



Universiteit
Leiden
The Netherlands

Efficient constraint multi-objective optimization with applications in ship design

Winter, R. de

Citation

Winter, R. de. (2024, October 8). *Efficient constraint multi-objective optimization with applications in ship design*. Retrieved from <https://hdl.handle.net/1887/4094606>

Version: Publisher's Version

License: [Licence agreement concerning inclusion of doctoral thesis in the Institutional Repository of the University of Leiden](#)

Downloaded from: <https://hdl.handle.net/1887/4094606>

Note: To cite this publication please use the final published version (if applicable).

Chapter 4

Empirical Design Optimization Approach

As described in the problem characteristics of Chapter 3, the quickest and most coarse way to estimate the main particulars of a new design is by using and learning from data of reference vessels. This already is part of the answer to research question 3: *How can data be used to find feasible Pareto efficient ship design solutions?* In the current chapter, a new method is introduced that uses reference data, machine learning algorithms, and an optimization algorithm to find suggestions for the main particulars and KPIs of new ship designs. With this new method, designers can make more informed decisions in the preliminary design phase where very limited information is available and decisions need to be made in a short amount of time. However, it is in the preliminary design phase where the most influential decisions are made regarding the global dimensions, the machinery, and therefore the performance and costs. In this chapter, it is shown that a machine learning algorithm trained with data from reference vessels is more accurate when estimating key performance indicators compared to existing empirical design formulas. Finally, the combination of the trained models with optimization algorithms proves to be a powerful tool for finding Pareto-optimal design solutions from which the naval architect can learn. Although the application domain of this chapter is ship design, the approach can also be transferred to other application domains.

4.1 Introduction

In the preliminary design stage, more knowledge should be used when making decisions. With the method described in this chapter, naval architects are better supported by data to make design decisions instead of relying only on instincts, knowledge, and experience. The decision support is in the form of machine learning models which can be used to validate ideas, assumptions, and design variations. This helps the naval architect avoid innovation risks and to find better design variations. On top of this, without much additional effort, the naval architect can use the trained machine-learning models in combination with an optimization algorithm. This optimization algorithm can then be deployed for searching advantageous and competitive design variations. The only requirement for the reference optimizer that is proposed in this chapter to work adequately is enough relevant ship data and a good design problem setup.

4.2 Data Description for Reference Studies

The solution proposed in this chapter utilizes the power of empirical design methods in combination with parametric optimization. However, for empirical design method to work properly, data is needed. Fortunately, a lot of data services have become available for the maritime industry. The most prominent ones are:

WORLD FLEET REGISTER is a ship data and intelligence platform from Clarksons Research with data about ship earnings, vessel parameters, and new-build data [40].

SEA-WEB collects static ship data of existing and even scrapped ships and tracks vessels worldwide [136].

BRL SHIPPING CONSULTANTS has a subscribers area where reports are available about the active fleet and about newly built vessels [26].

MARINE TRAFFIC is a platform that allows even without logging in to obtain the location of vessels plus general static ship data [98].

AISHUB is an AIS data sharing platform where you can get access to global AIS radar stations when you join with your own AIS antenna [2].

All this static and operational data has been collected and aggregated into more than 100 particular data fields per vessel and a large database with historical locations of ships. Examples of collected data fields are: Length, breadth, draft, block

coefficient, light ship weight, dead-weight, maximum continuous rating of the engines, maximum speed, but also more ship specific data fields for specific ship types such as: Bollard pull, passenger capacity of urban transportation vessels, number of car lanes, crane capacity for offshore vessels, hopper volume of dredgers, and ice class qualification.

This data can be used in a reference study in the preliminary ship design process since design trends can be visualized, design trends can be learned, and gaps and competitive advantages in the market can be found.

4.2.1 Visualizations

After the relevant parameters for a vessel type have been selected, the parameters can be summarized and plotted. When three parameters are relevant it is still possible to visualize them in two dimensions. As an example, the dead-weight and moulded breadth together with the Twenty-foot equivalent Unit (TEU) capacity is given for a set of container vessels in Figure 4.1.

However, it is often the case that more than three parameters are relevant in the preliminary ship design stage, which makes it challenging to visualize. To still be able to investigate a selection of ships or design variations with more than three parameters, parallel coordinate plots can be used from Section 2.5.2. In Figure 4.2, a parallel coordinate plot is made for several hundred container vessels with a length between perpendiculars between 175 and 200 meters.

Interpretation of Visualizations

As can be inspected from Figure 4.1, the moulded breadth has a maximum of 32.4 meters, a well-known maximum width for ships to still be able to pass through the Panama Canal. This maximum moulded breadth can also be seen in Figure 4.2. Moreover, it is now also possible to simultaneously see all other relevant parameters of the container vessels. For example, the limited draught for the majority of vessels in this selection is smaller than or equal to 12 meters, also an important Panama Canal dimension. Besides this, one can simultaneously see the conflicting relationship between block coefficient (C_b), and Maximum Continuous Rating (MCR) and their influence on service speed. The vessels with a high block coefficient, and small maximum continuous rating, also have a slow service speed and vice-versa. When designing new vessels, these plots can be very helpful for the designers.

4.2. Data Description for Reference Studies

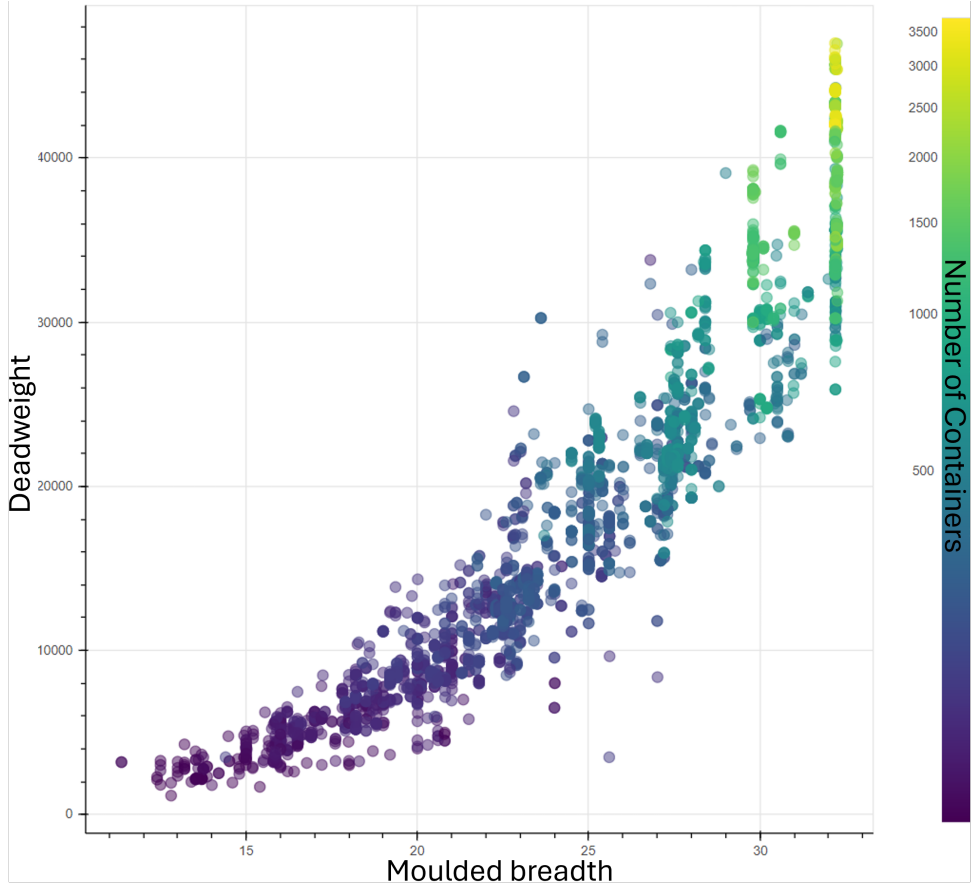


Figure 4.1: Container vessels color-coded by container capacity.

4.2.2 Data Pre-processing

Preliminary data analysis showed duplicate vessels and vessels which are very similar. To make sure that specific vessels are not over-represented, but still enough data is available, data pre-processing must be done. The pre-processing consist of three steps and is done so that machine learning algorithms can be trained with *cleaner* data.

1. All except for one of the vessels with exact duplicates must be deleted. Ships are considered to be duplicates if their *gross tonnage*, *length between perpendiculars (Lbp)*, *breadth overall (Boa)*, *draught (T)*, and *MCR* are equal.
2. All but one vessel out of a series of sister vessels are deleted. If the earlier

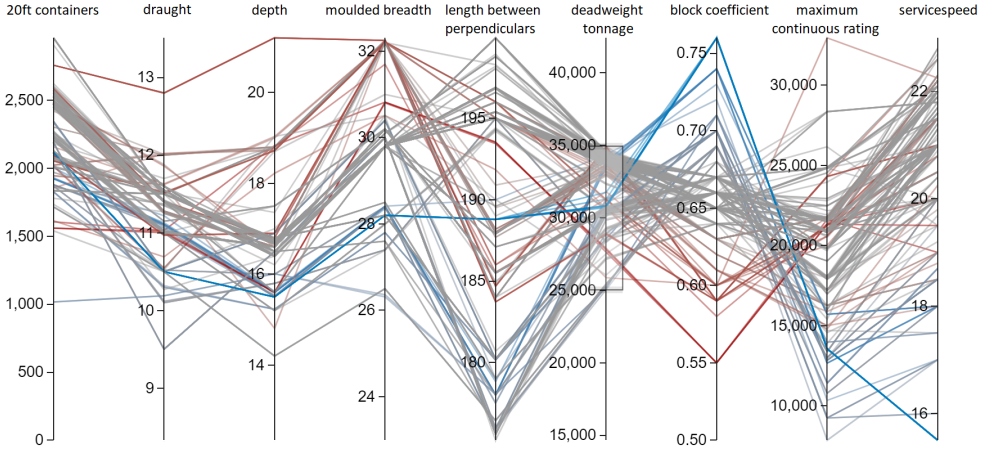


Figure 4.2: Parallel coordinate plot of container vessels color-coded by block coefficient, and a deadweight tonnage selection between 25000 and 35000 tonnes.

mentioned variables are all within 1 percent of each other, the vessels are marked as too similar.

3. A second degree polynomial and interacting features are created. The two degree polynomials and interacting features of the example $[a, b]$ would be: $[a, b, ab, a^2, b^2]$.

Reasons for deletion of duplicates and very similar vessels are to prevent the potential over-fitting of machine learning models. If a series of sister vessels would be present, the machine learning model would automatically put more weight on the sister vessels compared to one unique vessel. A second argument to delete sister vessels is, once a machine learning model has learned from a vessel, a second sister vessel does not add much knowledge but will only add computation and training time.

The second degree polynomial and interacting features are created to generate more potentially interesting features from the design parameters that are known. This way, the machine learning models have more features to learn from which potentially leads to more accurate results.

4.3 New Empirical Design Methodology

This section describes how the new empirical design method is used in combination with an optimization algorithm in the so called reference optimizer. As mentioned in the related work section, it is often the case that the empirical design equations are not available for a specific ship type or that the available equations are outdated. This is unfortunate since designing ships with wrong or outdated design equations will most likely not lead to optimal decisions. The empirical design equations are therefore replaced with machine learning models. These machine learning models make sure that it is no longer need to solely depend on predefined equations or the experience and knowledge of naval architects.

Machine learning models are used to learn the relationships, similarities, and trends between hundreds of data points. However, for machine learning models to work properly, the relationship between the dependent and independent variables need to be learned. The dependent and independent variables are chosen by the naval architect. The machine learning models learn the relation between the independent and dependent variables in the training phase. After the training phase, the trained machine learning models are coupled to an optimization algorithm that can exploit the trends learned and search for optimal design configurations that outperform the existing designs.

4.3.1 Setup Design Challenge

For the machine learning algorithm to work well a design challenge should be set up by the user. The design challenge consists of three parts, the design variables, the constraints, and finally the objectives as described earlier in Section 3.2.2.

DESIGN VARIABLES are set up by choosing the design parameters that have a significant influence on the final design and which are allowed to vary. The allowed variations in the variables are controlled with a user-defined lower and upper limit. However, the limit can not be smaller or larger than the smallest and largest ship in the collection. Examples of a set of design variables are: *Length between perpendiculars (Lbp)*, *draft (T)*, *Breadth overall (Boa)*, *block coefficient(Cb)*, and *service speed (V)*.

CONSTRAINTS are also set by the user. The design constraints are typically hard limitations or strong wishes for the to-be-designed ship. Examples of constraints

are *deadweight (DWT)* capacity of 30,000 tons, a *cargo capacity* of 2000 TEU, a *length overall* smaller than 180 meters, or a *draught* of not more than 12 meters.

OBJECTIVES of a ship design are usually the key performance indicators that deal with operational expenses and investments. Ideally, they are as low as possible, however, they most often do not go hand in hand and are most of the time conflicting. Examples of three objectives are: minimizing the *light ship weight (LSW)*, maximizing the *deadweight (DWT)* capacity, and minimizing the *Maximum Continuous Rating (MCR)* of the main engines.

Once the design variables, constraints, and objectives have been set by the user, the relationship between the variables and the constraints and objectives can be learned.

4.3.2 Random Forest Regression

A random forest regression model [25] can be used to learn the relationship between the features and one target variable. A random forest regressor is chosen because it is robust against outliers and overfitting and because it can deal with discrete parameters which comes in handy as the data used for training comes from existing ships and might not always be 100% reliable. In the new empirical design methodology the features are the design variables plus the polynomial features and the target variable is one of the constraints or one of the objectives. Therefore, for each constraint, and for each objective a new unique random forest regression model is trained.

The random forest regression model learns the relation between the features and the target by fitting a multitude of decision trees. One decision tree is fitted to learn the relation between a set of random selected features with the corresponding target values. The data with the random selected features is sequentially greedily split into two sub-samples based on one of the features until the number of samples in the nodes reaches a threshold value. Resulting in an upside-down tree with nodes, branches for splits, and leaves with similar target scores.

Once e.g. 100 decision trees have been trained with the 100 randomly selected feature sets, the random forest is done training. The trees in the forest can be traversed which makes a prediction of the target variable for an unseen combination of feature values possible for each tree. These 100 outcomes of the 100 decision trees are then averaged into a final prediction. The process of making a prediction is visualized in Figure 4.3. Because a multitude of trees are fitted, the random forest regression model is robust against outliers in the training data. However, due to the fact that the final

4.3. New Empirical Design Methodology

score depends on the average of all the trained trees, the random forest regression model is not capable of extrapolation.

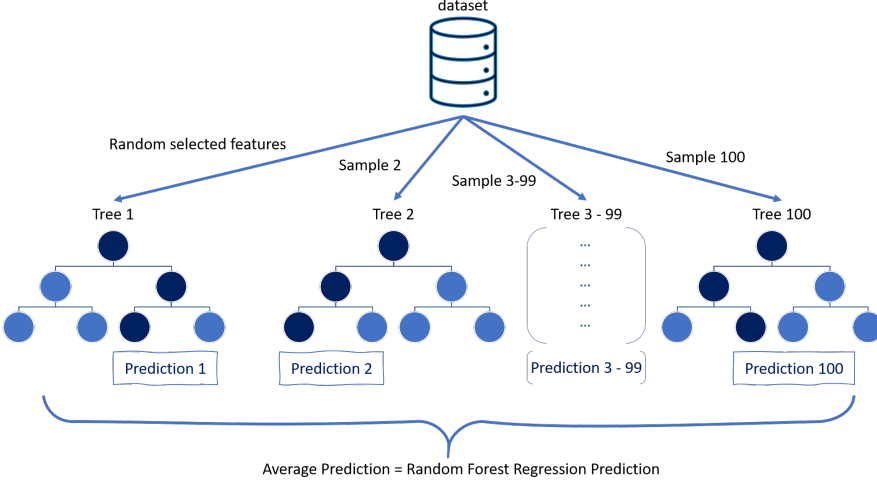


Figure 4.3: Random Forest Regression Model Illustration.

4.3.3 Isolation Forest

In the reference optimizer not only the user-defined constraints limit the search space. The search space is also limited by an anomaly detection algorithm. The anomaly detection algorithm used is named Isolation Forest [94]. Isolation forest is an unsupervised machine learning algorithm that tests how easy it is to isolate certain data points. It does so by recursively splitting the data by randomly selecting a variable and a random split value between the lower and upper limits. If a sample is easy to isolate by randomly splitting the data set, it is marked as an anomaly. A sample that is hard to isolate versus a sample that is easy to isolate is visualized in Figure 4.4.

In practice, this means that in case a design variation is unique and lies outside of the trend, or if the database contains a ship with length by accident reported in feet instead of meters, it is marked as an anomaly. When searching for a new design variation, design variations that are marked as anomalies by the isolation forest will no longer be considered. This is the case because they do not follow the pattern and therefore their prediction is probably incorrect. On top of this, the isolation forest will make sure that the design variations will not exceed the limits, so that the random forest regression model is not forced to extrapolate.



Figure 4.4: Illustration of Isolation forest with easy to isolate sample on the left and a hard to isolate sample on the right.

4.3.4 Design Problem Optimization

The reference optimizer searches for Pareto optimal designs that do not violate any of the constraints. This can be done with any multi-objective optimization algorithm that can deal with constraints but in the reference optimizer it is done with the Non-dominated Sorting Genetic Algorithm II (NSGA-II) [49]. NSGA-II optimizes the design challenge by modifying the design variables. NSGA-II is allowed to vary the design variables between the user-defined lower and the upper limit. The design variations that come out of NSGA-II are evaluated on the random forest regression models to predict the constraint and objective scores. The objective and constraint scores are then combined with the design variable values and tested to see if the combination can be easily isolated by the Isolation Forest. Once the isolation score, the objective score, and the constraint scores are evaluated they are given back to the NSGA-II algorithm. The NSGA-II algorithm includes the evaluated design variations in the population of previously evaluated solutions and then new solutions are generated with the non-dominated genetic sorting strategy. After NSGA-II has converged, the optimal designs are reported so that they can be inspected with a parallel coordinate plot and the objectives can be visualized on a Pareto frontier.

4.4 Empirical Design Experiments

To validate the models and the algorithms, different experiments are conducted. The first experiment is set up to test the predictive capabilities of the random forest regression models. In the second experiment, a set of ships is intentionally modified to see if the isolation forest is capable of identifying the newly created anomalies. Finally, in the last experiment, everything is connected and novel container ship variations are generated on a Pareto frontier.

For all the experiments 2538 container ships are used. 1219 of these vessels have the duplicate characteristics and are filtered out during the preprocessing phase as described in Section 4.2.2.

4.4.1 Random Forest Regression Experiment

The random forest regression models are intended to predict the performance and capital investment cost of the ships of the future. In this experiment, such a situation is mimicked. Three different KPIs are learned by the random forest regression models with data from 1019 ships built before 2005, and then the random forest regression models are tested with data from 96 ships built after 2010. By comparing the predicted values with the actual values it can be determined if the trained random forest regression model is good to use in practice.

The KPIs that are predicted in this experiment are LSW, MCR, and DWT. The KPIs are estimated with the random forest regressor and with empirical design equations for the specific KPIs. The design variables used to predict LSW are [*Lbp*, *Boa*, *T*, *Cb*, *MCR*]. The design variables used for MCR are [*Lbp*, *Boa*, *T*, *Cb*, *V*]. The design variables used to predict DWT are [*Lbp*, *Boa*, *T*, *Cb*].

Random Forest Regression Results

The accuracy of the random forest regression model is determined with the R^2 measure [104]. This measure compares the real KPI values with the predicted values and see how much variation of the dependent variable can be explained by the model. With R^2 scores of 0.93, 0.90, and 0.95 for LSW, MCR, and DWT, it can be confirmed that the random forest regressor is capable of capturing a lot of variance.

Empirical Design Equation Results

Light Ship Weight for container vessels can be predicted with the Empirical Design Equation of D'almeida [55]. The prediction of D'almeida for LSW is dependent on the *steel weight* (SW), *outfitting & equipment weight* (OEW), and *machinery weight* (MW):

$$\begin{aligned} LSW &= SW + OEW + MW \\ SW &= 0.0293 \cdot Lbp^{1.76} \cdot Boa^{0.712} \cdot T^{0.374} \\ MW &= 2.35 \cdot (MCR/0.745699872)^{0.60} \end{aligned} \tag{4.1}$$

The D'almeida equations use the same independent variables as in the random forest regressor model to calculate the LSW. However, the estimate of this empirical formulation only obtains an R^2 score of 0.84.

Maximum Continuous Rating can be estimated with the empirical formula named the *Admiralty constant* [129]. The admiralty constant C can be calculated by using the maximum continuous rating (MCR) and displacement values (Δ) from reference vessels and then plugging it in the following formula:

$$MCR = \frac{\Delta^{2/3} \cdot V^3}{C} \tag{4.2}$$

The mean admiralty constant (C) of the reference vessels is then stored so that it can be used in later approximations for MCR given different displacements. In the experiment, the *Admiralty constant* itself (C) is approximated with the reference vessels from before 2005. The mean C from the vessels before 2005 is used to make predictions for the container vessels after 2010. The R^2 score for this formula is 0.87, again a worse R^2 score compared to the random forest regressor.

Dead Weight Tonnage is dependent on the LSW of the vessel. The empirical formula for DWT is:

$$DWT = \Delta - LSW \tag{4.3}$$

but since the empirical equation for light ship weight has a worse R^2 score compared to the random forest regressor, it is no surprise that also for DWT, the R^2 score of 0.89 is lower compared to the R^2 score of the random forest regressor.

4.4. Empirical Design Experiments

4.4.2 Isolation Forest Experiment

In the isolation forest experiments, the isolation forest is trained with the data from the container vessels as described earlier. After this, two data fields per vessel are modified to create impossible design parameter/KPI combinations. All vessels are then evaluated by the trained isolation forest to see if they are marked as an anomaly or not. The modified design parameters/KPI values and the percentage of anomalies detected are presented in Table 4.1.

Modified Columns	Anomaly Percentage
no modification	15%
$Lbp/1.1, T \times 1.1$	36%
$Lbp/1.25, T \times 1.25$	56%
$Lbp/1.5, T \times 1.5$	88%
$Lbp/1.75, T \times 1.75$	99%
$Lbp/2, T \times 2$	100%
$MCR/2, V \times 2$	95%
$LSW/2, Cb \times 2$	97%
$Lbp/2, DWT \times 2$	97%
$Cb/2, Lbp \times 2$	100%
$T/2, Cb \times 2$	100%
$Cb/2, V \times 2$	100%

Table 4.1: Modified columns and classified anomaly percentage after this modification.

The experiments indicate that as the vessels undergo more significant modifications, the isolation forest identifies an increasing percentage of vessels as anomalies. The experiments also show that there is a small percentage of vessels that have been radically changed but have not been marked as an anomaly. This indicates that the anomaly detection algorithm does not detect all anomalies and that the naval architect should pay attention when analyzing the results and do a few integrity checks on the results.

4.4.3 NSGA-II Experiment

For this experiment, it is assumed that the random forest regression model and the isolation forest perform as intended so that the NSGA-II algorithm can be tested. If the NSGA-II algorithm can find feasible and realistic Pareto-optimal solutions it can be confirmed that the reference optimizer works as intended.

The reference optimizer is tested again on a container ship case. In this experiment NSGA-II was allowed to vary the main particulars of the vessel (LBP , Boa , T , Cb ,

V). The LSW and MCR are minimized, while the DWT capacity should be larger than or equal to 28000 tonnes. The results are visualized on the Pareto frontier in Figure 4.5 and Figure 4.6.

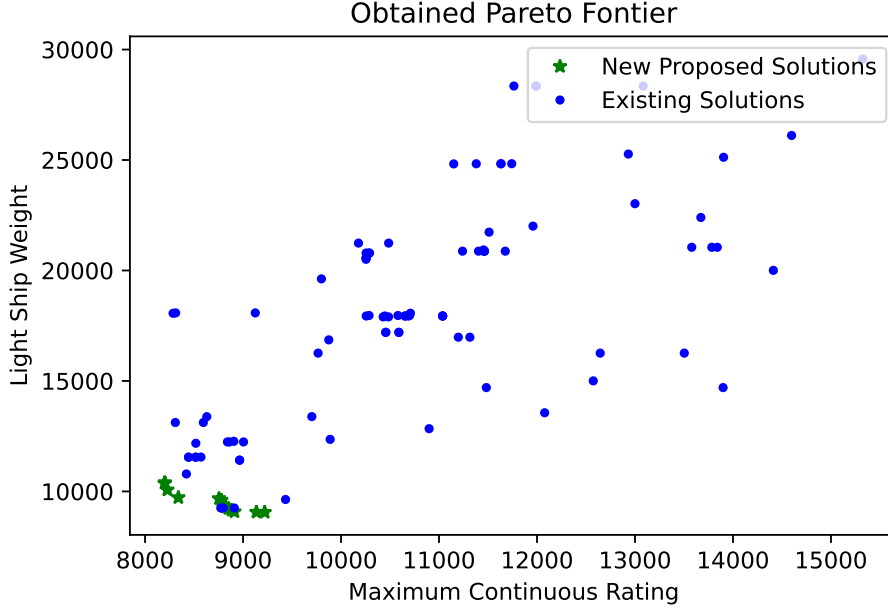


Figure 4.5: Obtained Pareto efficient solutions for test case. Blue solutions indicate existing vessels that do not violate any of the constraints while green solutions are the proposed solutions by NSGA-II.

NSGA-II in this experiment found 14 Pareto efficient solutions along the Pareto front interposed by Pareto efficient existing vessels. As previously described the algorithm only uses data and does not know any physics, it is the task of the naval architect to double-check the feasibility of the proposed solutions. In this experiment, the physical integrity of the proposed solutions are checked with the DWT Equation 4.3.

The weight balance of the vessel i.e. the sum of the DWT and LSW need to be in line with the corresponding displacement that can be calculated with: $Lbp \cdot Boa \cdot T \cdot Cb \cdot \rho$. Here $\rho = 1.025$ which is the water density of salt water. The found maximum deviation for existing vessels is 7% with an average of 0.2%. This indicates that data from the existing vessels is not always 100% accurate. The found maximum deviation for the proposed vessels by the reference finder is 2.6% with an average of 1.8%.

To give more details of the obtained Pareto efficient solutions, a parallel coordinate

4.5. Discussion

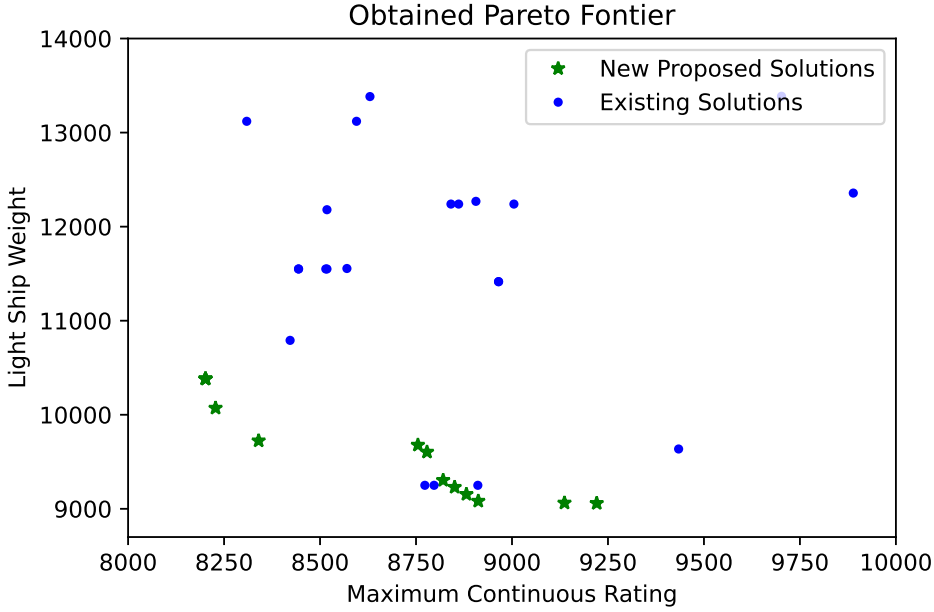


Figure 4.6: Zoomed in on Pareto frontier for test case. Blue solutions indicate existing vessels that do not violate any of the constraints while green solutions are the proposed solutions by NSGA-II. Inspection of this figure reveals that the majority of the existing solutions are dominated by the newly proposed solutions.

plot is made and presented in Figure 4.7. In the parallel coordinate plot, the main dimensions and the resulting performance indicators can be inspected and compared with the existing vessels.

4.5 Discussion

The reference optimizer as introduced and described in this chapter has two drawbacks that hold for any optimization process. The first drawback of the reference optimizer is that it needs a sufficient amount of good data for the random forest regressors to make accurate predictions. Good data without mistakes is important since otherwise, the random forest regressors will learn a wrong trend and the predictions will be off. Accurate and 100% reliable data is difficult to gather (especially for the less common vessel types) and sometimes only possible to obtain with expensive subscriptions.

A second drawback of the reference optimizer is that the design challenge should

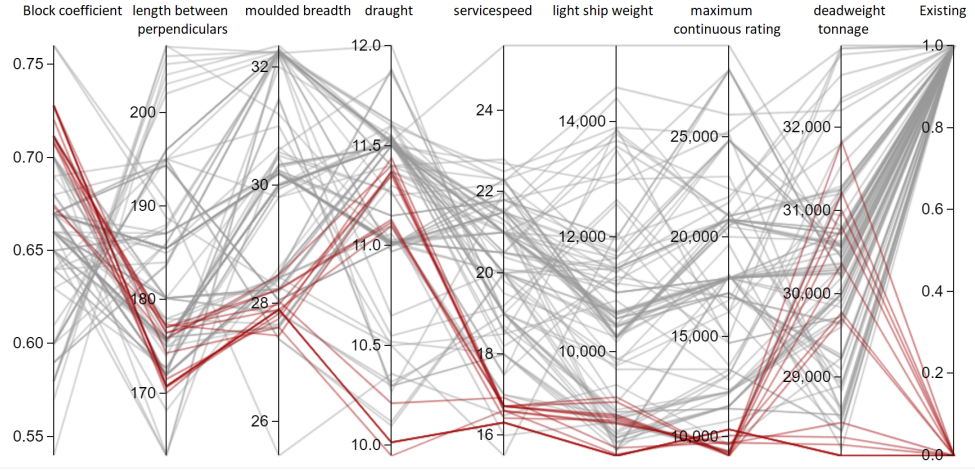


Figure 4.7: Parallel coordinate plot of the proposed solutions in red versus the existing vessels in grey.

be set up properly. For this, a naval architect needs to learn the basics of machine learning algorithms. During the training of naval architects at least the choice of what to choose as independent and dependent variables should be addressed in combination with different performance metrics.

4.6 Conclusion and Future Work

In this chapter, an alternative generic way is presented on how naval architects can use data to make preliminary design decisions by visualizing the data, learning from the data with machine learning algorithms, and finally finding optimal configurations with optimization algorithms.

The experiments in this chapter show that random forest regressors can give better estimations for light ship weight, dead weight, and maximum continuous rating compared to empirical design equations often used by naval architects. Besides a better estimation of key performance indicators, the random forest regressors are also capable of predicting key performance indicators for which no empirical design equations are readily available in the literature.

After training the random forest regressor and an anomaly detection algorithm, the models are coupled to a multi-objective optimization algorithm. This setup is capable of automatically generating optimal design configurations for preliminary ship

4.6. Conclusion and Future Work

design problems. As a practical use case, a container vessel design challenge has been executed. The setup proposed 14 new Pareto-efficient solutions. The preliminary designs consisted of the main particulars of the vessels plus the key performance indicators light ship weight, maximum continuous rating, and deadweight. After this, the preliminary designs have been validated with integrity checks to verify their feasibility.

For future work, it is intended to improve the performance of the machine learning models even further, train them with more accurate data, and integrate a more robust anomaly detection algorithm to detect obvious mistakes better.