smoothing
	MAE	12.2	13.4	12.68	11.65
	MSE	319.24	470.3	447.04	319.18
	MAPE	0.41	0.46	0.41	0.39
GONG		ARIMA	Holt-Winters'	Naive	Exp. Smoothing
	MAE	27.25	31.06	28.71	26.24
	MSE	1678.39	2526.93	2197.14	1645.14
	MAPE	0.45	0.47	0.43	0.42

Table 2: The Error Results of a Subset of Products

As seen, different forecasting techniques may give the best results for distinct products. On the other hand, exponential smoothing usually performs better than the others in terms of provided error metrics shown in the table.

Appendix D User Interface

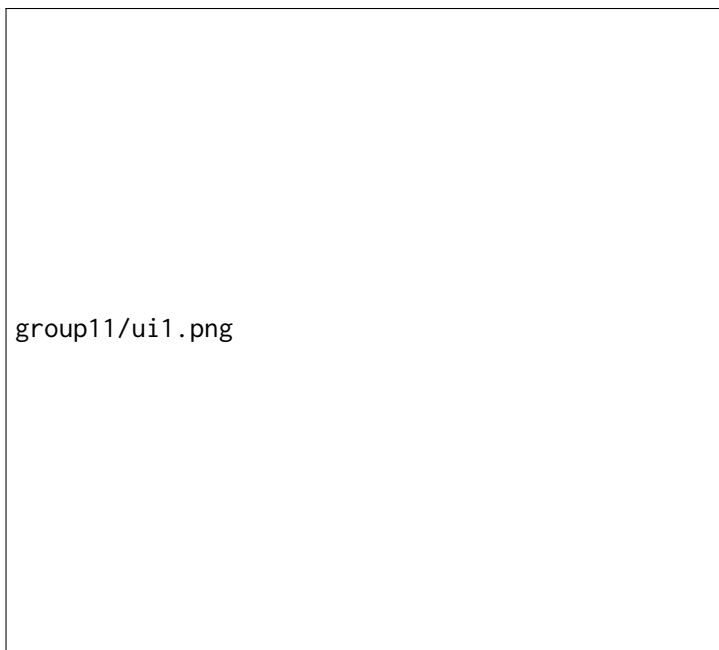


Figure 4: The Opening Screen of User Interface

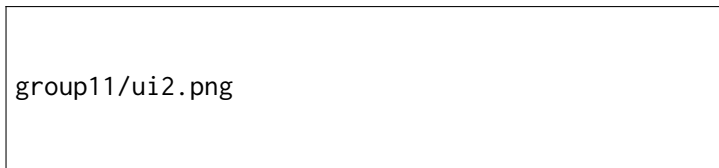


Figure 5: UI Showing Predicted and Targeted Sales and Predicted Growth Rate of Products for Multiple Months

Forklift Batarya Şarj İstasyonlarının Yerlerinin Eniyilenmesi

Arçelik A.Ş. Buzdolabı İşletmesi

group12/groupphoto.png

Proje Ekibi

Berke Atalay, Mürsel Başpınar, Ali Barış Bilen, Burcu Diana Erdoğan
Hasan Hüseyin Karadaş, Mutlu Yiğitcan Köse, Arda Yaşarlar

Şirket Danışmanı

Öyküm Özkök Büyüker
Endüstri Mühendisliği Takımı

Akademik Danışman

Prof. Dr. Oya Ekin Karaşan
Endüstri Mühendisliği Bölümü

ÖZET

Arçelik Buzdolabı İşletmesi Eskişehir fabrikasındaki forklift akü şarj istasyonlarının yerleşimi, forklift operasyonlarının ve buna bağlı olarak üretimin verimliliği üzerinde önemli bir rol oynamaktadır. Fabrika kurulurken hazırlanan yerleşim planı, zamanla değişen faktörler ile şirketi şarj istasyonlarının konumlarını sorgulamaya itmiştir. Bu projenin amacı, endüstri mühendisliği yöntemlerini kullanarak şarj istasyonlarının konumlarının eniyilenmesidir. Oluşturulan matematiksel model ile fabrikanın sağladığı veriler ışığında şarj istasyonlarının en iyi şekilde konumlandırılması ve projenin şirkete olan katkılarının hesaplanması hedeflenmiştir. Yapılan modelleme ile forkliftlerin şarj istasyonlarına gitmek için alması gereken mesafe %37,89 oranında azaltılmıştır. Bununla birlikte elde edilen faydanın, elektrik giderleri ve fabrikada operasyonel verimliliği arttırdığı göz önüne alınarak projenin uzun vadedeki faydası hesaplanmıştır.

Anahtar Kelimeler: tesis yerleşim problemi, forklift operasyonları, kapasiteli p-medyan problemi

Optimization of the Locations of Battery Charging Stations

1 Company Overview

Arçelik is a leading household appliances manufacturer of Turkey with more than 32,000 employees worldwide. It produces refrigerators, washing machines, dishwashers, dryers, ovens, cooking appliances, TVs, and others. The plant has a daily production capacity of 11,000 refrigerators and adopts make to order (MTO) production system. The facility currently has 7 production lines and most of them are fed by forklifts.

2 System Analysis

2.1 System Description

Currently, the company has 29 forklifts and 58 batteries and 28 charging lots available in 4 charging stations. A forklift is taken to the nearest station when its charge level drops to between 10 and 20 percent. When the forklift arrives at the charging station, the operator opens the battery cover, and gets into another forklift, which is assigned to that charging station and used for the battery replacement process. Empty battery is taken out via forklift and placed into the charging lot using a crane. After that, the full battery is placed into the immobile forklift and the station forklift is moved back to its parking area. Finally, the operator travels back to his workplace. The total time spent for this whole process is approximately 30 minutes.

2.2 Problem Definition and the Scope of the Project

In the current system, the main problem is long travel distances, which causes the decrease in the efficiency of forklift operations. The current travel time for forklifts' battery replacement process is measured as 25 minutes. It is more than desired and is believed that the underlying reason is the “non-optimal” placement of charging stations. Therefore, we decided to search for optimal locations of charging stations by using mathematical modelling and also a sensitivity analysis to be performed in case of various changes in the current system. Our main performance measures are decided to be improvement of distance and time spent during the travel. Finally, our deliverables are the final locations of battery charging stations and their corresponding capacities, forklift-station assignments, the model in Excel OpenSolver and an instruction manual explaining how to use and update the model for future use.

3 Solution Approach and Proposed Model

3.1 Modeling the Demand

In our problem, forklifts can be treated as customers to be served, while the charging stations are the facilities to meet their demand. Operations research scholars have developed several ways to handle such problems. In these problems, one

first needs to decide how she/he is going to formulate the demand. Upchurch and Kuby (2010) classify three different approaches to model the demand of customers; node-based demand models, arc-based demand models and path-based demand models. MirHassani and Ebrazi (2013) state that node-based models recognize the demand as aggregated in nodes where facilities can be placed anywhere within the network, including both nodes and the points on arcs. P-median, p-center and set covering problems are well-known representatives of this class. Arc-based models consider the demand (traffic flow) together with the length of the corresponding arc and try to determine the sufficient number of charging stations on an arc, such that the stations on the arc will cover the demands of its head and tail nodes. Thirdly, path-based models also take the traffic flow as the input for demand, however, in contrast with the arc-based models they locate facilities (stations) such that the flow capture will be maximized, where the flow is considered to be captured if it passes through a facility. Among the three approaches outlined above, we chose to adopt a node-based model approach for our problem. Since we do not have a reliable source of information regarding the traffic flow or the paths the forklifts follow, a node-based approach would be a more credible way to formulate the problem.

3.2 Problem Classification

There are many problem types in the literature that adopt a node-based approach, and we needed to pick the most appropriate one for our project. We were provided with the information that relocation of charging stations does not involve any capital investment and can be made at no cost, meaning that one will not incur facility opening cost. Company officers also confirmed that the amount of labor associated with the relocation of the charging stations is negligible. More important aspect of the problem is the connection cost, i.e. the distances between each demand node (customers) and facilities (charging stations). Thus, Simple Plant Location Problem (SPLP) as well as the Set Covering Problem (SCP) are clearly not the best candidates for our project. On the other hand, P-median problem matches perfectly with our expectations because as identified by Hakimi (1964), p-median problem deals with weighted sum of total distance travelled. We believe that optimal locations for the charging stations can be achieved by minimizing the total weighted distance between each forklift and its assigned station. Another reason for choosing the p-median problem is, as Upchurch and Kuby (2010) state, p-median models have quite simple data requirements. A distance matrix consisting of nodes and potential facility locations will be enough for our model because we will not be considering the cost of building a station. Besides the distance matrix, we also included a parameter for each forklift called "frequency" describing the number of times a forklift should replace its battery within a shift. The exact methodology will be explained in later sections. Moreover, we included a capacity constraint for the charging stations, which makes our problem a "Capacitated P-median Problem (CPMP)". CPMPs have not been well known as much as the generic p-median problem in the literature, however,

except a minor change in the constraint set, they are very similar. Ghoseiri and Ghannadpour (2007) introduce CPMP as a remarkable variation of the p-median problem where capacitated medians are placed such that the total weighted distance is minimized while meeting the capacity constraints. We had also sought some heuristic approaches and found that one meta-heuristic is widely used for CPMP which is Genetic Algorithm. Moghadam et al. (2013), and Ghoseiri and Ghannadpour (2007) solve CPMP using genetic algorithms on test problems and real-world cases and they both report good quality solutions. However, solving our model with genetic algorithm was the option we would consider only if the model did not scale well for realistic dimensions of the problem. The integer linear programming model that we have developed for the problem adopting CPMP could achieve an optimal solution within a reasonable amount of time.

3.3 Approach

By solving capacitated p-median problem, we aimed to find the optimal places of battery charging stations. Possible candidate locations of charging stations were determined, and each forklift is assumed to have a fixed location in the layout. In determining these fixed locations, we considered the designated work places of some forklifts, as they do not leave their work places frequently. We call them “fixed” forklifts. On the other hand, for the forklifts that carry loads through the entire factory floor, we used the following rule for determining the locations of each unfixed forklift: we located those forklifts considering the areas which the forklifts are observed the most frequently. The company currently has 4 battery charging stations and is not willing to make an investment for more. After reaching out the optimal locations for 4 charging stations, we considered the solutions for less than 4 charging station(s) to be able to understand the efficiency of the system and also evaluate the increased efficiency for more than 4 stations after reaching out optimal locations. In this way, we aimed to find the required number of charging stations and whether any investment for additional stations is worth it.

3.4 Model Development

Inputs:

- Forklifts’ designated work areas throughout the company using the factory layout.
- A distance matrix that includes the required locations for the forklifts in terms of x and y coordinates of the cartesian coordinate system. Distances are rectilinear since forklifts travel in a rectilinear fashion across the factory floor.
- Frequency details of the forklifts as a parameter. This parameter shows us how many times a forklift goes to a battery charging station in one shift. For instance, frequency will be 0.5 where a forklift goes to charging only once per two shifts.

- Other parameters (available data): 29 forklifts, 28 charging lots and 4 battery charging stations are currently available in the facility.

Outputs:

- Optimal locations of the battery charging stations and their capacities.

Table 1: Indices, Parameters and Decision Variables of the Mathematical Model

<i>Indices</i>	
$j \in S, i \in T$	
<i>Parameters</i>	
S	$\{1,2,3,...,175\}$ A set consisting of potential facility locations
T	$\{1,2,3,...,29\}$ A set consisting of demand locations
C	The total capacity (28)
d_{ij}	Distance between forklift i and station j
f_i	Frequency of demand of the forklift i in one shift
M	Sufficiently large number
<i>Decision Variables</i>	
x_{ij}	$\begin{cases} 1 & \text{if forklift } i \text{ assigned to site } j \\ 0 & \text{otherwise} \end{cases}$
y_i	$\begin{cases} 1 & \text{if site } j \text{ host a facility} \\ 0 & \text{otherwise} \end{cases}$
c_j	capacity assigned to station at site j

$$\min \sum_{j \in S} \sum_{i \in T} x_{ij} d_{ij} f_i \quad (1)$$

$$\text{subject to: } \sum_{j \in S} x_{ij} = 1 \quad \forall i \in T \quad (2)$$

$$\sum_{j \in S} y_j = 4 \quad (3)$$

$$\sum_{i \in T} x_{ij} f_j \leq c_j \quad \forall j \in S \quad (4)$$

$$\sum_{j \in S} c_j = C \quad (5)$$

$$c_j \leq M y_j \quad \forall j \in S \quad (6)$$

$$x_{ij} \leq y_j \quad \forall i \in T, \forall j \in S \quad (7)$$

$$x_{ij} \in \{0, 1\} \quad \forall i \in T, \forall j \in S \quad (8)$$

$$y_j \in \{0, 1\} \quad \forall j \in S \quad (9)$$

$$c_j \geq 0 \quad \forall j \in S \quad (10)$$

4 Model Validation

We have finalized our choice of model after verification and passed the validation phase in the middle of January. To validate the model constructed, we collected real data from the company. In our model, we have two main inputs; distance matrix and frequency. In order to construct our distance matrix so to convert the raw data to proper inputs of our data, we followed some steps. First, we have put the scaled layout on a grid (coordinate system) and identified the potential facility locations, designated locations of the forklifts and doors. (Doors are required for forklifts to get out of the factory to reach a charging station.) The layout on the grid and the figure showing the candidate locations are available in Appendix A and Appendix B, respectively. Pink points indicate potential locations distributed 8.33 meters apart from each other in a discrete manner. While reducing the continuous space to discrete candidate locations, we placed each of the candidate locations 8.33 meters apart from each other just for the sake of ease of calculation. Then, having outlined 175 potential locations to place charging stations, 29 forklifts and 18 doors, to find their (x,y) coordinates on the coordinate plane, we have utilized a tool called “GetData Graph Digitizer 2.26”. GetData Graph Digitizer is a publicly available tool and enables us to find exact coordinates (x,y) of any point we will work with. In our analysis, we settled with one unit length as the tolerance. This means that locations can vary at most one unit length from their real locations. We used a 82 x 82 grid (unit lengths) and it corresponds to 683.06 meter (one edge of the grid = $82 \times 8.33 = 683.06$). Thus, our analysis admits to $(8.33/683.06) = 1.22\%$ tolerance level. Thirdly, after obtaining all (x,y) pairs, we have constructed two distance matrices, one is showing rectilinear distances from each forklift to each door, the other is showing rectilinear distances from each door to each potential location. As one can see, the first matrix is a 29 x 18 matrix, and the latter is a 175 x 18 matrix. Finally, to construct our final distance matrix, which shows rectilinear distances from each forklift to each potential location. However, with this way, we allow forklifts to pass over factory walls, or any other place that is unavailable for forklift travel. We had used basic rectilinear distance calculation formula for this purpose but this method does not prevent this issue as its only duty is to find rectilinear distances no matter what. Therefore, our approach fails to find the correct distance value when a forklift needs to pass from two or more doors to reach that specific potential location. To fix this problem, we have identified each forklift/potential location pair such that our formula cannot correctly find the distance and we updated the final distance matrix manually. Excel screenshots which can clarify the subject on reader’s mind are available in Appendix C. After completing the distance matrix, we then decided for the second main input, frequencies. A couple of forklifts carry much heavier loads than the rest. We have set frequencies in accordance with these conditions, which are all confirmed by our industrial advisor. After data processing to model, we checked possible problems that some stations might be overloaded and some others might be underloaded, in terms of charging devices assigned to them. If it gives an unbalanced solution meaning that the difference between loads of the charging stations is high, we were going to introduce new constraints to prevent this from happening. We have discussed with our industrial advisor and concluded that all locations and assignments are sound and scale well for the realistic dimensions of the problem. Finally, to validate our model, we deleted all the candidate locations and forced model to place all four stations to their current locations in the factory. This model gave us a theoretical benchmark that we can compare with our optimal solution. The reason why we did that is, that it is practically impossible to obtain data regarding the total distance travelled in real life. We discussed with the company that the model’s result is well aligned with the current

setting and concluded that our model is valid.

Table 2: Station Placements

Potential Location ID	Coordinates	
	X	Y
25	25	35
40	9	7
116	30	61
172	55	35

Table 3: Forklift-station assignments

Forklifts	Assigned to the station (Pot. Loc. ID)
8,10,13,14,15,23,26,27	25
1,2,3,4,5	40
16,21,22,24,25,28	116
6,7,9,11,12,17,18,19,20,29	172

Table 4: Station-capacity assignments

Potential Location ID	Capacity assigned (#of charging lots(devices))
25	8
40	5
116	6
172	9

5 Project's Contribution to the Company

According to the information that is shared by the company, energy consumption for the forklifts is considered to be on average 6.6 kWh per unit. Also, at Arçelik Eskişehir Refrigerator Plant, the average speed of the forklifts is determined to be 8 km/h. Combining these parameters related to energy consumption and average speed of the forklifts, we can calculate that each forklift consumes 0.825 kWh of energy for traveling 1 km ($0.825 \text{ kWh} = 6.6 \text{ (kWh/h)} / 8 \text{ (km/h)}$). Therefore, we can conclude that energy consumption for each forklift can be stated as 0.825 kWh/km per unit. Additionally, the company shared with us that the cost of the energy they are using is 0.57 TL / kWh. After multiplying this price parameter with the energy consumption value that is derived for each forklift per 1 km, it can be stated that the electricity cost of each forklift is 0.47 TL for traveling 1 km ($0.47 \text{ TL/km} = 0.825 \text{ kWh/km} * 0.57 \text{ TL/kWh}$). The objective function of our mathematical model is defined to minimize the total distance traveled by the forklifts. As a result of our project, the improvements that have been made on the performance measures can be explained in the following manner:

We solved our mathematical model and got an objective value of 458.49-unit distance. In the grid of the factory layout that is used for this project, one unit distance corresponds to 8.33 meters. Then, we benchmark our model's result with another model result that has the current setting of the company by updating the locations of the battery charging stations. This theoretical objective value that we obtained in this benchmark model represents the total distance traveled with the current setting. The model gives us an objective value of 738.16. Compared to our proposed setting, this value is 1.61 times larger than 458.49. It can be concluded that the total improvement of our project in terms of total distance traveled by the forklifts is $738.16 - 458.49 = 279.67$ distance units (or $279.67 \times 8.33 = 2329.65$ meters). Therefore, the percentage improvement can be stated as 37.89% less travelled distance [$37.89\% = (738.16 - 458.49)/738.16$]. Considering that our model takes one shift (8 hours) as the model period but not three shifts (24 hours); this means that in one shift 279.67 less distance units will be covered. Moreover, our model takes only the distances that forklifts take to go to their assigned charging stations. We also have to account for the distances that forklifts should travel back to their workplaces. Thus, the total savings in terms of distance traveled can be calculated in the following manner: $279.67 \times 8.33 \times 2 = 4659.3$ meters in one shift. Considering three shifts per day and assuming 340 workdays in a year, **the proposed improvement in terms of distance** can be stated as $3 \times 340 \times 4659.3 = 4,752,486$ meters per year = **4752.49 km/year**. Finally, by utilizing all the above-mentioned information, **the financial improvement of this project** for a year can be calculated as $0.57 \text{ (TL/kWh)} \times 0.825 \text{ (kWh/km)} \times 4752.49 \text{ (km/year)} = \mathbf{2234.86 \text{ TL}}$

An improvement of 4752.49 km less traveled distance for the forklifts has been provided to the company and this corresponds to a total annual saving of 2234.86 TL. **Discounted savings for the next 10-year period** can be found as **9696.91 TL** using a Net Present Value Calculator (NPV (2234.86, $i = 0.19$, $n = 10$)) if 19 percent interest is taken into account.

Other than the improvements in performance measures, another benefit of our project to the company can be stated from opportunity cost point of view. In order to find it, **annual savings in terms of hours** can be calculated in the following manner: $4752.49 \text{ (km)} / 8 \text{ (km/h)} = \mathbf{594.06 \text{ hours}}$. The company currently has a plan of expanding Factory 6 in the near future, this process will require some new construction as well as the removal of unnecessary structures around the area. 594.06 hours of saved time for the forklifts mean that they can be utilized not only in the production system but also for different types of operations such as carrying construction materials to the site or removing scrap materials from the facility. In addition to that, the company will take into consideration that there has been an increase in the potential operational time of the current forklifts in the facility when there might be a situation for purchasing new forklifts. If scheduled properly, the new demand could be satisfied without even purchasing any extra forklifts or purchasing less forklifts than the planned number in the first place.

6 Implementation

In the implementation stage, during the meetings with the company officers, we are informed that some points are not suitable to place a station on. For example, as charging devices emit some amount of toxic gases during the charging, they have to be placed far from the places where humans are present, such as rest areas or rest rooms. After discussing with our industrial advisor, we eliminated such candidate locations and final-

ized the locations of the charging stations. After an updated settlement plan with new charging stations is provided, we analyzed the electricity infrastructure that should be implemented to the new locations. We informed the company that they should expand their electricity grid including the new locations that will be placed, in order to supply required electricity to the stations. Since charging stations must be located outside the factory, stations are vulnerable for external factors such as rain, snow, humidity or wind. Therefore, a protective structure should be constructed as a shell, in order to protect the charging devices, stations' forklift and the electricity supply from the external factors. We know that charging devices should be placed in a mold, in order to prevent any contact with water or any other liquid. In the current system, metal and canvas roofs are used in order to provide protection, so similar structure should be installed before the batteries are carried to the stations. Another requirement of a charging station is a parking area for the stations' forklift, which is used at battery replacement processes. After the electricity is provided and the protective structure is constructed then the crane and the battery charging devices should be placed according to the output of the model. Our model has already allocated 28 charging devices according to the number of forklifts that are assigned to a particular charging station. After the charging devices are placed into the stations, batteries will be carried to the allocated charging stations. With the batteries placed into the new stations, implementation process is completed. We have also provided an instruction manual to the company, where the model is explained step by step. This manual will be an useful reference when the model gives a solution and implementation is required . Recently, we are informed that due to the expansion operations held at the factory, Arçelik is not able to implement our project. However, by the time they complete the expansion process and finalized the factory layout, they will be able to run the model with new parameters and possible locations. The implementation process will be valid at that time and should be carried out as explained above. The instruction manual that we have prepared clearly explains how to use the developed model on Excel OpenSolver, which is an open source solver that the company may freely use. All the parameters and candidate locations can be updated as suggested in the manual, and solving the model will yield optimal locations for the charging stations. The interface we developed for using this model can be seen in Appendix 4. The users will use output screen to run the model and acquire the solutions as seen in Appendix F. The output screen shows the edges of the plant that will host charging stations, but the user has to check the exact locations by looking at the potential locations ((x,y) pairs) along with the grid.

7 Conclusion

As a result of our project, we concluded that the current locations for the battery charging stations are not optimal and our model helped us find the optimal locations with saving the total distance taken by the forklifts while changing their batteries. When our results are compared with the current application results which is in Appendix 2, it can be seen that 37.89% reduction in total taken distance is provided. Therefore, the energy spent for travelling will be decreased by this rate. Moreover, we expect operational time of these forklifts to be increased as well because their time spent while going to change their batteries will decrease and they will be able to spend more time in working their operational areas and they are expected to be more productive. If the company plans to buy new forklifts in the future they can decide the required amount of forklifts more appropriately using our model result and they will be saving from unnecessary costs. Company is very pleased with our results but as they decided to expand the factory

in near future, they will use our model with some new parameters and decide for the optimal locations upon the completion of the expansion process.

References

- K. Ghoseiri and S. F. Ghannadpour. Solving capacitated p -median problem using genetic algorithm. In *2007 IEEE International Conference on Industrial Engineering and Engineering Management*, pages 885–889, 2007. doi: 10.1109/IEEM.2007.4419318.
- S. L. Hakimi. Optimum Locations of Switching Centers and the Absolute Centers and Medians of a Graph. *Operations Research*, 12(3):450–459, June 1964. ISSN 0030-364X, 1526-5463. doi: 10.1287/opre.12.3.450. URL <http://pubsonline.informs.org/doi/abs/10.1287/opre.12.3.450>.
- S. A. MirHassani and R. Ebrazi. A Flexible Reformulation of the Refueling Station Location Problem. *Transportation Science*, 47(4):617–628, Nov. 2013. ISSN 0041-1655, 1526-5447. doi: 10.1287/trsc.1120.0430. URL <http://pubsonline.informs.org/doi/abs/10.1287/trsc.1120.0430>.
- A. M. Moghadam, H. Piroozfard, A. B. Ma'aram, and S. A. Mirzapour. Solving a Capacitated p -Median Location Allocation Problem Using Genetic Algorithm: A Case Study. *Advanced Materials Research*, 845:569–573, Dec. 2013. ISSN 1662-8985. doi: 10.4028/www.scientific.net/AMR.845.569. URL <https://www.scientific.net/AMR.845.569>.
- C. Upchurch and M. Kuby. Comparing the p -median and flow-refueling models for locating alternative-fuel stations. *Journal of Transport Geography*, 18(6):750–758, Nov. 2010. ISSN 09666923. doi: 10.1016/j.jtrangeo.2010.06.015. URL <https://linkinghub.elsevier.com/retrieve/pii/S0966692310000967>.

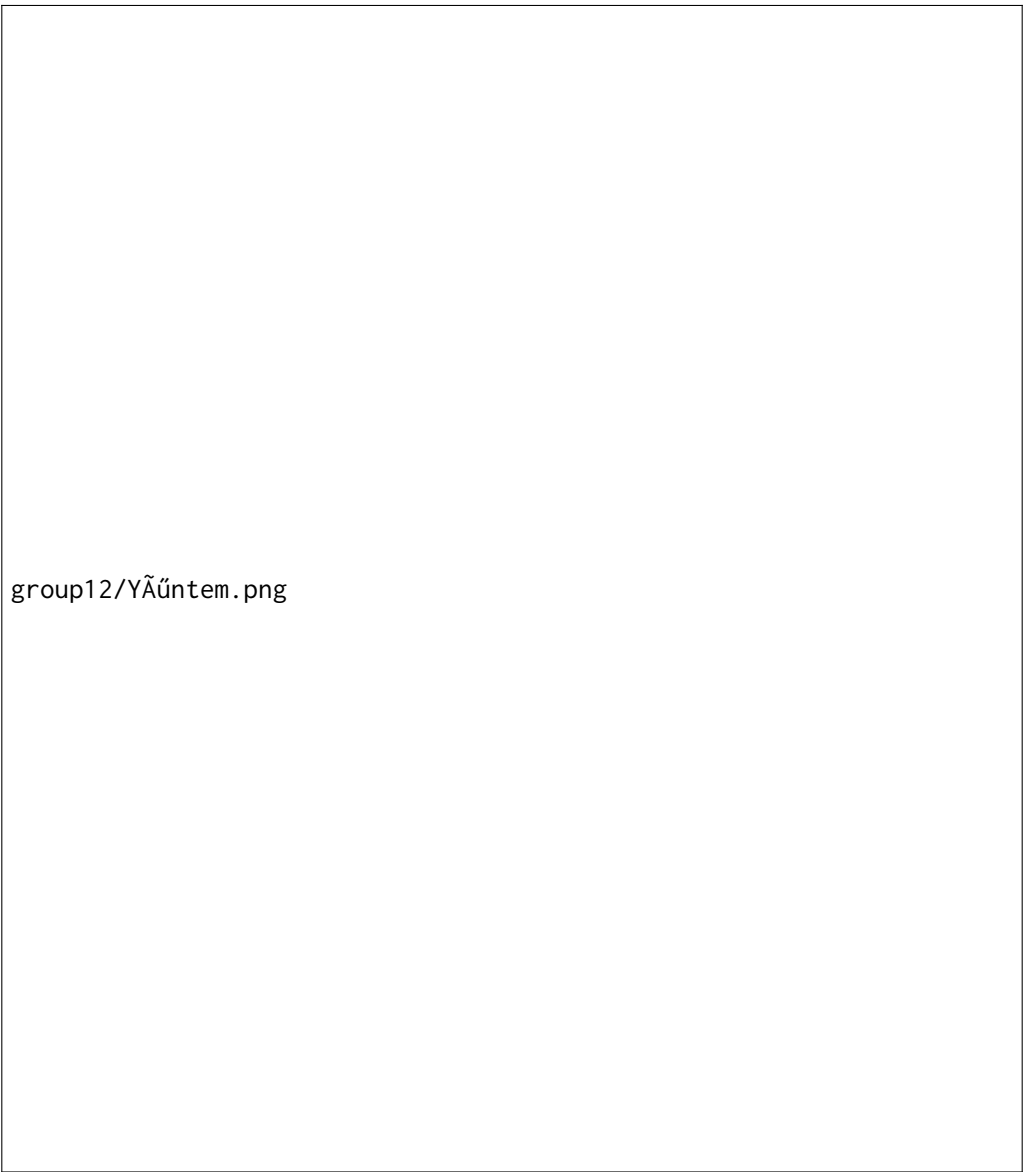
Appendix A Layout on the Grid



Appendix B Our Illustration Showing Potential Locations



Appendix C Excel Explanation

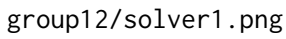


group12/YÃ¼ntem.png

Appendix D Currently Available Results

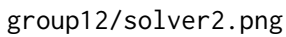
group12/Factory.png

Appendix E Excel Interface

The image shows the Excel Solver interface. It includes the Solver Parameters dialog box, the Solver Options tab, and the Solver Load/Save dialog box. The Solver Parameters dialog box is set to solve for the maximum of the objective function, with the variable cell range set to \$B\$1:\$B\$10. The Solver Options tab shows the Solver Engine set to GRG Nonlinear engine, Solver Options: Make Unconstrained Variables Non-Negative, and the Select a Solving Method checkbox checked. The Solver Load/Save dialog box shows the Solver Parameters file saved in the current workbook.

group12/solver1.png

Appendix F Excel Results

The image shows the Excel Solver Results dialog box. It displays the Solver Parameters, the Solver Options, and the Solver Results. The Solver Parameters dialog box is set to solve for the maximum of the objective function, with the variable cell range set to \$B\$1:\$B\$10. The Solver Options tab shows the Solver Engine set to GRG Nonlinear engine, Solver Options: Make Unconstrained Variables Non-Negative, and the Select a Solving Method checkbox checked. The Solver Results dialog box shows the Solver found a solution, and the Solver Results tab displays the Solver Parameters, the Solver Options, and the Solver Results.

group12/solver2.png

**Proje Bazlı Üretim Ortamında Raf Ömürlü
Malzemelerin Stok Planlaması İçin Karar Destek
Sistemi Tasarımı
Roketsan A.Ş.**

group13/group_photo.png

Proje Ekibi

İbrahim Emirhan Aydın, Gözde Çakmak, Defne Ekşioğlu,
Berk Meşe, Simay Pala, Elif Sadıkoğlu, Alper Yenigün

Şirket Danışmanı

Emre Aytac, Müdür
Saime Ceren Şar, Kıd. Uzm. Müh.
Duygu Soylu, Uzm. Müh.
Cansu Korkmaz, Uzm. Müh.
Cevat Enes Karaahmetoğlu, Müh.
Endüstri Mühendisliği Takımı

Akademik Danışman

Prof. Dr. Nesim K. Erkip
Endüstri Mühendisliği Bölümü

ÖZET

Roketsan, güncelleme testleri uygulayarak malzemelerin son kullanma tarihlerini uzatabilmektedir, ancak mevcut MRP sistemi, sipariş verirken malzemelerin raf ömürlerini uzatma olasılığını dikkate almamaktadır. Bu projedeki amacımız, son kullanma tarihlerini dinamik olarak güncelleyen ve sipariş miktarlarını sırasıyla sonuç veren, hurda oranını ve stok seviyelerini minimuma indiren bir metod geliştirip buna uygun verimliliği artırmak için bir envanter politikası oluşturmaktır. Oluşturulan bu metodoloji sayesinde materyallerin etkin raf ömürlerinin tahminlemesi de yapılmıştır. Geliştirilen sistem sayesinde hurda miktarları azaltılmıştır ve gereken sipariş miktarları seviyelerde belirlenmiştir.

Anahtar Kelimeler: Envanter, Güncelleme Testi, Sipariş Miktarı

Design of a Decision Support System for Inventory Planning for Materials with Limited Shelf Life in a Make-to-Order Environment

1 System Description

Roketsan A.Ş. is specialized in developing and manufacturing rockets and missiles for customers around the world. Their manufacturing system is a project-based discrete manufacturing with more than 40000 materials, where 8.45% of them have limited shelf lives. Company's planning horizon is 6 years, they use a periodic review system and roll their MRP monthly. The company uses Oracle with MES (Manufacturing Execution System) add-on as their ERP system customized for their manufacturing system.

2 Problem Analysis

2.1 Detailed Description of Current Operations

There are several keywords used by Roketsan which are referred to throughout the project report and they can be explained as the following:

- *Flexible Resource Materials* are the materials used commonly by different departments in several projects whereas *Non-flexible Materials* are not used commonly, they are ordered individually for different projects.
- *Efficient Shelf Life* is a term used for the expected usability of an item in the Roketsan inventory.
- *Updating Test* is applied to a material to extend its shelf-life. The number of tests applied on a material is different for materials, and the period of extension of shelf life is called the *Updating Period*.

Current system in Roketsan starts with materials' arrival to Roketsan. When ordered materials enter the Roketsan's system, they first go into the Quality Control Department and get checked whether their shelf lives are greater than 80% of their projected lifetimes or not. If it is greater than 80%, they are accepted into the company's inventory; if not the material can still be admitted under the condition that the Production Planning Department decides it is absolutely needed. The acceptance rule of 80% is determined by the Quality Control Department and it is calculated for each arriving material as stated in the Appendix A. Furthermore, it takes 3 weeks for Roketsan to request, propose and approve the orders, which is specified as "Preparation Period", and it takes 2 weeks for the Quality Control Process. Figure 1 represents the explained arrival system as a diagram and it can be found in Appendix A.

The second part of the current system is about the testing system of the materials in the inventory. The perishable materials are taken into updating tests at the end of their efficient shelf lives and if the result of these tests are successful, the

efficient shelf lives are extended with a specified amount of time for each material, and they stay in the inventory. Otherwise, the materials are scrapped. In figure 2, this testing system is represented in a diagram and it can be found in Appendix A.

Roketsan's test policy is to apply the updating tests only when the material is expired or very close to expiration. Also, in the current practice they currently hold 15% of the annual demand for each material as safety stock in order to eliminate stock outs.

3 Problem Definition

An inventory ordering and control decision system is designed to eliminate the following problems:

- Inventory Management Related Problems: Each project officer gives orders for the same material independently from the orders of other projects. Thus, excess orders are placed and eventually, scrap rate increases.
- Lead Time Related Problems: The lead times of the materials vary from order to order. Uncertainty of lead times causes exaggerated numbers of orders to prevent stock-outs.
- Efficient Shelf Life Related Problems: Roketsan can extend the expiration dates of the materials by applying updating tests, however the MRP system does not take into account the possibility of extending the shelf lives of materials when the orders are placed. As a result, the order quantities in ASCP appear to be higher than the required amount and eventually inventory level increases.

4 Objective, Scope, Constraints

In this project, our objective is to minimize the inventory level, scrap rate while having no stock-out. In this manner, the outcome of our methodology is the order quantities for each material and safety stock level with respect to the lead time uncertainty and reorder level. Safety stock determination and order decision for each material are made by considering the expiration date and updating tests of the materials. In this project, all procedures and our methodology were planned without cost constraints since we do not have any access to cost data.

5 Proposed System

5.1 Literature Review

According to the article written by Gonçalves et al. (2020), suggested safety stock calculation is modified to our case for overcoming the uncertainty caused by lead times and alteration of plans. In the article, it is suggested that safety stocks are strategies that help to deal with the lead-time, demand and supply uncertainty while diminishing the stock-outs. Wijngaard and Karaesmen (2007) suggested that

in some cases where lead times are positive and under a certain threshold value it is optimal to use “order base stock policy”. The model mentioned in the article is constructed upon production orders and related costs, which are not common with our project. However, in the following parts we used the ideas such as utilizing the order base stock model where we have planned to do replenishment orders when the inventory level falls below a certain level. In addition to this, the idea of using order base stock policy can also be supported by Gallego and Özer (2001) since he suggests that when cost parameters are not known it is better to use base-stock policy as our solution approach, and in our case cost parameters are unknown too. Moreover, Nahmias and Olsen (2015) and Silver et al. (2017) suggest that for the continuous inventory review systems, it is used (Q,R) policy, however, since we have stochastic lead time and no cost information we modified this policy to our case.

5.2 Proposed Methodology

The proposed system is a monthly rolling system that shows possible MRP scenarios according to the updates of the materials and outputs the resultant order quantity of each scenario. Therefore the proposed methodology can be evaluated as an order quantity decision support system and this system can not make an optimization since there is no available cost information.

MRP Table Construction

In order to better elaborate the proposed methodology, first the components of the MRP table in Appendix B are discussed which are *Starting inventory*, *Demand/MRP Output*, *Inventory level and position*, *Expired and updated inventory*, *Reorder and order up to level*, *Order quantity and order received* for which the detailed calculations and explanation can be found below.

1. **Starting Inventory:** The proposed methodology starts with the starting inventory. For the month in which the methodology is run, the starting inventory data is driven from the ASCP of Roketsan. For the remaining months it is calculated as:

$$\text{Inventory level at } (t-1) - \text{Expired Inventory at } (t-1) + \text{Updated Inventory at } (t-1) + \text{Order Received at time } (t)$$

2. **Demand/MRP Output:** Demand, MRP output, data is completely driven from the ASCP module of Roketsan every month.
3. **Inventory Level and Position:** Inventory level at time t shows the inventory at hand after the demand at time t is deducted and it can be calculated as:

$$\text{Starting Inventory at time } (t) - \text{Demand at time } (t)$$

While the inventory position at time t is calculated as follows:

$$\text{Inventory Position at time } (t-1) - \text{Demand at time } (t) + \text{Order Quantity}$$

at time (t) + Updated at time(t-1)– Expired at time (t-1)

For the starting month the *Inventory Position at time (t-1)* should be replaced by *Starting Inventory at time t*.

The difference between inventory position and level is that in inventory position the ordered quantity at time t is immediately added to the inventory position so that the system keeps up with the reorder level consistently, and gives more accurate order decisions. This can be considered as a virtual value, rather than a factual one in order to make necessary calculations.

4. **Expired and Updated Inventory:** Expired inventory row shows the expected number of materials that is going to expire on time t. This data is directly driven from the ASCP module of Roketsan every month. While the updated inventory row shows whether the expired inventory is updated or not in month t.
5. **Reorder Level Calculation:** The components mentioned so far are used to calculate the inventory position which is the key element for the order decision making. After those, the reorder level is calculated which is the level at which if the inventory position is smaller than that, the decision of giving an order will be made. This level differs every month when the system is rolled and it can be calculated as:

$$ROL_t = \left[\sum_t^{t+L+P} D_t \right] + [\Phi^{-1}(\alpha) * \sqrt{[(\mu_L + 1) * \sigma_D^2 + \mu_D^2 + \sigma_L^2]}] \quad (1)$$

The elements of the formula is as follows:

- D_t is the summation of demand row from time t to t +Lead time +1 period
- $(\Phi)^{-1}$ is the safety factor.
- μ_D is the average demand per unit time for a period of 12 months
- σ_D is the standard deviation of the demand D for a period of 12 months
- μ_L is the expected lead time in months for a period of 12 months
- σ_L is the standard deviation of the lead time L for a period of 12 months

All of the mentioned elements of the formula are calculated directly from the data that is provided by Roketsan except the safety factor. Safety factor is different for each material because it is based on the ABC class that the material belongs to. Roketsan uses class A for the most important materials which are harder to obtain, require special permissions and are expensive. Whereas class B stands for the medium important materials and C stands for the remaining materials which are commonly used and are not expensive.

6. **Order Up To Level Calculation:** By using the reorder level the order decision is made and by looking the order up to level the order quantity decision is made. For each period like reorder level, order up o level is different and it is calculated as:

$$OL_t = \left[\sum_t^{t+L+P} D_t \right] + [\Phi^{-1}(\alpha) * \sqrt{[(\mu_L + P) * \sigma_D^2 + \mu_D^2 + \sigma_L^2]}] \quad (2)$$

In this calculation, the parameters except P will be the same with the ones that are shown in reorder level calculation. P is the number of periods between two orders. The proposed methodology reviews the system every month and gives orders every P months. Since no cost information is provided to us, the same order frequencies as Roketsan has been used as long as possible for consistency.

7. **Order Quantity and Order Received:** By using the necessary components for order decision making explained above, order quantity is calculated. The order quantity stands for the amount of the order that should be given at time t and while calculating the order quantity the minimum order quantity and package quantities that are driven from the data are also considered and it is calculated as:

Order Quantity at time t:

*If Inventory Position at time (t-1) > Reorder level at time (t-1) then
OQ=0*

Else, OQ = MAX(Minimum Order Quantity, Order Up To Level at time (t-1) -Reorder Level at time (t-1))

After an order given it appears in the order quantity row at time t. The inventory position at t+1 is updated accordingly. Furthermore in order to update the starting inventory, the arrival of the material to the system should be considered and it can be found in the Order Received row where it equals to:

Order Quantity at time (t - lead time)

The necessary estimations of the alpha, $(\Phi)^{-1}(\alpha)$, standard deviation and P values for the calculation of order quantity can be found in Appendix C.

Success Rate and Expected Shelf Life Calculation

In addition to the essential components of the MRP table mentioned above, the success rate and expected shelf life calculation are significant to evaluate possible scenarios.

1. Conditional Success Rate Calculation The conditional success rate of the materials refers to the possibility of passing the i^{th} test and it is calculated as follows:

- (a) The lots of the same material are separated.
- (b) For each lot, raw data taken from Roketsan is analysed and by using “Test Completion Date”, whether the i^{th} test of each lot is successful or not is found.
- (c) For each test index i , probability p_i of passing that the i^{th} test is calculated by dividing successful lots to total number of lots having the i^{th} test.
- (d) Then, starting from the first test, conditional probabilities are calculated and success rate is found. For example, success rate for the first test is p_1 , for the second is $p_1 p_2$.

2. Expected Shelf Life Calculation

- p_i : probability passing the i^{th} test of the material.
- T : updating test period of the material which is driven from the data.
- SR_i : success rate of the i^{th} test of the material (found by conditional probability as explained above, $\prod_{i=1}^I p_i$)

By using the parameters above, expected shelf life extension that shows the amount of months that the material's shelf life will extend according to its expected shelf life is calculated. This calculation is used in the total expected shelf life calculation and in scenarios that are mentioned in the following sections and it as follows:

$$ExpectedShelfLifeExtension = \sum_{i=1}^I SR_i * T \quad (3)$$

By considering the above formulation the Expected Shelf Life is calculated as follows:

$$ExpectedShelfLife = S + \sum_{i=1}^I SR_i * T \quad (4)$$

Where S is the initial shelf life of the material that is driven from the data. The efficient shelf life which is raw data obtained from Roketsan, gives us a value that shows the total expected shelf life of a material without considering the updating tests.

Working Principle and Related Outputs of the Methodology

As it mentioned earlier the proposed methodology elaborates possible scenarios in which the expired inventory at every month is either passed or failed. Due to the stochastic nature of test probabilities there happens to be 2^n scenarios where n is the number months that have expired inventory. Since a large number of MRP tables is really hard to elaborate and decide from, the proposed methodology

will only consider 3 of these scenarios which can be denoted as “best”, “worst” and “expected shelf life” cases for which the explanations are as follows and the exemplary tables are shown in Appendix B.

1. **“Best case”** refers to the scenario in which all the materials located in the expired inventory row for each month will be updated so there is no change in the inventory.
2. **“Expected shelf life case”** refers to the scenario where the expected shelf life of the material is extended in accordance with the previously calculated “expected shelf life extension”. This extension is shown by writing the expired amount of material to the extended period of “Expired Inventory” row. This case is the same as the best case generally, as the expected shelf life of the material is generally longer than the lead time. However as we roll the system, the updated material will appear on the expired row, unless the amount of the material is not used before its extended period. By this way, a difference may occur in the system, so it should be considered.
3. **“Worst case”** refers to the scenario where all the materials in expired inventory row will expire so the amount is subtracted from the starting inventory.

As a result of these three scenarios, every month there is a possible decision node for ordering. At that decision point, in order for the Roketsan team to give a better order decision, the resultant order quantities of that month according to those three scenarios are displayed and the final decision is up to Roketsan. An example for how these three scenarios can be seen in the Appendix B while the resultant order quantities and the decision point can be seen in the Appendix D.

6 Validation

In order to validate that our system is consistent with Roketsan’s ongoing system, we have modified our safety stock calculations and found order quantities with respect to them. Roketsan takes the 15% of annual demand and keeps as safety stock, therefore, the second part of our ROL calculation which is used for safety stock calculations was changed to 15% Roketsan. The comparison of our methodology’s Worst and Best Case to Roketsan’s current system is summarized in Appendix E.

A sample of 15 materials from classes of A, B and C was created by adding frequently used materials on top of the ones suggested by Roketsan for validation. After running the methodology, we observed that our system yielded similar results in terms of order quantities, frequencies and times. When comparing these values of Roketsan system with our methodology’s Worst Case scenarios, for the majority of the materials in this sample, they remained the same. However, there was a slight difference for 3 materials. For CH1352, where our methodology suggests an order frequency of 5, Roketsan’s system suggests 3. For AD1010, Rocket-

san has an order frequency of 3, whereas our methodology has 4. For CH1352 and AD1010, the total order quantities in both systems remained the same. For CH1318, where Roketsan only orders once, our methodology orders twice with an additional order quantity of 15025 grams and it increases the total order quantity by 80.3% on the annual basis.

When comparing the order quantities, frequencies and times of Roketsan with our methodology's Best Case scenarios, it is observed that in our methodology, the majority of the samples have the same values as Roketsan. Only the two of the materials in the sample have a difference. For AD1010, Roketsan's system proposes an order frequency of 1 but our methodology proposes ordering 2 times with an additional quantity of 1872 grams which corresponds to a 12.4% increase in order quantities. For CH1318, where Roketsan has an order frequency of 1 for 18711 grams, our methodology orders twice with an additional quantity of 15025 grams which corresponds to an 80.3% increase in order quantities. This increase is calculated on an annual basis.

In order to find whether Roketsan's current system orders according to our methodology's Best or Worst Case scenarios, we compared them and found that for the majority of the samples, Roketsan orders according to the Worst Case Scenarios. Roketsan's current system and our methodology are yielding similar results and our methodology proposes decisions that are relevant and consistent with the Roketsan's current system for almost each material in the sample as expected, since our Worst Case simulates the Roketsan's current system. Therefore, our methodology is validated and works as it should be. However, in this analysis, the model is not validated for materials that enter the updating tests and extend their shelf lives because Roketsan's system does not consider this situation and we cannot compare our results with Roketsan for that reason. However, with our MRP tables validated according to the current system of Roketsan, it can be ascertained that our methodology will yield similar results for materials that enter the updating tests and will be consistent with Roketsan's system.

7 Benchmarking

For benchmarking, performance measures of scrap rate and inventory level of the current system of Roketsan and proposed system were compared and analyzed to observe the improvement. Since Roketsan does not have access to past MRP data, the measurements are completed for the period between November 2020 and March 2021. From the given data, the scrap rates and inventory levels of samples were calculated for the current system. Then, the proposed solution was implemented and run based on the given MRP data of November 2020. The scrap rates and inventory levels of the proposed system were measured for the same period. In benchmarking, the previously used sample was used by changing 3 of the materials, since 12 materials in the sample did not change due to the limited time horizon. Therefore, 3 new materials (R94321701-3, AD1007-1, AD1022-1) which expire within a specified time horizon and have short extension life, were

added to measure the effectiveness of the proposed system more accurately.

The scrap rate of a material is calculated by finding the ratio of the expired amount within the time horizon to initial amount of stock at the beginning. We took the “expected shelf-life” scenario into consideration. Inventory levels were also calculated by taking the average of inventory levels within the time horizon. For instance, for the recently added material R94321701-3, we first calculated scrap rate of the current system by finding actual expired amounts of material within November 2020 and March 2021, and calculated inventory level by taking the average. The scrap rate was found as 63.3% within the period. After the measurement of the current system, we ran our methodology by taking November 2020 as the starting month in order to obtain the scrap rate and inventory level between November 2020 and March 2021, which shows the output of “expected shelf life” case if Roketsan had used our proposed methodology. The scrap rate for the proposed system was also calculated by taking the ratio of expired amount to initial amount of the lot. The new scrap rate of the material was found 61%. The scrap rate decreased in our system because the proposed system uses 2.3%, which equals 1.27 KG, more than the current system of Roketsan. Also, average inventory decreases from 71,7 to 68 KG due to decreasing order quantities.

Comparison of scrap rate and average inventory level of the current and proposed system for each material can be found in table in Appendix F. In order to find the average improvement of scrap rate and inventory level, the formulas given in Appendix F are used. As a result, between November 2020 and March 2021, our proposed system yields a 4% inventory increase while decreasing scrap rate by 3% for the given sample. In the long term, it can be suggested that the scrap rate and average inventory level would decrease because since we do not have access to past MRP data, it was not possible to compare the results in a longer time horizon than the interval of November 2020 and March 2021.

8 Implementation and Integration to the Company

To implement the proposed system, the user interface and the decision support system have been created by using Excel VBA for the production planner team of Roketsan. In order to run our model, the respective inputs given as separate Excel sheets (for all materials combined) are copied into our worksheet once without changing their formats. Our Methodology automatically extracts the related information, constructs MRP and calculates order frequency and success rates.

The user interface in Appendix G, consists of three list boxes which are “Material”, “Starting Date” and “Ending Date” respectively. The end-user selects the material that he/she wants for planning and dates are selected to determine the planning horizon of the MRP output. After the selection, the end-user clicks the “Submit” button that runs VBA code. The algorithm behind the interface, works and outputs the order quantities and times that are suggested by the scenarios and

the efficient shelf life of the material. At the end, results of order quantities and success rate of material for each test are printed on the “Results” pop-up (Appendix D) and also the Roketsan team is able to see best and worst case MRP outputs after closing the pop-up screen.

In order to implement this methodology, the only thing that Roketsan should do is to copy the necessary data that they provided to us into excel worksheets. Because the code works in accordance with the exact format of the data they provide, they won’t encounter a problem. Also, the methodology is coded in a dynamic way that when new material’s data is included in the spreadsheets, it will automatically understand and perform the necessary calculations. After the discussions with Roketsan, upon their request, in order to ease the process of implementation a video that explains how to copy the data and create a new material for search is provided to Roketsan. Furthermore, necessary explanations on the excel sheets are made as comments for them to easily follow the methodology.

9 Project’s Possible Contribution

Our solution will provide several benefits to the company. First of all, as it mentioned earlier with the new methodology they will decrease the scrap rate and inventory levels. Secondly, the methodology will calculate the success rates of the materials by using past data and calculate expected shelf life of the material. In addition to this, it also overcomes the weakness of the current MRP system which is ignoring the possibility of lifetime extension of materials by evaluating the “best, expected shelf life and worst case” scenarios. Although the final decision is always up to Roketsan, we recommend them to evaluate scenarios as follows:

1. If all three scenarios suggest not the order then no order should be placed to prevent possible increase in scrap rate and inventory level.
2. If only the “worst case” scenario suggests ordering. Then the “expected shelf life case” and the lead time should be checked.
 - (a) If the material has a long lead time and “expected shelf life case” does not cover extended demand, then an order should be placed. Even if the expected shelf life case covers enough demand, due to long lead time, to be on the safe side, an order can be placed.
 - (b) If the material does not have a long lead time, then although expected shelf life case” does not cover extended demand, since it can be placed and obtained in a short time, it is suggested to not give an order.
3. If the “worst case” and the “expected shelf life” case scenario suggests ordering but the best case scenario does not, then an ordering should be made and to be on the safe side, the order quantity should be determined according to the worst case.
4. If all three suggest giving an order, then an order should be placed and the

quantity should be equal to the worst case scenario.

Thirdly, by using a more consistent approach to holding an inventory which is calculating the Reorder Level (ROL) we not only offer a better Safety Stock calculation than the current system but also satisfy the company's specification of not stocking out or backlogging, because several extreme cases in the system are considered and it is shown that in none of these cases there was a problem of stocking out. Furthermore, deciding order quantities by looking at the Order Up to Level (OL) helps us to overcome the fluctuating demand data and uncertainty of lead time by encountering their standard deviation and mean of them over a 12 months period. Moreover, the user interface that is created helps the company to reach any information about a material faster and quicker since the data of a material is more organized than before. Finally, from the societal and global perspective, our system has a positive impact on environmental pollution and carbon emission by enabling more efficient use of materials since it results in a decrease of scrap rates.

In the final meeting with Roketsan, they stated that they are not only happy with the output of the project but also with the contributions of it. All the necessary requirements they specified before are satisfied except running the code for all of the materials at the same time. However, this does not affect the working principle and the amount of materials that can be implemented to our project, it only affects the running time. Therefore, despite that, they mentioned that they are highly satisfied with the outcome of this project.

10 Conclusion and Recommendations

Roketsan can extend the expiration dates of the materials by applying updating tests, however their current MRP system does not consider the possibility of extending the shelf lives of materials when they place their orders. To satisfy the expectations, we come up with a methodology that works based on scenarios in which the expiration dates and the extension possibilities of materials are dynamically considered. As output, the system provides order quantities, and the expected shelf life of the materials for that month.

In order for this methodology and the process in general to work more effectively, it is recommended to Roketsan to control this data accumulation process regularly and in a more organized manner. Because as it mentioned earlier in the benchmarking and success rate calculation, there is a lack of past data and this deficiency affects the estimation of data driven parameters, past data analysis and order quantity decision directly.

Furthermore, in the current system of Roketsan, accepting the materials that have 80% or more efficient shelf life may cause inventory accumulation and hence scrap rate increases. As this check is performed independent of the material planned demand and safety stock, it is also recommended for Roketsan to evaluate whether they need the material or not for the following months at that point rather than

accepting according to that condition.

It is recommended to look at the demand that can be covered by the efficient shelf life of the material by using the expected shelf life information which is calculated by using success rate or scenarios. If the worst case scenario suggests not to give an order to the point that the accepted materials efficient shelf life will end, then accepting that material at the beginning will only increase the inventory level. If the only worst case suggests placing an order, then before the material is accepted, the lead time of the material should be checked. If the lead time of the material is long and the current "efficient shelf life case" covers just enough for the demand then materials can be accepted to the Roketsan inventory to be on the safe side. But, if the material does not have a long lead time, then in order to decrease the possible increase in scrap rate it can be rejected. If all three or "worst" and "efficient shelf life cases" suggest giving an order for that time, then the materials should be accepted. Therefore, considering the 80% accepting condition can improve the system, however, at the end the decision is always up to Roketsan.

In conclusion, we developed a methodology that considers expiration dates and the extension possibilities of materials and when comparing the current system of Roketsan to our methodology, we see a decrease in both the inventory levels and the scrap rates.

References

- G. Gallego and Özer. Integrating replenishment decisions with advance demand information. *Management Science*, 47(10):1344–1360, Oct. 2001. ISSN 0025-1909, 1526-5501. doi: 10.1287/mnsc.47.10.1344.10261. URL <http://pubsonline.informs.org/doi/abs/10.1287/mnsc.47.10.1344.10261>.
- J. N. Gonçalves, M. Sameiro Carvalho, and P. Cortez. Operations research models and methods for safety stock determination: A review. *Operations Research Perspectives*, 7:100164, 2020. ISSN 22147160. doi: 10.1016/j.orp.2020.100164. URL <https://linkinghub.elsevier.com/retrieve/pii/S2214716020300543>.
- S. Nahmias and T. L. Olsen. *Production and Operations Analysis*. Waveland Pr, Long Grove, Ill, 7. ed edition, 2015. ISBN 9781478623069. OCLC: 935795578.
- E. A. Silver, D. F. Pyke, E. A. Silver, and E. A. Silver. *Inventory and Production Management in Supply Chains*. CRC Press, Taylor & Francis Group, Boca Raton, fourth edition edition, 2017. ISBN 9781466558618.
- J. Wijngaard and F. Karaesmen. Advance Demand Information and a Restricted Production Capacity: On the Optimality of Order Base-Stock Policies. *OR Spectrum*, 29(4):643–660, July 2007. ISSN 0171-6468, 1436-6304. doi: 10.1007/s00291-006-0076-x. URL <http://link.springer.com/10.1007/s00291-006-0076-x>.

Appendix A Current System in Roketsan

Acceptance Rule = [the material's expiration date - the material's arrival date to Roketsan]/[the material's shelf life].

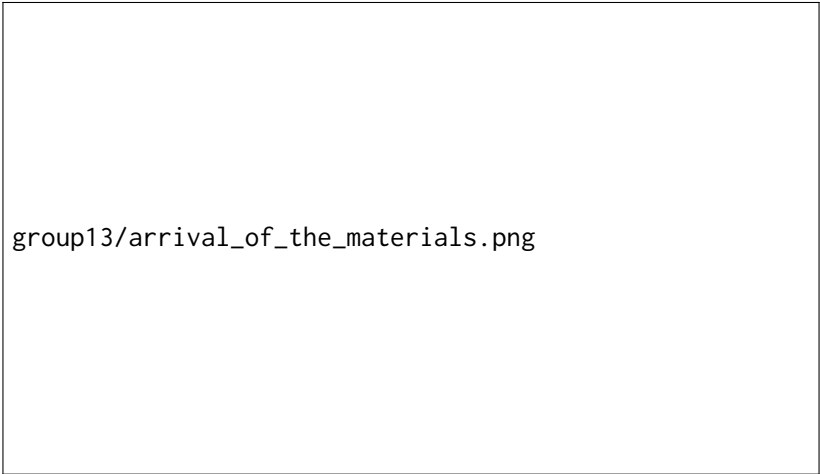


Figure 1: Arrival of the Materials

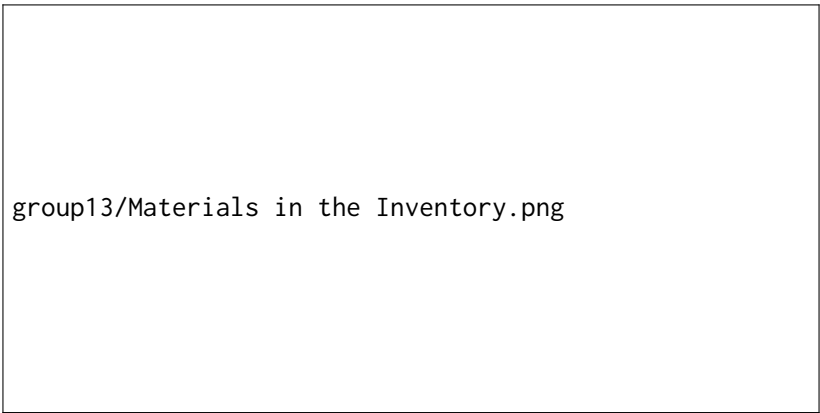


Figure 2: Material's in the Inventory

Appendix B Estimations

Alpha = Represents the service values. The service values of the A, B and C classes are determined by Roketsan.

Z = The z values corresponding alpha values are adapted from the literature and they depend on the normal distribution.

Standard Deviation of Lead Time = The standard deviation of the lead time is estimated from the “variable lead time” data while the mean of the lead time

is estimated from the “constant lead time” data that Roketsan provided.

Standard Deviation of Demand = It is calculated by finding the standard deviation of the demand data from time t for a period of 12 months.

P = P is the time between two orders. Order frequency of the materials is found by analyzing the past data and it is used to calculate the P value (periods between two orders) for each material. P is calculated as $12/\text{Order Frequency}$ of the material. For instance, if the order frequency of the material is 2 orders per year, then P value is equal to 6 months.

Appendix C Output of Results

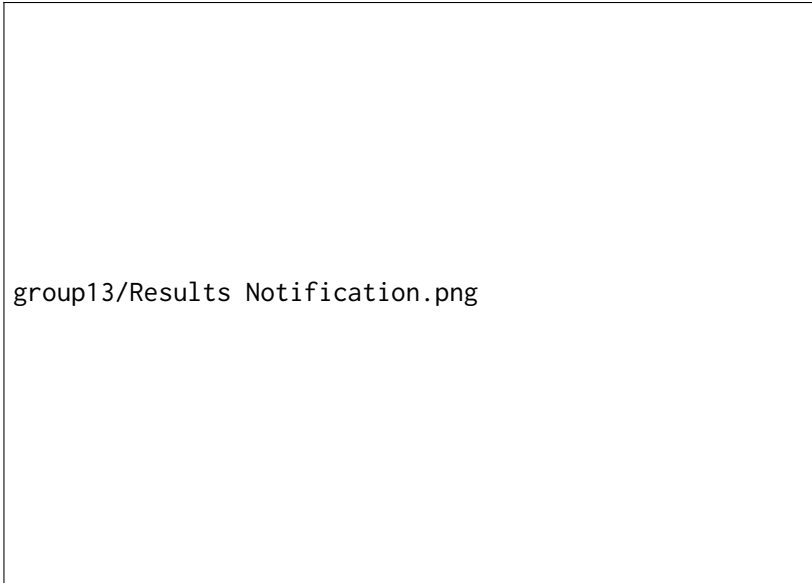


Figure 3: Results notification

Appendix D Benchmark

D.1 Benchmark Formulas

- ***Decrease in Inventory*** = $(\text{Average Inventory Level in Roketsan} - \text{Average Inventory Level in Our Methodology}) / (\text{Average Inventory Level in Roketsan})$
- ***Decrease in Scrap Rate*** = $(\text{Average Scrap Rate in Roketsan} - \text{Average Scrap Rate in Our Methodology}) / (\text{Average Scrap Rate in Roketsan})$

Appendix E User Interface



Figure 4: User Interface

Eti'nin Esenyurt Deposu için Hibrit Bir Palet Adresleme Sistemi Tasarımı ve Uygulanması

Eti Gıda Sanayi ve Ticaret A.Ş.

group14/group_photo.png

Proje Ekibi

Zeynep Arslan, Göksun İşliel, Pınar Şavkay Oymaner, Can Özgünsür,
Adnan Emre Selçuk, Eylül Yurtseven, Rabia Yücel

Şirket Danışmanı

Hüseyin Gençer
Mikro Lojistik Planlama Birim
Yöneticisi

Akademik Danışman

Yrd. Doç. Dr. Özlem Karsu
Endüstri Mühendisliği Bölümü

Endüstri Mühendisliği Takımı

ÖZET

Eti'nin Esenyurt deposundaki güncel palet adresleme sisteminde depoya gelen paletler depoda kendi ürün türlerine göre ayrılmış saklama alanlarına shopping alanına en yakın olacak şekilde yerleştirilir. Mevcut palet adresleme sistemi depo içi lojistik hareketin artmasına neden olmaktadır. Bundan ötürü şirketin ve deponun kaynakları verimsiz bir şekilde kullanılmış olur. Bu projenin amacı depo içinde ürün taşımak için kat edilen yolu azaltabilen bir dinamik palet adresleme yöntemi sunmaktır.

Anahtar Kelimeler: dinamik adresleme, depo yönetimi, depolama planlaması

Keywords: dynamic addressing, warehouse management, storage policy, material handling operations

Design and Application of a Hybrid Pallet Addressing System for Eti's Esenyurt Warehouse

1 Company Information

Eti is an Eskişehir based company that produces various types of food such as chocolates, biscuits, and cakes. Eti has eight production plants in total. Their largest production plant is in Eskişehir. Goods are directly transported from Eskişehir Logistics Center (ELC). Each city reserves at least one distributor with a total number of 120 country wide. Large-scaled locations such as İstanbul usually contain multiple distributors. Esenyurt Warehouse, which is in İstanbul, is the largest warehouse of the company and it contains almost every type of product that the company produces. The goods that are stored in these distribution centers later get delivered to the demand points.

2 System Analysis

There are different types of storage areas based on product types and their designated temperature inside the Esenyurt Warehouse. In addition to these product-type-based storage areas, there is another storage area called the shopping area which can be seen in Figure 1.



Figure 1: Illustration of the warehouse.

Pallets with bar-codes that contain product details arrive at the warehouse. Using these bar-codes, the relevant data gets checked and these pallets are addressed to the nearest available cells in their own storage areas. Then pallets are transferred to the shopping area to ease the shipping processes.

Generally, there are one or two pallets for each Stock Keeping Unit (SKU) in the shopping area, each pallet consists of a certain number of parcels. When the number of parcels falls below a certain limit, a refill takes place and a new pallet gets moved from the product base storage area to the shopping area. The customer orders are usually in parcel units and they are picked from the shopping area manually by an operator. Products with closer expiration dates than others are offered for sale first. This strategy is named as FEFO, which means first

expired first out. The collected products are placed in the haulage vehicles from the shipment area and get transported to their pre-defined destinations. The warehouse keep track of all these product flows and processes with their ERP system.

3 Problem Definition

When we examine material handling processes inside the warehouse, the problem starts when addressing the received pallets to their own storage areas. The current addressing system of the warehouse does not consider the characteristics of the products such as sales rate and material handling times while assigning them to a storage cell inside the warehouse. Current system looks at the empty storage cells and addresses the pallets to the nearest available cells within that storage area. Hence, none of the products have any specific or predetermined storage space in the warehouse.

This addressing system increases the total time spent for material handling inside the warehouse. To illustrate, if we assume that the warehouse is empty at the beginning, and the first products that arrive at the warehouse are the ones with low sales rates, then the current system will assign these products to the nearest cells to the shopping area in the warehouse. The current system works such that when products with higher sales rates arrive, they must be stored at the remaining cells which can sometimes be further away from the shopping area, hence, may require longer distances to be covered. That being the case, to satisfy the demand on products with high sales rates, the warehouse staff would sometimes have to pick up the products from the furthest cells of the storage area multiple times. This is a result of the fact that products with low sales rates may sometimes occupy the storage areas that are near the shopping area. Overall, the current system that is utilized at the product-type-based storage areas may cause the workers to cover longer distances, hence time spent on these processes increases. The deliverable of this project is to develop a new pallet addressing system which would decrease the total material handling time inside the warehouse.

3.1 Literature Review

To define our problem and form a systematic approach for order picking optimization and related performance problems we need a clear understanding of the warehouse processes and functions. Eti uses a manual storage system which employs a high-level picker-to-parts or picker-to-stock sub-system. In this system workers manually pick the items using logistic equipment (Manzini, 2012).

Eti uses a Unit Load Storage System based on pallets. These pallets are stored with a random storage policy in the main storage areas with changing stock levels. Random storage policy assigns the products to the closest location in the warehouse to minimize logistic cost (Manzini, 2012). A good approach for our problem includes the item activity profile to optimize the logistic costs. Item activity profile gives us product information such as: activities of an item and its

volume, daily demand variation, seasonality, handling characteristics etc., these information are very useful when making decisions for the warehouse and finding the proper slotting method.

There are different models for addressing policies. For example the “Individual Unit Load Storage Model” finds the best location for a single product while minimizing the logistic costs. In this model, product units cannot share the same location in a time period and one unit cannot be in more than one location (Manzini, 2012). It is proven that turn-over based systems for single command operations optimizes the time spent for the order picking processes (Park, 2012).

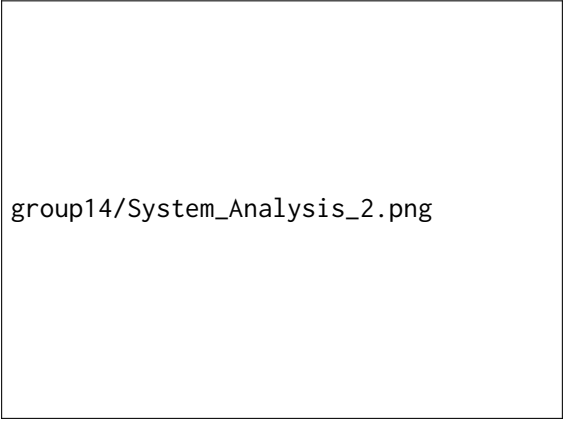
Warehouse operation optimization models includes elimination of non-value-added activities like unnecessary forklift and package movements, time and distance traveled for retrieval as well as storage (Sharma and Shah, 2015).

4 Solution Approach

4.1 Method

A hybrid addressing model which is a mixture of both dynamic and static models will be constructed to minimize the material handling times in the warehouse. This model will address the products to the most desired cells among their designated storage area. In this case the most desired cells are the ones that are closest to the shopping area.

At first, considering the receiving and shipping frequency of each product, we have created two different product segments: fixed and non-fixed. The fixed product segment is consisted of the products that have high receiving and shipping frequency. The remaining products automatically belong to the non-fixed segment. Both of the segments can include all product types. As a result of this segmentation a certain number of cells will be needed to store every fixed product according to their turnover rates and average stock levels. After creating the list a general area to store all fixed products will be reserved in their designated storage areas. The segmentation will be performed separately for different product types. Cells that are reserved for fixed products will only be allowed to store fixed products. However, if all the reserved cells are occupied the fixed products will be addressed in the most desired non-reserved cells. The decision making process regarding the segmentation of products and cell reservation process are explained in detail in Section 4.2. After the segmentation, an algorithm is designed to address the incoming pallets according to the designated segment of the products that are inside of them. The algorithm is explained in the Section 4.3. The illustration regarding this addressing policy can be seen in Figure 2.



group14/System_Analysis_2.png

Figure 2: Illustration of the proposed addressing policy.

4.2 Determining the List of Fixed Products

The algorithm reserves the best cells (least amount of time to reach) in the warehouse to the “fixed” segment by using a list consisting of the SKUs of the products in this segment, their types (ambient or chocolate) and number of cells to be reserved for each of them. When creating this list, we thought the warehouse as a “queuing system” where arrival times of customers are the arrival times of pallets to the warehouse and average waiting time is the amount of time pallets spent in the warehouse waiting to be shipped. We calculated the Poisson hourly arrival rate (λ) of pallets by utilizing the dataset regarding the product receiving which was provided by the company. This dataset was grouped according to the pallet receiving date and hour. Then average of hourly pallet arrival $\bar{X}$ is calculated for every product to estimate the λ parameter. To calculate the time that each pallet spends in the storage areas, turnover rates and average stock of each SKU were used. The average time a pallet spends in the warehouse (W) is calculated by the dividing turnover rate of products to their number of stocks. After calculating both λ and W we applied the Little’s Law ($L = \lambda W$) (Little and Graves, 2008), to find average amount of pallets stored in the warehouse (L). L is the number of cells that will be reserved for the fixed segment. For example a product has arrival rate of $\lambda = 1$ pallet per hour and occupancy time for $W = 2$ hours in the storage area, in this case Little’s Law will tell us to reserve $L = 2$ cells. After finding the 2 values of each product, products were sorted according to the frequency of their receiving and shipping operations. To decide on the number of products that will be in the “fixed” segment we used our algorithm mentioned in Section 4.3. “Fixed” lists with different number of products given to the algorithm as an input to find the best list that decreases the total material handling time of the system. We did this experiment to decide the number of cells to be reserved for the fixed products in order to decrease the material handling time. Now this experiment will be explained in detail. For this purpose, assume that there are 10 different Type A products (SKUs). For each product, the previously mentioned L values

are calculated. Then, all 10 of these products are sorted in the descending order with respect to their frequencies of receiving and removing operations. After that, starting from the top of the sorted products' list, the cumulative values of L , i.e., the total number of cells is calculated. This list is available in Table 1.

Table 1: List of candidate Type A fixed products.

SKU	Type	L	Total Cells
7	A	4	4
8	A	3	7
6	A	2	9
3	A	1	10
5	A	3	13
1	A	2	15
4	A	3	18
10	A	4	22
9	A	1	23
2	A	2	25

Now, all 10 of the products in this list are candidates to be included in the fixed segment and there are no products in the fixed segment yet. From this list, we select the first row, look at its value of "Total Cells", reserve that number of most desired cells in the Type A storage area for Type A's fixed segment. Now, only the SKU of the first row will be included in the fixed segment and other products will be labeled as non-fixed. Then, we run our replication model, obtain the total time spent for the material handling operations in that trial and proceed to the second row of the list. For the second row, we look at its value of "Total Cells", reserve that number of most desired cells in the Type A storage area for Type A's fixed segment and update the fixed segment. The fixed segment will now include the SKUs of the first two rows. After that, we run our replication model, obtain the total time spent for the material handling operations in that trial and apply the same iterations until we reach the end of the list. Once we reach the end of the list, we detect the row that yielded the lowest material handling time. Then, we go to that row, we look at its value of "Total Cells", reserve that number of most desired cells in the Type A storage area for Type A's fixed segment and update the fixed segment. The fixed segment will now include the SKUs up to that selected row. For example, when we consider the list in Table 1, if the lowest material handling time was observed in the fifth row, we reserve 13 cells in the Type A storage area for the Type A fixed segment and we select product 7, 8, 6, 3, and 5 as the fixed products. The remaining of the products are now non-fixed products and cannot not be addressed inside of these 13 reserved cells. This decision is supposed to be the final decision regarding the fixed segment of Type A products.

Moreover, we did this experiment with historical data. Assume that only two product types will be segmented: product type A and B. The algorithm was run separately for type A and type B products. For the type B list the number of cells reserved for type A products was kept as a constant number. The result of this first experiment is available in Appendix A, Figure 3. As a result of the experiment we decided to reserve 8.8% of the Type B storage area as fixed shown in the third trial.

After finding the best number of cells reserved for type B pallets, same process is applied for the type A fixed list while keeping the number of cells reserved for type B products a constant number. The result of this second experiment is available in Appendix A, Figure 4. The result of this experiment showed that, reserving 13.5% of two storage areas improves the current system by 1.45%.

4.3 Algorithm

The main goal of the algorithm is to determine in which cell that an incoming pallet should be addressed and from which cell an outgoing pallet should be removed so that the overall time spent for the material handling operations is decreased as much as possible. It is important to state that this algorithm does not include any addressing decisions regarding the addressing positions of the products at the shopping area. It only performs addressing decisions regarding the product-based-storage areas. Moreover, we will refer to our main algorithm as Hybrid Pallet Addressing System, HPAS in short. Sets and parameters of HPAS are available in Appendix B.1.

Moreover, the algorithm has three distinct procedures: "Reserve Fixed Cells", "Pallet Removal", and "Pallet Addressing" which are available in Appendix B.2. All three of these procedures run only when they are called to run. In addition, the algorithm does the separation of the product types at the very beginning and proceeds performing the operations only considering the product-based-storage area in which the relevant product belongs to. Regardless of a product's type, the same functions are going to be utilized. So, from now on, one should assume that the product type separation has already been completed and the operations are performed only regarding the relevant product-based-storage areas.

Reserve Fixed Cells

Each time the "Reserve Fixed Cells" function is called, it is assumed that none of the cells are reserved for fixed products. When it is called, the "Reserve Fixed Cells" function takes the first product (SKU) from the set F and sets the number of cells that are going to be assigned for that fixed product as f_p . Then, starting from the most desired cell, it checks whether a cell has already been reserved as a fixed product cell. If the current cell is not reserved, then, the function reserves this cell for fixed products. If the current cell is already reserved as a fixed product cell, the function simply proceeds to the next most desired cell in the layout. Once it completes reserving f_p cells, the function proceeds with the next product in the

set F and repeats the same process until all of the fixed products are considered or the value of cell variable reaches to the total number of pallet cells in the relevant product-type-based storage area.

Pallet Removal

The main goal of the "Pallet Removal" function is to make sure that the chances of a more desired cell being available is much higher than a less desired cell. The "Pallet Removal" function can take a single operation of pallet removal at a time. To be more specific, it takes a single SKU (p) as an input, which belongs to the product that is desired to be removed from its product-type-based-storage area. Then, it takes the total number of pallets that are desired to be collected for that specific product, which is the n parameter. Once these two parameters are provided, the function takes the sets C and K as inputs so that it can check whether a cell is reserved as a fixed product cell and which cell contains which SKU at the moment. After taking the total number of cells in the product-based-storage area for the last input, it starts scanning the storage area starting from the cell located at the most desired location of the warehouse and proceeds to the next most desired location until it scans the whole warehouse. During the scanning process, once a pallet of this product is found, it gets removed from its cell, the value of the *count* variable is increased by one, and the whole process is repeated until the value of *count* to the total number of pallets that are going to be removed, which was the parameter n . After the value of count is equal to n this individual pallet removing operation finally gets completed.

Pallet Addressing

The main goal of the "Pallet Addressing" function is to make sure that an incoming pallet is addressed in the most desired cell available in the relevant product-based-storage area. The "Pallet Addressing" function can take a single operation of pallet addressing at a time. To be more specific, it takes a single SKU (p) as an input, which belongs to the product that is desired to be addressed in its product-based-storage area. Then, it takes the total number of pallets that are desired to be addressed for that specific product, which is the n parameter. Once these two parameters are provided, the function takes the sets F , C and K as inputs so that it can check whether p is a fixed product, whether a cell is reserved for fixed products and which cell contains which SKU at the moment.

After taking the total number of cells in the product-based-storage area for the last input, it starts scanning the storage area starting from the cell located at the most desired location of the warehouse and proceeds to the next most desired location until it scans the whole warehouse. There is an important point. The function does not let a non-fixed product to be addressed inside a cell that is reserved for fixed products even if that cell is currently empty. So, the function first checks whether p is a fixed product or not. If p is a fixed product and the function finds an empty cell that is reserved for fixed products, it addresses that pallet inside that cell. If the function cannot find such a cell, it addresses the pallet inside the

cell that is the most desired one among the other available cells which are not reserved as a fixed product cell, i.e., the cells for the dynamic products. Notice that if the relevant product is not a fixed product, the function will eventually be forced to go through the later process. After an appropriate cell is found and a pallet of this product is addressed, the value of the count variable gets increased by one, the cell counter (*cell*) is reset to zero and the function restarts scanning the same product-based-storage area from the very beginning. The whole process is repeated until the value of count is equal to the total number of pallets that are going to be addressed, which was the parameter n . After the value of count is equal to n , this individual pallet addressing operation finally gets completed. Moreover, please note that the procedures should be performed each time they are called.

5 Validation

In the system, the main issue that has potential to make our approach invalid is maxing out the warehouse capacity, thus not having any empty cells for new coming pallets. In the past they have never experienced such scenario. However, as our approach reserves a particular portion of the storage areas for the fixed segment products, the portion of the storage areas that can contain non-fixed segment products gets decreased. So, even though there are empty cells in fixed segment, the non-fixed products cannot be addressed if all of the non-fixed segment cells are occupied. When we observed the results of the replication model, which is explained in Section 6 in more detail, we have seen that the occupancy ratio of the non-reserved areas was varying between 60% to 70%. In conclusion, with the utilization of historical data in our replications, it can be said that our approach is not imposed to the risk of being invalid considering that the physical limits of the warehouse were never violated.

6 Implementation

For the implementation, we were not able to implement our solution approach inside the actual warehouse due to COVID-19 restrictions. So, we have created a model to artificially replicate the relevant material handling operations inside the warehouse. For the inputs of the replication model, the company provided us previous data regarding the actual pallet addressing and removal operations performed in the product-type-based storage areas for 11 consecutive months. This dataset included the SKUs of the addressed/removed pallets, the exact time that each operation was performed, and location of the cells that each of these pallets were addressed/removed. The company also provided us with the layout plan of the warehouse. By utilizing the layout plan, we have replicated the entire warehouse inside our replication model. Once our replication model was constructed, it was ready to run and perform the pallet addressing/removal operations. For the comparison of our approach's results with the current system's results, we have run the same replication model with the same inputs once for both approaches. In order to do that, we have written a second algorithm which performs the pallet

addressing/removal operations based on decisions given by the current system. Note that we already had the algorithm of our own approach, HPAS. Once both algorithms were ready, we had to decide on the starting condition of the warehouse before running the replication model. For that, the company provided us the current state of the warehouse cells in the very beginning of the same 11-month period. In that state, approximately 36% of the total product-type-based storage area cells were occupied. The occupied cells were dispersed inside the warehouse. We acted greedy and decided to try the worst-case scenario by assuming that only the top 36% of the cells were occupied instead. Meaning that when we order the cells from the most desired to the least desired, the first 36% would be occupied and we would not be able to reserve these cells for the fixed products in the early stages of the replication. After every requirement was satisfied, we have finally run the replication model separately for both addressing approaches. Note that the list of fixed products was provided to HPAS before we ran its replication. In the replication, we have kept track of the time spent for each individual pallet addressing/removal operation and added each of them to the total material handling time in both separate occasions. After both of the replications were completed, the model provided us the total time spent for material handling operations for both cases. When we looked at the final results, we have seen that HPAS was able to record an approximate of 1.45% improvement compared to the total material handling time when the current approach is utilized. Thus, by replicating the system, we were able to predict what would have been the outcome if they had utilized HPAS during the same 11-month period and observed that there could have been around 1.5% improvement in the total material handling time. Also, note that this improvement was yielded when the initial condition of the warehouse is in the favor of the current approach as it is in the worst possible starting condition for HPAS. So, the improvement percentage can expected to be higher in real life applications. Moreover, the company can implement the proposed system in two ways, either by removing all of the non-fixed product pallets that are currently located inside the cells that should have been reserved for the fixed products or by keeping the current state of the warehouse as is and letting the warehouse adapt the system gradually. First option is not feasible since the warehouse cannot fit such operation to their schedules, either they have to pause all other warehouse operations or work over-time. Second option is more attainable however there will be an adaptation process. As mentioned, HPAS algorithm takes the fixed product list as an input and aims to address these products at the cells that are reserved for fixed products and to address the non-fixed products at the remaining cells. In the adaptation process to the system, they may not be able to observe any improvements right after initialization. This happens because of the current state of the warehouse as some of the cells that are closer to the shopping area are occupied with non-fixed products. As they keep utilizing the algorithm, these cells will be emptied and recently received fixed product pallets can finally be addressed inside these cells which are closer to the shopping area. We will provide the list of fixed products and the total number of cells that will

be reserved for fixed products. However, the fixed product list and the number of cells that are going to be reserved for fixed products can be updated by the user whenever they need. They specifically requested such a flexible system so that they can adapt themselves quickly in cases of limited products, new releases, and campaigns when the related products' sales increase cause frequent material handling operations. Note that although the system has not been implemented yet our results with historical data demonstrate the improvement potential of the suggested approach over the current system.

7 Interface and Other Tools

HPAS algorithm requires inputs such as product's SKU, amount of pallets entering or leaving the warehouse at a certain time, and the operation type. There are two types of operations: single product operation and multiple product operation. First one is for addressing or changing the location of a single pallet. Second operation is for truck shipment with multiple pallets, so that the user can enter or remove multiple pallets at once. These operation interfaces are available in the Appendix C. The interface is designed considering potential errors users can make and guides them to prevent mistakes. For instance, entering an SKU that does not exist or wanting to remove a product more than its available amount in the storage area. In such cases, the system gives feedback to the user. The algorithm is a Python script however, Eti is not familiar with Python tools and environments. We created a VBA Macro to run the Python script from the Excel. The interface was prepared via VBA. An Excel VBA macro executes the Python script in the back and the user will only see the output on the newly created Excel Sheet.

8 Expected Contributions and Conclusion

The main contribution of the project is to decrease the total time spent for material handling operations inside warehouse. As mentioned before, HPAS algorithm has provided an approximate of 1.45% improvement on this metric when compared with the performance of the current system. On the other hand, our measurements are strictly limited to calculating total time spent on relevant material handling operations. Beyond that, we cannot confidently mention any other specific contribution. However, considering the operating logic and performance of HPAS, we can predict other potential outcomes. For instance, the decrease in total material handling time is highly correlated with the decrease in the total distances covered by the forklifts inside the warehouse. Hence, we can expect a decrease in the fuel consumption of the forklifts that operate inside the warehouse and anticipate a potential economic benefit. Moreover, there can also be several non-quantitative benefits such as improved customer relations and considering that the decrease in material handling time can result in quicker deliveries.

References

- J. D. C. Little and S. C. Graves. *Little's Law*, pages 81–100. Springer US, Boston, MA, 2008. ISBN 978-0-387-73699-0. doi: 10.1007/978-0-387-73699-0_5. URL https://doi.org/10.1007/978-0-387-73699-0_5.
- R. Manzini. *Warehousing in the Global Supply Chain*. 01 2012. ISBN 978-1-4471-2273-9. doi: 10.1007/978-1-4471-2274-6.
- B. Park. *Order Picking: Issues, Systems and Models*, pages 1–30. 01 2012. ISBN 978-1-4471-2273-9. doi: 10.1007/978-1-4471-2274-6_1.
- S. Sharma and B. Shah. A proposed hybrid storage assignment framework: A case study. *International Journal of Productivity and Performance Management*, 64: 870–892, 07 2015. doi: 10.1108/IJPPM-04-2014-0053.

Appendix A Experiment

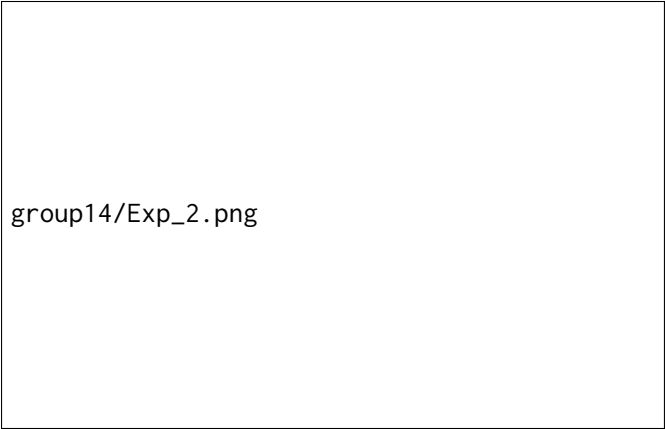


Figure 3: Experiment results regarding cell reservation for Type B fixed products.

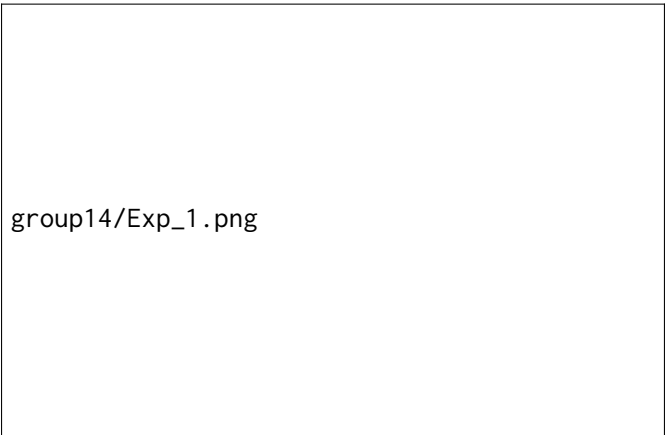


Figure 4: Experiment results regarding cell reservation for Type A fixed products.

Appendix B Algorithm

B.1 Sets and Parameters

Sets

F : set of SKUs of the products that will have fixed addresses, where F_i denotes the i^{th} SKU of the i^{th} fixed product, recall that the products are already sorted with respect to their sales rates in the descending order

C : set that denotes the products inside each storage cell, where C_i denotes the SKU of the product that is stored in the i^{th} cell, C_i is zero if that i^{th} cell is not assigned to any product, note that the cells are listed in the ascending order with respect to their time values, i.e., the first cell in the list ($i=1$) is the most desired one and the last cell in this list is the last desired one

K : Set that denotes whether a cell is reserved for fixed products or not, where K_i is a binary variable and $K_i = 1$ if the cell is reserved for fixed products, $K_i = 0$ otherwise

Parameters

tc : total number of cells in the product-type-based storage area

f_p : the total number of cells that are going to be allocated for product p given that product p is included in the list of fixed products, i.e., F

n : number pallets that are going to be addressed inside the cells or taken out from the cells

B.2 Pseudo-codes

Input: F, f_p, tc, K, C

```
for  $p$  in  $F$  do
   $f = f_p$  and  $cell = 0$  and  $count = 0$ 
  while  $count < f$  and  $cell < tc$  do
     $IsFixed = False$  while  $IsFixed = False$  and  $cell < tc$  do
      if  $K_{cell} = 0$  then
        |  $IsFixed = True$ 
      else
        |  $cell = cell + 1$ 
      end if
       $K_{cell} = 1$ 
    end while
    end while
  end while
end for
```

Algorithm 1: Reserve Fixed Cells

Input: p, n, C, K, tc
 $cell = 0$ and $count = 0$ and $same = False$;
while $count < n$ and $cell < tc$ **do**
 while $same = False$ and $cell < tc$ **do**
 if $C_{cell} = p$ and $K_{cell} = 1$ **then**
 | $same = True$
 else
 | $cell = cell + 1$
 end if
 if $same = True$ **then**
 | $C_{cell} = 0$
 | $same = False$
 | $cell = cell + 1$
 | $count = count + 1$
 end if
 end while
end while
 $cell = 0$ and $same = False$
while $count < n$ and $cell < tc$ **do**
 while $same = False$ and $cell < tc$ **do**
 if $C_{cell} = p$ **then**
 | $same = True$
 else
 | $cell = cell + 1$
 end if
 if $same = True$ **then**
 | $C_{cell} = 0$
 | $same = False$
 | $cell = cell + 1$
 | $count = count + 1$
 end if
 end while
end while

Algorithm 2: *Pallet Removal*

```

Input:  $F, p, n, C, K, tc$ 
 $cell = 0$  and  $count = 0$ 
while  $count < n$  and  $cell < tc$  and  $p$  in  $F$  do
     $empty = False$ ;
    while  $empty = False$  and  $cell < tc$  do
        if  $C_{cell} = 0$  and  $K_{cell} = 1$  then
             $empty = True$ 
        else
             $cell = cell + 1$ 
        end if
        if  $empty = True$  then
             $C_{cell} = p$ 
             $empty = False$ 
             $cell = cell + 1$ 
             $count = count + 1$ 
        end if
    end while
end while
 $cell = 0$ 
while  $count < n$  and  $cell < tc$  do
     $empty = False$ 
    while  $empty = False$  and  $cell < tc$  do
        if  $C_{cell} = 0$  and  $K_{cell} = 0$  then
             $empty = True$ 
        else
             $cell = cell + 1$ 
        end if
        if  $empty = True$  then
             $C_{cell} = p$ 
             $empty = False$ 
             $cell = cell + 1$ 
             $count = count + 1$ 
        end if
    end while
end while

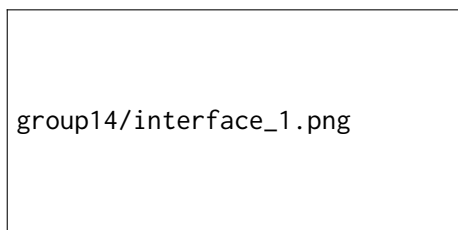
```

Algorithm 3: *Pallet Addressing*

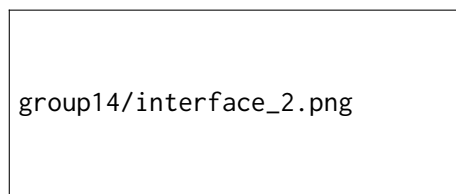
Appendix C User Interface

group14/interface_3.png

Figure 5: Snapshot of the main user interface.



(a) Single pallet operation



(b) Multiple pallet operation

Figure 6: Snapshots of the user forms.

Tahminleme Sistemi ile Hatalı Pişme Sayısının Azaltılması

**Brisa Bridgestone Sabancı Lastik Sanayi ve Ticaret
A.Ş.**

group15/group_photo.png

Proje Ekibi

Merve Ayyıldız, Hümeysra Buğçe Bali, Kağan Çalikoğlu,
Yunus Doğukan Çalışkan, Muhammed Ferman Dağ, Berfin Korhan, Erdil Şen

Şirket Danışmanı

Emre Vardar
Pişirme Geliştirme Mühendisi
BRISA Bridgestone
Endüstri Mühendisliği Takımı

Akademik Danışman

Yrd. Doç. Dr. Çağın Ararat
Bilkent Üniversitesi
Endüstri Mühendisliği Bölümü

ÖZET

Brisa Bridgestone Sabancı Lastik Sanayi ve Ticaret A.Ş.'ye ait Kocaeli fabrikasında lastik üretimi sırasında platen kaynaklı olarak gerçekleşen hatalı pişmeler, firmanın mevcut sistemleriyle bu hatalı pişmeler gerçekleşmeden önce tespit edilememektedir. Bu tespitin yapılamaması sonucunda hurda lastik sayısı artmakta; bu hatalı pişmeyi değerlendirme süreci ise iş gücü kaybına yol açmakta, maliyetleri arttırmakta ve sürdürülebilirlik açısından olumsuz bir etki yaratmaktadır. Tahminleme sistemi ile hatalı pişme sayısının azaltılması projesinin amacı, yaklaşmakta olan hatalı pişmeyi henüz bu hatalı pişme gerçekleşmeden öngörecektir bir alarm sistemi oluşturmaktır. Bu doğrultuda pişme işleminin gerçekleştiği platen sıcaklık verilerini kullanarak gerçek zamanlı çalışan ve veriden öğrenen bir tahminleme algoritması geliştirildi. Bu algoritma veri üzerinde değişim noktası tespit eden bir başka istatistiksel yöntem ile ikili kontrol yaparak hibrit bir sistem oluşturacak şekilde tasarlandı. Oluşturulan sistem arayüzü ile fabrikada uygulanmaya uygun hale getirildi.

Anahtar Kelimeler: Hatalı pişme, tahminleme, erken uyarı sistemi

Miscure Reduction by Early Prediction Systems

1 Introduction

The Bridgestone Corporation commenced service in Turkey with an equal share partnership with Sabancı Group on November 1, 1988. As a result of this partnership, Bridgestone's name was changed to Brisa Bridgestone Sabancı Tire Industry and Trade Inc. In 1990, Brisa started serial production in its Kocaeli factory. The mission of Brisa is to offer high-quality values to public with sustainable improvements. Brisa has seven significant values for its business: work safety, sustainability, innovation, customer centricity, team spirit, and work perfection. Today, Brisa is the pioneer tire manufacturer in Turkey and saves its place as the sixth in Europe. Although the main product of Brisa is tire, the company has a wide range of product scales such as industrial rubber, golf balls, tennis products, and bicycle equipment. Also, the tire products of Brisa consist of various types such as car, truck, bus, dozer and crane tires. In Brisa Bridgestone, there are actively working 3,112 employees. 2,930 of these employees work in the factories of Brisa, which are located in Kocaeli and Aksaray. Kocaeli factory has the denotation of one of the biggest tire factories in the world with 361,000 m² closed area. In this project, we carry out our work to be implemented in the Kocaeli factory.

2 Process Analysis and Current System

The curing process is one of the most significant tire manufacturing steps in which the final shape, thread pattern, and desired physical properties are given to the tire (Appendix A). The main parameters of the curing process are temperature, time, and pressure. The values of these parameters are predetermined for each tire type to gain the physical properties needed by the tire. The curing process occurs inside the presses. They have three heat sources: platens (top and bottom), bladder, and jacket. Hot water steam is supplied to the tire at the determined pressures from these sources to keep the main parameters within the desired tolerances in terms of temperature. There are also tolerances defined for each parameter, and these tolerances have upper and lower limits. If the tires are produced within the desired tolerances, then the process is called the "optimum cure". When the parameter values are below the specified tolerances, cross-linking does not fully engage and this situation is named as the "undercure". When the parameter values are above the specified tolerance, the desired physical properties start to be gained in the opposite direction and this is named as the "overcure". In total, these two cases are termed as "miscure". Besides, miscure also occurs when the tire is exposed to a temperature that exceeds tolerances for more than a certain time. Additionally, the term "cycle" is used in order to define the time interval in which a tire is being cured by the presses. The starting and ending time of a cycle can be understood from the pressure's fluctuations. In order to separate the cycles and idle times, we have used the pressure data as well.

Currently, to understand the problem related to miscures, there are several dif-

ferent alarm systems assisting the presses within the factory. Each of these alarm gives warning only after a miscure occurs. Since the present system provides a reactive solution, it becomes too late to rescue the tire before it is deemed miscure. Our objective is to create a model that gives the level of miscure risk for future tires in advance by examining the data in the presses. This proactive approach will eventually help reducing the tire scraps that stem from miscures while gaining the previously-wasted work-hours.

3 Problem Definition

Miscure is a problem that does not add value in production and is desired to be eliminated because it creates both tire waste and requires a considerable amount of labor time. When a miscure is detected, the tire entails many processes such as loading, handling and inspection to determine whether the tire is scrap or not. A detailed analysis is made to see how important the time wasted with this miscure judgement process (L. Li and Ni, 2009). The distribution of the miscured products as "scrap" or as regular tires can be seen in Appendix B. It is concluded that a remarkable amount of labor time is wasted due to miscure problems. Secondly, another detailed analysis is carried out to see how we can reclaim the wasted time on the evaluation process by converting that into production. We have calculated how many tires would be produced additionally, and what the contribution of this would be. Calculations have shown that this waste (both for the scrap tires and the man-power perspectives) actually causes great damage to the company. Hence, the analyses emphasize the significance of having an early miscure detection system. The detailed information about the loss is not stated specifically due to the confidentiality.

The variety of factors that cause miscures makes the problem more complicated. Miscures, which arise from platens create approximately 25% of the total reasons alone, highlighting the fact that platens are the major cause of miscures. Therefore, it is more logical to focus on platen miscures at the first stage, so it is determined as the scope of the problem. Additionally, when other causes of miscures are examined, some reasons contain human factors which make them difficult to intervene. Instead of dealing with them, working on the platen miscures is more preferable since they are more mechanically-sourced. As platen is one of the heat sources in the presses, it causes temperature changes. These temperature change is the main factor, resulting in the temperature parameter to go out of tolerance in the curing process. Hence, the main parameter that we need to work on is temperature. We can also say that our problem is statistically analyzable since the presses have a system that measures and records the platen temperature between equally divided time intervals. Thus, working on the platen miscures will ensure a significant improvement for the factory and provide crucial information against overall miscure problems.

Platen-sourced miscures have three different miscure types which are termed as CA, CP and Acute. CA miscure occurs when the temperature is close to the upper

or lower tolerance values, CP miscure occurs when the temperature fluctuates between the upper and lower tolerance values, and Acute miscure occurs as a result of sudden deviations in temperature. While 65% of platen-caused miscures are CA and CP type miscures that we aim to detect, 35% of them are caused by Acute type miscure.

4 Solution Approach

4.1 Apriori and Double Exponential Smoothing

To convey the idea of being proactive, we developed a forecasting model that works with the historical data (P. Prithiviraj, 2015). The algorithm observes the frequencies of cycles that precede miscure formations. By establishing if-then relations between cycles' mean temperature differences and the desired set values, the probability of leading to a miscure is aimed to be predicted (Montgomery, 2013). The forecasting of the next cycle's mean temperature value is calculated with the double exponential smoothing method (Nahmias, 2005). The historical data feeding the algorithm are taken from each press individually. When the algorithm receives the training data, it first identifies the cycle means in terms of bins. We conduct this additional step to generalize the effects of specific data points. For example, if 78.9% is the cycle mean index of a miscure-causing cycle, then it is classified as "B5", which corresponds to index of cycle means between 75% and 90%. Thus, the prediction for the upcoming cycles does not have to be exactly equal to the value in the training data, but it just needs to be within the range of that index bin to be alerted by the association. Appendix C illustrates the categorization of means according to bin ranges.

By the means of this process, the bins are divided into lists containing up to five items to control the relationship of a possible reoccurring trend (Toivonen, 1996). The Apriori algorithm operates to find association rules for these lists of bins (Agrawal, 1993). The rules convey the risk of a bin that might cause a miscure in the following cycles. At the end of the training process, the intake for the live data starts. Cycle means are calculated at the end of each cycle. Accordingly, a forecast for the next cycle is calculated and then categorized in terms of bins. If the bin of the forecasted value is compatible with the predetermined rules, then an alert pops up to warn about the miscure probability. Since we make a forecast about the next cycle and since the rule warns about the miscure-causing bins, we alert the miscure formation possibility at least two cycles before. Therefore, this algorithm can be defined as a way of a self-learning mechanism. Yet, if there becomes a failure to notice, then the algorithm re-runs itself with the new data including the uncaptured miscure to come up with new rules. In this way, even if the data characteristics change in time, the algorithm stays up to date. This method was initially written in Python to observe its performance. Then the transformation of it to an executable file (.exe) is completed.

4.2 Two-Sample t-statistics

For our alarm system that aims to predict miscues before they happen, we worked with a change point detection method in addition to the association rule mining (Apriori Algorithm and Double Exponential Smoothing), which is based on the Two-sample t-statistics. The method of two-sample t-statistics is one of the most useful variations of the change point models when both μ and σ parameters are unknown. Also, this method is able to detect the shifts in real time data sets (D. M. Hawkins, 2003). We constructed our model with the assumption that none of the parameters is known and benefited from "two-sample t-statistic" between the left and right sections of the sequence, maximized across all possible change points. As data, we used the increments between platen temperatures measured in equally-divided time intervals to be independently distributed and normality assumption is reasonable. Afterwards, the model will show the t-statistic value at a specific time by comparing the means of two different samples. Hence, by checking the critical value against t-statistic, we can make inference to determine whether there is a change in the process or not. If the critical value is exceeded, it can be concluded that there has been a change. See the Appendix D for better understanding.

In this approach, we distinguished the observed variations in data as positive and negative changes. A change can be classified as a negative change if temperature values of platen gets closer to the tolerance levels. However, positive change is defined as a change of observed variation towards the desired set temperatures values. Therefore, dangerous t-test alarms are the negative change alarms. Also, we aim to demonstrate the acceptability of data for our approaches, regarding the normality of platen temperature increments. For this purpose, probability plots and statistical tests are used in the R-Language. The selected two statistical tests, Shapiro - Wilk Test and Kolmogorov - Smirnov Test, are used to ensure the normal distribution (Oztuna D, 2006; HJ., 2002; DJ., 2007). According to the Shapiro - Wilk test, we found a p-value that is greater than 0.05 ($0.63 > 0.05$) and the data is accepted as normally distributed. Also, the Kolmogorov - Smirnov test accepted that the increments are normally distributed since p-value is greater than 0.05 ($0.85 > 0.05$).

As mentioned before, we used the increments between platen temperatures. While T_i is the i^{th} platen temperature measured in equally-divided time intervals, we indicated Y_i as the i^{th} increment: $Y_i = T_{i+1} - T_i$. For a given change point $j \in \{1, 2, \dots, n-1\}$, $\bar{Y}_{jn}$ is the sample mean of the first j observations, and $\bar{Y}_{jn}^*$ is the sample mean of the remaining $(n-j)$ observations. $V_{jn} = \sum_{i=1}^j (Y_i - \bar{Y}_{jn})^2 + \sum_{i=j+1}^n (Y_i - \bar{Y}_{jn}^*)^2$ is the residual sum of squares. $\hat{\sigma}_{jn}^2 = V_{jn}/(n-2)$ is the usual pooled estimator of σ^2 . Finally, a conventional two-sample t-statistic for comparing the two segments is: $T_{jn} = \sqrt{\frac{j(n-j)}{n}} \frac{\bar{Y}_{jn} - \bar{Y}_{jn}^*}{\hat{\sigma}_{jn}}$ (D. M. Hawkins, 2003).

Pseudocode for two sample t-statistic can be found below.

initialization

calculate increments between platen temperatures measured every
equally-divided time and store them in a Dataset

for each data in Dataset **do** calculate Y_{jn} (the mean of the first j
observations)

calculate Y_{jn}^* (the mean of the remaining $(n - j)$ observations)

calculate $\hat{\sigma}_{jn}^2 = V_{jn}/(n - 2)$ (the usual pooled estimator of σ^2)

calculate $T_{jn} = \sqrt{\frac{j(n-j)}{n} \frac{\bar{Y}_{jn} - \bar{Y}_{jn}^*}{\hat{\sigma}_{jn}^2}}$

if $T_{max,n} \geq h_n$ **then** the change is detected **end if**

end end for

next

Algorithm 4: Pseudocode for Two sample t-statistic

The generalized likelihood ratio test for the presence of a change point consists of finding $T_{max,n} = \max_{1 \leq j \leq n-1} |T_{j,n}|$, $j^* \in \arg \max_{1 \leq j \leq n-1} |T_{j,n}|$.

The j^* gives the MLE of the true change point, and the corresponding $\bar{Y}_{j^*}$ and $\bar{Y}_{j^*}^*$ are MLE's of the unknown means μ_1 and μ_2 . If $T_{max,n}$ exceeds some critical value h_n which is calculated as:

$$h_{n,\alpha} = h_{10,\alpha} (0.677 + 0.019 \ln \alpha + (1 - 0.0115 \ln \alpha)/(n-6))$$

where $\ln ()$ is the natural log function.

After that, we conclude that there was indeed a shift. Otherwise we state that there is not sufficient evidence of a shift.

4.3 Proposed Model: Hybrid System

The final model is designed as a combination of the two methodologies described above. While Apriori Algorithm and Double Exponential Smoothing give a prediction for the upcoming cycles related to being miscure or not, the second approach is focused only to detect the changes by showing how far the observed values are from the set values, and it does not directly alert for the probability of a miscure formation. These two methods are integrated with a robust, four-level algorithm. The integrated system would create the hybrid methodology where the number of false alarms would be significantly reduced while the detection of miscures would increase. These levels are described with colors. To visualize the system, Figure 1 will help demonstrating the weights of the methods.

There are four outcomes which are likely to happen in this system:

1. Neither of the methods detect anything and the system continues working. Then, there would be no color-warning.
2. There would not be any alert from the Apriori + Forecasting Algorithm after the detection of change by the t-statistic method. Thus, this is an example of low miscure risk so that the blue color would turn on.



Figure 1: The Hybrid Methodology

3. Only Apriori + Forecasting Algorithm gives an alert about the miscure even if no change is detected. Since Apriori + Forecasting Algorithm is more direct for detecting miscures, its warnings are expected to have more weights. Therefore, an orange color would turn on.
4. Both algorithms give an alert about the possible miscure formation. This basically means that there is an obvious problem detected by both of our algorithms. Hence, the risk is higher than usual. At this point, the red color would turn on.

4.4 Validation

We used randomly chosen presses from different production lines to test our model that we coded in Python. We tried to choose long-range data to be able to observe the number of false alarms. After our model gives an alarm, we expect to observe a miscure within two cycles. Despite, there may be some cases that miscure does not occur within two cycles but occurs close to two hours. Therefore, we have categorized them as "false / early alarms". Table 1 displays the performance of the hybrid model on different tests.

Results show that our success rate is $14/39 = 35.9\%$ which means that we are catching 35.9% of the upcoming miscures at least two cycles before they occur. The correct alarm rate is $14/26 = 53.84\%$ and it shows that 53.84% of the alarms that our model gives are correct. In terms of the trust in the model, this result is found to be satisfactory.

4.5 Discussion

In order to get more beneficial results from the hybrid model, there are some diagnostics of our model so that for the future plans of the company and the proposed system, they might be significant to consider. While achieving promising results in general for Apriori model, the model has not been able to achieve the desired efficiency on acute type miscures. Therefore, the first goal can be

Tests	Test 1	Test 2	Test 3	Test 4	Test 5	Test 6
Duration	6 months	3 months	3 months	3 months	1 month	1 week
Total Miscure	32	1	0	0	1	5
Total Alarm	14	1	1	0	1	9
Correct Alarm	9	1	0	0	1	3
False Alarm	5	0	1	0	0	6
Uncaught	23	0	0	0	0	2
Choice	Random	Random	Random	Random	Planned	Planned

Table 1: Test Results

settled as increasing this efficiency. Additionally, if the model succeeds on platens, the system will be made applicable for another heat source, Jacket. With that purpose, jacket sourced miscures can be also eliminated. Finally, by analyzing the current situation of input data (temperatures received from platens), we have seen that there are lots of different characteristics of miscures. Therefore, it would be more useful to categorize the presses according to their behaviors through time. Some factors having impact on the different behaviors of different presses can be:

- The age of the press
- The change frequency of the cured tire type
- Bladder
- Type of the tire being cured

These and other factors may cause the different patterns and the miscure preventative solutions should be derived accordingly for better results. In Appendix E, some behavior types of the presses can be seen.

5 Implementation

Since the company has many presses that contain two platens that run simultaneously, our model is expected to work for all of the presses at the same time while conducting the calculations for both platens as one. Therefore, we further evolved our model as an independent .exe that can intake multiple presses as a single input. To open the file, one needs to log in to the system with a pre-assigned username and password due to confidentiality reasons. As a group, we tried to

create a user interface as convenient as possible for the workers on the field. By simply pushing the "run" button, the system works with the ability of saving the data history including the alarms and the allocated time slots without any outside help. Also, workers do not need to check the system constantly to see if there is a miscure coming. Instead, our application sends an e-mail an adequate time before, and the e-mail contains the information of the alarms' time slots alongside the graph of the last 15 cycles in order to observe the trend that led our program to notify an alarm (See Appendix F).

Our system is designed to work non-stop without losing any execution performance while continuously learning about the environment. In other words, as far as the data comes from SQL without any interruption, our model will catch miscures. In order to work with the live data coming from the presses, we programmed our .exe file to work in synchronize with the company's database. As the sensors within the presses read temperature values and write them onto the SabancıDX database, our program reads necessary inputs, such as temperature value, set temperature and pressure value for every equally divided time intervals. These necessary inputs are then passed to our algorithm to continue calculations. The Python code for reading SQL database can be found in Appendix G.

6 Outcome and Deliverables

6.1 Deliverables

While implementing the system, there are some key points that are taken into account. First, the desired system's alarm numbers should not be unrealistically large. It may cause stopping the production line a lot and reduce the production volume of the company. Moreover, if it is observed that the alarms that we give are actually not successful overall, they may ignore the successful alarms as well. Therefore, the false alarm rate is an important measure to detect the system. Second, after the alarms, time loss should not exceed the current time loss spent while detecting the miscures and judging them to determine whether they are scrap. Thus, the production should start fast after the alarms. Lastly, the past data is extremely important to measure the desired system. The new system's alarms and accuracy will be used on the company's past data and compared to actual results. For delivering of the system, the following tasks have been also completed:

- We have made it possible to have input files as SQL documents like the company database is restored. This will ease the transaction from installation to implementation.
- We have written the executable file (.exe) of the codes so as to make it easier to work with.
- One of the possible situations in which our code had inaccuracy was the case of having multiple miscure alarms at the same time for different presses.

This was the hardest part for us because we had coded and tested the methods for single press cases. Therefore, we had made our system capable of controlling all of the presses (almost 300 presses) in the company by adjusting our code.

- In order to develop better calculations about the improvements, we have prepared a simulation. The settings and results are stated below.

6.2 Simulation of the System

To fully present the results of our model in the manufacturing field, we decided to create a simulation model to come up with answers for the unanswered performance measures (Appendix H). Accordingly, it can be clearly seen that the total cycles within the same time interval drastically changes when the model detects the miscures due to change in inspection time. This is not only an opportunity to produce extra tires but also a way to use the resources more effectively. E-mailing the last 15 cycle's means in the form of a line graph every time a miscure is detected increases the mobility of a worker to be more flexible with the presses. In other words, no one needs to waste their times checking the mailed presses when the hybrid model alarms but instead, it is planned to be enough to check their mails and act accordingly if needed. Unless it is needed, then it is a false alarm. Talking in terms of the performance indicators mentioned above:

- The number of stops of the production line due to miscure alarms is less than the current system. It is thanks to our mailing system. There should be no need to stop the line unless it is needed.
- The miscure judgement time drastically decreases thanks to the correct alarms. Unfortunately, false alarms do consume time even though it is quite small.

6.3 Benefits to the Company

As mentioned in the project description, the main expected outcome of the Brisa Project is the 10% enhancement in total miscure tires. The project's expected side impacts can be promoted as improved time-efficiency, workforce allocation, financial outcome, and sustainable environmental solutions in general. The solution will bring benefits to the company in three terms:

- Miscure stops need lots of long processes by engineers and operators. As a result of the quality examination processes to identify miscure tires, operators spend a significant amount of time. With this project, at least 10% improvement is expected on time management.
- In the mid term, the saved time from the miscure judgement process is expected to be utilized to produce more tires. This would naturally increase the profit of the company due to manufacturing and selling more tires .
- Talking about the long-term impact, it is related to the environmental issues

of miscure tires. The rapidly growing waste tire stocks are becoming a major environmental problem in our country as in the world and we expect that this project will decrease the number of waste tires by decreasing the number of miscures (Demir, 2015; T24, 2008).

In conclusion, we have identified three key performance indicators related to our alarm system besides from the financial and work-power improvements stated above:

1. The Number of Stops of the Production Line due to Miscure Alarms

Suppose we assume that the company stops the production for each alarm of ours, according to Table 1. In that case, we can compare the current system and the proposed model. Previously, they had to stop for each miscure alarm (39 times), while we suggest stopping the production when we give proactive alarms (26 times). Therefore, there are $(39 - 26)/39 = 33.3\%$ less stops. Simulation results also supported this improvement. The impact of mailing is important.

2. The Miscure Judgement Time

In terms of work-power, using the proposed methodology will save vital amount of minutes of labor time which could be utilized in other problems.

3. The Accuracy of the Alarms

According to Table 1, total correct alarm number is 14 while total alarm number is 26. Proportion of $14/26$ gives us $\approx 54\%$ accuracy which is much more than 10% and for a proactive method, quite promising.

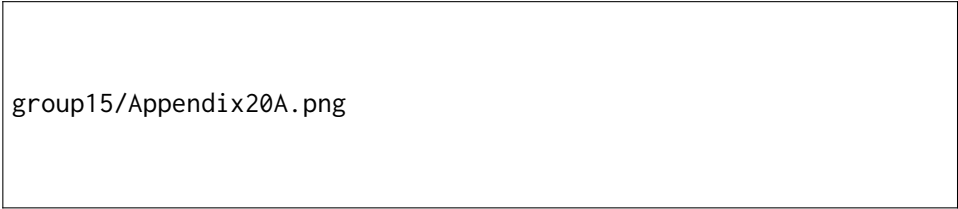
References

- T. I. A. S. Agrawal, Rakesh. Mining association rules between sets of items in large databases. *Acm Sigmod*, pages 205–275, 1993.
- C. W. K. D. M. Hawkins, P. Qiu. The change point model for statistical process control. *Journal of Quality Technology* vol. 35, no. 4, pages 355–366, 2003.
- P. T. Demir. Araç lastiklerinin geri dönüşümü Üzerine, 2015. <http://www.plastik-ambalaj.com/tr/plastik-ambalaj-makale/2452-arac-lastiklerinin-geri-doenuesuemue-uezerine-bir-derleme>.
- S. DJ. A cautionary note on the use of the kolmogorov-smirnov test for normality. *American Meteor Soc.* vol. 135, page 135, 2007.
- T. HJ. *Testing for normality*. Marcel Dekker, New York, 2002.
- Q. C. L. Li and J. Ni. Real time production improvement through bottleneck control. *International Journal of Production Research*, 2009. doi: 10.1080/00207540802244240.

- D. Montgomery. *Statistical Quality Control: A Modern Introduction*. Wiley, Hoboken, NJ, 2013.
- S. Nahmias. *Production and Operations Analysis*. McGraw-Hill Irwin, Boston, Mass, 2005.
- T. E. Oztuna D, Elhan AH. Investigation of four different normality tests in terms of type 1 error rate and power under different distributions. *Turkish Journal of Medical Sciences* vol. 36, no. 3, page 171, 2006.
- R. P. P. Prithiviraj. A comparative analysis of association rule mining algorithms in data mining: A study. *American Journal of Computer Science and Engineering Survey*, pages 122–144, 2015.
- T24. Hurda Lastikle Gelen Çevre Felaketi, 2008. <https://t24.com.tr/haber/hurda-lastikle-gelen-cevre-felaketi>, 22651.
- H. Toivonen. Sampling large databases for association rules. *The VLDB Journal*, pages 134–145, 1996.

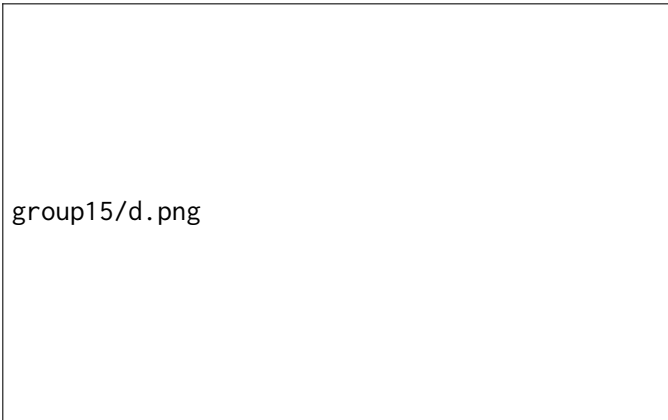
7 Appendix

Appendix A Tire production



group15/Appendix20A.png

Appendix B Miscure scrap rate

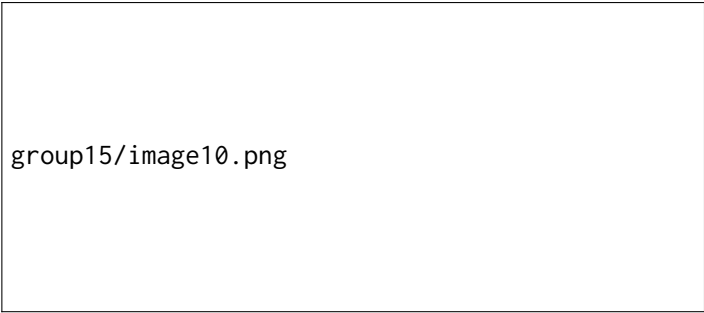


group15/d.png

Appendix C Cycle bins for Apriori Algorithm + Double Exponential Smoothing

Lower Bound	Upper Bound (including)	Bin
$-\infty$	-100%	-B7
-100%	-90%	-B6
-90%	-75%	-B5
-75%	-65%	-B4
-65%	-50%	-B3
-50%	-25%	-B2
-25%	0%	-B1
0%	25%	B1
25%	50%	B2
50%	65%	B3
65%	75%	B4
75%	90%	B5
90%	100%	B6
100%	∞	B7

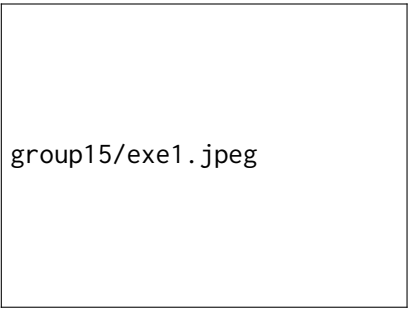
Appendix D Schematic representation of inputs and outputs of the two sample t-statistic method



Appendix E Examples of different press behaviors



Appendix F Userface and output of .exe





group15/exe2.jpg



group15/mail.JPG

Appendix G Python code for reading data from SQL server

Appendix H Simulation

group15/sim1.png

İ.D. Bilkent Üniversitesi Personel Servis Güzergahlarının Modellenmesi ve Optimizasyonu

İ.D. Bilkent Üniversitesi

group16/group16_group_photo.png

Proje Ekibi

Gökhan Akkaya, Aslı Sena Akyurt, Mustafa Kaan Çoban,
Hasan Kağan Gürbüz, Ezgi Can Kılıç, Batuhan Tavşan

Şirket Danışmanı

Tuncay İbiş
İ.D Bilkent Üniversitesi Genel Sekreter
Yardımcısı
Endüstri Mühendisliği Takımı

Akademik Danışman

Doktora Adayı Milad MalekiPirbazari
Endüstri Mühendisliği Bölümü

ÖZET

İ.D. Bilkent Üniversitesi Ulaşım Birimi sürücü sezgilerine ve tecrübelerine dayalı servis güzergahlarına sahiptir ve rotalama için bilimsel araç/bilgisayar sistemi kullanılmamaktadır. Dolayısıyla, mevcut uygulamada servis sürelerinin beklenenden daha fazla zaman aldığını varsayılabilir. Ayrıca pandemi nedeniyle belirsizliğin daha fazla olduğu günümüzde, planlama daha da zorlaşmaktadır. Bu nedenle, bu proje matematiksel modelleme ve modelin doğrulanması ile optimuma yakın çözümler sağlamak için hızlı ve güvenilir bir sistem geliştirmeyi amaçlamaktadır.

Anahtar Kelimeler: Okul Servisi Rotalama Problemi, matematiksel optimizasyon modeli, performans iyileştirme

Modeling and Optimization of Personnel Service Routes for İ.D. Bilkent University

1 Description of the System

The project that we carry out is “Modeling and Optimization of Personnel Service Routes for İ.D. Bilkent University”. These transportation services are provided by Bilkent Transportation Unit. The main purpose of the services provided by the Transportation Unit is to transport Bilkent University students, academic and administrative staff, between the university and the city. Considering the current transportation system, there are five different transportation programs, namely transportation programs of Main Campus and East Campus, Main Campus-East Campus Ring, Student Neighborhood Shuttle, and Academic-Administrative Staff District Shuttle. Main Campus as well as East Campus transportation programs provide transportation services between the university and Tunus/Sihhiye stops at certain hours of the day. Student Neighborhood Shuttle and Academic-Administrative Staff District Shuttle provide transportation to academic staff and students to various parts of the city. In addition to these, the Main-East Campus Ring program provides inter-campus transportation within Bilkent University. While providing all these transportation systems, the Transportation Unit has to prioritize its own economic benefits. Thus, all routes and number of buses should be determined in the most reasonable way to meet the economic concerns.

The Academic-Administrative Staff District Shuttle program provides transportation services for Bilkent personnel. The service in the morning takes them from their districts to the university and the afternoon service takes them from the university to the same district. There are usually 37 personnel buses available for transporting between 1100 and 1200 people, but as the number of active personnel is lower due to the pandemic (nearly 400), there are currently only 25 buses in operation. There are three types of vehicles, namely minibuses, midibuses, and buses. The routing is planned by considering the distance between the stops on the main streets. Moreover, as stated by the Transportation Unit, sometimes the drivers decide on which routes to take based on their experience. The Transportation Unit pays the daily cost of 543 TL+VAT, 303 TL+VAT, and 223 TL+VAT for each bus, midibus, and minibus, respectively.

Each morning, after picking up the personnel, all buses are required to enter the university within a 15 minutes interval, that is, they need to arrive at the university between 7:45-8:00 a.m. Moreover, in the afternoon, on the way back from school, all buses need to take the personnel and leave the university at 5:45 p.m.

Currently, there is no mobile application that matches the personnel to the stops, but there is an ID card reader on vehicles. Moreover, the vehicles have a GPRS-based tracking system. Thus, when the card is read, it is possible to determine the number of people who get on the bus from each stop. This system is especially

significant for the programs such as the Main Campus Transportation program. However, regarding the Academic-Administrative Staff District Shuttle program, the Transportation Unit has information on the number and the address of all personnel. Therefore, the tracking system may be used to access the amount of time each personnel spends on the bus.

2 Problem Definition

As mentioned earlier, the transportation system of Bilkent University has five different transportation programs. Since all five problems are interconnected, any improvement in one program affects the others. We chose a program which is used by the same personnel in the morning and the afternoon. These personnel use the buses every weekday, so the data is more stable compared to other programs. Academic-Administrative Staff District Shuttle is the most reviewable program to optimize and would result in improvements in the whole system.

Bilkent University Transportation Unit is currently planning the bus routes along with assigning passengers to specific buses just based on intuition and so far, no scientific way or tool has been used. So, we can assume that in the current practice, the travels take more time than expected. Furthermore, what makes the scheduling even more challenging is the need to adapt to rapidly changing conditions, especially in the present days with more uncertainty due to the pandemic. Therefore, in this project, we aim to develop a fast and reliable system to provide near-to-optimal solutions to the unit, based on mathematical programming. In this regard, we aimed to improve the efficiency of these schedules from different perspectives.

Firstly, since the number of vehicles to be used is affected by the service capacity, how many vehicles from each type is used is significant for the Transportation Unit. In general, we cannot always say that the number of vehicles used needs to be reduced by taking profit into consideration. However, the unit pays based on the number of buses used, not the total distance the buses travel, so the number of vehicles used should be kept at a minimum amount. Thus, reducing the existing number of vehicles may be needed.

Secondly, based on the start and end times of the current schedules, time spent on the buses may be a problem. Due to employing fewer numbers of buses than required or inefficient scheduling, some buses may use long routes which means some personnel need to stay longer in the buses until they arrive at their destinations. In that case, they have to start the day quite early to be able to catch the buses. In this case, the Transportation Unit encounters complaints of personnel due to early service hours. We find optimal routes such that it would also reduce the time spent on the bus, so personnel will get on the bus at the most convenient time which will probably be later than the current schedule.

Lastly, time between the first and the last stops are not determined in the current system. One of the routes can be seen in appendix A. Personnel need to wait at

the bus stop according to their previous observations and experience. Personnel might miss the vehicle or might wait a long time. This causes serious problems for the unit and the personnel.

3 Project Outcome and Deliverables

Our primary objective is to optimize the routing of the Academic-Administrative Staff District Shuttle program by minimizing the total cost corresponding to the number of buses used. Furthermore, we enforce the maximum time spent in the vehicles by each personnel to be less than a threshold in the constraint part of our model.

Our system provides an optimal solution which includes the number of buses and routes required. Assigning personnel to the specific bus stops that are convenient to them can also be provided. With this dynamic system, any changes regarding the buses or the number of personnel can be integrated into the system, and new optimal solutions will be obtained. For instance, due to the pandemic, most of the personnel do not use the buses. However, when the pandemic ends, the number of personnel that use the buses would increase and a new routing might be needed. With our system, when this situation occurs, the Transportation Unit can obtain the updated routes. The flowchart that explains the inputs and outputs of our model is shown in Figure 1.



Figure 1: Flowchart of the model's inputs and outputs

3.1 Mathematical Model

For the project, we describe the mathematical formulations of our SBRP that in some parts are similar to the formulation approach of Schittekat et al. (2006). In particular, we present two separate formulations for morning and afternoon pick-ups. The two models are very similar, but they have small differences. These differences can be found in Figure 2.

group16/group16_models_differences.png

Figure 2: Differences between morning and afternoon models

Mathematical model for morning pick-ups

Parameters and decision variables used in the morning pick-up formulation are reported in Table 1.

Table 1: Parameters and decision variables of the morning model

<i>Indices</i>		<i>Decision Variables</i>	
$b \in B, (i, j) \in A, i \in S$		$x_{i,j,b}$	$\begin{cases} 1 & \text{if bus } b \text{ traverses arc}(i,j) \\ 0 & \text{otherwise} \end{cases}$
<i>Parameters</i>		$y_{i,b}$	$\begin{cases} 1 & \text{if bus } b \text{ visits stop } i \\ 0 & \text{otherwise} \end{cases}$
S	Set of all potential stops (Stop 0 is the school)	$z_{i,p,b}$	$\begin{cases} 1 & \text{if personnel } p \text{ is picked up by bus } b \\ & \text{at stop } i \\ 0 & \text{otherwise} \end{cases}$
A	Set of all routes between all the stops	w_b	$\begin{cases} 1 & \text{if bus } b \text{ is used} \\ 0 & \text{otherwise} \end{cases}$
B	Set of available buses		
Cap_b	Capacity of bus b		
$t_{i,j}$	Time of traversing arc(i,j)		
c_b	Cost of using bus b		
MD	Maximum acceptable duration of each travel		
$q_{i,p}$	$\begin{cases} 1 & \text{if personnel } p \text{ can walk to stop } i \text{ in} \\ & \text{an acceptable amount of time} \\ 0 & \text{otherwise} \end{cases}$		

The mathematical model for morning pick-ups is provided below.

$$\min \sum_{b \in B} c_b w_b \quad (1)$$

$$\text{subject to: } \sum_{b \in B} y_{0,b} = \sum_{b \in B} w_b \quad (2)$$

$$\sum_{i \in S} y_{i,b} \leq w_b(|S| + 1) \quad \forall b \in B \quad (3)$$

$$\sum_{(i,j) \in A} x_{i,j,b} = \sum_{(i,j) \in A} x_{j,i,b} = y_{i,b} \quad \forall i \in S, \forall b \in B \quad (4)$$

$$\sum_{b=1}^B y_{i,b} \leq 1 \quad \forall i \in S \setminus \{0\} \quad (5)$$

$$\sum_{b=1}^B z_{i,p,b} \leq q_{i,p} \quad \forall p \in P, \forall i \in S \quad (6)$$

$$\sum_{i \in S} \sum_{p \in P} z_{i,p,b} \leq Cap_b \quad \forall b \in B \quad (7)$$

$$z_{i,p,b} \leq y_{i,b} \quad \forall i \in S, \forall p \in P, \forall b \in B \quad (8)$$

$$\sum_{i \in S} \sum_{b \in B} z_{i,p,b} = 1 \quad \forall p \in P \quad (9)$$

$$u_{i,b} - u_{j,b} + t_{i,j} \leq MD(1 - x_{i,j,b}) \quad \forall b \in B, \forall (i,j) \in A \quad (10)$$

$$0 \leq u_{i,b} \leq MD \quad \forall b \in B, \forall i \in S \quad (11)$$

$$u_{i,b} \geq 0 \quad \forall i \in S, \forall b \in B \quad (12)$$

$$x_{i,j,b} \in \{0, 1\} \quad \forall (i,j) \in A, \forall b \in B \quad (13)$$

$$y_{i,b} \in \{0, 1\} \quad \forall i \in S, \forall b \in B \quad (14)$$

$$z_{i,p,b} \in \{0, 1\} \quad \forall i \in S, \forall p \in P, \forall b \in B \quad (15)$$

$$w_b \in \{0, 1\} \quad \forall b \in B \quad (16)$$

The objective function (1) minimizes the total cost of using different buses for the whole morning travels. Constraint (2) ensures that all used buses come to school. Constraint (3) ensures that if bus b is not used no stops can be assigned to that bus. Constraint (4) ensures that if stop i is visited by bus b , then one arc should be traversed by bus b entering stop i and leaving stop i . Constraint (5) ensures that all stops are visited no more than once, except the stop corresponding to the school. Constraint (6) ensures that each personnel walks to a single stop they are allowed to walk to. Constraint (7) ensures that the capacity of buses is not exceeded. Constraint (8) ensures that personnel p is not picked up at stop i by bus b if bus b does not visit stop i . Constraint (9) ensures that all personnel are picked up once. Constraints (10) and (11) are Miller-Tucker-Zemlin constraints. Constraint (12) ensures that the time until arrival to school is more than or equal to zero. Constraints (13), (14), (15), (16) are integrality constraints. The earliest and the latest time that buses can enter the school were not used directly in the model. However, this information is considered in the coding part.

The table for parameters and decision variables used in the afternoon pick-up can be found in Appendix B. Moreover, the mathematical model for afternoon is provided in Appendix C.

3.2 Data and Solution Methods

Besides the academic resource, other resources are used in this project. These resources include the data provided by the Transportation Unit. To measure the distances between pick-up points, Python programming language, Google Maps and Google API were used. Since Google API has \$200 free usage monthly, it was enough for the project in order to

reach coordinates of the EGO stops and personnel addresses. The Transportation Unit did not need to do any extra payment since the group pays attention to use it in a way that does not exceed the free usage limit. Moreover, the group has informed the unit about the necessity of API for the progress of the project. Although the addresses of the personnel sent by Human Resources are missing, the coordinates of 178 addresses were transferred to a new Excel file. The step of getting the coordinates of the addresses can be seen in Figure 3. The Excel file where the coordinates are added is shown in Figure 4.

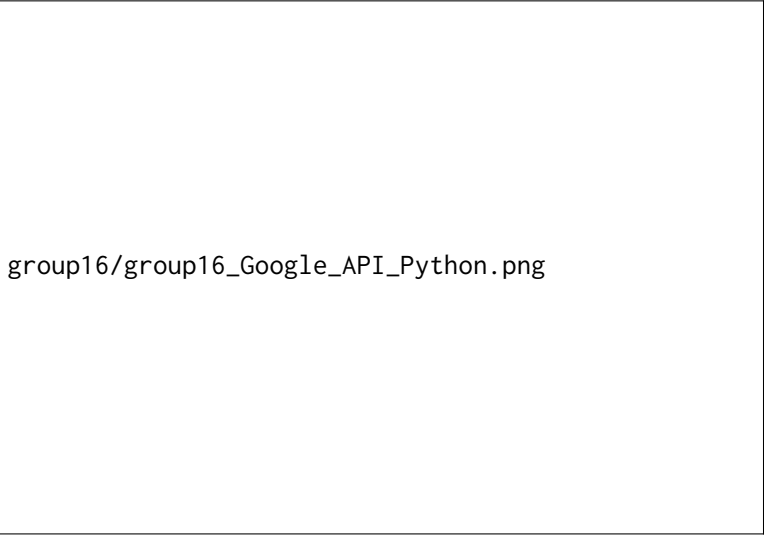


Figure 3: Getting coordinates using Google API in Python

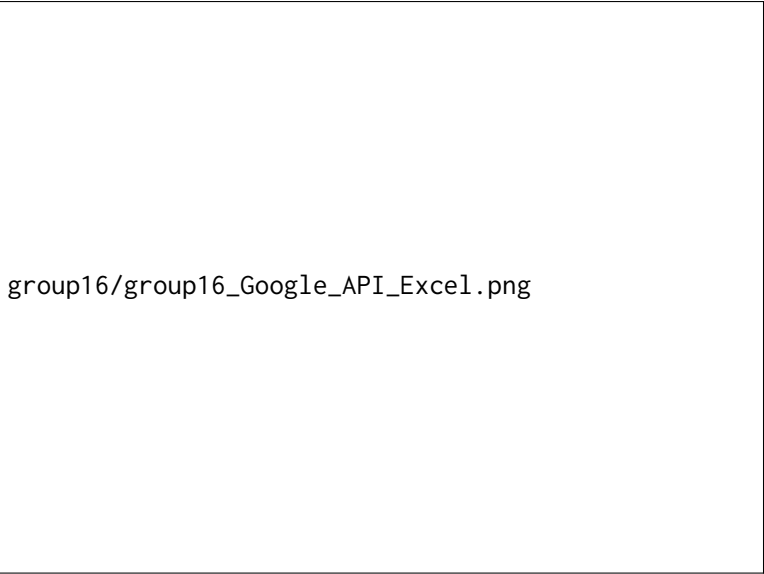


Figure 4: Excel file with addresses and the coordinates

4 Interface

While designing the interface, attention was paid to be as user-friendly as possible, and therefore we created a simple and easy-to-use interface. First, the user should select three Excel files with .csv extension. After the file selection is made, the name of the selected file and the date is written next to the button as shown in Figure 5.

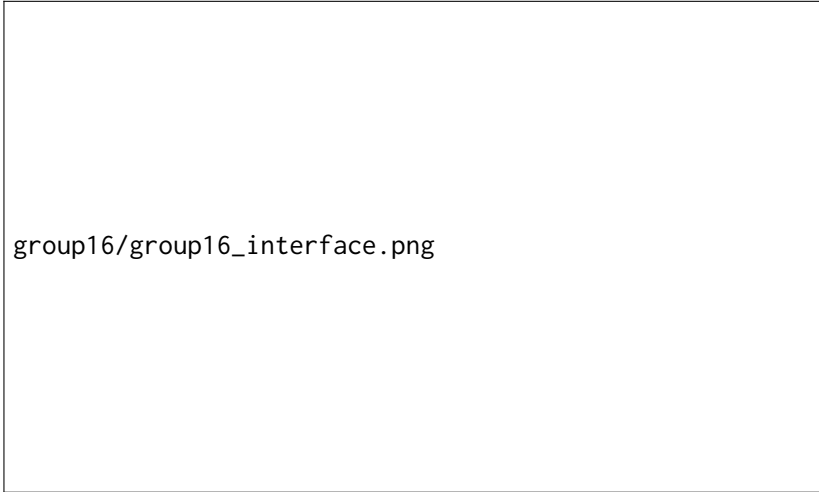


Figure 5: Display of the name of the selected file and the date

This way, the user can be sure that s/he has selected the correct file. At the same time, the user can see the preview of the selected file, if s/he wants. An example of such a preview of a small sample can be seen as in Figure 6.

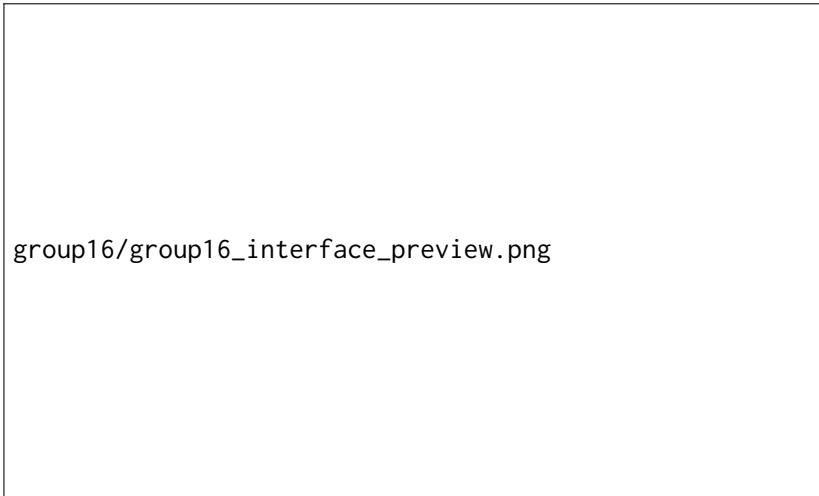


Figure 6: Preview of the selected file

5 Conclusion

When we examine the current system, three main issues are identified, as we mentioned above. These are the number of vehicles used, the time spent on the bus and when to get on/off the buses. Moreover, the Transport Unit does not use computer programs or any scientific methods, so the data is not reliable and some parts are missing. These issues can be solved as much as possible with the proposed approach.

For the stops, Ankara EGO stops were taken from the website via Google API, and the addresses of the personnel that are provided to us within the scope of Personal Data Protection Authority are used. It can be said that this is one of the most significant benefits for the unit and the personnel.

Although a reduction in the total number of buses does not always decrease the total cost, we can reduce the total cost depending on the number of buses. There may be a cost difference in using one large bus instead of more than one minibus. In other words, the distance a bus travels while picking up all the passengers is much more than two minibuses. However, sometimes it may be logical to use two minibuses. Thus, the model that can give the exact number of buses and types to minimize cost is prepared. Since, we know the cost of all three bus types and the total number of buses used, we can do a comparison between how much paid and how much could have been paid according to our solution.

Another benefit is that the proposed system gives the approximate arrival time for all the stops. Thus, all vehicles can be arranged to be at 8 a.m. at the school, and problems such as waiting at the station or missing the buses can be eliminated.

References

P. Schittekat, M. Sevaux, and K. Sorensen. A mathematical formulation for a school bus routing problem. In *2006 international conference on service systems and service management*, volume 2, pages 1552–1557. IEEE, 2006.

Appendix A Current System

group16/group16_Current_System.png

Appendix B Table for Afternoon Pick-ups

Table 2: Parameters and decision variables of the afternoon model

<i>Indices</i>		<i>Decision Variables</i>	
$b \in B, (i, j) \in A, i \in S$		$x_{i,j,b}$	$\begin{cases} 1 & \text{if bus } b \text{ traverses arc}(i,j) \\ 0 & \text{otherwise} \end{cases}$
<i>Parameters</i>		$y_{i,b}$	$\begin{cases} 1 & \text{if bus } b \text{ visits stop } i \\ 0 & \text{otherwise} \end{cases}$
S	Set of all potential stops (Stop 0 is the school)	$z_{i,p,b}$	$\begin{cases} 1 & \text{if personnel } p \text{ is dropped off by bus } b \text{ at stop } i \\ 0 & \text{otherwise} \end{cases}$
A	Set of all routes between all the stops	w_b	$\begin{cases} 1 & \text{if bus } b \text{ is used} \\ 0 & \text{otherwise} \end{cases}$
B	Set of available buses		
Cap_b	Capacity of bus b		
$t_{i,j}$	Time of traversing arc(i,j)		
c_b	Cost of using bus b		
MD	Maximum acceptable duration of each travel		

Appendix C Mathematical Model for Afternoon Pick-ups

$$\min \sum_{b \in B} c_b w_b \quad (1)$$

$$\text{subject to: } \sum_{b \in B} y_{0,b} = \sum_{b \in B} w_b \quad (2)$$

$$y_{0,b} = \sum_{b \in B} w_b \quad \forall b \in B \quad (3)$$

$$\sum_{i \in S} y_{i,b} \leq w_b (|S| + 1) \quad \forall b \in B \quad (4)$$

$$\sum_{(i,j) \in A} x_{i,j,b} = \sum_{(i,j) \in A} x_{j,i,b} = y_{i,b} \quad \forall i \in S, \forall b \in B \quad (5)$$

$$\sum_{b=1}^B y_{i,b} \leq 1 \quad \forall i \in S \setminus \{0\} \quad (6)$$

$$\sum_{i \in S} \sum_{p \in P} z_{i,p,b} \leq Cap_b \quad \forall b \in B \quad (7)$$

$$z_{i,p,b} \leq y_{i,b} \quad \forall i \in S, \forall p \in P, \forall b \in B \quad (8)$$

$$\sum_{i \in S} \sum_{b \in B} z_{i,p,b} = 1 \quad \forall p \in P \quad (9)$$

$$u_{0,b} = 0 \quad \forall b \in B \quad (10)$$

$$u_{i,b} - u_{j,b} + t_{i,j} \leq MD(1 - x_{i,j,b}) \quad \forall b \in B, \forall (i, j) \in A \quad (11)$$

$$0 \leq u_{i,b} \leq MD \quad \forall b \in B, \forall i \in S \quad (12)$$

$$u_{i,b} \geq 0$$

$$x_{i,j,b} \in \{0, 1\}$$

$$y_{i,b} \in \{0, 1\}$$

$$z_{i,p,b} \in \{0, 1\}$$

$$w_b \in \{0, 1\}$$

$$\forall i \in S, \forall b \in B \quad (13)$$

$$\forall (i, j) \in A, \forall b \in B \quad (14)$$

$$\forall i \in S, \forall b \in B \quad (15)$$

$$\forall i \in S, \forall p \in P, \forall b \in B \quad (16)$$

$$\forall b \in B \quad (17)$$

Personel Servis Filosu Rotalaması için Karar Destek Sistemi

Brisa Bridgestone Sabancı Lastik San. ve Tic. A.Ş.

group17/Group Photo.png

Proje Ekibi

Pınar Aktaş, Barış Aydın, Sıla Ayık, Ahmet Hakan Çolak
Ömer Faruk Düzgeç, Mehmet Can Gökçimen, Waleed Khaled Waleed Qutob

Şirket Danışmanı

Buğra Kılıç
Endüstri Mühendisliği Takımı

Akademik Danışman

Asst. Prof. Özlem Çavuş
Endüstri Mühendisliği Bölümü

ÖZET

Brisa Bridgestone Sabancı'nın Türkiye'deki iki tekerlek üretim fabrikasından biri olan Aksaray fabrikasında mavi ve beyaz yaka çalışanlarının ulaşımını sağlayan personel servis filosu mevcuttur. Ulaşım işlemleri, verimi düşük rotalama ve çalışan eşleştirmeleri sonucu geliştirilebilir ulaşım maliyetleri oluşturmaktadır. Projenin temel amacı verimsiz noktaların tespit edilerek düzeltilmesi ve uygulanabilir, kullanıcı dostu, maliyet düşüren bir karar destek sistemi tasarlamaktır. Optimizasyon modeline bağlı kurulan karar destek sistemi sayesinde güncel ulaşım maliyetleri üzerinde %10-20 arası düşüş hedeflenmektedir.

Anahtar Kelimeler: Heterojen Araç Rotalama Problemi, Araç Eşleştirmeleri, Çalışan Ulaşımı, Karar Destek Sistemi, Sezgisel Methotlar

Key Words: Heterogeneous Vehicle Routing Problem, Vehicle Assignments, Employee Transportation, Decision Support System, Heuristic Methods

Decision Support System for Personnel Service Fleet Routing

1 Introduction

1.1 Detailed Description of Current Operations

Brisa Bridgestone has a manufacturing plant in the Aksaray Organized Industrial Zone on a 952,000 m² area. The factory is equipped with strong assets to produce annually 4.2 million tires at full production capacity while employing 700 people.

In the Brisa Aksaray factory, there exists an ongoing employee schedule designed to collect employees of both blue-collar and white-collar workers of three shifts working plan. Due to geographical dispersion of employees, there are three different regions and the employee service system is divided into sub-regions.

The ongoing carried out operations for the service system only utilize mid-ranged travels (15-30 kms) for any vehicle. Accordingly, the sum of associated fixed costs for the traveling distance results in the overall cost. Detailed breakdown of fixed costs for incremental distances for vehicle types are provided by the company.

There are varying passenger capacities for each type of vehicle. Use of multiple vehicles may be necessary if the demand for any particular sub-region is beyond the capacity of any single vehicle. It is a side note from the company that, due to current necessities of the pandemic crisis, there are even more vehicles being used in order to limit the number of passengers in any vehicle to half of its capacity. For instance the "Temsä" bus with capacity of 50 seats is only allowed to collect 25 employees.

Some vehicles are used for all three shifts for both collecting and distributing the employees whereas some are used only for one shift and only to collect employees and arrive at the factory. There is also diversity in the number of vehicles used on weekdays and weekends considering changing total number of employees collected.

1.2 Analysis and Interpretation of Data

It is possible to divide the data into three headings which are the specifications regarding the ongoing service schedule, pricing information for any vehicle per incremental distance traveled and the open addresses for employees.

For the interpretation of data for specific addresses of the employees, the conversion to the latitude and longitude coordinates are obtained. The issue during the analysis of the open addresses is that most of the data were incomplete. That situation arose the need of extra work to either search manually the location or approximate location. The main objective for the utilization of coordinates will be highlighted in the methodology section but before that additional assumptions

are needed to be introduced during the data interpretation stage. For instance, the cases where open addresses are only left out with boulevard or street names; the midpoint of the roads are selected to be the candidate coordinates. For the cases where the data are even more scarce such as it is being left out with only district names, they are eliminated since it is not possible to give reasonable pinpoint coordinates.

Figure 1: Address Distribution of Employees



group17/Addresses.png

For the pricing information for any vehicle for distance traveled, it is an opportunity for the project to utilize that pricing information such that choosing different distance and vehicle combinations may yield greater benefit as opposed to current way of carried out operations. These costs which can be interpreted as the fixed costs and vehicle capacities are our project's input parameters.

2 Problem

In this project, the encountered problem can be generalized under Heterogeneous Vehicle Routing Problem (HVRP) for the employee transportation system of Brisa Bridgestone. This transportation system consists of three shuttle shift periods which carry employees from their homes to work on specific time intervals, and back to home from work after their shifts are done. There is a specific fixed cost incurred for each trip depending on the distance covered. The current system lacks routing optimization for each vehicle since bus stops are not correctly assigned according to the demand of each area and fleet capacities are not optimally used when assigning employees to the vehicles. Hence, it opens up an opportunity to overcome the excess costs with correct optimization methods. In order to minimize the cost of transportation, it is needed to find optimal routes for each vehicle by choosing which bus stops to be utilized. The bus stops which the optimal routes will pass through should include correct assignment of employees while keeping the walking distance and transportation time of each employee within agreed walking distance of 650 meters. Each vehicle should be assigned to a time period to depart

according to the starting and ending time of each shift.

Under the numerous constraints, the main performance measure will be the total cost which is provided as the sum of fixed costs by the contracted agency of Brisa. The routes and assignment of employees to bus stops and busses are not designed efficiently and can be improved. Any transition for a bus from medium-range total distance travel cost to a short-range total distance travel cost yields approximately 10-20% overall cost reduction. Therefore, the need for achieving a reduced cost can be validated through this assessment.

3 Proposed System

3.1 Inputs and Outputs of the Proposed System

The mathematical formulation of the proposed system will take inputs: employees' locations, possible bus stop locations, maximum walking distance, capacities of the available buses and the fixed costs of the buses. In accordance with the inputs, the model will give us the following outputs: the routes of the busses, assignment of employees to bus stops and busses, utilized capacities of the busses.

For the applicability measures of such a project demanded by the representatives, there has to be a fixation on the employee bus stop assignments which can stay as is through different shift cycles. Initial mathematical model was not restricted with the consistent bus stop assignments. Nonetheless, as the Table 1 below and the mathematical model illustrates, there is a shift index which serves this objective. During construction of the mathematical model, there were useful literature utilized in order to build the key elements for sets and constraints Schittekat et al. (2006), Baldacci et al. (2008).

3.2 Mathematical Model and Constraints in Detail

The objective of our model is to minimize total transportation cost. The objective function utilizes incremental cost groups due to way fixed costs are defined on the data provided. The step-wise function breakdown of the c_i , is divided into cost groups aa_i , bb_i and cc_i . The multiplications of q_{vi} , ss_{vi} and tt_{vi} binary decision variables with aa_i , bb_i and cc_i parameters let the correct association of incremental costs. The number of vehicles leaving the origin at a shift should be equal to the number of vehicles utilized at that shift (1). The number of vehicles arriving at the origin at a shift should be equal to the number of vehicles utilized at that shift (2). Constraint (3) implies that each employee should be assigned to exactly one bus stop. Constraint (4) demonstrates that each employee's walking distance should be less than maximum allowed limit of walking distance. Constraint (5) enforces the binary decision variable q_{vi} to be greater than or equal to ss_{vi} since if the distance pushes the ss_{vi} variable to be equal to 1, then intuitively q_i should be 1 as well. Similar logic is followed on the constraint (6) with the link between ss_{vi} and tt_{vi} . In a way, these constraints can be followed in reverse order since the hierarchy is established with this approach. For the constraint (7), those

three variables and the incremental distance intervals of 0-15-30-45 are matched with total distance traveled of any vehicle i being less than or equal to the right-hand side. More specific indication of the distance traveled by any vehicle i is represented in constraint (8), which returns the summation of distances between assigned bus stops of vehicle i . Constraint (9) satisfies the network balance for any vehicle. At any shift, if vehicle i arrives at bus stop a then it should also leave this bus stop. Constraint (10) satisfies the case of a vehicle i not visiting the bus stop a should not be taking any employees from that bus stop. For the constraint (11), in this case, if the vehicle i visits bus stop a , there has to be at least one employee assigned to vehicle i for that bus stop so that the employee collected at bus stop a , s_{ia} , is greater than or equal to the sum of arriving arcs to a .

Table 1: Sets/Parameters/Variables of the model

<i>Sets</i>		<i>Decision Variables</i>	
	I Set of vehicles K Set of employees A Set of bus stops V Set of shifts $\delta_{(+)}$ Arcs entering $\delta_{(-)}$ Arcs leaving	x_{viab}	$\begin{cases} 1 & \text{if vehicle } i \text{ uses arc } (a,b) \\ & \text{at shift } v, v \in V, i \in I, a \in A, b \in A \\ 0 & \text{o.w.} \end{cases}$
<i>Parameters</i>		e_{ka}	$\begin{cases} 1 & \text{if employee } k \text{ is assigned to} \\ & \text{bus stop } a, k \in K, a \in A \\ 0 & \text{o.w.} \end{cases}$
c_i	$\begin{cases} aa_i & \text{TL for } [0,16] \text{ km distance} \\ bb_i & \text{TL for } [16,31] \text{ km distance} \\ cc_i & \text{TL for } [31,45] \text{ km distance} \end{cases}$	q_{vi}	$\begin{cases} 1 & \text{if vehicle } i \text{ is utilized at shift } v \\ & v \in V, i \in I \\ 0 & \text{o.w.} \end{cases}$
f_{vk}	$\begin{cases} 1 & \text{if employee } k \text{ is working at} \\ & \text{shift } v, v \in V, k \in K \\ 0 & \text{o.w.} \end{cases}$	ss_{vi}	$\begin{cases} 1 & \text{if distance traveled by } i\text{th vehicle at} \\ & \text{shift } v \text{ is greater} \\ & \text{than 15 km, } v \in V, i \in I \\ 0 & \text{o.w.} \end{cases}$
w_{ka}	walking distance of employee k to bus stop a , $k \in K, a \in A$	tt_{vi}	$\begin{cases} 1 & \text{if distance traveled by vehicle } i \text{ at} \\ & \text{shift } v \text{ is in } [31\text{km}, 45\text{km}], v \in V, i \in I \\ 0 & \text{o.w.} \end{cases}$
M	used as considerably large number (Big-M)	s_{via}	the number of employees that will be taken from bus stop a via vehicle i at shift v , $a \in A, i \in I, v \in V$
p_{ab}	distance between the bus stops a and b , $a, b \in A$	y_{viab}	number of employees transported through arc (a,b) via vehicle i at shift v , $i \in I, a \in A, b \in A, v \in V$
W	max allowed walking distance per employee	u_{via}	the visiting sequence of vehicle i for the bus stop a at shift v , $i \in I, a \in A, v \in V$
O_i	predetermined constant for capacity utilization of vehicle i , $i \in I$	d_{vi}	distance traveled by vehicle i at shift v , $i \in I, v \in V$
D_i	capacity of vehicle i , $i \in I$		

$$\begin{aligned}
\min \quad & \sum_{v \in V} \sum_{i \in I} (q_{vi}aa_i + (bb_i - aa_i)ss_{vi} + (cc_i - bb_i)tt_{vi}) \\
& \sum_{i \in I} \sum_{a \in \delta(0-)} x_{via0} = \sum_{i \in I} q_{vi} & \forall v \in V(1) \\
& \sum_{i \in I} \sum_{b \in \delta(0+)} x_{vi0b} = \sum_{i \in I} q_{vi} & \forall v \in V(2) \\
& \sum_{a \in A - \{0\}} e_{ka} = 1 & \forall k \in K(3) \\
& \sum_{a \in A} w_{ka}e_{ka} \leq W & \forall k \in K, \forall a \in A - \{0\}(4) \\
& ss_{vi} \leq q_{vi} & \forall i \in I, \forall v \in V(5) \\
& tt_{vi} \leq ss_{vi} & \forall i \in I, \forall v \in V(6) \\
& d_{vi} \leq 15(ss_{vi} + tt_{vi} + q_{vi}) & \forall i \in I, \forall v \in V(7) \\
& \sum_{a, b \in A} x_{viab}p_{ab} = d_{vi} & \forall i \in I, \forall v \in V(8) \\
& \sum_{b \in \delta(a+)} x_{viab} = \sum_{c \in \delta(a-)} x_{vica} & \forall i \in I, \forall v \in V(9) \\
& M \sum_{b \in \delta(a+)} x_{viab} \geq s_{via} & \forall i \in I, \forall a \in A - \{0\}, \forall v \in V(10) \\
& \sum_{b \in \delta(a+)} x_{viab} \leq s_{via} & \forall i \in I, \forall a \in A - \{0\}, \forall v \in V(11) \\
& \sum_{k \in K} e_{ka} = \sum_{i \in I} s_{via} & \forall v \in V, \forall i \in I, \forall a \in A - \{0\}(12) \\
& \sum_{a \in A - \{0\}} s_{via} \leq D_i q_{vi} O_i & \forall v \in V, \forall i \in I(13) \\
& x_{viab} + x_{viba} \leq 1 & \forall i \in I, \forall v \in V, \forall a, b \in A - \{0\}(14) \\
& u_{via} - u_{vib} + 1 \leq (1 - x_{viab})|N| & \forall i \in I, \forall v \in V, \forall a, b \in A(15) \\
& 1 \leq u_{via} \leq |N| & \forall i \in I, \forall v \in V, \forall a, b \in A(16) \\
& q_{vi} \in \{0, 1\} & \forall i \in I, \forall v \in V(17) \\
& x_{viab} \in \{0, 1\} & \forall i \in I, \forall v \in V, \forall a, b \in A(18) \\
& e_{ka} \in \{0, 1\} & \forall a \in A, \forall k \in K(19) \\
& s_{via} \in \mathbb{Z}_+ & \forall v \in V, \forall i \in I, \forall a \in A(20) \\
& y_{viab} \in \mathbb{Z}_+ & \forall v \in V, \forall i \in I, \forall a, b \in A(21) \\
& u_{via} \in \mathbb{Z}_+ & \forall v \in V, \forall i \in I, \forall a \in A(22) \\
& d_{vi} \in \mathbb{Z}_+ & \forall v \in V, \forall i \in I(23) \\
& ss_{vi} \in \{0, 1\} & \forall i \in I, \forall v \in V(24) \\
& tt_{vi} \in \{0, 1\} & \forall i \in I, \forall v \in V(25)
\end{aligned}$$

When constraint (10) is redundant, constraint (11) satisfies that vehicle i takes

at least one employee from bus stop a . Constraint (12) emphasizes that for any shift, total number of assigned employees to bus stop a will be equal to total number of people collected with all the buses from bus stop a . Constraint (13) shows that at any shift, number of employees transported through that arc (a,b) via vehicle i should be always less than the capacity of utilized vehicle i with a predetermined gap which is satisfied through multiplication of capacity with slack O_i . The logic behind that slack constant multiplication is to enforce the vehicles to keep an amount of empty space rather than full utilization in case of undesired scenarios. Constraint (14) implies that a vehicle cannot travel both on arc (a,b) and (b,a) . Constraints (15) and (16) are MTZ constraints for subtour elimination. Constraints (17) to (25) are the domain constraints.

3.3 Heuristic Approach

Heuristic approach for the quicker run time of the decision support system is needed at the point where administrative aspect of the project asked for a daily output performance. Mathematical model itself was not capable of providing the outputs needed for each shift within a day of run time. To overcome this obstacle, Greedy Edge Cover Heuristic is utilized, Dharmarajan and Ramachandran (2019), with the help of IE 303 course material and additional literature review Gramm et al. (2006).

In our heuristic approach, edges are formed between employee and bus stop coordinates within the constraint of 650 meters of walking distance. With the edges formed, each bus stop is assigned with the same weight of 1 and sigma values are computed as follows:

$$\sigma_a = \frac{\# \text{ of edges formed for bus stop } a}{\text{weight}}$$

Result: Demand for Selected Bus Stops

initialization Set to collect demand and bus stops $B = \emptyset$

Compute $\sigma_a, \forall a \in A$

while *There are uncovered edges* **do**

 Pick the bus stop a^* with largest σ_a

$B = B \cup a^*$

$s_{ia^*} = \# \text{Number of edges erased}$

end while

Algorithm 5: Greedy Edge Cover for Demand Determination

Each bus stop is ranked with respect to the value and assignments are followed accordingly with the upper limit of 20 employees. 20 employee upper limit is decided with the minimum capacity inside the fleet so that uneven allocation is eliminated during routing. Assignment procedure is repeated step by step as one bus stop and its assignments are completed, edges are deleted and sigma values are recomputed until no edge exists. These demand numbers for each bus stop

is used for the mathematical model to be taken as inputs so that routing can be determined.

Benefit behind the heuristic approach as a facilitator for the mathematical model is that, run time significantly reduces to a point where daily run capability is achieved. To illustrate the incurred cost differences for the month of March 2021 on morning shift, ongoing cost is known as 35,452 TRY. Mathematical model optimal solution provides the cost of 28,883 TRY and heuristically adjusted system provides the cost of 30,196 TRY. Scenario analyses of two paths yield 17.4% and 22.7% cost reduction accordingly. Both cases satisfy the objective of 10-20% cost reduction whereas heuristic approach additionally lets the administrative aspect of the decision support system to be highly functional. Validation tests for both of the options are also administered before applying bench-marking and scenario analyses which is explained in detail below.

3.4 Validation

In order to validate our model, the mathematical model is coded in Python IDE and solved using the solver Gurobi to provide the output data. With the actual data from the ongoing routes and employee assignments, one by one arcs are formed to create the network referring to variable x_{viab} . For each region and sub-region order of travels are entered manually with pulling coordinate-bus stop index values of 98 bus stops. Each bus stop demand is entered manually to specify how many employees will be collected from each bus stop. This refers to the decision variable s_{ia} where vehicle indices are again manually set with locating which bus stop a it relates to by matching coordinate from the data set. Lastly, how many and which buses are used for any route are manually entered. Associated fixed cost for each of these travels, c_i 's are matched with the data provided. Overall, fixed variables can be listed as x_{viab} , q_{vi} , s_{via} , u_{via} . The data utilized for all these purposes are taken from the March 2021 employee shift schedule and daily cost verification is multiplied with number of working days of the month separately for white and blue collar workers. Due to current capacity utilization constraint of 50% because of safety measures, setup for the mathematical model is adjusted. The actual cost for that month was again reported as 359,734 TRY. Comparing the result with the validation result from the model resulted in 3.47% difference with 372,217 TRY. Possible minor deviations of the result can be explained as unplanned changes on employee transportation such as from day to day several employees may be out of office, prefer different transportation method and within month shift adjustments. All these fluctuations create the possibility of new demand for a bus stop and accordingly change on which bus to choose or what total distance traveled become.

Since validation step is the test for accuracy on actual costs with the use of manually implemented inputs, both heuristic approach and mathematical model separately calculated the exact number which is mentioned above.

3.5 Integration

The medium for delivering the ideal decision support system is decided as Python GUI as the mathematical model and heuristic approach both structured in Python IDE as well. For the users of the decision support system, required login credentials are introduced. Within the decision support system, multiple functionalities to maintain sustainable usage are constructed such as employee list changes, fleet changes and fixed cost adjustments. The output generated on the decision support system for Brisa Bridgestone illustrates which Employee ID is assigned to which bus stops on a consolidated list. In the case of new employee additions to the system or oppositely removals, coordinate list of employee addresses is adjusted accordingly to be introduced to the decision support system if needed. Representative design of the final design for the decision support system with the mentioned characteristics can be observed in Appendix A with output mapped in Appendix C.

All in all, decision support system is designed to help Brisa Bridgestone generate the outputs and manipulated within single medium. However, a consolidated list for Employee ID and their coordinates to pull upon demand via Excel VBA Macro is included as well which can also be seen in Appendix B.

4 Conclusion

In the scope of this project, which basically aims to optimize vehicle routing problem for employee transportation of Brisa Aksaray factory, cost reduction objectives are achieved by developing a vehicle routing decision support system, that is supported by a mathematical model and heuristic approach with the resulted cost for that month as 304,335 TRY. 15.6% cost reduction is observed as opposed to current way of carrying out operations.

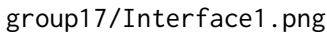
There are several steps that can be taken forward to make this project more comprehensive and sustainable in the future. The proposed system does not consider potential new bus stops which was studied but eliminated due to inapplicability concerns. In a case with new bus stops total transportation cost may be decreased. Additionally, distribution of employees to their shifts from one cycle to another is not possible to foresee. Shift schedule can be redesigned in a way that it can be introduced as a likely forecast so that it can be used in decision support system. Hence, this would provide further improved results with greater accuracy.

References

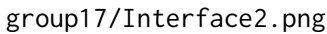
- R. Baldacci, M. Battarra, and D. Vigo. Routing a heterogeneous fleet of vehicles. In *The vehicle routing problem: latest advances and new challenges*, pages 3–27. Springer, 2008.
- R. Dharmarajan and D. Ramachandran. A modified greedy algorithm to improve bounds for the vertex cover number. *arXiv preprint arXiv:1901.00626*, 2019.

- J. Gramm, J. Guo, F. Hüffner, and R. Niedermeier. Data reduction, exact, and heuristic algorithms for clique cover. In *2006 Proceedings of the Eighth Workshop on Algorithm Engineering and Experiments (ALENEX)*, pages 86–94. SIAM, 2006.
- P. Schittekat, M. Sevaux, and K. Sorensen. A mathematical formulation for a school bus routing problem. In *2006 international conference on service systems and service management*, volume 2, pages 1552–1557. IEEE, 2006.

Appendix A User-Interface



group17/Interface1.png

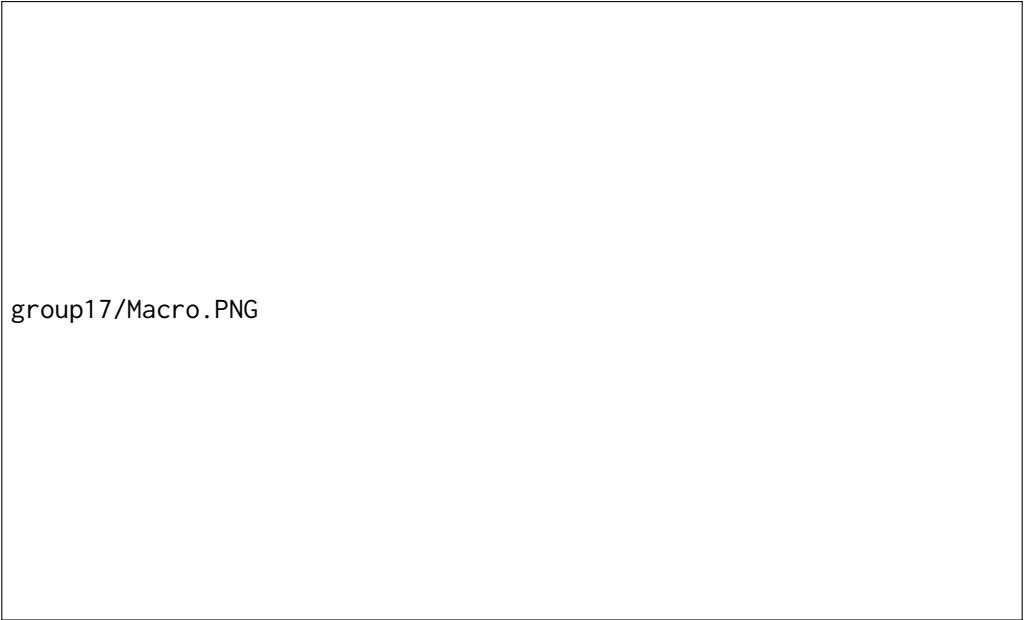


group17/Interface2.png



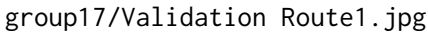
group17/Interface3.png

Appendix B Employee-List

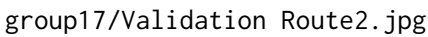


group17/Macro.PNG

Appendix C Routing-Instances

A diagram showing a routing instance configuration. It includes a box labeled "group17/Validation Route1.jpg".

group17/Validation Route1.jpg

A diagram showing a routing instance configuration. It includes a box labeled "group17/Validation Route2.jpg".

group17/Validation Route2.jpg

Şişecam Mamul Ambar Yönetimi Karar Destek Sistemi

Şişecam

group18/group_photo.png

Proje Ekibi

Mirza Benek, Ceren Bozkaya, Göktuğ Doğrul,
Berkay İnal, Ozan Mert Kaya, Çise Kuvan, Ekin Tepe

Şirket Danışmanı

Kemal Eker
Endüstri Mühendisliği Takımı

Akademik Danışman

Yrd. Doç. Dr. Emre Nadar
Endüstri Mühendisliği Bölümü

ÖZET

Şişecam Polatlı Fabrikasında, sevk edilmeyi bekleyen mamüller depodaki sehpalara yerleştirilmektedir. Hangi ürünün hangi sehpayaya yerleştirileceği ABC analizi ışığında RFID sistemi ile belirlenmektedir. Bu projede sehpalardaki ürünlerin, çıkış noktası olan rampalara uzaklıklarının azaltılmasını amaçlayan bir eniyileme modeli geliştirilmiştir. Depo yerleşim planında belirlenen pilot bölgelerde ürünlerin yerleşiminde sevkiyat bilgileri de göz önünde bulundurularak mevcut adresleme algoritması daha dinamik ve verimli hale getirilmektedir.

Anahtar Kelimeler: Depo Yönetimi, RFID, Adresleme, ABC analizi, Pilot Bölge

Decision Support System for Finished Goods Warehouse Management in Şişecam

1 Introduction

1.1 General Information about Şişecam

Şişecam is a global company that produces float glass, glassware, glass packaging, and glass fiber in all key areas of business including soda and chromium compounds. Şişecam continues its activities as an international company with more than 22 thousand employees, production activities spanning 14 countries, and sales in over 150 countries. Due to its leading position in the industry, the company's products are in high demand. This high demand brings along large quantities of production. Stock and warehouse optimization are some of the top priorities of the company due to the high volume of production and glass being a delicate product.

1.2 Current System

The factory has a daily volume of 1650 tons of glass plates. Finished goods are directly transferred to the warehouse which contains 674 storage units and 12 ramps where shipping occurs. Generally, while approximately 30-40% percent of the daily production is shipped on the same day, approximately 60-70% percent is stocked in the warehouse. RFID systems are being used to address the products that enter the warehouse. RFID addressing is made according to the FIFO rule and ABC analysis of the company that is based on production volumes and updated every 3 months with the latest orders. Polatlı factory has currently 373 product types and it continues to increase. These products are categorized as Jumbo Size A, B, C and Machine Size A, B, C. In the current system category A products are placed closer to the shipment areas, while category C products are placed further.

1.3 Problem Definition

The current RFID algorithm does not take into account the order information of the products. It makes suggestions based on what is already produced but not on what will be produced. When the data and layout are examined, this approach reduces the handling efficiency. Some category A products may block the stands located closer to the ramps even if they will not be shipped earlier than some category B or C products. The existing system precludes the products in categories B and C from passing beyond the predetermined borders for these categories. This way of locating the products may increase the handling time and transportation to the shipping ramps of the products. The priority of category A products leads to an unnecessary increase in the total distance traveled to the shipping areas since they gum up the products that should be shipped soon in other categories.

2 Solution Methodology

2.1 Proposed System

Rather than addressing the products according to the ABC analysis for all over the warehouse, it is decided to choose two regions for each product size that have enough capacity to contain the products with the available shipment data for the next week. These regions are chosen from stands that are close to ramps to lower the total distance. The pilot regions are highlighted with red in the warehouse layout where the orange regions are the shipment areas. The layout can be seen in Appendix A. Since category A products are already being placed to stands that are close to ramps, this method does not make a significant difference for products in category A but it helps eliminate the unnecessarily long distances traveled by categories B and C products with the closer shipment date. It is observed that addressing products in categories B and C that will be shipped in a week in certain regions of categories B and C causes unnecessary handling, and this proposed solution has the potential to reduce the total distance traveled to the shipment areas.

2.2 Optimization Model

There are two types of stands in the warehouse for each product size. Thus, two identical mathematical models are constructed. The mathematical models find the nearest available place to the shipment area in the pilot regions. For both models, the set of stands, N , contains the stands only in our restricted region. As mentioned earlier, these models take into account orders within a period of 1 week. The parameter o_i indicates the ordered amount of product i that will be placed to stands. Also, parameter c_n stands for the remaining capacity of the stand n at the time the model will be run.

Table 1: Variables of the model

<i>Indices</i>		<i>Decision Variables</i>	
$i \in I, n \in N$		x_{in}	$\left\{ \begin{array}{l} \text{fraction of the amount} \\ \text{of product } i \text{ that will} \\ \text{be placed to stand } n \end{array} \right.$
<i>Parameters</i>		y_{in}	$\left\{ \begin{array}{l} 1 \quad \text{if } x_{in} \geq 0 \\ 0 \quad \text{otherwise} \end{array} \right.$
c_n	remaining total capacity of the stand n		
o_i	order amount of the product i		
r_i	remaining day to the shipment of product i		
d_n	distance of the stand n to the ramps		

Mathematical model:

$$\min \sum_{i \in I} \sum_{n \in N} d_n y_{in} o_i (1/r_i) \quad (1)$$

$$\text{subject to: } \sum_{i \in I} o_i x_{in} \leq c_n \quad \forall n \in N \quad (2)$$

$$\sum_{n \in N} x_{in} = 1 \quad \forall i \in I \quad (3)$$

$$x_{in} \leq y_{in} \quad \forall i \in I, \forall n \in N \quad (4)$$

$$0 \leq x_{in} \leq 1 \quad \forall i \in I, \forall n \in N \quad (5)$$

$$y_{in} \in \{0, 1\} \quad \forall i \in I, \forall n \in N \quad (6)$$

To explain models in more detail, their objective (1) is to minimize the total distance traveled to the ramps by considering the order quantities and shipping dates of the products. Constraint (2) ensures that the amount of the products placed into any particular stand does not exceed its capacity. Constraint (3) ensures that each product is assigned to a stand. Constraint (4) relates the decision variable x_{in} to the decision variable y_{in} .

3 Verification and Validation

The proposed method aims to reduce the total distance glass packages have to travel in the warehouse. We conducted a simulation study in order to investigate if the traveled distances in the current system can be reduced with the proposed method under uncertainty. To analyze the accuracy of simulation models, we calculated standard errors for each product by running 25 replications, which can be also seen in Appendix B. In both of the simulation models of the current system and proposed system, the production quantities were generated according to the real production volumes of the products obtained from the company data. We have considered the volume of those products in a certain month from the data that was provided by the company and have taken the percentages proportional to their production volumes. We also set the arrival rate of products as EXPO(40) so that it would both be realistic and we would not exceed the entity limit of the Arena version we were using. Simulation results for the average distance traveled by each product can be seen in Appendix B. The first two letters of the product codes in the results show the size and category of the products (JA = jumbo-sized, category A and so on) and the rest is the product code. We observe that the distances traveled by the B and C products drop more significantly than the A products. This result is intuitive since we choose the region as close to the ramps as possible and B and C products could never be placed there with the current system while an A product can be addressed there. Although the volumes of the B and C products are lower than that of the A products, our proposed method helps to reduce the total distance traveled by the products. Also, utilization percentages of the stands were examined and it was seen that with the new pilot areas, utilizations of stands are increased. These utility changes per stand can be seen in Appendix C. We have analyzed the outcome of our simulation model in Arena Output Analyzer. We have compared the mean values of total distances traveled for each product in both of our simulations for every replication. From this comparison, it can be seen that the means of these simulations are different and therefore there is a significant difference between them. This means that the proposed system gives better results than the current one. Results of this output analysis can be seen in Appendix D.

4 User Interface and Implementation

The mathematical models were coded on Excel VBA and a user-friendly interface was designed. The interface for the Excel VBA can be seen in the Appendix E. Logic behind our Excel program is as follows: There is a predetermined pilot area for each product size and there is an order list for products that were produced today and will be shipped in 7 days. We place those products to those stands according to our model. There are 4 buttons in the interface. First button allows the user to add in an order to the order list manually. Second button automatically extracts the information of orders that fits to our criteria from the production plan sheet provided by the company. Sheet received from the company was used in the demo study. However, this interface can be used automatically by the company after implementation. Using SAP, the company will be able to use the interface by easily pulling the data into Excel from the system. After the order list is ready, the third button runs the model and places the packages to stands. Finally, the fourth button deletes packages from the system when they get shipped to the customer. Each time the model is run, it takes the current capacity of the stands into consideration so our model can run incrementally. An example for location assignments of each product can be seen in Appendix F.

5 Benefits to the Company

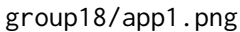
The purpose of the developed models is to minimize the distance of priority products to the shipment area. The current system determines the priority of products based on the analysis of past data only. In the proposed system, the future shipment data is also taken into consideration. The remaining days until the shipment, the volume of the shipment, and some other criteria that are found suitable can be used from the future data to set the priority scores of the products. As stated above, the simulation results indicate that especially B and C products are placed in much closer stands to ramps. Results demonstrate that products will be placed more efficiently with the proposed system: The total distance taken to the ramps in the proposed system is predicted to be lower by approximately 19% than that of the current system. The company will also benefit from savings in terms of time and use of labor thanks to the improvements in handling of the B and C products. Furthermore, usage of transportation machines will be reduced since the total distance traveled will be much lower. Consequently, the service life of the transportation machines will be increased too. Considering all of the benefits stated, it is expected to have a significant improvement in the handling and transportation cost of the company in the warehouse.

References

- C. W. Bulinski, Jerzy and P. Buraczewski. Utilization of abc/xyz analysis in stock planning in the enterprise. *Annals of Warsaw University of Life Sciences-SGGW. Agriculture 61 Agric.*, 2013.
- R. P. S. Kelton, W. David and N. B. Zupick. Simulation with arena.
- E. Kemal. Interview with kemal eker, Oct 2020.
- S. Nahmias and T. L. Olsen. Production and operations analytics. 2020.
- W. L. Winston and J. B. Goldberg. Operations research: applications and algorithms. 3, 2004.

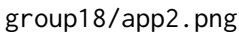
R. G. Zhekun, Li and B. S. Prabhu. Applications of rfid technology and smart parts in manufacturing. *Proceedings of DETC*.

Appendix A Warehouse Layout

A rectangular box containing the text 'group18/app1.png', which likely represents a placeholder for a warehouse layout diagram.

group18/app1.png

Appendix B Proposed and Current System Simulation Results

A rectangular box containing the text 'group18/app2.png', which likely represents a placeholder for simulation results.

group18/app2.png

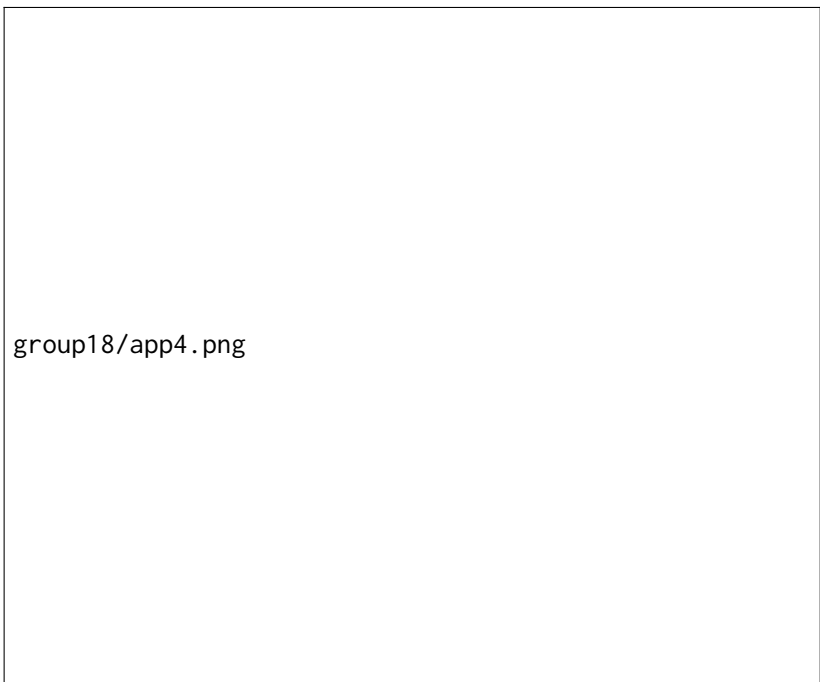
Appendix C Utilization Percentages of the Stands



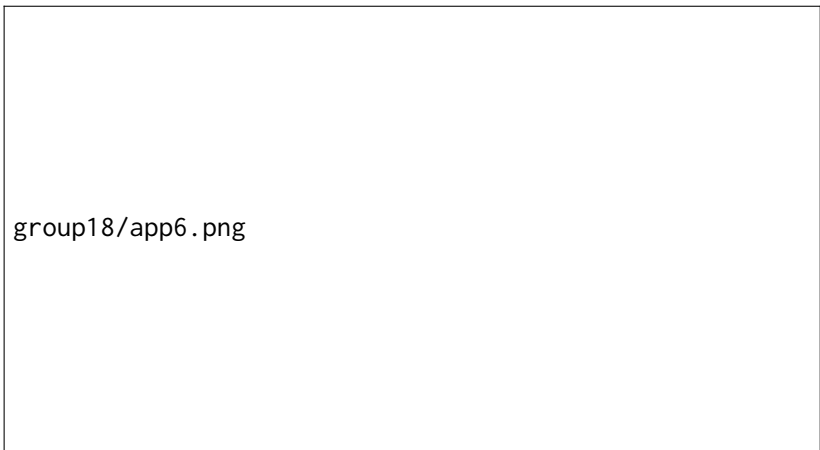
Appendix D Output Analysis of the Simulation



Appendix E User Interface



Appendix F Example of Location Assignments



Üretim Kafile Büyüklüğünü ve Sıklığını Belirlemek İçin Karar Destek Sistemi

Ortadoğu Rulman Sanayi (ORS)

group19/5.png

Proje Ekibi

Mehmet Oğuzhan Ayla, Ezgi İrem Bora, Elif Çağlar
Arda Cem Gürkan, Tuğba Betül İmre, Eyad Kajjoun, Defne Kaya

Şirket Danışmanı

Dr. Alptekin Demiray
Üretim Planlama Müdürü
Endüstri Mühendisliği Takımı

Akademik Danışman

Doç. Dr. Alper Şen
Endüstri Mühendisliği Bölümü

ÖZET

Ortadoğu Rulman Sanayi (ORS) mevcut durumda ürünlerinin hangi sıklıkla ve ne kadar üretileceğine tecrübelerine ve talep geçmişlerine göre karar vermektedir. Bu planlama yöntemi; bazı senaryolarda talebe karşılık verememe, envanter yoğunluğu ve bunların yol açtığı müşteri memnuniyetsizliğine sebep olmaktadır. Şirketle yapılan görüşmeler ve geçmiş verilerin incelenmesi sonucunda en düşük maliyet ile en yüksek talep karşılama oranına ulaşmayı hedefleyen karar destek sistemi oluşturulmuştur. Bu karar destek sisteminin, üretim planlama sürecini kullanıcı dostu bir arayüz ile sistematik bir hale getirmesi ve kolaylaştırması hedeflenmiştir.

Anahtar Kelimeler: Karar destek sistemi, Ana Üretim Programlama, Genişletilmiş EOQ modeli, Üretim Sıklığı, Kafile büyüklüğü

Decision Support System for Dynamic Determination of Manufacturing Frequencies and Lot Sizes

1 System Description and Problem Definition

1.1 System Description

Bearings are important machine components used in many devices in our everyday lives that we use directly or indirectly. It is found in all industrial and consumer machines whose working principle is based on rotation. The first and largest producer of bearings in Turkey is Ortadoğu Rulman Sanayi (ORS). ORS manufactures tapered rollers, cylindrical rollers, deep groove ball bearings, self-aligning, four-point contact bearings, as well as replacement parts such as rings, rollers, and bushings. In total, 750 different main types and 10,000 different versions of bearings are manufactured. The global demand for bearings is over \$ 50 billion, with ORS accounting for around 1 % of the world's market. ORS's market share is about one third in Turkey.

The production of bearings requires several stages. In the first step, metal bars are heated, forged and cut to manufacture rings of different sizes. Then, these rings go through a turning process. After the rings reach their ideal shapes, they are strengthened by heat treatment using different chemicals. After this stage, grinding and stamping processes are applied. In the final and most crucial manufacturing stage, rings and other components are assembled into different bearings. The bearings are then prepared for the final inspection and packaged.

In ORS, production is carried out in large batches until the assembly stage. The assembly is obviously specific for each bearing and is done in smaller batches. Finished bearings that go through all the operations are stored in a warehouse. This warehouse has an automatic storage and retrieval system. The operation chart is shown in Appendix A.

1.2 Description of Operations

As it is mentioned in Section 1, the production of the bearings is made in large batches until they arrive at their assembly stage. These batches are produced in 50 different production lines and there are 750 different types according to their dimensions. The assembly stage differentiates the products into final versions sold to the customer. How much of each version is produced is decided by taking inventory policies, customer needs, setup costs, setup times and inventory obsolescence risks into consideration.

ORS Bearings make their long-term plan for 12 months and develop a master product schedule at a monthly level for each bearing size after deciding the output demand of each product annually. In Appendix B, the production planning flow chart of ORS is available. Some products are produced every month in this master production schedule, while some other products can be produced bi-monthly or even less frequently.

1.3 Problem Definition and Needs Assessment

Decisions made in ORS are currently made on an ad-hoc basis and based on previous experiences. This leads to many problems. For example, if a particular bearing is not available in stock, a customer ordering that bearing must wait until the next production batch. This problem leads to customer frustration (that may affect customer loyalty and future orders) or even the loss of the current order. Demand uncertainty and the cost of not being able to fulfill demand on time should be quantified and objectively integrated into frequency decisions. Also, the current production plan in ORS is made according to the product types, but there are many versions of a product type with very different demand quantities and obsolescence probabilities. In this case, the production plan is not effective because version details are ignored.

ORS reports that it observes consumer concerns regarding long lead times for bearings that are not manufactured every month. At the same time, ORS also complains that excess inventory of many of its bearings are becoming obsolete. Finally, ORS also notes that their effective capacity is lower due to repeated setups, and they have difficulty fulfilling customer demand due to insufficient capacity.

They have difficulties determining the output frequency when they plan to manufacture a new type of product as there is no decision support mechanism inside the factory. They are currently deciding on the basis of past data and experience. Since they do not make these decisions based on a mathematical model, they are unable to know if their choices are optimal or even sufficiently good. The principal concerns of this company are these, and the project is to address these concerns.

1.4 The Objectives and Scopes of the Project

The main objective of the project is to develop a methodology and a decision support system for the company to help determining its production frequency of each product and their lot sizes. While addressing this problem, the company's fill rates are considered as a reference point. We tried to understand the current production status by examining the company's fill rates (Appendix C) and observed that the fill rates are already at a sufficiently high level. So, at the end of this project, the aim is to reduce costs by at least keeping fill rates at existing level. If it is possible, improving the fill rate levels can also be aimed. Since the fill rate is one of the indicators of customer satisfaction, we asked ORS to provide the fill rate levels and how they compute them operationally. In ORS, fill rates are being calculated according to the order amount, customer request date and confirmed delivery time. As a result, orders that cannot be delivered on time decrease the fill rate, since according to ORS, fill rate is equal to the orders delivered on time divided by the total order amount. Calculation of the fill rates is demonstrated in Appendix D. According to the data given by ORS, it is seen that fill rates are sufficiently high. However losing even a fraction of these customers can be quite

costly for the company. In ORS's production system, product types are divided into multiple versions at the assembly stage. To clarify, there is more than one version of a type. In line with the conversations made with the company, it was decided to make production planning on the basis of versions. At the end of the first semester, a production plan was made for the versions accordingly. Again, in line with the company's requests and recommendations, regulating the common production frequency of versions that belong to the same type is left to the company.

2 Proposed System (Methodology)

2.1 Literature Review

In the project, we are asked to find a frequency for production. We mainly planned to use the EOQ (economic order quantity) model and its extensions. The textbook of the IE375 course, *Production and Operation Analysis* Nahmias and Olsen (2021), was our main source for understanding the basics of the EOQ model. Especially Chapter 4 of the book, *Inventory Control Subject to Known Demand* is relevant and we used this chapter as a guide. Also, we found an article written by Doll and Whybark Doll and Whybark (1973) that can be helpful in this project. In the article, an iterative procedure is applied to find near optimal frequencies of products and a base for production scheduling while reducing the total cost. Also they explained the basic period policy (BP) which is a model that considers different production frequencies for different products by relaxing common cycle policy. Another one of our main sources is the article "Lot sizing and scheduling - Survey and extensions" by Drexel and Kimms Drexel and Kimms (1996). The article explains the current practices of different lot sizing and scheduling models, and also exposes their shortcomings. In this respect, we found it very useful for understanding different models and comparing them with each other to find the most applicable one for our project. Also, this article focuses on the MRP II model in the chapter 'Current Practice' (222). Hence it was useful for us to refresh our memory from what we have learned in the IE376 course. For that purpose we used the textbook of IE376, *Manufacturing Planning and Control for Supply Chain Management* Vollmann (2005), for better understanding of the subject MRP. In the project, one of the limitations that we faced was obsolescence. In order to understand the concept of obsolescence, we used the article by Cobbaert and van Oudheusden Cobbaert and Van Oudheusden (1996). This paper explains three different cases related to obsolescence. In our project, the risk of sudden obsolescence is an important sub-problem and this paper helped us figure out how to handle this problem. Another paper was found which contained a rather complex multi-item production planning model. The paper provided some important ideas on how to tackle the obstacle of extreme customization and product variety while planning the production Rakhshani et al. (2020).

2.2 Analysis of Inputs and Outputs

As it is mentioned in 1.3, the main problem of ORS is facing stock-outs and obsolete inventory. In brief, being stock-out means not being able to satisfy the demand fully, and facing obsolete inventory means producing more than required and not being able to sell these products for a long period of time (or not being able to sell them at all). The possibility of these two scenarios happening for a product depends on the production frequency of that product, which is not decided using a scientific methodology as explained in Section 1. In order to solve these problems, we established a decision support system by using an extended version of the Economic Order Quantity (EOQ) model. The reason why we used an extended version of EOQ is because of the uncertain demand. For this purpose, we assumed that the demand for a version is Normally distributed in the extended version of EOQ. We calculated the standard deviation and the mean of the demand from the sales data given for the last 4 years. Our system has both main inputs that directly be plugged into our model and sub inputs which provide us opportunities to calculate main inputs. All of these input values are specific to each version. The main inputs of our model are the monthly production period, monthly demand, set-up cost, obsolescence cost, stock-out cost and the critical ratio. A flowchart that briefly illustrates our system and the relations between inputs is available in Appendix E. Explanations of our decision variables, main inputs and related sub inputs are as follows.

Production Period, n

This is our decision variable which can take the integer values between 1, 2, 4, 8 month(s) and represents the time (in months) between consecutive production runs of a particular bearing version. In other words, this value is the cycle time of a particular bearing version in months and restricted to be powers of two in order to easily assign a common setup for the versions that belong to the same bearing type. To clarify, although we determined the optimal n value according to the product versions, since the versions belong to a certain type, we need to be able to combine production frequencies at a common time, because a setup for producing the bearing types must be done before branching out these types into versions. For this reason, we chose the possible n values so that the LCM (least common multiple) of different versions' cycle time equal to the cycle time of the version that has the least frequent production period (greatest n value), if these versions belong to the same product type.

Monthly Demand, D

Based on the analysis of data and the information that we were given by the company, we consider demand as a normally distributed random variable with parameters mean μ and standard deviation sigma. We calculated μ and σ according to the sales data of the last 4 years. In order to calculate the probabilities that represent the level of demand, we are using μ and σ as sub inputs.

Setup Cost, K

ORS does not consider the setup cost as the cost of necessary adjustments made to produce different types/versions of bearings. They consider this cost as a loss of production capacity due to the time that that setup takes. Accordingly, our system calculates this input as the multiplication of production quantity per shift, unit product cost, setup time that has scale unit as shifts and a ratio that is given to us by the company, represents the other expenditures and termed the fixed cost ratio. This fixed cost ratio is equal to 45% for the whole production processes, but in line with company's advices, our model only considers this ratio for the assembly stage which is approximately a quarter of the whole ratio. In the data set that we have received from the company, setup times were given for the product type. According to IA's opinion, we split the setup times according to the versions of a type in order to make our model sensitive to accurately optimize the production frequency of product versions instead of product types. To clarify, product versions that belong to the same type have considerable differences in their demands (i.e. one's demand is equal to ten times the other's). So, we split the total setup time to product versions according to their demand in order to make the marginal effect of setup time to each product version equal. The calculation method we use when distributing the setup time for a type to versions (setup time for type) * (μ of version) / (total μ).

Visible Months, v

This input was requested by the company because they claimed that within a specific period, they can estimate the demand very accurately but outside of this period, their estimates may be heavily inaccurate. If n is larger than v , then we assume that the demand in the first v periods will be certain and the demand in the remaining $(n-v)$ periods will be random.

Obsolescence Cost, h

This input is one of our main reasons to use an extension version of EOQ model instead of the original one. To clarify, obsolescence cost can be associated with the traditional holding cost but the logic behind it differs. Generally, holding cost can be considered as a cost of carrying a product and providing the necessary conditions to store it in the warehouse. However, ORS considers this cost as obsolescence cost which can be calculated by multiplying the raw material cost and the obsolescence probability. Moreover, probability of obsolescence were given to us as a function and for calculating the expected obsolete product quantity, we used a selection formula which can be illustrated as $E[\max Q * -D, 0]$ where Q^* is the optimal production quantity and D is the normally distributed random variable for demand over the cycle length explained in the previous section. Also, the obsolescence cost is a function of n because the more n means more uncertainty so that the obsolescence probability may differ for each n . So, the obsolescence cost may be denoted as $h(n)$. The obsolescence probability function was given as a logarithmic function of n by the company and we generated probabilities for

each n values. And the formula of obsolescence probability is $[(n - visibility) * 0.015 * \ln(n)]$.

Stock-out cost, π

This input can be associated with the traditional shortage cost. In ORS's production system and sector, there is not a specific unit cost to penalize unsatisfied demands. Therefore, we calculated this value by multiplying the expected value of stock-out level with the probability of being stock-out and some parameters given to us by the company about the cost and profit loss by taking company policies into consideration. Similar to the previous section, we calculated the expected value of stock-out level by means of the adjusted selection formula; $E[\max\{D - Q^*, 0\}]$.

Critical Ratio, z

We determine to use this by analyzing the traditional news-vendor problem's solution. It is calculated by dividing the underage cost by the sum of underage and overage costs in the news-vendor solution approach. Therefore, we implement this method to our model by dividing the stock-out cost π by the sum of stock-out cost π and obsolescence cost h . We are using this ratio in order to find the optimal lot size Q^* by means of the features of normal distribution.

2.3 Model

As it was mentioned before, we used the extension of the EOQ since we have the obsolescence cost instead of the inventory cost and more importantly, we have uncertain demand for multiple products. In addition to the standard EOQ, we have visibility input in our model. So, we added the visibility input and by this input, we claim that for visible months, the demand is considered as μ and after the visible months, the demand is a random variable for each month. In order to calculate μ and σ we used the past 48 months of demand data. As a result, final demand random variable which has mean $n\mu$ and the standard deviation as $\sqrt{n - v}$ was used, where n is the number of months in one cycle (a decision variable) and v is the number of visible months. The reason why we use such a standard deviation formula is that since there is no uncertainty for the demand of the visible months, we subtracted them, and we had remaining $(n-v)$ month that has uncertainty in demand. After this point, the expected cost over the cycle time (n months) for specific version of product was calculated by the formula in Appendix E. After the calculation for each n , we will divide the expected total cost ($TC(n)$) over the cycle time to n so that we have the average monthly cost for each version of a product. By using inputs that was mentioned above, our system analyzes all the possible n [1,12] values and chooses the n -value with minimum total expected cost as the optimal production frequency which is our system's first output (choosing the minimum $TC(n)/n$). The other output of our system is the optimal lot size Q^* which is calculated by means of the critical ratio that is explained in the previous section. To conclude, our system utilizes all inputs

and makes regarding calculations for the purpose of providing the optimal policy statement; “Produce Q^* units in every n month” for each bearing version.

$$TC(n) = K + h(n)(\sigma L(z)\sqrt{n-v} + z\sigma\sqrt{n-v}) + \pi(\sigma L(z)\sqrt{n-v}) \quad (1)$$

2.4 Verification

For the verification part, the main purpose was the testing of certain scenarios and observing if they give a logical result or not. For our model, we tested the model to see how the setup cost affects the production frequency. For instance, we considered differences between n and visibility values for the obsolescence probability. If the beyond-visible period in n (cycle time) is short, the obsolescence probability became less than the cases where the beyond-visible period is long. For example, assuming the n (cycle time) values are 4 and 8 for the same visibility value and the same version, the obsolescence probability is greater for $n = 8$. In addition, we tried our model for different types of products that branch into different versions. For example, we analyzed some extreme products which have high demand and low material cost, and have low demand and high material cost. According to the data, our model generates expected results. To clarify, when the setup cost is too high, our model chooses the largest n (cycle time) value and vice versa. Another data might be the number of months in a cycle (n). Additionally, we observed how our model reacts for the same version of the product in $n=1$ and $n=8$. After those observations, we observed that our model gives logical results. There are 823 different versions of 63 different product types in the given data set. When cycle time (n) is 1,2,4 and 8, the number of versions to be produced according to n is shown in Appendix E. When we examined the optimal cycle time periods, we realized that the versions with small obsolescence probability coincided with the versions with cycle times of 8. At the same time, while the versions with a cycle time of 1 have high obsolescence probability, versions with a cycle time of 8 have a small obsolescence probability as expected. When we examined the optimal cycle times over setup costs, as expected, the cycle times of versions with less setup costs are 1 and 2, while the cycle times of versions with higher setup costs are 4 and 8.

2.5 Validation

In order to validate our results, we separate the demand data for the past 48 months into two. The data for the first 36 months is used for training (determining the frequencies and lot sizes) and the data for the last 12 months is used to test. The results of the last 12 months are used to calculate the fill rates for 3 types and 10 version. For these products, resulting fill rate was 97.17

2.6 Benchmarking

In the benchmarking phase of our project, we primarily thought of using the fill rate. For this reason, we got the information from the company about the fill rates for 2020 and how they calculated the fill rate. Fill rates in ORS are determined

based on order amount, customer request date, and delivery time. As a result, orders that cannot be delivered on time reduce the fill rate, which is equal to the number of orders delivered on time divided by the total order amount, according to ORS. Also average fill rate in 2020 was 93.01 %. As it can be seen that ORS has already a good value. Therefore with our project, we planned to increase this number in order to benchmark. When we did a pilot study with the ORS's data, we were able to see the fill rate around 97.17 % . So, it is highly possible to improve ORS's fill rates using the decision support system we develop.

Additionally, the inventory turnover ratio can be used for benchmarking. It is a measure of how well a company generates sales from its inventory. Since the inventory turnover ratio is a measure of how good a company satisfies sales from its inventory, we thought that it would be suitable for benchmarking. As we know, ORS's competitors are NTN, SKF, NSK and SCHAEFFLER. By dividing the cost of goods to the average inventory, we calculated these companies' inventory turnover ratio for 2020. The inventory turnover ratios for NTN, SKF, NSK and Schaeffler were found to be 3.03, 0.94, 4.43, and 3.27, respectively. According to the IA, ORS's inventory turnover ratio of 2020 is 5.88(100M/17M). As IA stated that the reason why 2020's ratio is 5.88 is because that the company has decreased its stocks due to pandemic. Since the ratio of NSK and ORS are similar, it would be fair to say their their operational performance, in particular, inventory management performance are similar. Considering that NSK is in the 3rd place in market share, we can think that this is a good similarity. We can say that ORS is doing well in comparison to its competitors. With our project we aimed to keep ORS's turnover ratio at these levels.

3 Benefits to the company

This decision support system can provide many benefits to ORS with determining optimal production frequency and quantity. A reduction in the total cost of ORS is expected as our system can minimize setup, stock-out and obsolescence costs. Apart from the cost, we plan to increase ORS's customer satisfaction. We did this by minimizing the unmet demand with the optimum production frequency. Even if the company does not have a capacity issue and holding cost for production, when the product cannot be sold, it becomes an obsolescence cost for the company. This decision support system makes this problem disappeared. Finally, with this automated decision support system, we expect the phase of determining the frequency and quantity of production to speed up. Also, the main aim of this decision support system is increasing the fill rate or decreasing the total production cost while keeping the fill rate constant. The average fill rate of 2020 in ORS was 93.01 %. After using this system on the sample data that we were given, the fill rate became 97.17 %.

3.1 Implementation Plan

We planned to create a user-friendly interface which will provide convenience to the user. We decided to use Microsoft Office Excel for our project because when we asked our industrial advisor, he stated that this tool will be used by the employees who are from different work levels so it must be simple and clearly understandable. So, we used Excel VBA. We decided to have 3 different sheets in the excel workbook which are named as Start and Guide, Data Entry, and Sales. In the Start and Guide sheet, we have the instructions in order to run the program. In the Data Entry sheet, we have purple columns which are needed to be filled by user. The last sheet is Sales and it contains the past sales data. After entering the data, user will run the program and the result page will be shown. There will be type, version, optimal production quantity, and the optimal production frequency outputs for each product in the Result sheet. Also, there will be Multiplication Page sheet that will be shown after running the program and it contains all the combination of production frequencies. All the sheets that are used in our program can be seen in Appendix F.

In order to implement the project, we created a Gantt chart which can be seen in Appendix G. Before submitting the project to the company, we ran our model on our own computer for 63 types of sample data provided by the company and tested whether we achieved reasonable results. Once we confirmed that the results are reasonable, on April 18, we delivered the model to our industrial advisor to be tested on company computers. IA ran the model on the company computer for 63 types of sample data and there was no problem. Since we have successfully completed this stage, they will test the model for all types and versions produced in the company. If there is no problem with the model at this stage, the project and user manual will be delivered to the ORS employee who will use this program in the company. As the last step of the implementation plan, IA will carry out the final acceptance test, which is the procedure that the company applies in all projects. Finally, if it is successful in the final acceptance test, our project will be ready for use.

References

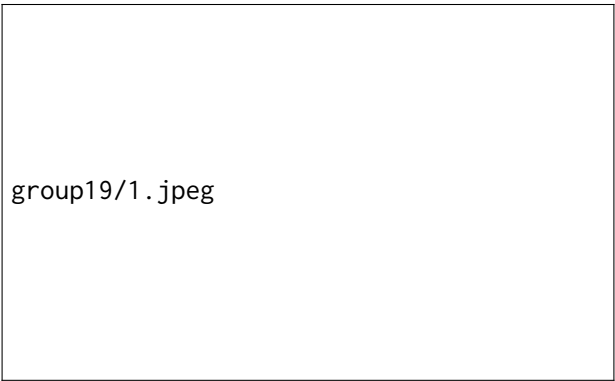
- K. Cobbaert and D. Van Oudheusden. Inventory models for fast moving spare parts subject to “sudden death” obsolescence. *International Journal of Production Economics*, 44(3):239–248, 1996. doi: 10.1016/0925-5273(96)00062-x.
- C. L. Doll and D. C. Whybark. An iterative procedure for the single-machine multi-product lot scheduling problem. *Management Science*, 20(1):50–55, 1973. doi: 10.1287/mnsc.20.1.50.
- A. Drexl and A. Kimms. *Lot sizing and scheduling: survey and extensions*. Inst. für Betriebswirtschaftslehre, 1996.
- S. Nahmias and T. Olsen. *Production amp; operations analytics*. Waveland Press,

Inc., 2021.

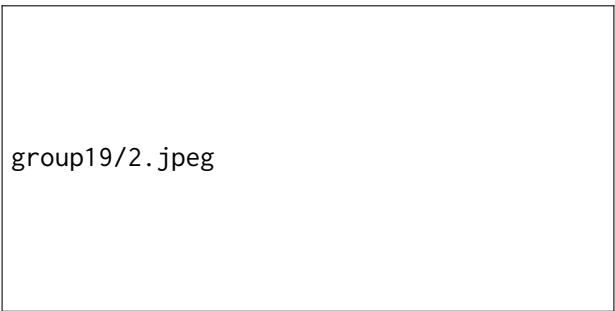
E. Rakhshani, H. Mehrjerdi, and A. Iqbal. Hybrid wind-diesel-battery system planning considering multiple different wind turbine technologies installation. *Journal of Cleaner Production*, 247:119654, 2020. doi: 10.1016/j.jclepro.2019.119654.

T. E. Vollmann. *Manufacturing planning and control systems for supply chain management*. McGraw-Hill, 2005.

Appendix A Operation Chart



Appendix B Planning Flow Chart



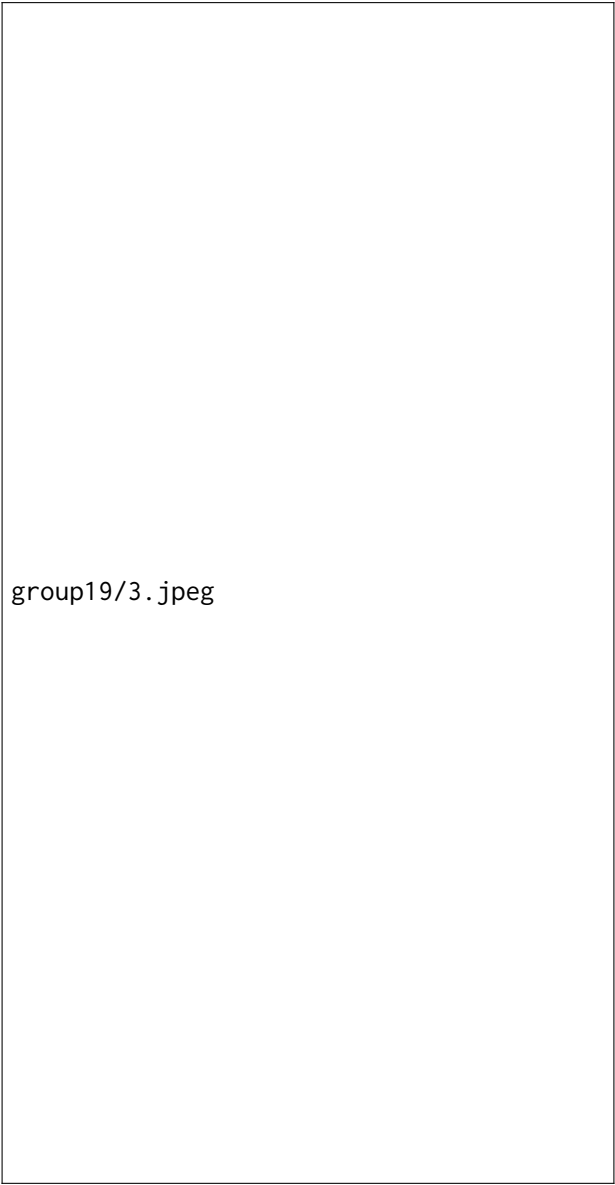
Appendix C Fill Rate Table For the year 2020

Date	Rate
2020-01	93 %
2020-02	95 %
2020-03	94 %
2020-04	92 %
2020-05	93 %
2020-06	92 %
2020-07	93 %
2020-08	93 %
2020-09	93 %
2020-10	95 %
2020-11	91 %

Appendix D Total Number of Versions According to Cycle Times

Cycle Time	Number of Versions
n=1	335
n=2	217
n=3	162
n=4	109

Appendix E Sheets of the program



group19/3.jpeg

Envanter Elleçleme ve Önceliklendirme Projesi

Aselsan

group20/group_photo.png

Proje Ekibi

İrem Naz Başar, Eda Nur Doğan, Kaan Tuna Erkman
Ayşe Secde Gençler, Kaan Kubilay Gökakın, Halis Kerem Özkılıç, Derya Şahin

Şirket Danışmanı

Mühendis Yasin Sarı
Endüstri Mühendisliği Takımı

Akademik Danışman

Dr. Emre Uzun
Endüstri Mühendisliği Bölümü

ÖZET

Bu proje, Aselsan Rehis yerleşkesinde bulunan depo içerisindeki günlük olarak elleçlenen malzeme sayısını en büyüleyecek bir karar destek sistemini grafiksel kullanıcı arayüzü ile sunmayı amaçlamaktadır. Karar verme süreci, başlıca iş emirlerinin seçilmesi, partilenmesi ve daha sonrasında seçilen iş emirlerine belirli kıstaslar uyarınca öncelik verilmesi ve sonra en yüksek malzeme ortaklığı ve dinamik adam-saat kısıtlaması göz önünde bulundurularak gruplanmasına odaklanmaktadır. Belirtilen süreç için karışık tamsayılı progama modeli ve tabu arama sezgisel tekniği çözüm olarak sunulmuştur. Her iki yöntemin doğrulama sonuçları karşılaştırılmıştır. Tabu arama sezgisel modeli, kullanıcı arayüzü tasarımına entegre edilerek karar destek sistemi oluşturulmuştur.

Anahtar Kelimeler: envanter, seçme, partileme, toplama

Inventory Handling and Prioritization Project

1 Introduction

Aselsan is a company of Turkish Armed Forces Foundation and it is established in 1975 in order to satisfy the needs of Turkish Armed Forces. Aselsan is the largest defense electronics company in Turkey. The company has 1680 blue collar employees and 6691 white collar employees. It is an indigenous product exporting company which invests in international markets through various cooperation models with local partners. Furthermore, It has been listed as one of the top 100 defense companies in the world (Defense News Top 100). The product/capability portfolio of the company consists of communication and information technologies, radar and electronic warfare, electro-optics, avionics, unmanned systems, land, naval and weapon systems, air defense and missile systems, command and control systems, transportation, security, traffic, automation and medical systems. The 74.20 % of the shares are owned by the Foundation as the remaining 25.8 % runs in Istanbul Borsa stock market. Aselsan Gölbaşı facility is the place which the Inventory Handling and Prioritization Project will be conducted.

2 Problem

The main focus of the Inventory Handling and Prioritization project is about the storage units, therefore, first, it should be considered carefully. There are thousands of materials stored in the warehouse, either on racks or on special stock keeping units which are known as vertical storage units. The racks are not automated stock keeping units; the workers need to manually retrieve the materials from the racks using forklifts. Usually, materials that are larger in size are stored in the racks. On the contrary, vertical storage units are automated; they have computer interface that the workers in the warehouse enter the location of the desired material, and the vertical storage unit brings that material to the worker automatically. The materials in these units are mostly smaller in size. The address location of a material indicates whether the material is stored in a vertical storage unit or in a rack. For example, a material address R8-04-3501 refers to a rack since it starts with the letter “R”. If an address starts with a letter other than “R”, then it often refers to a material that is stored in the vertical storage unit. At the beginning of a day, the warehouse engineer receives all the work orders that are issued that day and transferred from previous days. Afterwards, they select the work orders according to several criteria. These criteria are the priority of the work order (P), the statistical delivery date of the work order (SDD) and lastly the sub material assembly date of the work order (ATME). The P criterion refers to the work order which take precedence over the other work orders; these are the work orders that the engineer must select them within the production limits of that day. The SDD is the criterion which indicates the company’s estimated delivery date for the work order. The ATME date shows the date on which the subparts that are required for a work order are completed. These three are the selection basis of the engineer, but they do not use a systematic approach.

Therefore, the first objective will be to come up with a systematic approach which will select the work orders according to this selection basis. After the selection, the engineer groups the selected orders into batches. Therefore, the second objective in this process will be based on batching the orders according to the locations of the materials requested by the work orders and then the material similarities of work orders. These two considerations will be analyzed together in order to come up with a better batching planning. Then, the materials will be collected from their locations by the workers. They will collect the materials of the work orders of the batch. This collection will be done with forklifts. The number of employees in this section is dynamic, which means that it can be determined differently. Workers have to work for 8 hours each day. However, it is possible to work overtime if needed. Thus, each day has different man-hour constraints. Lastly, the materials will be put into the checking area and “reconstructing the work orders and checking” will be completed in this station at the end of each batch. In this process, the workers mainly check all the materials of a batch and reconstruct all the materials of a work order accordingly.

3 Solution

3.1 Mathematical Model

Sets

W = Set of all work orders that are received that day

S_i = Set of materials required for work order i

T = Set of all materials

Parameters

p_i = priority of work order i , $p_i \in \{0, 1\}$, $\forall i \in W$

d_{Si} = the difference between today and SDD for work order i , $\forall i \in W$

d_{Ai} = the difference between today and ATME for work order i , $\forall i \in W$

M = Available man-hour in the warehouse that day

E_{max} = Maximum available time to spent without any operations

$$\beta_{ij} = \begin{cases} 1 & \text{if } p_i > p_j , \forall i \neq j \in W \\ 0 & \text{otherwise} \end{cases}$$

$$\theta_{ij} = \begin{cases} 1 & \text{if } p_i = p_j \text{ and } d_{si} > d_{sj} , \forall i \neq j \in W \\ 0 & \text{otherwise} \end{cases}$$

$$\alpha_{ij} = \begin{cases} 1 & \text{if } p_i = p_j \text{ and } d_{si} = d_{sj} \text{ and } d_{Ai} > d_{Aj}, \forall i \neq j \in W \\ 0 & \text{otherwise} \end{cases}$$

$$C_{ij} = \frac{|S_i \cap S_j|}{|S_i \cup S_j|}$$

D_m = distance of material m to the warehouse collection location, $\forall m \in T$

N_{im} = number of material m included in work order i

r_i = retrieval time of work order i

By definition, d_{Ai} is always greater than or equal to zero, whereas d_{Si} can be negative or positive. α_{ij} is an $n \times n$ matrix and each entry can take the values one or zero. If the priority of two work orders is the same, and their d_{Si} values are also the same, and the ATME date of work order i is larger than work order j , the entry takes value one. It is equal to zero otherwise. Similarly, β_{ij} is an $n \times n$ matrix and each entry can take the values one or zero. The entries of this matrix denote the priorities of each work order compared to each other work order. Similarly, the $n \times n$ matrix θ_{ij} was constructed. Its entries compare the SDD of each work order with one another, in the case that their priorities are the same. C_{ij} is the commonality matrix, it is also an $n \times n$ matrix. The entries denote the commonality of each pair of work orders. These values are between zero and one and found by using the formula above. In this formula, we take the ratio of the number of materials that are the same for work orders i and j , and the combined number of materials that are required for these work orders. D_m is an $l \times n$ matrix which denotes the needed distance made while collecting material m , starting from the checking location, and going back to there. This assumption is made to estimate the total distance made during the collection of a batch and hence, the total time needed for that specific batch. Although this is an over-estimation, as one batch can be collected in fewer tours than the total material number, this assumption is the most efficient way of calculating the needed time. Also, this estimation creates an upper limit which makes the model feasible. Finally, r denotes the retrieval time which is a constant parameter of getting an object from their warehouse location, for example cutting the required length of wire or required number of any other material. With the help of this constant and the distance matrix, we can obtain the total time needed to of a specific material from the warehouse, considering both the required distance made in the warehouse and the time spent retrieving the material from the racks or vertical storage systems.

Decision Variables

$$L_{mk} = \begin{cases} 1 & \text{if material } m \text{ is collected in batch } k, \forall m \in T, k \in K \\ 0 & \text{otherwise} \end{cases}$$

$$X_{ik} = \begin{cases} 1 & \text{if work order } i \text{ is grouped into batch } k, \forall i \in W, k \in K \\ 0 & \text{otherwise} \end{cases}$$

M_k : available man-hour for batch k that day, $\forall k \in K$

$$Z_{ij} = X_{ik}X_{jk}, \forall i \neq j \in W, k \in K$$

E : time spent without doing any operations during the work hours

The nature of the problem requires the usage of several binary variables. For example, the model should determine whether to include a certain work order in a certain batch or not. Moreover, there is another decision variable that aims to determine the man-hour needed to be allocated to a particular batch. This is a decision variable because it depends on the size of the batch, and the size of the batch is determined by the binary decision variables. This man-hour decision variable is continuous. Also, the resulting number of man-hours will be interpreted by the company based on their employee work hours policies. Thus, it does not need to be an integer.

The Model

$$\text{maximize } \sum_{i,j \in W} Z_{ij}C_{ij}$$

subject to;

$$\sum_{i \in W, m \in T} (X_{ik}N_{im}r + D_mL_{mk}) = M_k, \quad \forall k \in K \quad (1)$$

$$\sum_{k \in K} M_k + E = M \quad (2)$$

$$B \geq M_k \geq 0, \quad \forall k \in K \quad (3)$$

$$\beta \left(\sum_{k \in K} X_{ik} - \sum_{k \in K} X_{jk} \right) \geq 0, \quad \forall i \neq j \in W \quad (4)$$

$$\theta \left(\sum_{k \in K} X_{ik} - \sum_{k \in K} X_{jk} \right) \geq 0, \quad \forall i \neq j \in W \quad (5)$$

$$\alpha \left(\sum_{k \in K} X_{ik} - \sum_{k \in K} X_{jk} \right) \geq 0, \quad \forall i \neq j \in W \quad (6)$$

$$X_{ik}N_{im} \leq L_{mk}N_{im}, \quad \forall i \in W, k \in K, m \in T \quad (7)$$

$$\sum_{k \in K} X_{ik} \leq 1, \quad \forall i \in W \quad (8)$$

$$Z_{ij} \leq 0.5(X_{ik} + X_{jk}) \quad \forall i \neq j \in W, k \in K \quad (9)$$

$$\sum_{k \in K} X_{ik}N_{im} \geq L_{mk}, \quad \forall m \in T, k \in K \quad (10)$$

$$E \leq E_{max} \quad (11)$$

- The objective of the model is to maximize the commonality of the work orders in the batches while considering the following constraints:
- Constraint 1 ensures that the total man-hour needed for the collection of one batch is equal to the determined man-hour for that specific batch.
- Constraint 2 ensures that the sum of man-hour needed for all the batches in that day is equal to the allocated man-hour parameter given by the company for that day.
- Constraint 3 is giving upper and lower bounds to the M_k decision variable to have a less volatility among the decided manhours for batches.
- Constraint 4 and 5, 6 make sure that the priority, ATME and SDD requirements are all considered during the batching process.
- Constraint 7 ensures that all materials included in work order i should also be included in the batch that contains work order i .
- Constraint 8 ensures that a work order cannot be batched into more than one time.
- Constraint 9 is a linearization constraint. By including this constraint, we are ensuring that Z_{ij} will take value one if and only if X_{ik} and X_{jk} are both one, and it will be zero otherwise. Normally, we would need two additional constraints in order to linearize the multiplication of two binary variables. However, because of the nature of our model, this one constraint is sufficient. This is mainly caused by the objective function. Since it will try to maximize the total commonality, it will try to assign one to X_{ik} for all i and k . Moreover, X_{ik} and X_{jk} do not need to be considered separately in this case, since i and j belong to the same set, W .
- Constraint 11 makes sure that total nonoperational time is less than user defined parameters.

3.2 Heuristic Algorithm

After running the model with a larger data set, a portion of the real data of Aselsan that includes one hundred work orders, reasonable results were obtained. Even though the results were satisfactory, we realized that the model had a lengthy run time, most of which was actually caused by the generation of parameter matrices mentioned above. Because of the complexity of the pairwise comparison of each work order, matrix generation proves to be a tedious process. Moreover, Aselsan deals with as much as two thousand work orders each day. With a data set this sizable, the model seems impractical. We implement a heuristic algorithm in order to solve the problem. The algorithm operates in two parts. The first part acts as a filtering mechanism for the work orders and the second part includes a grading function and a tabu search function (The algorithm can be found in Appendix A).

The Filtering Preprocess

The filtering preprocess is the initial operation which is conducted on the work orders in the Excel sheet that is provided by the company. The sheet includes all work orders, with their priorities (being 1 if the work order has priority, and 0 otherwise). Another column has the value of the current date minus the SDD for each work order. The next column has similar values but instead, ATME date is used. Other work order information is also present in the sheet, such as the number of materials included in each work order, etc. An Excel VBA code is written in order to sort the data in the sheet. The data is sorted firstly based on the priorities of each work order. The second and third sorting criteria are the SDD column and the ATME column, respectively. For example, the first work order in the sorted data has priority, the SDD is relatively old, and so is the ATME date. Thus, by doing that, we sorted the work orders based on their eligibility for selection. As we move down the sorted data, eligibility of work orders for selection decreases. Then, we choose all work orders that are due to be delivered in the next three months. This corresponds to approximately one third of all work orders. By doing the filtering process, we manage to eliminate the work orders that would most likely not be selected and batched in a given day. By decreasing the number of work orders that will be used in the grading function and the tabu search, we greatly reduce the run time.

The Grading Function

After the filtering preprocess is completed, the filtered data is fed into the second part of the algorithm. In this part, each possible solution is evaluated using the grading function and modified using the tabu search algorithm. The tabu search algorithm will be explained in further detail in the following section. Initially, each batch is filled with ten work orders, starting from the top of the sorted data. That means, first ten work orders in the data set are put in the first batch. These batches are stored in a list of lists object called solution. This object is a list that contains n lists, where n is the maximum possible number of batches plus one. So, each list in solution denotes a batch, except the last one. The extra batch can be considered as an overflow batch and it includes the work orders that are not selected that day. Then, the solution is graded. There are two different grading options in the function and the user is free to choose among them.

Time-Based Grading

In this option, each solution has a numerical grade. This grade is calculated by adding several weights to the grade that represent the requirements provided by Aselsan. As mentioned before, work orders that have priority must be selected on any given day. That is why, the grading function adds w_1 to the grade for each selected work order that has a priority. Similarly, numbers $w_2 \times (\text{current date} - \text{SDD})$ and $w_3 \times (\text{current date} - \text{ATME date})$ are added to the grade for each work order. Since we want to include work orders that have older SDD's and ATME dates, the function captures this by increasing the grade in such a way. Moreover, the

time-based grading function includes a proximity evaluator. This part calculates the address location proximities of each work order in each batch and returns a numerical value. Then, $w_4 \times \text{proximity value}$ is added to the grade, since we want to maximize the address location proximities of work orders in each batch. In this option, the selection and batching of work orders are constrained by time limitations. Each batch must be collected within a specific time and the workers in the warehouse cannot work more than a specific hour during any day. These specific times (batch time and total time is used as a jargon in the project) are determined by the user. If the solution at hand does not satisfy these constraints, the function returns a negative number which represents infeasibility.

Material-Based Grading

In this option, w_1, w_2 and w_3 remain the same as they were in the time-based grading. However, the proximity calculation is done in a different fashion. Instead of calculating the address location proximity of each work order in each batch, material commonality is calculated for each batch. This is done in the following way: for each batch, the number of common unique materials is calculated, then this number is divided to the total unique materials in that batch. This corresponds to a numerical value that is between zero and one. After that, w_5 commonality ratio is added to the grade for each solution. Furthermore, feasibility conditions are also different from the time-based grading. In the material-based grading option, the constraints that determine feasibility are in terms of number of materials. The number of materials that are collected in each batch cannot exceed a specific number. Similarly, there is an upper limit to the total number of materials to be collected during the day. Again, these numbers are determined by the user.

The Grading Function Parameters

After extensive testing and analyzing, reasonable values for the parameters in the grading function are determined. Since priority is the most important attribute of a work order, w_1 is the largest weight. When calculating the grade of a solution, w_1 needs to dominate the other parameters. Between SDD and the ATME date of a work order, SDD is more important. In order to reflect this relationship, w_2 is twice as big as w_3 . This is because w_2 needs to affect the grade more than w_3 while not completely dominating it. Considering these and testing the parameters several times, having w_2 twice as big as w_3 shows to be the best possible relationship between ATME date and SDD. Furthermore, w_4 and w_5 symbolize similar constraints in their respective grading systems. That is why, they have the same value. They do not have to dominate each other nor them having the same value affects the grade in any way, since they are never used together. Both w_4 and w_5 are significantly larger than w_2 and w_3 , but still smaller than w_1 .

Tabu Search

Tabu search part of the algorithm is where we modify the solution on hand. This part starts with the initial solution mentioned above. The function randomly chooses a work order and either swaps it with a different work order or removes it from the current batch and places it into another. After each operation, a new solution is obtained. Then, the new solution's grade is recalculated using the grading function. If the operation increases the grade of the solution, this specific operation is stored in a tabu list. This operation cannot be done again for the duration of tabu tenure, which is a number that can be changed on will. Thus, tabu search algorithm increases the grade of the solution by modifying the solution using the mentioned operations. This part of the algorithm runs iteratively until the specified running time is reached. The user determines for how long the algorithm will run. It is noticed that longer run times result in better solutions.

3.3 Verification and Validation

Verification of the heuristic and the model ensured that both model and heuristic functions properly and provides us solutions we look for. For validation, we obtained updated and organized data from Aselsan. We used this real data to run the heuristic algorithm and the mathematical model. In addition to that, we obtained sorted and formulated batches by responsible engineer, using same data. In the process of implementation, the first stage was the comparison of the model and the heuristic by using 100 work orders. For this comparison we used material based grading. We observed that, solution obtained by the model has 2,365 materials collected daily while the solution of heuristic contains only 2,046 materials which show us that there is a 13.49% gap. Since the outcome of the heuristic algorithm depends on the number of trials in a run, the output can be improved simply by extending the run time. On the other hand, the grade increased by 43.37% because in the heuristic algorithm verbal constraints can be described better. We also compared outputs from the heuristic run with real data, to the batches Aselsan is currently formulating with the same data in order to observe the improvement done by the proposed heuristic algorithm. Both time based grading and material based grading was used during comparison. When we used material based grading the grade of the solution obtained by the heuristic increased by 30.04% compared to the grade of the current formulated batches. As it was stated before these results can be improved by running the heuristic multiple times and increasing run time. Also, the number of materials collected during the day in the current system was 1,315 while the solution obtained by heuristic contains 1,360 materials which again shows that there was a 3.31% increase. This part depends mainly on our three fundamental constraint; priority, statistical delivery date and sub-material completeness date. If data has large number of prioritized work orders or work orders can be sorted almost perfectly chronological. Next, when we used time based grading and compared the heuristic to the current used system, improvement in the grade was 35.19%. Similarly,

we observed that collected daily materials increased from 1,315 to 1,380 which is an increase of 4.71%. We can see there is a noticeable difference and it is proven that the system is efficient and will increase the daily amount of collected materials.

3.4 User Interface

In the phase we built the user interface, we tried to keep it as simple and useful as possible in order to provide the most user-friendly application we could. The interface which we have built via Python, consists of three different pages; on the first page the user can provide the inputs, and also we ask them to search for the input excel file from the computer (as it can be found in Appendix B). An important remark is that if the user chooses “Materyale Göre”, the inputs related to “Zamana Göre” becomes inaccessible and same is valid for the reverse situation. After the user enters inputs, they push “ÇALIŞTIR” button, after they push the first page closes and the second page comes. On the second page, we have a “BAŞLAT” button which makes the process start when pushed. After pushing “BAŞLAT”, the heuristic becomes its running process in the background. While the heuristic runs, the progress bar which we put as an animation also becomes running with respect to the running time of the code. There are also two push buttons, “ÖNCEKİ” makes the user go to previous page and stops the code in order to take new inputs, and “SONRAKİ” button is pushed after the heuristic is done and progress bar comes to %100 to move forward to the third page. On the last page, we have the guidance sentence which says that the user can find their final plan in the file. There is one push button in the bottom right of the page which is “ANA MENÜ”, when it’s pushed the third page closes and the first page comes for the user to enter the inputs again. We added that option in case of situations that if the user is not satisfied with the outputs or the plan, they can change the parameters and run the heuristic again.

4 Conclusion

As the contributions to the project, by using tabu search algorithm and a grading algorithm which we built to be able to pick the work orders of the day to be handled that day with respect to their Priority, ATME and SDD status, we provided a new, rational and automated system while picking the work orders to be handled that day. By using our algorithm, ASELSAN will be able to end this process with an increased rate of efficiency and decreased rate of labor spent by assigned engineer.

Appendix A Heuristic Algorithm



Appendix B User Interface



Off-Site ATM Temizliđi Rota Planlaması

TEPE Servis A.Ş.

group21/group_photo.png

Proje Ekibi

Eylül Akkaş, Ahmet Selman Korkmaz, Emin Gökalg Öz,
Enes Özen, Cafer Eren Şimşek, Furkan Yasin Taksim, Enes Uğur

Şirket Danışmanı

Tolga Türkeroglu
Tepe Servis ve Yönetim Sistemleri San.
A.Ş.
Bölge Müdürü
Endüstri Mühendisliđi Takımı

Akademik Danışman

Halil İbrahim Bayrak
Bilkent Üniversitesi Endüstri
Mühendisliđi Bölümü

ÖZET

Tepe Servis ATM temizlik hizmetleri aylık 13.000 ATM ve 120.000 temizlik işinden sorumludur. Her temizlik çalışanı kendisine atanan günlük işleri takip etmeleri için Ontime adı verilen bir mobil uygulama kullanılır. İşçilerin tamamlaması gereken bir dizi görevi vardır ve zamanında tamamlanmayan işler için önceden belirlenmiş ceza ücreti ödenir. İşçilerin tamamlayacakları görevi seçmek için takip edebilecekleri rehber yoktur. Şirketin şikayetleri; sistematik olmayan iş seçimi ve tamamlanmamış görevleri ertesi güne devredememesidir. Bu proje, tamamlanmamış iş sayısını ve dolayısıyla ödenen ceza maliyetlerini en aza indirmeyi ve tamamlanmamış işleri bir sonraki döneme aktaran bir sistem geliştirmeyi amaçlamaktadır. Yöntem olarak Araç Yönlendirme Probleminin bir uygulaması kullanılacaktır. Geliştirilen algoritmanın geçmiş dönem bilgileriyle karşılaştırılmasıyla Mart 2021'de yapılan silimlerin bizim algoritmamızla %54 daha az yol katederek %36 daha fazla başarı oranına ulaşacağı gözlemlenmiştir.

Anahtar Kelimeler: İş Planlama, Araç Yönlendirme Problemi, Off-Site ATM Temizliđi

Planning Routes for Off-Site ATM Cleaning

1 Company Information

Tepe Servis and Management Inc. was established in 2008 by Bilkent Holding as a cleaning services company. Today, they continue their journey as an integrated facility management company (Tepe-Servis). Tepe Servis provides several services such as Facility Management Services, Cleaning Services, Special Integrated Projects, and Call Center Services (Tepe-Servis).

One of Tepe Servis's Cleaning Services is Off-Site ATM Cleaning Service. Tepe Servis operates 5 different banks' Off-Site ATM cleaning in 11 cities of Turkey. Moreover, they supervise remaining cities and KKTC served by the subcontractors. This service consists of 13.000 ATMs and approximately 87.000 cleaning jobs per month. Automobile or pedestrian teams provide the service in approximately three minutes per job. Besides ATMs, the company offers cleaning services for signboards and branch windows. Extra services take forty-five minutes on average.

2 System Analysis

2.1 Current System

In the current system, Tepe Servis signs a contract to settle the service details while settling with a bank. Visit frequencies for each ATM are specified in this contract. The ATMs can change depending on the contracts with the customers. Sometimes, the customers request branch windows and signboards to be cleaned as well. Besides the jobs settled in the contract, the customers may request urgent jobs for ATMs that are not on the schedule. These additional tasks and the time windows of the urgent demands are also stated in the contract. Unfinished jobs during the period cause penalty cost to the company. Penalty cost is another detail settled during onboarding, which may vary from 100% to 500% of a cleaning job cost per incomplete job.

The company uses a mobile application called Ontime to assign jobs to teams. Tepe Servis has formed ATM pools in each region heuristically using their expertise. The teams receive job assignments out of these pools. They can follow the jobs assigned to them for a given period via the app and mark the tasks completed. Each period, which usually consists of 3 days, the system automatically assigns jobs in a non-ordered, non-scheduled fashion. The teams are free to choose in which order and on which days they complete the cleaning jobs. The workers take a photograph of the clean ATM and upload the image to the system via the Ontime application.

2.2 Problem Definition

Unsystematic selection of jobs causes longer routes traveled, fewer jobs completed, and a high amount of fuel consumed. Furthermore, incomplete jobs do not transfer

to the next period. Since new jobs are assigned each period, incompleting jobs cannot be integrated into the program of new tasks. This issue increases the number of incompleting jobs and hence, the amount of penalty issued.

The performance measure is the number of monthly incompleting jobs. Both of the issues addressed lower the performance. Another metric is the cost per job. Parameters such as the number of crews working, distance traveled, number of hours and shifts worked, etc surge the cost. Since penalty costs are proportional to the number of incompleting jobs, the main objective of this project is to plan routes for the cleaning teams. Thus, facilitate workers to complete more tasks than they do in the current system and decrease the cost per job.

The company's completion rate is close to 100% on most months. However, since the workers decide which routes to follow without a systematic approach, high completion rates come with a high cost. Hence, the company has set two main goals for our project team:

- Decrease the number of incomplete jobs.
- Reduce the cost per job.

3 Proposed Tools and Methodology

This section defines the problem as well as all the tools and methods we are using to solve it. We give most of our literature review and preliminary benchmarks here.

3.1 Solution Description

Cost per job is proportional to the length of the routes. Shorter routes mean less fuel consumption and depreciation as well as more completed jobs. Hence, the ideal solution is to come up with an algorithm that will lead to shorter routes.

The routing problem is a variant of the Vehicle Routing Problem (VRP). The objective of VRP is to minimize the total route cost for a given interconnected node network and vehicles (Steinbach et al. (2018)). The emphasis on the definition is "given node network". While we know every node, ATM, in a given month and their required visit frequencies, it is uncertain on which day or shifts a node should be visited. It is possible to model this problem using Vehicle Routing Problem With Time Windows (VRPTW). VRPTW introduces time windows to each node in which they have to be visited, and possibly a visit duration as well. VRPTW is conventional in the modeling of problems with distinct nodes such as customers with varying service needs (Tan et al. (2006)). Then, this approach yields another problem: How do we assign time windows to identical nodes? Instead of assigning hourly intervals to the nodes, we would have to specify a day, or shift, in which that node has to be visited. That is the equivalent of solving a VRP for each day with different scheduled nodes. This approach would yield more but smaller prob-

lems compared to VRPTW. VRP is said to be an NP-complete problem, that is, time to solution increases outstandingly as the problem gets larger. Hence, more but smaller problems are favourable to one big problem (Karkula et al. (2019)). Then, it is unsuitable to use VRPTW when it may be possible to achieve the same results with much smaller problem sizes. A disadvantage to this approach is by separating the problems, we are forgoing possible solutions because they do not exist in the same space anymore. In an effort to keep the desirable routes in our global solution pool, this separation of jobs must be done carefully. That is to say, in a singular problem, we should preserve the nodes that would be scheduled together and forego the ones that would be eliminated by the solver. This yields a primary problem to solve before optimization: assignment of jobs to days.

To recap, we define the problem as the assignment and routing of jobs repeating every week. Each period, the data will be updated, and incomplete jobs will be rescheduled.

3.2 The Clustering Algorithm

To schedule nodes, we will first cluster the close data points together and make sure they are appointed to the same days. This helps us eliminate unfavorable routing options that the solver would have checked and waste time while restricting our search to a set of favorable ones.

Our literature search yielded two alternative clustering algorithms: k-means clustering and Density-based spatial clustering of applications with noise (DBSCAN).

- **k-means Algorithm:** "k-means clustering is a prototype-based, partitional clustering technique that attempts to find a user-specified number of clusters (k), which are represented by their centroids" (Bostel et al. (2008)). This algorithm is used to divide N observations into k clusters by minimizing variance. Due to the initial seeding mechanism it implements, k-means algorithm tends to over-represent nodes that have many nodes around, which in turn causes some nodes to not appear in clusters" (Boeing (2018)).
- **DBSCAN:** DBSCAN clusters nodes according to the Haversine distance between nodes and the minimum sample size by minimizing geodetic distance. Using Haversine distance over Euclidean distance used in the k-means algorithm essentially yields more realistic distances by taking Earth's curvature into account. The algorithm helps reduce distortion on points far from the equator. DBSCAN allows classifying some nodes as noise as well, however, we want to schedule each node (Borah and Bhattacharyya (2004); Boeing (2018)).
 - **Haversine distance:** Haversine distance calculates the shortest distance between two geolocation by taking Earth's curvature into account (Scikit). The formula can be found in Appendix A.

Due to the above, we concluded that DBSCAN is a superior method when it

comes to the classification of spatial latitude-longitude data. After some trial and error, we set our maximum distance constraint to 500 meters and used ball tree space-partitioning data structure to calculate great-circle distances between the points (Marius (2020)).

3.3 The Solver

The problem requires re-solving the model regularly every week for each pool's each workday. That accounts to initiating the solver thousands of times each month. Since we are trying to deliver a cost-free solution, we decided to use an open-source software as the solver. We conducted a literature search to find possible tools and to construct our criteria (Karkula et al. (2019)). We considered VROOM, JSPRIT, Google OR-Tools, OptaPlanner, and reinterpretcat's A Vehicle Routing Problem Solver (Builuk (2021)).

Our criteria for tool choice was:

- **Speed:** Since the model will run every week, the number of hours to initial solution is a crucial criterion and needs to be as low as possible.
- **Output quality:** The project's aim is to decrease the number of incomplete jobs and will be achieved by improving the routes. Hence, the solution quality is important.
- **Model flexibility:** We need a tool that could solve various flavors of VRP as well as allow us to introduce custom constraints and/or an objective function.
- **Documentation quality:** We need a tool that would make it easy to build and update a solution according to changing requirements. To achieve this flexibility and to fully utilize the tool, a comprehensive documentation is important.
- **Tool's programming language:** The software part of the project is not the objective but the instrument. Therefore, we wanted a tool written in a language that is easy to write, fast, and easy to integrate into the existing infrastructure in the form of a web server or a GUI application.

At the end of our investigation, we decided to use Google OR-Tools because while its engine is written in C++, it is possible to interact with it using Python, making it a fast and easy-to-use solver (Surana (2019)).

- **Google OR-Tools:** "OR-Tools is an open-source software suite for optimization, tuned for tackling the world's toughest problems in vehicle routing, flows, integer and linear programming, and constraint programming" (OR-Tools).

3.4 The First Solution Strategy and the Local Search Metaheuristic

The next step after choosing the solver is choosing solution algorithms. OR-Tools offers 12 different First Solution Strategies and 6 different Local Search Metaheuristics. Thus, we have 72 different combinations on top of 12 different first solution strategies with no metaheuristic, 84 choices in total. Fifty-two of these configurations are suitable to the problem at hand and were examined in detail.

- **First solution strategy:** A first solution strategy is an algorithm used by the solver to find a not necessarily optimal initial solution that satisfies all of the constraints. The solution is often suboptimal for large problems ($n > 10$).
- **Local search metaheuristic:** Local search metaheuristics sit on top of the strategy and alter its behavior (its objective function) to find the best solution out of candidate solutions by moving locally in the search space until a predetermined time limit is reached or a predetermined number of iterations made.

We examined existing benchmarks, conducted ours using different algorithm combinations on various sample sets provided by Tepe Servis, and a literature review to confirm that our benchmark results fit the theoretical expectations as well as other experiments in the literature.

OR-Tools Benchmarks

The first step was to select timeout settings for the metaheuristic regardless of which algorithm we use. Timeout is the maximum duration the algorithm will look for better local solutions in each iteration. We've experimented with different timeout metrics for all of the metaheuristics. The solution time and average total weekly route distance data are given in Table 1.

We determined that a metaheuristic with a large enough timeout returns good results compared to strategies without metaheuristics. Our computational results suggest that a time limit beyond 90 seconds has little impact on the algorithm's performance, and hence, we set it to 90 seconds. Moreover, we have found that OR-Tools' VRP solver found better solutions than best known solutions for various benchmark problems (Karkula et al. (2019)).

Table 1: Time and Average Performance Comparison of Metaheuristic Local Search Timeout Settings

Timeout (secs)	Solution Time (mins)	Total Distance (kms)
10	1.02	213.6
30	3.02	200.0
90	9.02	196.3
120	12.03	197.3
-	53.97	192.1

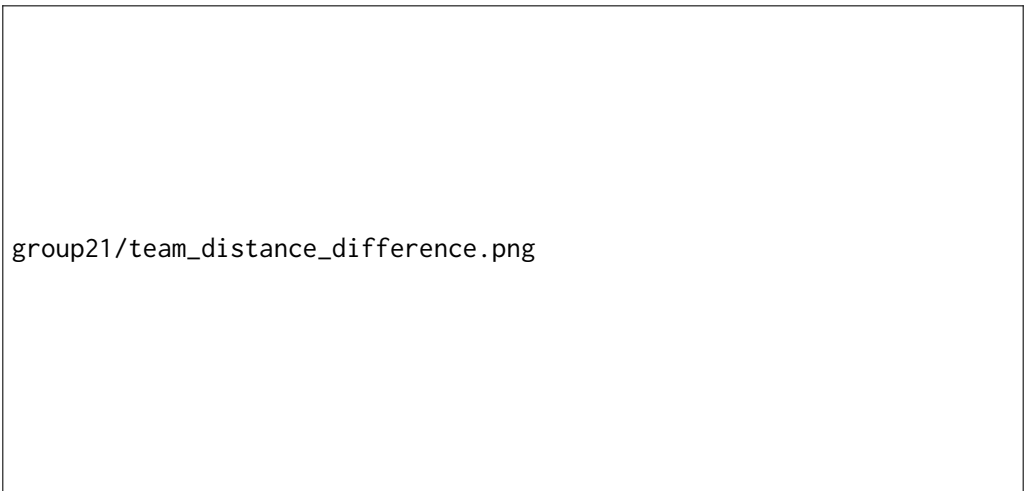
Next, we choose the algorithms. While there is no best First Solution Strategy - Local Search Metaheuristic combination, our benchmarks overwhelmingly performed better with Local Cheapest Insertion strategy - Guided Local Search metaheuristic combination. Figure 1 shows the performance comparison in terms of cumulative weekly total distances of all teams in Ankara. Each group of histogram represents a First Solution Strategy and each column in a histogram represents a Local Search Metaheuristic.

group21/algo_performance.png

Figure 1: Total cumulative distances comparison (km)

As visible, the Local Cheapest Insertion strategy is superior to others. Among the 5 metaheuristics applied, 3 of them performed similarly in terms of the total distance. Hence, we examined another metric: the deviation in team weekly loads. In this experiment, there were two teams, therefore taking the difference of their total distances would be enough to compare different configurations. Figure 2 demonstrates how the absolute value of team loads differs across algorithms.

While the Global Cheapest Arc strategy yields a slightly lesser difference between team loads, Local Cheapest Insertion’s total distance performance makes up for it. Among the 3 metaheuristics that performed similarly, Guided Local Search is the



group21/team_distance_difference.png

Figure 2: Difference Between Team Loads (km)

best in both metrics. Hence, the selected configuration is Local Cheapest Insertion strategy with Guided Local Search metaheuristic. We will examine these further in our literature review and confirm that they are suitable for the problem.

Literature Review

- **Local Cheapest Insertion:** "Iteratively build a solution by inserting each node at its cheapest position; the cost of insertion is based on the arc cost function" (OR-Tools).

The arc cost function is the function that calculates the total route length. In our problem, we are not just looking for the shortest routes but balanced routes in terms of the number of nodes assigned to each vehicle. Hence, our arc cost function also increases a significant amount whenever a node is added to a route. In other words, there is a high extra cost for inserting a node into a route. This makes sure that the algorithm does not stack too many close nodes into one vehicle's schedule. The absolute value of the differences between the number of nodes assigned to each vehicle can be found in Appendix B. Due to our clustering algorithm and modified arc cost function, the difference is always one since the number of nodes is odd and there two vehicles.

The even distribution is important because the cleaning jobs also take time, hence, shorter routes by themselves do not necessarily mean more completed jobs.

- **Guided Local Search:** GLS applies penalties to the objective function according to selected features and their assigned penalty values. A feature in our case is if two nodes are directly connected. Whenever the strategy reaches a local optimum, GLS applies the penalties to these unfavorable fea-

tures. If a penalty is applied, the penalty value for that feature is increased to keep the algorithm from penalizing the same features over and over again. Doing so, GLS focuses its search on areas that have favorable features instead of focusing on a cost-effective local area (Voudouris C. (2003)). For the mathematical representation of GLS, see Appendix C.

Some of the visualized routes of a 50 node, 2 vehicle sample problem can be found in Appendix D.

4 Final Deliverable and Conclusion

4.1 The Deliverable

On top of the scheduling and optimization algorithms, the deliverable requires a connection to Tepe Servis' database and a web server. These are also delivered along with a comprehensive benchmarking module that can be used to measure performance if modifications are made to the algorithm in the future. Ontime servers will be connecting to ours to get the weekly routes and display our output to the users. So, there is no user interaction required, thus there is no Graphical User Interface. Ontime can specify more than one pool to calculate the routes for. In that case, we merge the nodes and the vehicles in that pool into one problem.

We do not cluster all nodes together but we split them by weekly frequency first because the weekly frequency of a node affects which days they are going to be scheduled to and we want to schedule closer nodes (clusters) together. Our assignment algorithm loops through each cluster and looks for the most suitable day that the cluster can be assigned to. The most suitable day is defined to be the day with the least amount of nodes assigned to it. Depending on the weekly frequency of the cluster, different possibilities are considered while looking for suitable days. For instance, the nodes do not get cleaned two days in a row, or nodes that will be visited twice a week will be visited at least once by Thursday. Once a cluster is assigned to a day, all the nodes in it are added to the day's schedule. The algorithm evenly apportions nodes to days while ensuring that the nodes that are assigned to a day are close to each other. Once the assignment is done, each day's distance matrix is created. There is an extra dummy node in each matrix that has 0 distance to all other nodes as the depot. This means that start and endpoints do not affect the final solution. This is because Tepe Servis does not operate from a depot. Finally, these matrices are fed to the solver and the formatted routes are sent in an HTTP response.

4.2 Results and Simulation

Simulation Environment

As we have already discussed, company's job completion rate is already close to 100% and our main goal is to decrease the cost per job while getting rid of the occasional incomplete jobs. Since we cannot test our routes in the real world,

we created a simulation environment to test the performance of the algorithm. Some of the past routes the crews took are also simulated in the same environment. In this environment, we define the random parameters given below and take the routes as the input. There are three routes that we simulate on Ankara-3 pool:

- The routes followed by Tepe Servis in December 2020
- The routes followed by Tepe Servis in March 2021
- The optimized routes

We run the simulation a thousand times and average the results. The success measure is the percentage of visited nodes and the total distance travelled.

The parameters of the simulation environment are:

- **Speed:** The speed of the vehicle between each node is a uniformly distributed random variable. We have tested several intervals between (25,45) to (40,60). Our tests revealed that the improvement rate is not heavily dependant on the speed. Therefore, we chose the interval (25, 45).
- **Total Time:** Tepe Servis' typical workday is 8 hours. However, to allow teams a slack time, we run the simulation for 7 hours. The rest is reserved for unforeseen conditions. The simulation always finishes the last job even if it exceeds the time limit.
- **Delay Time:** On top of overall slack time, we introduce random delays throughout the day. There may be delays up to 4 times, again distributed uniformly, 15 minutes each. These delays represent the time spent looking for parking space, waiting for bank customers to finish using the ATM, traffic etc.
- **Cleaning Time:** Cleaning time is set as 3 minutes on average by Tepe Servis. However, we also take it as a uniform random variable between 3 and 5 minutes to account for possible slight delays.
- **Lunch Break:** There is a lunch break everyday that takes a value between 45 and 75 minutes randomly.

As a result of the randomness introduced, we are measuring the performance of the model for 4.75 to 6.25 hours workdays instead of 8.

Results and Analysis

The simulation results given in Table 2 indicate that the proposed routes result into more completed jobs and shorter traveled distances. More specifically:

- 57.6% fewer kilometers travelled and 32.0% more completion rate compared to December.

- 54.53% fewer kilometers travelled and 36.3% more completion rate compared to March

Table 2: Simulated Performance Comparison of the Generated and Historic Routes

Followed Route	Completion Rate (%)	Total Distance (kms)
December	73.7	13695.7
March	71.4	12779.6
Optimized	97.3	5812

4.3 Further Suggestions

Below are some suggestions that may be considered to further improve the model.

- A dynamic database to properly update monthly cleaning jobs left to make sure incomplete jobs do not stay that way throughout the month. Given this data our scheduling algorithm can already implement dynamic computation.
- Revision of the pools. We suggest merging pools that do not have operational contradictions to get better results out of the algorithms. As mentioned above, our current model can merge some pools if the company specifies it.
- Usage of a distance matrix API. Our current distance calculation methods while widely used in academia and projects, are not sufficient for a business-focused model. We recommend using a paid API to create real distance matrices that account for parameters such as city structures and traffic.

References

- G. Boeing. Clustering to reduce spatial data set size. *SocArXiv Papers*, pages 1–7, 2018.
- B. Borah and D. Bhattacharyya. An improved sampling-based dbscan for large spatial databases. *International Conference on Intelligent Sensing and Information Processing, 2004. Proceedings of*, pages 92–96, 2004. doi: 10.1109/ICISIP.2004.1287631.
- N. Bostel, P. Dejax, P. Guez, and F. Tricoire. Multiperiod planning and routing on a rolling horizon for field force optimization logistics. In B. Golden, S. Raghavan, and E. Wasil, editors, *The Vehicle Routing Problem: Latest Advances and New Challenges*, pages 503–525. Springer US, Boston, MA, 2008. ISBN 978-0-387-77778-8. doi: 10.1007/978-0-387-77778-8_23.
- I. Builuk. A new solver for rich Vehicle Routing Problem, 2021. URL <https://doi.org/10.5281/zenodo.4624037>.

- M. Karkula, J. Duda, and S. Iwona. Comparison of capabilities of recent open-source tools for solving capacitated vehicle routing problems with time windows. *AGH University of Science and Technology*, pages 1–6, 2019.
- H. Marius. Tree algorithms explained: Ball tree algorithm vs. kd tree vs. brute force. *Towards Data Science*, pages 1–11, 2020.
- OR-Tools. Google OR-Tools Routing Optimization Suite, <https://developers.google.com/optimization/routing>, Accessed: 2020-12-01.
- Scikit. Haversine Distance, https://scikit-learn.org/stable/modules/generated/sklearn.metrics.pairwise.haversine_distances.html, Accessed: 2021-03-15.
- M. Steinbach, T. Pang-Ning, A. Karpatne, and V. Kumar. *Introduction to Data Mining (Second Edition)*, pages 495–535. Pearson, 2 edition, 2018.
- P. Surana. Benchmarking optimization algorithms for capacitated vehicle routing problems. *Master’s Projects*, pages 1–53, 2019.
- X. Tan, X. Zhuo, and J. Zhang. Ant colony system for optimizing vehicle routing problem with time windows (vrptw). In D.-S. Huang, K. Li, and G. W. Irwin, editors, *Computational Intelligence and Bioinformatics*, pages 33–38. Springer Berlin Heidelberg, 2006.
- Tepe-Servis. Tepe Servis ve Yönetim A.Ş., <http://www.tepeservis.com.tr>, Accessed: 2021-04-14.
- T. Voudouris C. Handbook of metaheuristics. *International Series in Operations Research & Management Science, vol 57*. Springer, pages 185–188, 2003. doi: 10.1007/0-306-48056-5_7.

Appendix A Haversine Distance Formula

$$D(x, y) = 2 \times \arcsin \sqrt{\sin^2 \frac{x_1 - y_1}{2} + \cos x_1 \times \cos y_1 \times \sin^2 \frac{x_2 - y_2}{2}}$$

Appendix B Team Load Comparison in Terms of Number of Nodes

group21/team_load_difference.png

Figure 3: Team load comparison (number of nodes)

Appendix C Guided Local Search

$g(s)$: Objective function value for a candidate solution s

F : Set of features defined

p_i : Penalty value of a feature $i \quad \forall i \in F$

$I(s)$: Whether or not candidate solution s exhibits feature $i \quad \forall i \in F$

λ : Algorithm penalty tune parameter

$util_i(s_*) = I_i(s_*) \times \frac{c_i}{1+p_i} \quad \forall i \in F$: The utility of penalizing feature $i \quad \forall i \in F$

c_i : The cost of feature $i \quad \forall i \in F$

$h(s) = g(s) + \lambda \times \sum_i (p_i \times I_i(s)) \quad \forall i \in F$: The new objective function altered by the algorithm

Appendix D Sample Problem Routes



Figure 4: All nodes to be scheduled and routed

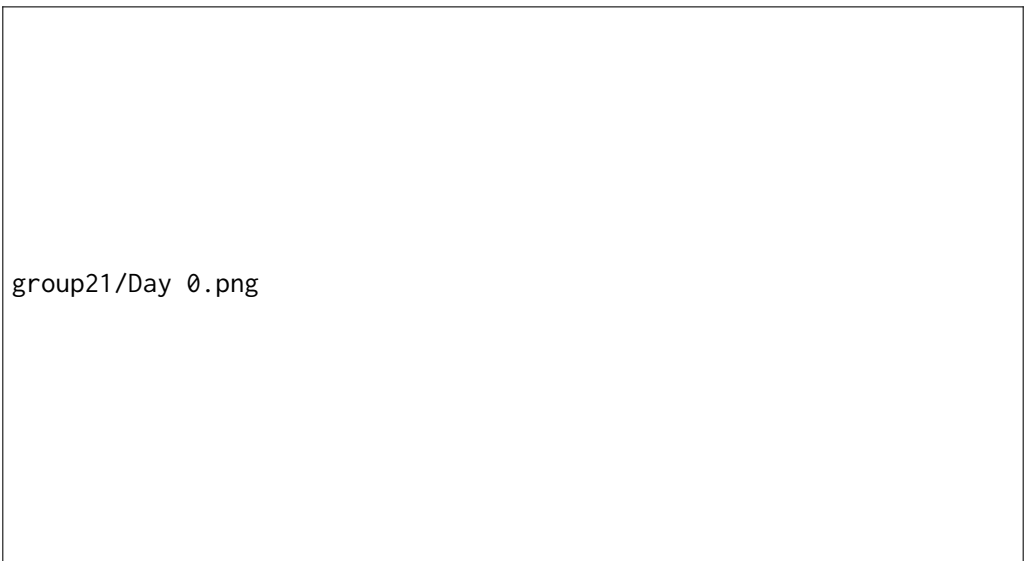
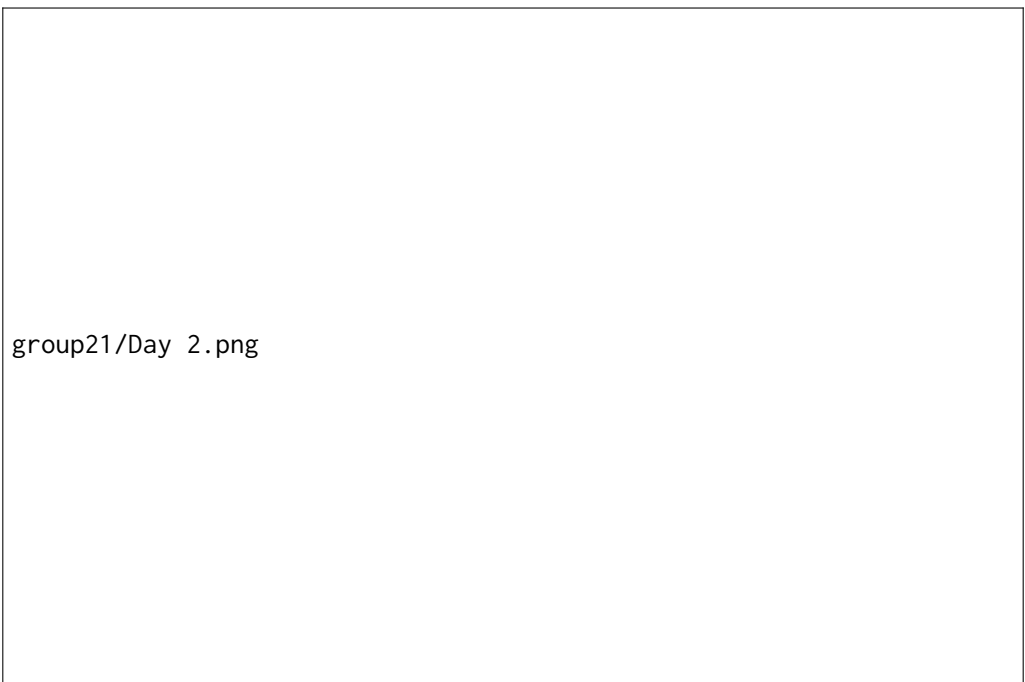


Figure 5: First day routes of Team 1



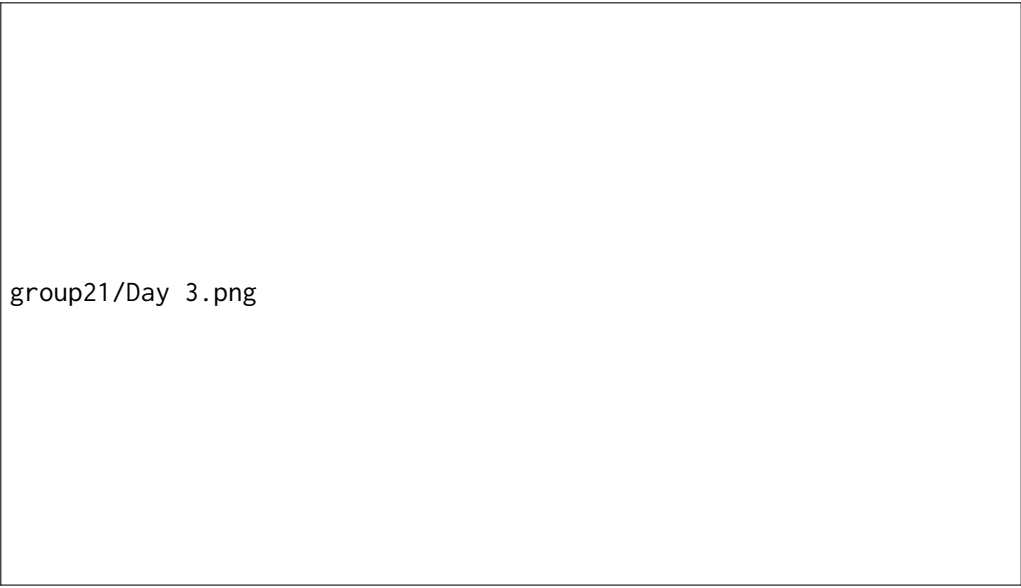
group21/Day 1.png

Figure 6: Second day routes of Team 1



group21/Day 2.png

Figure 7: Third day routes of Team 1



group21/Day 3.png

Figure 8: Fourth day routes of Team 1