

This is a public version. All sensitive data was removed.

Final Public Version

9/9/2013

Preventing Line Congestion at Aebi Schmidt Nederland

Predicting Cycle Times and Work in Process using Queuing Network Analysis

By Thomas Klijnsma, student Industrial Engineering and Management

Preventing Line Congestion at Aebi Schmidt Nederland

Predicting Cycle Times and Work in Process using Queuing Network Analysis

Author

Thomas Klijnsma (s1015362)

t.klijnsma@student.utwente.nl

Supervisor

Dr. Ir. Ahmad Al Hanbali

Second Supervisor

Dr. Jan-Kees van Ommeren

External Supervisor

Fred Harbers

University of Twente

School of Management and Governance

Educational Program

Industrial Engineering and Management

Aebi Schmidt Nederland

Handelsweg 6, 7451 PJ, Holten, The Netherlands

Tel. +31(0)548 370 000

Management Summary

This report is about preventing line congestion at the welding and surface process at Aebi Schmidt Nederland. Aebi Schmidt is active worldwide in traffic area cleaning and maintenance; the production facility in Holten, the Netherlands, produces snow clearance and de-icing machines.

The manufacturing process at hand is divided into six main processes: Tack welding, seam welding, blasting, primer coating, baking, and final paint coating. Products are baked in the oven after each paint coating. The different product types are categorized as top parts, bottom parts, tail pieces, frames, and integral designs.

Frequently, long queues arise at the welding and surface treatment process. In an effort to explain and provide possible solutions for this problem, the welding and surface treatment process is modeled as a queuing network. With queuing network analysis it is possible to analyze the current situation in terms of performance measures such as cycle time and work in process.

An Excel tool was developed to perform the queuing network analysis for both the current situation and for several scenarios, based on data collected from time registrations and inquiry. The following conclusions were drawn:

- The time registrations are prone to many errors and inconsistencies, which makes the data collection in this research more time consuming than necessary.
- <<Not available in public version>>
- Of the six tested scenarios, adding a blasting cabin was the most impactful but also the most expensive solution. Increasing seam welding work hours and decreasing seam welding process times were the most effective solutions.

The following recommendations were made:

- Improve time registration, by implementing forms or a tracking system.
- <<Not available in public version>>
- Attempt to decrease seam welding process times, or increase work hours at welding.
- Attempting to decrease outages at blasting will help prevent congestion, and is probably not a very expensive option. It is recommended to implement preventive maintenance or other methods to reduce variability in mean time to repair and mean time to failure.

Table of Contents

Management Summary	2
Table of Contents	3
Chapter 1: Introduction	5
1.1 Aebi Schmidt Nederland	5
1.2 Problem Description	6
1.3 Research Questions	6
1.4 Scope and Limitations	7
1.5 Methodology	8
1.5.1 Research Design	9
1.5.2 Data Collection	10
1.5.3 Data Analysis	10
1.6 Outline of the report	11
1.7 Conclusion Chapter 1	11
Chapter 2: Case Information	12
2.1 The Products	12
2.2 The Process	13
2.3 Conclusion Chapter 2	15
Chapter 3: Theoretical Aspects	16
3.1 Measuring Manufacturing Performance	16
3.2 Quantitative Models	19
3.2.1 Statistical analysis	19
3.2.2 Simulation	20
3.2.3 Queuing Theory	20
3.3 Observation	28
3.4 Conclusion Chapter 3	28

Chapter 4: Data Collection.....	29
4.1 Process Parameters	29
4.1.1 Data Cleaning	29
4.1.2 Results.....	30
4.2 Observations	30
4.3 Conclusion Chapter 4.....	30
Chapter 5: Queuing Network Analysis	32
5.1 Modeling the welding and surface treatment line	32
5.1.1 The Input.....	32
5.1.2 Modeling different shift hours	33
5.1.3 Single-class G/G/c Queuing Network and Disaggregation	34
5.1.4 Excel Tool	35
5.2 Current Situation Results.....	36
5.2.1 Aggregation <<Not available in public version>>	36
5.2.2 Single-class G/G/c Open Queuing Network Analysis <<Not available in public version>>	36
5.3 Validation	37
5.4 Sensitivity Analysis	37
5.5 Conclusion Chapter 5.....	41
Chapter 6: Conclusion & Recommendations.....	43
Bibliography.....	46
Appendix A: Data cleaning of the Time Registrations	47
Appendix B: Queuing Networks	48
Appendix C: Causes of Variability	50
Appendix D: Disaggregated performance measures	52

Chapter 1: Introduction

About the company Aebi Schmidt Nederland, the problem at hand and research methodology

In March 2013 3 other students and I started the second year Industrial Engineering and Management course Project 2: Organization and Process Analysis. We were assigned a case at Aebi Schmidt Nederland in Holten. The goal: Develop a method to predict cycle times.

Filled with enthusiasm we worked on the case for 2 months, and developed a Matlab simulation that was able to effectively predict cycle times for very specific states of the welding and surface treatment facility.

This study aims to improve upon this already performed research. Whereas a simulation was especially useful at the operational level of the plant, it does not actually solve the problem: the frequent occurrence of long queues. In this study, a tactical level analysis will be made of the welding and surface treatment manufacturing system at Aebi Schmidt Nederland.

1.1 Aebi Schmidt Nederland

Aebi Schmidt Nederland is a part of the Aebi Schmidt Group, the *“leading system provider of innovative technical solutions for the cleaning and clearing of traffic areas”* (Aebi Schmidt). Aebi Schmidt Nederland is mostly active in winter maintenance technology.

The winter maintenance facility in Holten was founded in 1949 as Universal NV, producing simple snow clearance tools. The first line of de-icing equipment was produced in 1955. In 1972 Universal NV fused with Nido, establishing Nido Universal Machines; even now the facility in Holten is commonly known as Nido. In 1983 Nido Universal Machines was adopted in the international Schmidt Group. After the fusion between the Aebi Group and the Schmidt Group in 2007, Aebi Schmidt was founded as it is known today.

Due to its historical association with snow clearance and de-icing, the production site in Holten comprises the winter maintenance section of the Aebi Schmidt Group. At the production site in Holten a broad range of spreading and spraying machines are produced. Generally speaking, the manufacturing process at Aebi Schmidt consists of welding metal parts, surface treatment, and assembling the machines.

1.2 Problem Description

The welding and surface treatment process is currently quite complex. The production steps are subject to intricacies such as human capacity constraints at tack welding, non-uniform size process batching in the oven, and material dependent production routes. Combined with a lot of different product types, this poses problems in the predictability of the production process, which in turn complicates production planning.

Demand for winter maintenance equipment is heavily influenced by seasonality. Production peaks during the summer, and then steadily declines to a minimum in the winter. For this reason popular products are welded to stock during the production low season. This flattens out the seasonality of the production rate, but further complicates the production planning, since customer orders should be given priority to stock orders.

The combination of a complex production process, a high product variety and heavy seasonal influence on demand makes production planning very difficult. Determining measures such as expected cycle time of new orders or work in process can be a complicated matter. Frequently new arrivals are admitted into the system while certain production stations are already highly utilized, which inevitably leads to a long queue. Since orders are of varying priority, queues have to be actively monitored to prevent late orders, which is time consuming and frustrating for both manufacturing staff and management.

The problem perceived by Aebi Schmidt is ‘unexplained long queues’: Too frequently long queues arise, not always at the same workstation. The research conducted in this study aims to assess the production process in terms of performance measures. This does not only make the current production process more insightful, it also opens up the possibility of finding the most impactful and realistic way to increase the efficiency of the line. It is intended to perform the analysis in such a way that it is easily repeatable; this is achieved by automation of the analysis of the collected data.

1.3 Research Questions

More concisely stated the perceived problem by Aebi Schmidt is as follows:

The production line that encompasses the welding and surface treatment is too often congested.

The research question is thus posed as follows:

How can long queues in the welding and surface treatment line be prevented?

In order to answer this question, several sub-questions are formulated:

1. *What are the characteristics of the current welding and surface treatment process?*
2. *What are the product characteristics, and how can different products be categorized?*

The first two sub-questions deal with the current characteristics of the process and the products. The intricacies of the manufacturing process are explained first. Second, the different types of products and product parts are defined and categorized in data groups.

3. *Which performance measures need to be determined to measure the current manufacturing performance?*
4. *How can these performance measures be determined using a quantitative model?*
5. *What are the input parameters of the model, and how can they be measured?*

The next sub-questions deal with the theoretical aspects. In order to improve the current situation, it is necessary to establish a set of measures that describe the performance, and a framework to determine these measures. Additionally, the input parameters for the model need to be clear.

6. *What are the current input parameters for the welding and surface treatment processes?*

This sub-question pertains to the data collection. Several data manipulation techniques are required to eventually estimate current input parameters.

7. *What are the current performance measures of the welding and surface treatment processes?*
8. *Can these results be validated by Aebi Schmidt?*

This sub-question will be answered by applying the model to the determined input parameters, and checking whether the found results are realistic.

9. *What scenarios could improve the performance of the process the most?*

The last question builds toward a conclusion. By changing the input parameters according to several scenarios, recommendations can be made to improve overall performance.

1.4 Scope and Limitations

This study aims to find an overall solution to prevent line congestion. It is not our goal to find solutions for specific situations that might occur in the production process, but rather to decrease line congestion in the first place, thereby raising performance. It is therefore at a 'tactical level' instead of an 'operational level'.

The goal of this study is to find solutions to prevent long queues in (parts of) manufacturing systems, in this case for the welding and surface treatment process at Aebi Schmidt. Preventing long queues leads to lower work in process, and hence less capital tied up in inventory. The benefits of reducing unnecessary work in process are self-explanatory. Analyzing the whole manufacturing process, including engineering and assembly, would be too elaborate. The analysis therefore considers only the welding and surface treatment process.

The scope of this research is limited to the science of operations management; other areas of research which might have significant impact on the manufacturing performance such as human resource management or strategic management are not covered in this study. Solutions are sought within the current situation and strategy of the company.

Although this study focusses on the situation of Aebi Schmidt, this study is meant to be generic, and can be used for other manufacturing cases. The methodology is therefore general. To allow easy repeatability of the analysis, the analysis will be performed using a computerized tool.

Due to time constraints it is necessary to choose a specific course of research. In this study, the welding process will be considered an independent system from the other facilities at the production site. That way data collection and analysis are manageable within current time constraints. In reality, products and parts will flow through the entire manufacturing system which includes engineering and assembly, but assuming independent systems allows reasonable accurateness while upholding manageability of the study.

1.5 Methodology

An important aspect of methodology in this study is repeatability: the research method employed for the current case should be extendable to other cases. Without loss of generality, the research is conducted with the standard research cycle (Heerkens & Van Winden, 2012):

- Problem statement
- Research questions
- Research design
- Data collection
- Data analysis
- Conclusion

Problem statement and research questions have already been covered. This paragraph will elaborate on the research design, data collection and data analysis.

1.5.1 Research Design

Concretely stated, the research design for this study encompasses the following:

- Analyze the manufacturing system in its current state in terms of predefined performance measures.
- Validate the used model for determining these measures with the company.
- Observe the production process.
- Make recommendations to alleviate the undesirable impacts of long queues based on the used model and observations of the production process.

This is the general approach for answering the main research question. The current study follows these steps. More specifically geared towards the case company, the research is designed as follows:

- Collect necessary data to measure performance measures: fill in missing data, filter out unusable data, and group data in relevant data groups. Technically Aebi Schmidt produces over 700 different products, but many products are very similar as far as welding and surface treatment are concerned. Because such a large product variety needlessly complicates an analytical model, it is necessary to group different products into general product groups.
- Use an analytical model to determine the performance measures
- Develop an analytical tool to calculate the performance measures, built so that recalculation for different parameters is relatively easy.
- Observe the production process for multiple days.
- Validate the tool with Aebi Schmidt.
- Use the tool to find ways to decrease the odds of long queues arising.
- Analyze the observations to identify which parameters of the manufacturing system are most easily changed.
- Recommend the most impactful ways to decrease the odds of long queues

The necessary knowledge to perform these steps include at least:

- A detailed description of the characteristics of the welding process as it is performed now.
- A concise overview of all classes of products that undergo the welding process, including (quantitative) information on process times and arrival rates.

- A balanced set of performance measures that accurately estimate manufacturing performance.
- A theoretical, mathematical framework to calculate performance measures based on the line characteristics and product information.

1.5.2 Data Collection

The required data can be classified in product data and line data: The first pertains to relevant data that differs per product class. Examples of this kind of data are process times, arrival rates and product routings throughout the line. The second pertains to relevant data that is specific for the manufacturing system at hand. Examples of this kind of data are the amount of processes, the amount of servers per workstation, and the typology of the production process.

The line data is well known at Aebi Schmidt. Collection of this type of data has already been performed in earlier reports (Dröge, De Jong, Klijnsma, & Reimert, 2013). From earlier visits to Aebi Schmidt it is known that a lot of product data is tracked: especially the welding process is documented fairly accurately. However, some data, especially some process times for surface treatment, are currently not tracked. For this data estimations will have to be made, based on interviews with experienced production personnel.

Gathered data will have to be cleaned, by eliminating incorrect and double data. Missing data will have to be eliminated or completed, depending on the nature of the record with missing data. Rare events in the data will have to be handled properly. Finally, data will be grouped into a manageable amount of product classes. The data is then ready for the data analysis.

Besides quantitative product and line data, this study aims to identify the characteristics of the production process more qualitatively. This is necessary to strengthen reliability of the recommendations; it is hardly possible to formulate a worthwhile recommendation without ever closely observing the manufacturing process. Qualitative data regarding the production process will be acquired by observations.

1.5.3 Data Analysis

The collected data will be analyzed by a mathematical model. The nature of this model will be described in Chapter 3. To make the analysis easily repeatable, the mathematical model will be programmed in a suitable programming language. The output of this program should be the predefined performance measures, calculated with the input determined in the data collection.

The second step of data analysis, is finding possible improvements for the current situation. By adjusting input from the data collection, and comparing outcome, it is possible to determine a

set of possible ways to reduce the risk of long queues. This set of possible improvements will then be compared with the outcome of the observations; that way a recommendation can be formulated that is both impactful (based on the model) and realistic (based on the observations).

1.6 Outline of the report

The report is structured according to the standard research cycle (Heerkens & Van Winden, 2012). In each chapter, one or more sub-questions are answered.

In Chapter 2, the case at hand is more elaborately described, mapping out the production process and specifying the different products. This effectively answers sub-questions 1 and 2.

The theoretical aspects are treated in Chapter 3. In this chapter, a set of performance measures is determined and an analytical model is chosen. The sub-questions answered in this chapter are sub-questions 3, 4 and 5.

Chapter 4 briefly describes the techniques applied to find the input parameters for the model, followed by displaying the definitive input parameters determined from the data. This answers sub-question 6.

The analysis is performed in Chapter 5. First, the program used to perform the analysis is discussed. Second, the current situation is analyzed based on the input parameters determined in Chapter 4. Finally, several realistic scenarios are proposed to improve manufacturing performance. Sub-questions 7, 8 and 9 are answered.

The report is finalized with a conclusion. Previous chapters are shortly summarized and recommendations are made according to the most effective scenarios.

1.7 Conclusion Chapter 1

In this chapter, the case company was shortly introduced. The problem at hand was described more elaborately: Aebi Schmidt faces long queues as a result of many complexities. Several research questions were posed, in order to answer the main research question: How can long queues in the welding and surface treatment line be prevented?

The scope of the study was explained: the study aims to find ways to improve manufacturing performance and decrease the undesirable impact of line congestion at the welding and surface treatment facility of Aebi Schmidt. Solutions are sought within the domain of the science of operations management. The methodology was explained by shortly describing the research design, the data collection methods and the data analysis. Finally, an outline of the report was given.

Chapter 2: Case Information

Information about the products and processes at Aebi Schmidt Nederland

To understand the mathematical analysis and observations, it is necessary to have a clear overview of what the welding and surface treatment facility at Aebi Schmidt comprises, and what kind of products are made. First products and product parts are categorized; data grouping (as described in Chapter 4) is based on this categorization. Second, the production process will be described in detail.

2.1 The Products

The facility in Holten produces snow clearance and de-icing equipment and machines. The variety of the product is nearly endless, but certain generalizations are possible. Only the different products and product parts that are treated at the welding and surface treatment facility in Holten are dealt with.

Three main distinctions can be made between different machines: small design, large design, and integral design (see Figure 1).



Figure 1: From left to right: The large design Stratos, the small design Syntos, and the integral design Galeox

The large Stratos and the small Syntos design are both modular and consist of four parts: A bottom part ("onderbak"), a top part ("bovenbak"), a tail piece ("staartstuk") and a frame ("frame") (see Figure 2). These four parts come in different sizes, and have many customizable features. The integral design consists of, as the name implies, only one part, which comprises all metal plating of the integral design. In this study, the different product classes that will be considered are bottom parts, top parts, tail pieces, frames, and Galeox units (the integral product); variability in these products are treated as inherent variability.

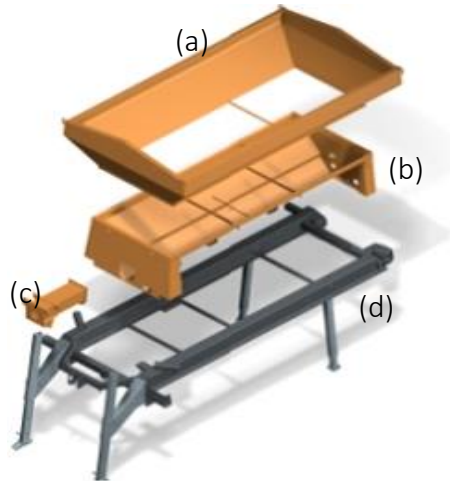


Figure 2: The top part (a), the bottom part (b), the tail piece (c) and the frame (d)

2.2 The Process

In this report, six different processes are distinguished: tack welding, seam welding, blasting, primer coating, baking, and final paint coating. Every product follows the same route: from tack welding to seam welding, from seam welding to blasting, from blasting to primer coating, from primer coating to baking, then from baking to final paint coating, and from paint coating once again to baking; the product is then finished. Because capacity of the surface treatment process is limited by machines, the capacity of the surface treatment line is lower than the capacity of the welding line. This is compensated for by working longer. It is important that this is accounted for in the analysis. The process is schematically drawn as a network in Figure 3.

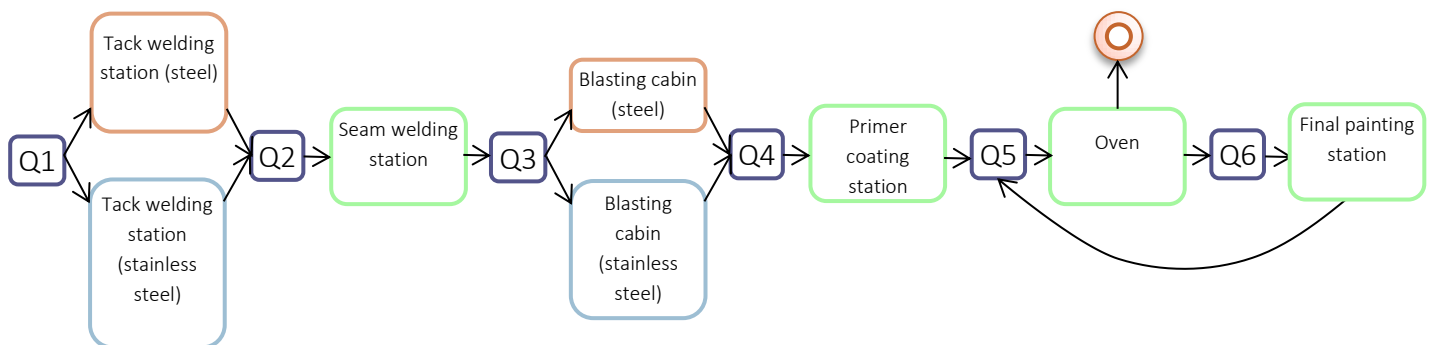


Figure 3: Schematic overview of the production process

Since each process has its own intricacies, the processes are explained shortly.

Tack Welding

Tack welding is the assembly of separate metal parts in its final form. After tack welding, the product part is not yet water tight and not coated; it only has its form. Separate metal plates are placed into a mold to position them, and are then welded together. Since each product type requires a different mold, every article can only be tack welded at a specific location.

Tack welding capacity is usually upper bound by labor constraints rather than number of tack welding locations. Generally speaking, there are about <<censored>> welders active, although this number may be increased during production high season with temporary workers.

Increasing capacity or decreasing variability at tack welding is not trivial: welders have certain skills and traits, and are therefore not available for every type of work. Temporary forces are generally poorly trained; management is already satisfied if the temporary force performs adequately at one task.

Seam Welding

At seam welding the seams in the metal construction are closed. This makes the metal part water tight, a necessary feature considering the parts are designed to contain liquids. Seam welding used to be done exclusively manually, but since April 2012 welding robots were purchased. During the course of this study, these machines were still in the early stages of operations, and most seam welding was still done manually. Once the robots are utilized to their potential, it is estimated that process times will decrease with <<censored>>%.

Blasting

Blasting is the first of the four processes that comprise surface treatment. At this process, the product is blasted to clean it thoroughly before being coated. There are two cabins for blasting: one for steel, and one specifically for stainless steel. Each material type is restricted to its own cabin.

Primer coating

The products are first coated with a primer coating before their final color is applied. When steel is exposed to moisture it will oxidize. A primer coating provides an additional layer of protection against direct exposure to moisture. Furthermore it increases the adhesion of the final paint coating. There is only one primer coating cabin available at Aebi Schmidt; it is expected this will lead to a capacity problem when the welding robots are fully operational.

Oven

Once a layer of painting is applied, it is necessary to 'bake' the layer. The duration that the product is baked depends on the maximum thickness of the sheet metal of the product, and the type of paint.

Final paint coating

At this process the final color is applied to the product. Like the primer coating process, only one painting cabin is available.

There are some aspects of this process that need to be addressed when a queuing network analysis is performed. Changing the paint color requires some setup time and setup costs, but this might be negligible. Furthermore there are indications that this process has a significant rework rate.

Additional considerations regarding surface treatment

<<Not available in public version>>

2.3 Conclusion Chapter 2

In this chapter, the product characteristics and the process characteristics were described. The process encompasses 6 main processes: tack welding, seam welding, blasting, primer coating, the oven, and final paint coating. The product parts that flow through the manufacturing system can be divided into 5 main categories: top parts, bottom parts, tail pieces, frames and integral designs. By describing the characteristics of products and processes, sub-questions 1 and 2 were answered.

Chapter 3: Theoretical Aspects

Review of literature on manufacturing system analysis and performance measures

In Chapter 1, an overview was given regarding the necessary information to conduct the designed research: A set of performance measures that accurately describe the performance of a manufacturing system, a theoretical, mathematical framework to determine these measures, and the necessary input from the case to calculate these measures. In this chapter these theoretical aspects are determined, starting with a balanced set of performance measures.

3.1 Measuring Manufacturing Performance

Before determining the necessary indicators for performance for the study at hand, it is necessary to get a grip on what is meant with the term 'performance' in the first place. Hopp and Spearman (2001, p. 289) state:

"...management of any system begins with an objective. The decision maker manipulates controls in an attempt to achieve this objective and evaluates performance in terms of measures."

A performance measure should measure to what extent a company is achieving its goals. It has been long known that simply measuring productivity, usually in terms of capital but also in terms of labor, material or overhead, can lead to inadequate performance assessment (Kaplan, 1983). Performance depends not solely on return on capital (or labor, material or overhead, for that matter), but also on less tangible measures such as flexibility and quality. According to Kaplan (1983) measuring performance depends on numerous factors, such as the maturity of the product, the extent of automation that is present in a production line, and the scale of the production.

There have been numerous attempts at building a general framework for assessing manufacturing performance. A noteworthy attempt was made by Son & Park (1987): They expressed an Integral Performance Measure as a weighted average of productivity, quality and flexibility, and operationalized all indicators necessary to compute values for these measures (Figure 4).

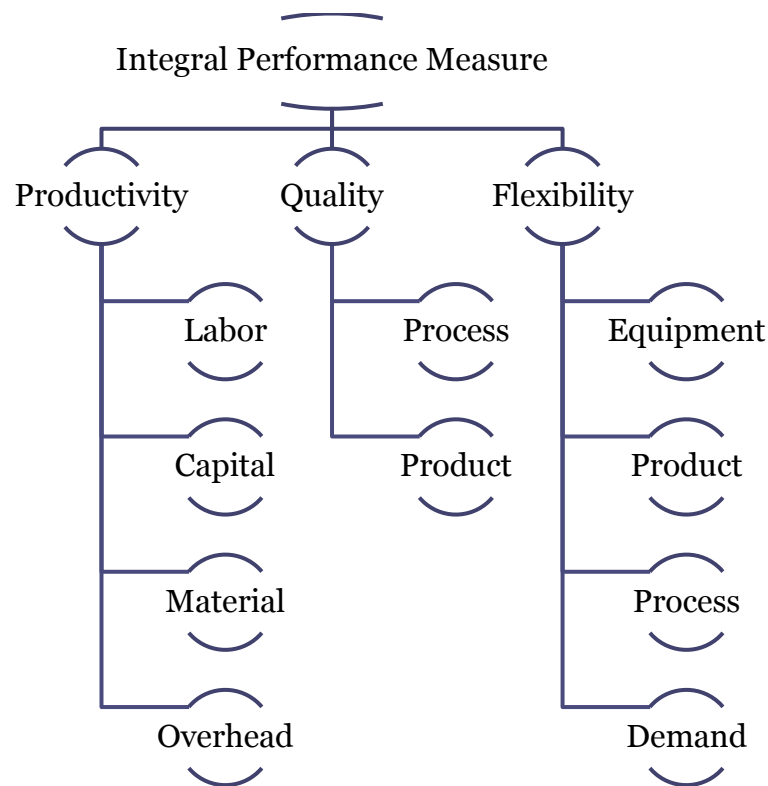


Figure 4: Hierarchy of Performance Measures and their cost elements (Son & Park, 1987).

More aimed at production line performance, Hopp & Spearman (2001, pp. 291-293) came up with seven efficiencies that measure the performance of a single-product line, but are easily extended to a multi-product line by aggregating flows and inventories:

- **Throughput Efficiency.** Ideally, throughput is equal to the demand rate. Throughput efficiency is defined as:

$$E_{TH} = \frac{\min(TH, D)}{D}$$

D being the demand rate and TH being the throughput, which means it is only less than 1 if demand is not fully met. Since inventory efficiency is measured elsewhere, it is not necessary to penalize unnecessary inventory building here.

- **Utilization Efficiency.** Ideally, all machines are fully utilized. Utilization efficiency is thus defined as:

$$E_u = \frac{1}{n} \sum_{i=1}^n \frac{TH_i}{r_i}$$

With n the number of machines in the line, TH the throughput, and r the optimal production rate of the machine.

- **Inventory Efficiency.** Ideally a line only has the minimal required inventory: No raw material inventory, only as much work in process as is needed for the line to operate at the desired throughput, and no finished goods inventory. Inventory Efficiency is measured by:

$$E_{inv} = \frac{\sum_i \frac{TH_i}{r_i}}{RMI + WIP + FGI}$$

With TH the throughput, r the optimal production rate of the machine, RMI the raw material inventory, WIP the work in process, and FGI the finished goods inventory.

- **Cycle Time Efficiency.** Ideally average cycle time equals the shortest possibly achievable cycle time. Cycle Time Efficiency is hence measured as follows:

$$E_{CT} = \frac{T_0^*}{CT}$$

With T_0^* the shortest possibly achievable cycle time, and CT the actual average cycle time.

- **Lead Time Efficiency.** Penalizing late deliveries is done in the measure for customer service efficiency below. Lead time efficiency is only concerned with setting lead time as short as possible in a make-to-order system (in a make-to-stock system, it is not a sensible measure). It is measured as follows:

$$E_{LT} = \frac{T_0^*}{\max(LT, T_0^*)}$$

With T_0^* the shortest possibly achievable cycle time, and LT the lead time quoted to customers. This measure is 1 when quoted lead time is equal to or lesser than the shortest possibly achievable cycle time. Obviously, when quoted lead time is shorter than shortest possibly achievable cycle time, every delivery will be late.

- **Service Efficiency.** Service efficiency is the fraction of demand that is served on time. It is measured as follows:

$$E_s = \begin{cases} \text{fraction of demand filled from stock in a make to order system} \\ \text{fraction of order filled within lead time in a make to stock system} \end{cases}$$

The fraction of demand filled from stock is also known as the *fill rate*.

- **Quality Efficiency.** Although the term quality is generally used broadly, operational quality is only concerned with the fraction of products that do not require any form of rework. Hopp & Spearman define quality efficiency as:

$$E_Q = \text{fraction of jobs that go through the line with no defects on first pass}$$

This particular study is concerned with a production line that is up and running, and has been for a long time. An integral performance measure would not only take too much time to compute, it

would also be of limited value, since it would take many measures into account that do not seem to be directly linked with Aebi Schmidt's line performance. It seems sensible to minimize the amount of indicators that constitute the manufacturing performance measure.

Any marketing activities undertaken by Aebi Schmidt are outside the scope of this study; hence manufacturing performance will in this study be assumed independent of lead time efficiency and service efficiency. Also utilization, although an important measure, is difficult to directly link to performance: welding is very labor intensive, not very capital intensive.

Therefore, in this study, it is attempted to express manufacturing performance in terms of inventory, cycle time and quality. The indicators for these variables will be work in process, average cycle time per station, average waiting time per station, and whenever possible, the fraction of products that require rework.

3.2 Quantitative Models

At this point it is necessary to evaluate the different options of computing the set performance measures: average cycle time, average work in process, and quality. Here treated are three different methods of calculating these measures: statistical analysis, which builds on the science of using historical results to predict the present; simulation, which essentially negates the requirement for cumbersome mathematics; and queuing theory, which attempts to approximate reality through an analytical model.

3.2.1 Statistical analysis

The main application of statistical analysis in the science of operations management is to understand historical results, and forecast accordingly. Two main methods can be distinguished: causal statistical analysis and statistical analysis of time series (Hopp & Spearman, 2001, p. 415). Causal statistical analysis aims to describe a parameter that is difficult to measure directly as a function of other parameters that are measurable. Although important in areas such as marketing, it is of limited use in determining manufacturing performance measures such as average cycle time or work-in-process. Statistical analyses of time series are more useful when a parameter is not easily expressed as a function of other parameters, and historical data is assumed to be a good indicator. This is also the main weakness of time series analysis: the parameter data from the future will inherently not be taken from the same population as historical data, since it is usually impossible to keep all other system parameters equal. Its main use is thus to describe the trend of a parameter over time, but it is difficult to predict the behavior of a parameter when the situation is altered.

The aim of this study is to prevent line congestion in the future; although statistical analysis of time series might be of some use, it is difficult to use statistical analysis to find impactful and realistic ways to prevent line congestion.

3.2.2 Simulation

When manufacturing systems get very complex, it is often hard to establish an analytical model that is not either unsolvable or too simplistic. The most prominent feature of a simulation is that many simplifying assumptions that are necessary for an analytical model simply do not have to be made; it is that much more straightforward (Winston, 2004, pp. 1145-1146). Simulation techniques can be classified as discrete or continuous, static or dynamic, and stochastic or deterministic. For applications in manufacturing simulations are generally discrete, i.e. states of the system only change at discrete points in time (in contrast with continuous simulation, in which variables change continuously over time). Deterministic simulation has some major advantages: nearly any process, no matter how atypical, can be simulated deterministically. Stochastic simulation is already considerably more time consuming. In an earlier research a deterministic simulation was made for the welding and surface treatment line at Aebi Schmidt (Dröge, De Jong, Klijnsma, & Reimert, 2013). Although the simulation has great potential for enhanced production scheduling, its primary weakness was that variability in process times and arrival rates was not accounted for. Incorporating variability measures increased simulation time beyond limits for practical use.

Simulation thus has two main disadvantages. The first is that simulation, while arguably easier than solving an analytical model, is still complex and time consuming to set up. Second, simulation is not very suitable for optimization purposes, especially not when variability is involved (Winston, 2004, p. 1145). Optimizing involves repeating the simulation with different parameters until a conclusion can be drawn regarding the optimal solution; this is, in nearly every practical case, too time consuming.

3.2.3 Queuing Theory

The general goal of queuing theory is to compute performance measures by the application of a queuing model. In this study, the performance measures are average cycle time, average waiting time, and average work in process. The required input for these models are the arrival rate, which is the amount of jobs per unit time that arrive at the server, and the process time, which is the time required to process a job. By definition, process time is the inverse of service rate, which measures the rate at which jobs can be processed by a server per unit time. When it is necessary to take into account the variability in the arrival process and process times (as is often the case, since a deterministic model is usually not very representative), some measure of

variability in these input parameters is necessary; it is found that taking into account only the first moment of variability, the standard deviation, generally results in sufficiently accurate models (Whitt, 1982). This is especially the case when the interest is on average cycle times or queue length.

The queuing system, on which queuing theory is applied, is usually described in Kendall-Lee notation (Winston, 2004, pp. 1060-1061):

$$A/B/C$$

A: Specifies the distribution of interarrival times (the time between two subsequent arrivals). The most common distributions are the exponential distribution (denoted with the letter *M*), a deterministic distribution (denoted with the letter *D*), an Erlang distribution (denoted with the letter E_k , k being the shape parameter) and a general distribution (denoted with either the letter *G* or the two letters *GI*; in this report the letter *G* will be used). From the distribution of interarrival times, it is possible to calculate the distribution of the amount of jobs that arrive during a specific time interval. Most noteworthy is that exponential interarrival times leads to a distribution for $N(t)$, the number of arrivals in time period t , that is Poisson; hence an arrival process with exponential interarrival times is called a Poisson arrival process. *Poisson Arrivals See Time Averages*, or in other words, every arrival has a probability of $\pi(n)$ finding the queuing system in state n , where $\pi(n)$ is the steady-state probability for state n (Wolff, 1982) (Winston, 2004, pp. 1051-1144). This property makes processes with Poisson arrivals easily solvable.

B: Specifies the distribution of service times, with the same symbols that are used to specify the distribution of interarrival times.

C: Specifies the number of servers at the station.

The extended Kendall-Lee notation also includes characteristics that describe the maximum number of jobs in the system, the queuing discipline and the size of the population that jobs are drawn from when they enter the queuing system. Since none of these characteristics are of vital importance in the study at hand, it is assumed that the line does not 'block', i.e. production is never (partly) shut down because there is no more room for products, the queuing discipline is general and the population orders are drawn from is infinitely large.

An important requirement to allow any meaningful analysis of queuing systems is the stability condition. On average, the rate of arrivals into the system should be strictly less than the capacity of the system (with the sole exception of completely deterministic systems). If arrival

rate is equal to or exceeds the capacity in the long term, work in process grows infinitely large. Mathematically, the stability condition can be summarized as follows:

$$u = \frac{\lambda}{c\mu} < 1$$

In which u is the server utilization, λ the arrival rate, c the number of parallel servers and μ the service rate (Zijm, 2012).

Queuing models

The analytical model tends to become increasingly more difficult as simplifications are dropped. The most basic application of queuing theory computes the performance measures for an M/M/1 queuing system: a single server with exponentially distributed service times, processing jobs that arrive with exponentially distributed interarrival times. This basic example of queuing system has been extensively treated (Winston, 2004, pp. 1072-1077) (Hopp & Spearman, 2001, pp. 267-270). Its performance measures are as follows:

$$WIP = \text{Average Work in Process} = \frac{\lambda}{\mu - \lambda} = \frac{u}{1 - u}$$

λ here denotes the arrival rate of jobs (which should be, on average, equal to the throughput in the long term), $t_e = \frac{1}{\mu}$ the mean service time, and $u = \frac{\lambda}{\mu} = \lambda \cdot t_e$ the utilization (the fraction of time the station is active).

$$CT = \text{Average Cycle Time} = \frac{WIP}{TH} = \frac{1}{\mu - \lambda} = \frac{\frac{1}{\mu}}{1 - u} = \frac{t_e}{1 - u}$$

λ here once again denotes the arrival rate of jobs, $\frac{1}{\mu}$ and t_e the mean service time, and $u = \frac{\lambda}{\mu}$ the utilization.

$$WIP_Q = \text{Average Queued Work in Process} = \frac{\lambda}{\mu(\mu - \lambda)} = \frac{u^2}{1 - u}$$

$$CT_Q = \text{Average Waiting Time} = \frac{\lambda}{\mu(\mu - \lambda)} = CT - t_e = \frac{u}{1 - u} t_e$$

These performance measures are exact for this model, but the model is not very useful: it only describes one station, with only one server, with exponential interarrival and service times.

Applicability can be greatly improved by allowing general interarrival and service times. The PASTA property (*Poisson Arrivals See Time Averages*) is violated for a G/G/1 queuing model, and it is not always exactly solvable. However, there are multiple approximations that still allow fairly accurate calculation of the performance measures. The most widely used approximation in G/G/1 queuing models is Kingman's formula (Kingman, 1961):

$$CT_Q = \frac{c_a^2 + c_e^2}{2} \cdot \frac{u}{1-u} \cdot t_e$$

In which c_a is the coefficient of variation of interarrival times and c_e the coefficient of variation of (effective) service times. This approximation is most accurate at high traffic intensity. A notable feature of this approximation is that it reduces to the performance measures of the M/M/1 queuing model when interarrival times and service times are exponentially distributed, i.e. $c_a = c_e = 1$.

The other performance measures follow logically:

$$CT = CT_Q + t_e = \frac{c_a^2 + c_e^2}{2} \cdot \frac{u}{1-u} \cdot t_e + t_e$$

$$WIP = TH \cdot CT = \frac{c_a^2 + c_e^2}{2} \cdot \frac{u^2}{1-u} + u$$

$$WIP_Q = TH \cdot CT_Q = TH \cdot CT = \frac{c_a^2 + c_e^2}{2} \cdot \frac{u^2}{1-u}$$

The next step is to extend this queuing model to a multi-server model: the G/G/c model. A widely used approximation of the waiting time of the G/G/c model is (Whitt, 1983):

$$CT_Q = \frac{c_a^2 + c_e^2}{2} \cdot CT_Q(M/M/c)$$

In which $CT_Q(M/M/c)$ is the well-known waiting time for an M/M/c queuing model. Although solvable exactly, it is convenient to use the approximation by Sakasegawa (1977), for computational purposes and because it is easier to work with in analyzing queuing networks:

$$CT_Q(M/M/c) = \frac{u^{(\sqrt{2c+2}-1)}}{c(1-u)} t_e$$

$$CT_Q(G/G/c) = \frac{c_a^2 + c_e^2}{2} \cdot CT_Q(M/M/c) = \frac{c_a^2 + c_e^2}{2} \frac{u^{(\sqrt{2c+2}-1)}}{c(1-u)} t_e$$

From this result the other performance measures follow logically. This concludes the theory on single-station single-class queuing models.

Single-Class Queuing Networks

The single station models described before can be expanded by allowing jobs to flow from one station to the other, establishing a queuing network. Once again the simplest case of a queuing network consists of solely M/M/1 or M/M/c queues. These networks are called open Jackson-networks, and are treated in Appendix B.

The open network of M/M/c queues can be expanded to an open network of G/G/c queues: Coefficients of variation of interarrival times (c_{0j}^2) and service times (c_{sj}^2) are no longer assumed to be equal to 1, and are now necessary input. The main difference compared to the analysis of Jackson-networks is that combining and splitting job flows in the system is no longer a trivial matter of adding up. The squared coefficient of variation of the arrival rate per station, c_{aj}^2 , is determined by solving the following linear equation (Whitt, 1983):

$$c_{aj}^2 = a_j + \sum_{i=1}^M c_{ai}^2 b_{ij}$$

In which a_j and b_{ij} depend on the input. Once these squared coefficients of variation are known per station, the stations are treated as independent queuing systems, at which point performance measures can be computed.

In this report, superposition and splitting of flows will be approximated using the formulas found in Whitt (1983). The linear equation above is then described by:

$$a_j = 1 + w_j \left((Q_{0j} c_{0j}^2 - 1) + \sum_{i=1}^M Q_{ij} ((1 - P_{ij}) + P_{ij} u_i^2 x_i) \right)$$

$$b_{ij} = w_j P_{ij} Q_{ij} (1 - u_i^2)$$

The used variables to compute these two parameters all depend on the input:

$$x_i = 1 + c_i^{-0.5} (\max[c_{si}^2, 0.2] - 1)$$

$$Q_{ij} = \frac{\lambda_{ij}}{\lambda_j}$$

Q_{ij} is like “the receiving end” of the routing matrix, which is defined as $P_{ij} = \frac{\lambda_i}{\lambda_{ij}}$. It “denotes the proportion of the arrival flow of station j originating from station i ” (Zijm, 2012, p. 44). v_j , w_j and x_i are once again determined by the input. They do not have a clear physical meaning, but are convenient for computational purposes:

$$v_j = \left(\sum_{i=0}^M Q_{ij}^2 \right)^{-1}$$

$$w_j = \left(1 + 4(1 - u_j)^2 (v_j - 1) \right)^{-1}$$

To calculate the performance measures, the stations are assumed to be independent and are therefore treated as separate G/G/c queues. Computation of these measures is remarkably trivial considering the complexity that was necessary to superpose job flows.

$$CT_{Q,i} = \frac{c_{ai}^2 + c_{si}^2}{2} \cdot \frac{u_i^{(\sqrt{2c_i+2}-1)}}{c_i(1-u_i)} \cdot \frac{1}{\mu_i}$$

$$CT_i = \frac{c_{ai}^2 + c_{si}^2}{2} \cdot \frac{u_i^{(\sqrt{2c_i+2}-1)}}{c_i(1-u_i)} \cdot \frac{1}{\mu_i} + \frac{1}{\mu_i}$$

$$WIP_i = \lambda_i CT_i$$

$$WIP_{Q,i} = \lambda_i CT_{Q,i}$$

Multi-Class G/G/c Queuing Network

The final step is to allow more than one type of job into the system; a very necessary feature to accurately analyze the welding and surface treatment line at Aebi Schmidt. There are multiple methods available for analyzing these types of queuing networks; in this report, the ‘complete reduction method’ will be used (Zijm, 2012). This method consists of three general steps:

- Aggregate the separate classes to one single class.
- Use the aggregated class to analyze a single-class queuing network.
- Disaggregate to determine performance measures per class.

The aggregation of interarrival times and service times is done by simply taking a weighted average:

$$\frac{1}{\mu_j} = \frac{1}{\lambda_j} \sum_{r=1}^R \lambda_j^{(r)} \cdot \frac{1}{\mu_j^{(r)}}$$

Here $\frac{1}{\mu_j^{(r)}}$ is the mean service time of a job of class r at station j and $\lambda_j^{(r)}$ the arrival rate of jobs of class r at station j , from both outside the system and from stations in the system (including itself, in case of rework).

$$\lambda_j^{(r)} = \lambda_{0j}^{(r)} + \sum_{i=1}^M \lambda_i^{(r)} P_{ij}^{(r)}$$

Here $\lambda_{0j}^{(r)}$ is the arrival rate of jobs of class r from outside the system at station j , and $P_{ij}^{(r)}$ the probability a job of class r flows to station j once it has left station i .

$$\lambda_j = \sum_{r=1}^R \lambda_j^{(r)}$$

The aggregated squared coefficient of service times is more complex:

$$c_{sj}^2 = \frac{\mu_j^2}{\lambda_j} \sum_{r=1}^R \lambda_j^{(r)} \left(\frac{1}{\mu_j^{(r)}} \right)^2 \left((c_{sj}^{(r)})^2 + 1 \right) - 1$$

$c_{sj}^{(r)}$ is the squared coefficient of variation of service times of jobs of class r at station j .

Superposing the job flows of different classes from outside the system allows calculation of the squared coefficient of variation of interarrival times (Zijm, 2012):

$$c_{0j}^2 = w_j \sum_{r=1}^R Q_j^{(r)} (c_{0j}^{(r)})^2 + 1 - w_j$$

Here $Q_j^{(r)}$ is the proportion of the arrival flow of jobs of class r from outside the system to station j , and $c_{0j}^{(r)2}$ is the squared coefficient of variation of the interarrival times of jobs of class r from outside the system to station j .

Finally, the routing matrices for the individual classes are aggregated:

$$P_{ij} = \frac{1}{\lambda_j} \sum_{r=1}^r \lambda_i^{(r)} P_{ij}^{(r)}$$

Which concludes the aggregation; with these results it is possible to calculate the performance measures with the method described for a single-class queuing system. Disaggregation is relatively straightforward. The performance measures per class, in terms of input parameters and aggregated performance measures are computed as follows (Zijm, 2012):

$$WIP_{Q,j}^{(r)} = \frac{\lambda_j^{(r)}}{\lambda_j} WIP_{Q,j}$$

$$WIP_j^{(r)} = WIP_{Q,j}^{(r)} + u_j^{(r)}$$

$$CT_j^{(r)} = \frac{V_j^{(r)} WIP_j^{(r)}}{\lambda_j^{(r)}}$$

At which point the performance measures per class are computed. This concludes the multi-class queuing network analysis; this is the framework proposed for this study.

Additional considerations in queuing theory

The variability measures that are input for the queuing network analysis are a measure of natural variability in process times, and additional factors that cause variability. The most prominent factor are batching, which inherently increases the squared coefficient of variation of arrival rates, and outages, which inherently increases the squared coefficient of variation of process times. Their influence on variability is more elaborately treated in Appendix C.

Rework

Accounting for rework in queuing network analysis is remarkably easy: the routing can be adjusted so that there is a probability that jobs flow directly to station they came from, effectively requiring a second process time. In practice however, the rework time is significantly shorter than the whole processing time. Instead of the actual rework rate (number of products that require rework divided by total number of products), a relative rework rate is used:

$$\text{relative rework rate} = \frac{\text{rework time}}{\text{process time}} \cdot \frac{\text{number of rework instances}}{\text{number of product instances}}$$

3.3 Observation

Besides performing a queuing network analysis, the welding line is qualitatively observed. The goal of these observations is to provide a 'reality check' for the queuing network analysis. Although it is observed how workers spend their time, the fraction of spent time on different activities is not exactly measured; this would require extensive and elaborate measurements, which are far beyond the scope of this study. Alternatively, it is attempted to find some recurring inefficiencies in the process, which might cause variability or a longer than necessary process time. If an observed possibility for improvement coincides with a sensitive parameter of the queuing network analysis, a solution is found which is both impactful and realistic.

3.4 Conclusion Chapter 3

It was concluded that statistical analysis is not very suitable for the study at hand. A simulation of the welding and surface treatment line was already performed in an earlier study. During this study, some difficulties were encountered: simulation time to find effective solutions took prohibitively long.

In this study, the cycle time and work in process will be determined using queuing network analysis. To model the welding and surface treatment line as accurately as possible, it is chosen to model the line as a network of G/G/c queues with multiple classes. This way job variety can be maintained, and routings can be taken into account. Calculation of performance measures is very quick; a moderate contemporary computer should be able to calculate performance measures in a matter of seconds.

The queuing network analysis will take into account outages and rework when necessary. Batches are not specifically accounted for; however, some high squared coefficients of variation are expected for some classes.

By choosing a set of performance measures and an analytical model to determine these, sub-questions 3, 4 and 5 were answered. Average cycle time, work in process and waiting time will be determined using queuing network analysis.

Chapter 4: Data Collection

About the methods and choices of collecting and selecting data

This chapter will elaborate on the applied methods to collect necessary data. Recapitulating previous chapters, the data consists of product and line data (quantitative) and observations (qualitative). The product and line data was to be acquired partly by time registration Excel sheets from Aebi Schmidt, and interviews with specialists regarding the production process. This two-pronged approach was necessary because not all process times were recorded in time registrations.

4.1 Process Parameters

Of the six processes at the welding and surface treatment facility (tack welding, seam welding, blasting, primer coating, oven, final paint coating), only tack welding, seam welding, primer coating and final painting coating were recorded. There are no time registrations for the blasting cabin, but there is a log file of outages; the process times of blasting are therefore determined by inquiry, and the variability is calculated by analyzing the outages. First, treatment of the time registrations is discussed.

4.1.1 Data Cleaning

Process data was recorded by the welders and painters themselves. The computers used for data collection were located near the actual processes. Data was recorded in Excel sheets.

Unfortunately, data was not collected with the use of forms. Although the data is largely uniform, many entries are not properly recorded, or contain remarks that require further inspection before being used for data analysis. This data cleaning is a very laborious and tedious task. Since both averages as well as frequencies need to be determined it is not possible to simply omit the faulty entries; measures of frequency would then be grossly underestimated.

For this reason, over 13.000 entries have been checked manually for inconsistencies. Whenever entries were deemed too inconsistent to represent reality, their end time and process time were removed. This way, the entry is not counted in computing the average process time, but it is counted in determining the frequency of the product. A more elaborate description of how data was handled can be found in Appendix A.

Analysis of the outages of the blasting cabin has been performed by first cleaning up faulty entries. The log file indicates that whenever a failure occurs, often another failure occurs the same day; this is because reparation is not a trivial task and often requires some adjustments.

Log entries made on the same day are therefore not independent. To calculate squared coefficient of variation of effective service time, multiple entries on the same day were considered one outage (time to repair was added for these outages). Additionally, some outages were very short, and hardly disruptive for the process. These outages were ignored, in an effort to compute a representable coefficient of variation. The used formula to calculate the squared coefficient of variation of the process times at blasting is (assuming an otherwise non-variable process): $c_e^2 = c_0^2 + (1 + c_r^2)A(1 - A)\frac{m_r}{t_0} = \frac{\langle censored \rangle}{t_0}$, in which c_0^2 is the natural squared coefficient of variation of process times, c_r^2 the squared coefficient of variation of time to repair, m_r the mean time to repair, A the availability and t_0 the process time excluding outages. See Appendix C for a more elaborate treatment.

t_0 was determined by simply asking the worker responsible for blasting. Since blasting is for the most part automated, assuming $c_0^2 = 0$ is not a very unrealistic estimate, and t_0 is very close to the programmed times of the blasting robot.

4.1.2 Results

For the application of the primer coat and the final paint coat, there is a significant amount of rework. The average rework time and the rework rate are displayed after the computed average process time and the arrival rate.

The process times were determined by averaging all process times in a data group. The corresponding squared coefficient of variation was determined by dividing the variance of process times by average process time squared. Arrival rate was determined by counting the number of arrivals per day, and corresponding squared coefficient of variation by dividing variance in arrivals per day by average arrivals per day. N (Clean) is the number of entries on which average process times are based; these are all data entries with a valid process time. N (Total) is the number of entries on which arrival rate is based; these are all data entries with a valid starting date. Since arrivals only occur at tack welding, only the arrival rate at tack welding is shown.

<<Numerical results not available in the public version>>

4.2 Observations

<<Not available in the public version>>

4.3 Conclusion Chapter 4

In this chapter, the input parameters for the analytical model were determined. Tack welding, seam welding, primer coating and final paint coating data was determined by analysis of time

registrations; at both primer coating and final paint coating, a significant rework rate was found. Input parameters for blasting and the oven were determined by inquiry; blasting is subject to frequent outages, which have been accounted for in the final data collection. The oven is modeled as a deterministic process; variability is only present due to different oven times primer coating and final paint coating. By establishing the input parameters, sub-question 6 has been answered.

Chapter 5: Queuing Network Analysis

Explanation of the developed tool, the results of the queuing network analysis, and sensitivity analysis

In this chapter, the queuing network analysis is performed. First the necessary assumptions and choices are discussed.

5.1 Modeling the welding and surface treatment line

A general solver for a multi-class G/G/c queuing network was created to analyze the data collected in the previous chapter. The solver was programmed in Visual Basic for Applications in Microsoft Excel; although this programming language is arguably not the most flexible, it is fast enough for queuing network purposes and most importantly: Excel is widely used, so nearly everyone will be able to use the solver. In this paragraph the solver will be discussed in the same order as queuing network analysis: The input, the aggregation, the single-class Open Queuing Network analysis, and the disaggregation.

5.1.1 The Input

The program currently does not use all collected data. Process times per product group per station and the squared coefficient of variation hereof are all used as input; for the arrival rate, only the arrival rate and corresponding squared coefficient of variation of tack welding is used as input. The different product groups were numbered as follows:

r	Product group
1	Bottom Part
2	Top Part
3	Tail Piece
4	Frame
5	Integral Design

Table 1: Product group numbering

For the queuing network analysis, it is assumed the arrival rate at tack welding represents the long term arrival rate of products. Since jobs can only enter the line at tack welding, the arrival rate vector $\lambda_0^{(r)}$ is only non-zero for tack welding. The vector with corresponding squared coefficients of variation in arrival rate $(c_{0j}^{(r)})^2$ is also only non-zero at tack welding.

Not taking rework into account rework, all products have the same routing matrix:

$$P_{ij}^{(r)} = \begin{bmatrix} 0 & 1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0,5 & 0,5 \\ 0 & 0 & 0 & 0 & 1 & 0 & 0 \end{bmatrix}$$

In which the last column represents a node 'outside the system', to which finished products are sent.

For primer coating and final paint coating, a significant rework rate was found. However, quality problems are not properly represented by simply dividing occurrences of rework by the occurrences of the corresponding product, because the required time to rework the product is generally shorter than the initial process time. Therefore, a relative rework rate was calculated.

This relative rework rate was incorporated in the routing matrix. Since rework only seemed to be a significant factor at primer coating and final paint coating, the routing matrix per product was adjusted:

$$P_{ij}^{(r)} = \begin{bmatrix} 0 & 1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & rrw & 1 - rrw & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0,5 & 0,5 \\ 0 & 0 & 0 & 0 & 1 - rrw & rrw & 0 \end{bmatrix}$$

rrw being the relative rework rate per product class. It is possible to add rework rates to other stations as well; the program decreases other probabilities proportionally.

5.1.2 Modeling different shift hours

At this point, there is one obstacle that makes it very difficult to model the queuing network as a single network: the welding process is active <<censored>> hours per day, while the surface treatment process is active for <<censored>> hours per day. Using the fact that all jobs from seam welding flow to blasting, the line is split between seam welding and blasting. The analysis is performed as follows:

- Queuing network analysis is performed on only tack welding and seam welding.
- The departure rate and squared coefficient of variation of seam welding are used as the direct input for blasting.
- Queuing network analysis is performed on the rest of the stations.

The analysis is thus actually performed twice. This is the most accurate way to deal with different shift hours, but it is only possible under the assumption that all departures from seam welding go to blasting. It is difficult to generalize dealing with different shift hours; the current solver has lost some generality to account for it.

5.1.3 Single-class G/G/c Queuing Network and Disaggregation

In this part of the program it is important to think about how to translate the real line to representable model that is solvable by queuing network analysis. There are three problem areas that need to be dealt with: Tack welding, which is primarily labor constrained; seam welding, in which welding robots are currently being tested; and the oven, which allows multiple products to be baked at once.

At tack welding, it is chosen to assume that the number of active welders equals the number of parallel servers, and that each server can handle each product. The number of active welders is, on average, about <<censored>> workers (average calculated over one year), but this number is increased during production high season. This is a simplification in multiple ways: First, welders are better at some products and tasks than others, and are sometimes not even capable of performing a very specific task that is usually done by someone else. The model thus assumes a uniformly trained workforce, which in reality is slightly more rigid. Second, jobs arriving at tack welding require a specific mold. Although for most product groups there is more than one mold available, simultaneous production of some products is limited by the amount of available tack welding stations.

Current seam welding data is based on manual welding. Welding robots are being installed now, which supposedly have much shorter process times once they are fully operational. This is not yet the case; the majority of products is still seam welded manually. Therefore only manual seam welding is analyzed in this report. For robotic seam welding, Aebi Schmidt estimates that process times can be <<censored>> percent of the corresponding manual process time. The squared coefficient of variation is assumed to be equal to the manual case. Taking the average over one year, the average number of seam welders is approximately <<censored>>.

The oven is simply a heated space in which products can be placed. It is difficult to model the oven exactly, because there are products that take up the whole oven (e.g. the integral designs) and products that can be baked six at a time with ease (e.g. tail pieces and frames). Oven times are relatively short compared to other processes. Baking is not a variable process, but there is some variability introduced because oven times for the first baking and the second baking are different.

Initially, the oven was modeled as <<censored>> separate servers that could each handle one entire product. This way capacity to bake large products is overestimated, and the capacity to bake small products is underestimated. After validation with Aebi Schmidt, it was concluded that this overestimates the capacity too much. The oven is now modeled as <<censored>> separate servers.

With these assumptions, the number of parallel stations is as follows:

Process	Tack Welding	Seam Welding	Blasting	Primer	Oven	Final Paint
c	<<Not available in public version>>					

Table 2: Number of parallel servers per station

The final step of the single-class queuing network analysis is to compute the aggregated performance measures. After that, the performance measures per class can be determined using the formulas in Chapter 3.

5.1.4 Excel Tool

The queuing network analysis is performed entirely in Excel. The development of the tool was based on three principles: Generality to allow the tool to be used in various situations, responsiveness (here defined as quickness of execution) to make the tool usable in practice, and usefulness to make the tool a valuable aid to decision-making.

The actual code to perform the queuing network analysis is completely general. However, to account for shift hours, the current tool performs the queuing network analysis according to the method described in paragraph 5.1.2, so in its current state it is specifically designed for the welding and surface treatment process at Aebi Schmidt. Adding more classes or processing is all possible.

The responsiveness of the program is very good; performing a single analysis is done within milliseconds. Because of this high speed of execution, it was possible to automate the sensitivity analysis, which requires performing the analysis many times.

The program allows easy adjustment of the input parameters. There are also options on how to deal with different shift hours, and it is possible to adjust the shift hours at welding or surface treatment. Because the program executes quick enough, a sensitivity analysis feature was created. This feature allows the user to test different scenarios: what would happen to average cycle time if process times were to decrease, or another blasting cabin was installed? The ability to give an indication of the impact of different solutions makes the tool highly useful.

5.2 Current Situation Results

The queuing network analysis was run with data collected in Chapter 3. First, the aggregation is shown, followed by the single-class open queuing network analysis. The disaggregation can be found in the Appendix.

5.2.1 Aggregation <<Not available in public version>>

	<u>Welding</u>			<u>Surface treatment</u>				
	Tack Welding	Seam Welding		Blasting	Primer	Oven	Final Paint	
Arrival rate (/day)								
SCV Arrival rate								
Service time (min)								
SCV Proces tijd								
Routing	Tack Welding	Seam Welding	Leave System	Blasting	Primer	Oven	Final Paint	Leave System

Table 3: Aggregation of the input parameters

Since there is no rework in tack welding, seam welding and blasting, the arrival rates at these stations are equal to the throughput. The oven is visited twice, but has no rework; hence the arrival rate at the oven is exactly twice the throughput. At primer coating and final paint coating the arrival rate is slightly higher, which is due to rework.

The variability in interarrival times approaches 1 as the system progresses. This is expected, since process time variability is at most stations quite low.

5.2.2 Single-class G/G/c Open Queuing Network Analysis <<Not available in public version>>

	Tack Welding	Seam Welding	Sum	Blasting	Primer	Oven	Final Paint	Sum
Utilization								

CTQ (minutes)								
CT (minutes)								
Visit ratio								
WIP								
WIPQ								

Table 4: Aggregate performance measures of the single-class open queuing network

Total average waiting time	
Total average cycle time	
Total average work in process	
Total average work in queues	

Table 5: Summed performance measures

5.3 Validation

<<Not available in public version>>

5.4 Sensitivity Analysis

Based on the analysis performed, several scenarios are now proposed, that would improve manufacturing performance. A set of scenarios is tested that are likely to have a significant impact in decreasing average cycle time and work in process. It is beyond the scope of this study to determine costs of implementing these scenarios quantitatively.

- Add capacity to seam welding by hiring more seam welders
- Reduce seam welding process times
- Add capacity to blasting by installing more cabins
- Reduce outages at blasting, thus reducing both average process times and variability
- Increase work hours at welding or surface treatment

The program is able to deal with other kinds of scenarios as well. Input parameters such as capacity, process times, arrival rates and squared coefficients of variation hereof, can all be tested for impact on average cycle time. Additionally, the impact of adjusting shift hours can be determined.

Add capacity to seam welding by hiring more seam welders

It is possible to increase the number of seam welders, although this is not an easy solution. Generally, number of welders is already adjusted slightly throughout the year to reduce seasonality influences. The effects of increasing seam welders are explored here.

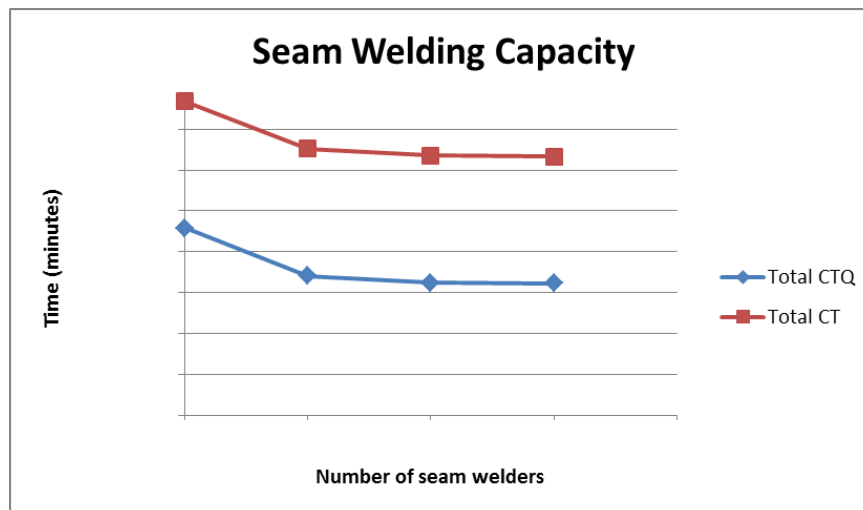


Figure 5: Total average cycle time and total average waiting time vs. number of seam welders

The decrease in waiting time is tremendous when one seam welder is added. Waiting time at seam welding alone drops by over <<censored>> hours. Adding another welder decreases cycle time even further, but the impact is clearly not as large.

This scenario may not directly lead to a clear recommendation, but it is a valuable insight that having more seam welder can have tremendous impact on the performance of the entire system.

Reduce seam welding process times

From the observations it was clear that welding times, at both tack welding and seam welding, could be decreased if the right steps are taken. The impact of lowering process times at seam welding are shown below. Performance measures were determined for different percentages of the current seam welding times, e.g. 90% of the current process times at seam welding shows the impact of a 10% process time decrease.

The observations indicate that there are possibilities for improvement at both tack welding and seam welding. Assuming it is possible to decrease process times at seam welding, it is interesting to know the impact of decreasing process times by a certain amount. The performance measures were determined for continuously decreasing process times, measured as a percentage of the current seam welding process time.

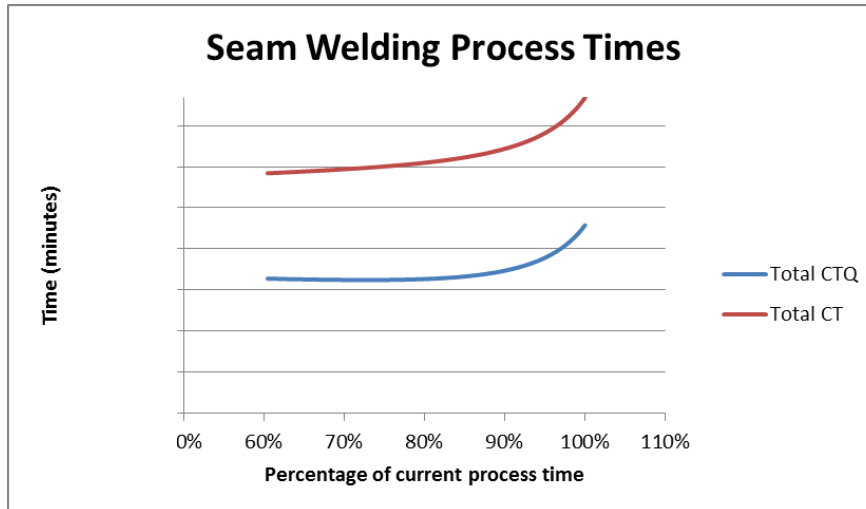


Figure 6: Total average cycle time and total average waiting time vs. percentage of current seam welding process times

The most radical decrease of cycle times is within the first ten percent decrease; reducing cycle times by ten percent would decrease total average cycle time by over <<censored>> hours, given the current characteristics of the system. It is noteworthy that a ten percent decrease in cycle times improves the system approximately as much as adding another seam welder.

Add capacity to blasting by installing a second cabin

There is a <<censored>> blasting cabin, but it is currently underutilized because it is only used for stainless steel. With some investments in equipment, it is possible to add another blasting cabin. The impact of having a <<censored>> blasting cabin was calculated.

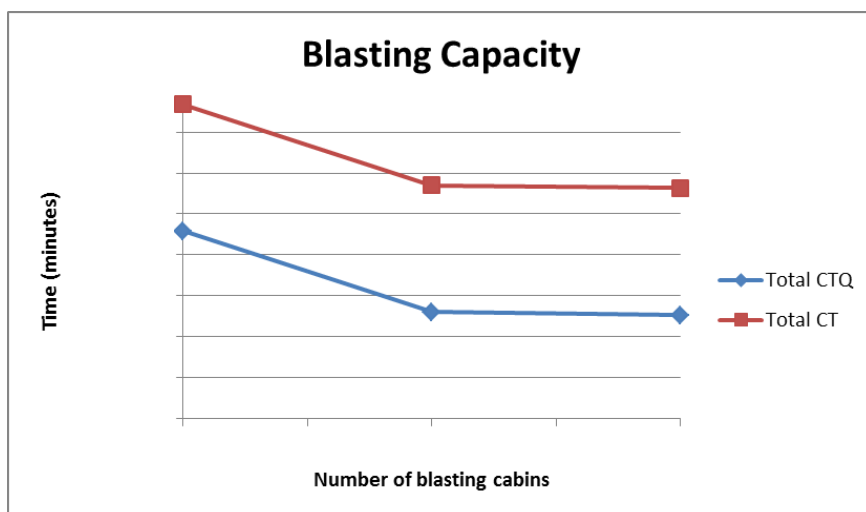


Figure 7: Total average cycle time and total average waiting time vs. number of blasting cabins

Although the capital investments are likely to be large, adding a <<censored>> blasting cabin reduces the waiting time drastically. It should be noted that allowing steel to be processed in the stainless steel blasting cabin introduces some setup costs. Furthermore, the current results are slightly overestimated: capacity is now considered twice as high, which might not be true for the smaller stainless steel cabin. Another interesting find is that adding a <<censored>> blasting cabin would hardly influence cycle times at all.

Reduce outages at blasting

There four possible ways of reducing the impacts of outages at blasting: reducing mean time to repair, increasing mean time to failure, decreasing the variability in repair times, and decreasing variability in time to next failure. Of these methods, reducing variability in repair times is the most interesting approach, since it was found that the squared coefficient of variation of repair times is quite high.

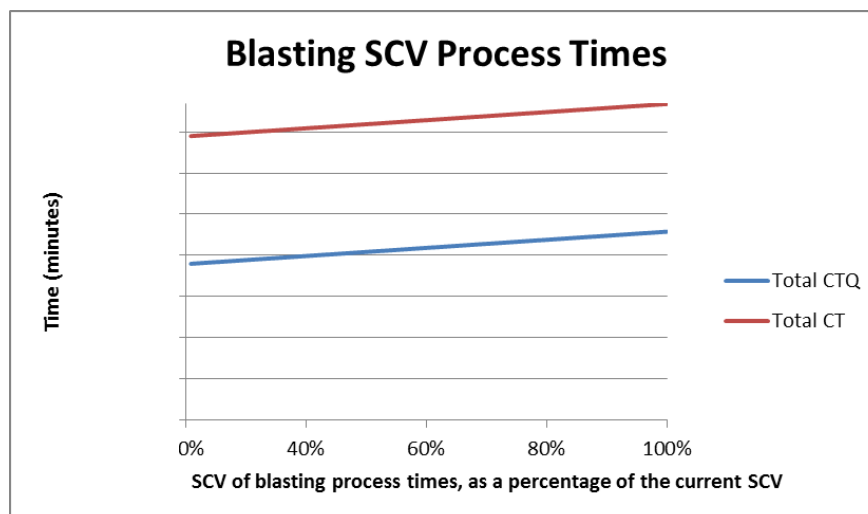


Figure 8: Total average cycle time and total average waiting time vs. percentage of current squared coefficient of variation of blasting process times

Although cycle times would clearly benefit from a less variable process at blasting, the impact is not as drastic as previously mentioned improvements. It is however noteworthy that the impact is not altered much for different blasting process times. Reducing variability in blasting process times is therefore best suited after increasing either blasting capacity, or decreasing blasting process times.

Increase work hours

The last explored option is to increase work hours. First the impact of increasing work time at welding is determined.

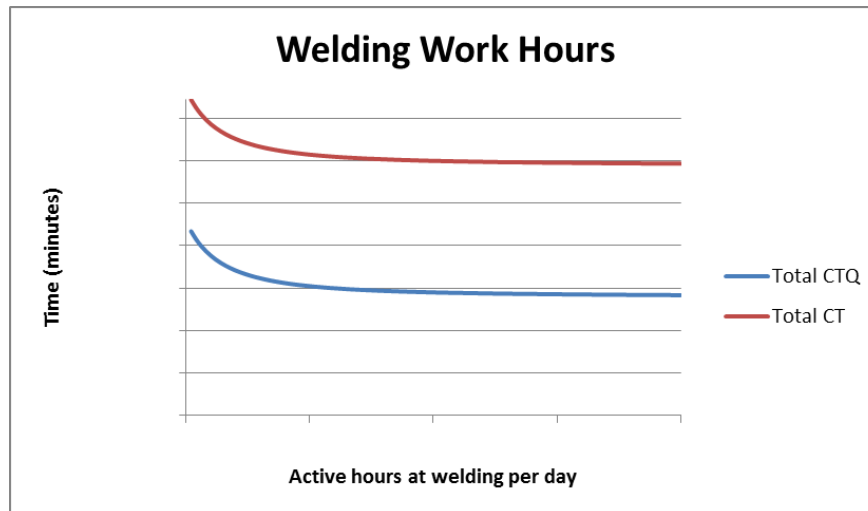


Figure 9: Total average cycle time and total average waiting time vs. welding work hours per day

By adding only two work hours to the current welding time per day, waiting times at welding are decreased significantly. The remaining waiting time is then primarily due to surface treatment.

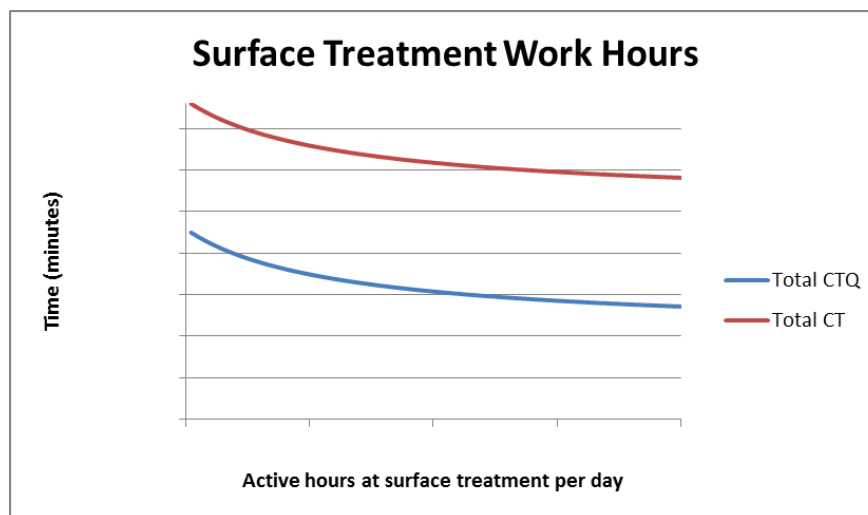


Figure 10: Total average cycle time and total average waiting time vs. surface treatment work hours per day

The impact of adding more work hours at surface treatment is not as big as increasing work hours at welding. Still, large improvements are possible by adding a third shift at surface treatment.

5.5 Conclusion Chapter 5

In this chapter, the details of the model were described, the current situation was analyzed, and sensitivity analysis was performed. In building the model, several assumptions had to be made:

the oven was modeled as <<censored>> parallel servers, the number of welders was based on the average number of welders throughout the year, and the average number of welders was assumed to represent number of parallel servers at welding. The current situation analysis was performed using the input established in Chapter 4. It was found that <<censored>> are the bottleneck processes. Finally the sensitivity analysis revealed several methods to increase performance measures, although a definitive assessment of the best option is difficult to establish.

Chapter 6: Conclusion & Recommendations

Main findings and targeted recommendations to prevent line congestion

In this report the complex manufacturing system at the welding and surface treatment facility of Aebi Schmidt Nederland was analyzed as an open queuing network. The thesis is structured along the lines of the standard research cycle. By answering a sub-question in each chapter, the main research question can be answered: How can long queues in the welding and surface treatment line be prevented?

In Chapter 2 the characteristics of the different classes of products and of the process itself were explained. Six processes were distinguished: tack welding and seam welding encompass the welding process, and blasting, primer coating, the oven, and final paint coating encompass the surface treatment process. Five different product types were distinguished: A top part, bottom part, tail piece, frame, and an integral design.

In Chapter 3 the performance measures that most intrinsically linked with line congestion are determined. Secondly, the analytical framework that is used to analyze the manufacturing system is explained. By using queuing network analysis, the cycle time and work in process can be determined analytically.

In Chapter 4 the required data was collected. The majority of the data was collected by analyzing time registrations. Additional data was acquired by inquiry. The input required for the queuing network analysis consisted of process times per product type per process, and arrival rate per product. Since arrivals only enter the system at tack welding, only the arrival rate at tack welding was determined. Additional input were measures of the statistical first moments of these quantities: the squared coefficient of variation of process times and arrival rates.

In Chapter 5 the analysis was performed using the input parameters of Chapter 4. First the assumptions and methods of modeling the process were discussed: several difficulties, such as determining the capacity of stations and dealing with different work hours were treated. Second, the current situation was analyzed. It was found that <<censored>> are the main contributors to waiting times. Finally, by altering the current situation, a sensitivity analysis was performed to test the impact of different process improvements. Because the cost of implementation is not known, it is difficult to assess the best option.

Final Conclusions

Regarding the method and data collection, the following conclusions are drawn:

- The time registrations are prone to many errors and inconsistencies, which makes the data collection in this research more time consuming than necessary.
- The analysis can be performed with a computer tool with great generality, responsiveness and usefulness.

Regarding the analysis, the following conclusions are drawn:

- <<Not available in public version>>
- Improvements at seam welding and blasting are generally more impactful than improvements at other stations.
- Six different scenarios were tested. Quantitative cost estimation of each scenario is beyond the scope of this research, but a qualitative classification of scenarios reveals the following results:

Scenario	Impact	Implementation Cost
Increase seam welding capacity	Medium	Medium
Decrease seam welding process times	High	Medium
Increase blasting capacity	Very high	High
Decrease blasting SCV of process times	Low	Low
Increase shift hours at welding	High	Medium
Increase shift hours at surface treatment	Medium	Medium

Table 6: Different scenarios and their impacts and implementation costs

Recommendations

Based on the conclusions above, several recommendations are made.

- Improve time registration. The cheapest and easiest way to do this is to use forms, and making sure workers use a uniform way of recording their work. A more elaborate but ultimately more efficient approach is to implement a tracking system. Using a system with barcodes that allows tracking of the product throughout the process and automatically records process times would greatly improve the validity of the research and the manageability of the method in general.

- The two processes that contribute most to average waiting time are <<censored>>. It is recommended to further investigate the possibilities of increasing capacity or decreasing process times of these stations.
- The most impactful scenario is to add another blasting cabin, but this is also the most expensive. Based on the sensitivity analysis, it is recommended to look for solutions that decrease seam welding process times, or to increase the number of work hours at welding.
- Attempting to decrease outages at blasting will help prevent congestion, and is probably not a very expensive option. It is recommended to implement preventive maintenance or other methods to reduce variability in mean time to repair and mean time to failure.

Bibliography

- Aebi Schmidt. (n.d.). Retrieved July 20, 2013, from Aebi Schmidt Website: <http://www.aebi-schmidt.com/en/home>
- Dröge, R., De Jong, A., Klijnsma, T., & Reimert, M. (2013). *Final Report Project 2: Organization and Process Analysis*.
- Heerkens, H., & Van Winden, A. (2012). *Geen Probleem: Een aanpak voor alle bedrijfskundige vragen en mysteries*. Nieuwegein: Van Winden Communicatie.
- Hopp, W., & Spearman, M. (2001). *Factory Physics*. New York: Irwin/McGraw-Hill.
- Kaplan, R. (1983). Measuring Manufacturing Performance: A New Challenge for Managerial Accounting Research. *The Accounting Review*, 58(4), 686-705.
- Kingman, J. (1961). The single server queue in heavy traffic. *Mathematical Proceedings of the Cambridge Philosophical Society*, 57(4), 902-904.
- Pinedo, M. (2009). *Planning and Scheduling in Manufacturing and Services* (2nd ed.). New York: Springer.
- Sakasegawa. (1977). An Approximation Formula $L_q \approx \alpha \rho^3 / (1 - \rho)$. *Annals of the Institute of Statistical Mathematics*, 29(1), 67-75.
- Son, Y., & Park, C. (1987). Economic Measure of Productivity, Quality and Flexibility in Advanced Manufacturing Systems. *Journal of Manufacturing Systems*, 6(3), 193-207.
- Whitt, W. (1982). Approximating a Point Process by a Renewal Process, I: Two Basic Methods. *Operations Research*, 30(1), 125-147.
- Whitt, W. (1983). The Queueing Network Analyzer. *The Bell System Technical Journal*, 62(9), 2779-2815.
- Winston, W. (2004). *Operations Research: Applications and Algorithms* (4th ed.). Belmont: Brooks/Cole, Cengage Learning.
- Wolff, R. (1982). Poisson Arrivals See Time Averages. *Operations Research*, 30(2), 223-231.
- Zijm, W. (2012). *Manufacturing and Logistic Systems Analysis, Planning and Control*. Enschede: University of Twente.

Appendix A: Data cleaning of the Time Registrations

<<Not available in public version>>

Appendix B: Queuing Networks

The single station models can be expanded by allowing flows between single stations, establishing a network. Here the performance measures for open Jackson-networks are explained.

The basic queuing network is an open queuing network with exponential process times, and arrivals from outside the network according to a Poisson process. This type of network is called an open Jackson-network, and it is exactly solvable (Zijm, 2012).

The required input for a single Jackson-network is the arrival rate of jobs from outside the system at each station (since it is not necessary that jobs only arrive from outside the system at the first station) and the process (or service) times. Additionally a routing matrix is required, which describes the flow of jobs through the system. Obviously, the number of stations and the number of servers per station are also required.

The total arrival rate per station is calculated as follows:

$$\lambda_i = \gamma_i + \sum_{j=1}^M \lambda_j P_{ji}$$

In which λ_i is the total arrival rate of jobs at station i , γ_i is the arrival rate of jobs at station i from outside the system, P_{ij} is the probability a job leaving station i will go to station j , and M the number of different stations in the network, often referred to as 'nodes'. This equation can be written in matrix form, with the vector λ , the routing matrix P , and the vector γ :

$$\lambda = \gamma + \lambda P$$

Since γ and P are input, solving this equation is trivial, given that λ exists (Zijm, 2012).

The next quantity of interest is the visit ratio per station, which is a measure of how frequently a job visits a station. It is defined as:

$$V_i = \frac{\lambda_i}{\sum_{i=1}^M \gamma_M}$$

Open Jackson-networks are quite easily solvable, because the Poisson arrival process allows one to calculate performance measures for each node independently (Zijm, 2012). Network throughput is given by:

$$TH = \sum_{i=1}^M \gamma_i$$

Each station behaves as a M/M/c queuing model, so computing performance measures per station is relatively straightforward:

$$CT_{Q,i} = \frac{u_i^{(\sqrt{2c_i+2}-1)}}{c_i(1-u_i)} \cdot \frac{1}{\mu_i}$$

c_i being the number of parallel servers at station i . From this formula all other performance measures can be determined by using Little's Law and the visit ratios.

Appendix C: Causes of Variability

Batching and Outages

When the squared coefficients of variation of interarrival times are determined, it is helpful to realize some inherent causes of variability. In this appendix the influence of serial batching and preemptive outages on variability is shortly discussed.

Hopp & Spearman (2001, p. 264) give the following example:

“... suppose a forklift brings 16 jobs once per shift (eight hours) to a work station. Since arrivals always occur in this way with no randomness whatever, one might reasonably interpret the variability and the CV [coefficient of variation] to be zero.

However, a very different picture results from looking at the interarrival times of the jobs in the batch from the perspective of the individual jobs. The interarrival time for the first job in the batch is eight hours. For the next 15 jobs it is zero. Therefore, the mean time between arrivals t_a is one-half hour, and the variance of these times is given by:

$$\sigma_a^2 = \frac{1}{16} \cdot 8^2 + \frac{15}{16} \cdot 0^2 - t_a^2 = 3,75$$

The arrival SCV [squared coefficient of variation] is therefore

$$c_a^2 = \frac{3,75}{0,5^2} = 15$$

In general, if we have a batch size k , this analysis will yield $c_a^2 = k - 1$ ”

The question is now what the real squared coefficient of variation is: zero, or 15. Hopp & Spearman reason that the system will behave somewhere in between. It is known that some product parts are produced in serial batches at Aebi Schmidt, although there is no standard batch size. The example illustrates that this may lead to high measures of variability in arrival rates for some products.

A second prominent reason for variability is preemptive outages, or ‘breakdowns’. Hopp & Spearman (2001, pp. 256-258) identify four characteristics of a breakdown: the mean time to failure (denoted with m_f), the mean time to repair (denoted with m_r), and the coefficients of variation of both these numbers (denoted with c_f and c_r respectively).

Breakdowns directly influence cycle time through availability: a machine is not available when it is broken. Availability is calculated as:

$$A = \frac{m_f}{m_f + m_r}$$

The effective mean service time is then longer than the mean service time if there had not been breakdowns:

$$t_e = \frac{t_0}{A}, \quad r_e = Ar_0$$

r_e is the effective rate, which is lower than the rate without breakdowns, r_0 .

c_f is often assumed to be 1; this assumption is generally valid since it portrays a memoryless property of the mean time to failure. Since in practice equipment is usually a combination of old and new components, and outages can occur in both new and old components, the memoryless property of the exponential distribution seems realistic (Hopp & Spearman, 2001, p. 257).

The squared coefficient of variation of the effective service time, assuming $c_f = 1$, then equals (Hopp & Spearman, 2001, p. 257):

$$c_e^2 = c_0^2 + (1 + c_r^2)A(1 - A)\frac{m_r}{t_0}$$

It is noteworthy that the difference between the squared coefficient of variation of the process with and without breakdowns depends linearly on the mean time to repair; this is the reason it is usually preferred to have short, frequent breakdowns instead of long, infrequent breakdowns.

Appendix D: Disaggregated performance measures

The disaggregated performance measures for the current situation

<<Numerical data not available in public version>>

Name	Bottom Part		Blasting	Primer	Oven	Final Paint
	Tack Welding	Seam Welding				
CT (minutes)						
rho						
WIP						
WIPQ						

Name	Top Part		Blasting	Primer	Oven	Final Paint
	Tack Welding	Seam Welding				
CT (minutes)						
rho						
WIP						
WIPQ						

Name	Tail Piece		Blasting	Primer	Oven	Final Paint
	Tack Welding	Seam Welding				
CT (minutes)						
rho						
WIP						
WIPQ						

Name	Frame		Blasting	Primer	Oven	Final Paint
	Tack Welding	Seam Welding				
CT (minutes)						
rho						
WIP						
WIPQ						

Name

CT (minutes)

rho

WIP

WIPQ

Integral Design

Tack Welding

Seam Welding

Blasting

Primer

Oven

Final Paint
