

HOW DO DIGITAL ENGINEERING AND INCLUDED AI BASED ASSISTANCE TOOLS CHANGE THE PRODUCT DEVELOPMENT PROCESS AND THE INVOLVED ENGINEERS

Bickel, Sebastian; Spruegel, Tobias C.; Schleich, Benjamin; Wartzack, Sandro

Friedrich-Alexander-Universität Erlangen-Nürnberg (FAU)

ABSTRACT

Current trends in product development are digital engineering, the increasing use of assistance tools based on artificial intelligence and in general shorter product lifecycles. These trends and new tools strongly rely on available data and will irreversibly change established product development processes. One example for such a new data driven tool is the plausibility check of linear finite element simulations with Convolutional Neural Networks (CNN). This tool is capable of determining whether new simulation results are plausible or non-plausible according to numeric input data. The digitalization and the increased use of data driven tools employing algorithms known from Artificial Intelligence also shifts the roles of many involved engineers. This paper describes and highlights this transition from current product development processes to a data driven / simulation driven product development process. Particularly, the shifts and changes of different roles and domains are illustrated and an example for changing roles in the design and simulation department is described. Furthermore, required adjustments in the design process are derived and compared to the current status.

Keywords: Digital Engineering, Machine learning, Data Driven Design, Organisation of product development, Artificial intelligence

Contact:

Bickel, Sebastian
Friedrich-Alexander-Universität Erlangen-Nürnberg
Engineering Design
Germany
bickel@mfk.fau.de

Cite this article: Bickel, S., Spruegel, T.C., Schleich, B., Wartzack, S. (2019) 'How Do Digital Engineering and Included AI Based Assistance Tools Change the Product Development Process and the Involved Engineers', in *Proceedings of the 22nd International Conference on Engineering Design (ICED19)*, Delft, The Netherlands, 5-8 August 2019. DOI:10.1017/dsi.2019.263

1 INTRODUCTION

Current trends in product development, such as digitalization, digital twins, shorter product lifecycles, and digital engineering (Trage *et al.*, 2018), imply a radical change of established product development processes and engineering workflows. In this regard, for example assistance systems for the knowledge-based setup and plausibility check of finite-element-simulations (Kestel *et al.*, 2016 and Spruegel *et al.*, 2018) strongly re-shape the established simulation and design verification processes. The application of such new data-driven design tools will transform the product development process as we know it today to a digital engineering process. This also means that the roles of many involved engineers will change in the future.

The current paper systematically highlights this transformation exemplarily for the finite-element simulation process and carves out important shifts and changes of different roles and domains in the design and simulation department. Furthermore, required adjustments in the design process are derived and compared to the current status.

The paper is structured as follows. First, a typical current product development process is highlighted. After that, the current state of the art in Data Mining and data-driven methods for knowledge-based simulation is presented. Afterwards, a future product development process and the transition of different domains and the involved people is described. Finally, a conclusion and an outlook are given.

2 A TYPICAL CURRENT PRODUCT DEVELOPMENT AND TESTING PROCESS

The characteristics of new products must be predicted as accurate as possible before the delivery to the customer. Otherwise the probability of complaints and damages rises. The most common ways to evaluate new products are simulation and technical testing. In general, a distinction can be made between customer testing and self-testing. Usually customer testing is used for single-piece production e.g. big machines or large-sized gears whereas self-testing is applied for series and small series production. This typical sequence for self-testing is shown in Figure 1. (Ehrlenspiel and Meerkamm, 2013)

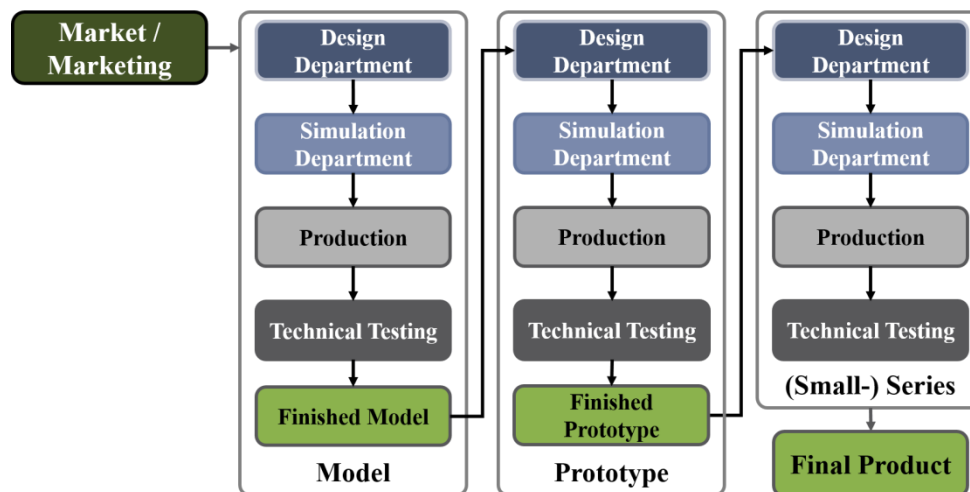


Figure 1: General procedure of product creation with self-testing acc. to (Ehrlenspiel and Meerkamm, 2013)

In general, the process of developing new parts needs various iterations to be completed successfully. This is shown in the flowchart of Figure 1 through the three different iterations Model, Prototype and (Small-) Series. In this example, the start of the whole process is the market or marketing which provides the need or idea for a new design task. After that step, the design and development department generate the first concepts, which are then checked by simulation. The next actions are the production and technical testing of the model. Once these tasks have been completed, the Prototype-phase can be initiated. The involved departments and steps are the same as in the first run, but the result is a functional prototype of the new product. The acquired knowledge is used to complete the last stage, the (Small-) Series. The outcome of this iteration is the final product ready for series production.

3 OVERVIEW OF THE CRISP-DM PROCESS AND DATA MINING METHODS

This section firstly introduces the reader to a standard process for the application of data mining to various use and business cases and describes established data mining methods. After this, a brief introduction to assistance and support systems for knowledge-based simulation is given.

3.1 CRISP-DM process

A standardized process for the application of Data Mining is the “Cross Industry Standard Process for Data Mining” (CRISP-DM). This process was developed in 1996 and supported by an EU project (funded under FP4-ESPRIT 4, Grant agreement ID: 25959). Participants in this research project included Daimler-Benz, ISL (now SPSS) and IBM. The result of this research project was the development of a proven method to guide Data Mining tasks. (Chapman *et al.*, 1999) and (Wirth and Hipp, 2000)

CRISP-DM includes a methodology and a process model for Data Mining problems. The methodology describes the typical phases of a project and the work to be carried out in them. The process model in Figure 2 provides an overview of the Data Mining cycle, which comprises six phases.

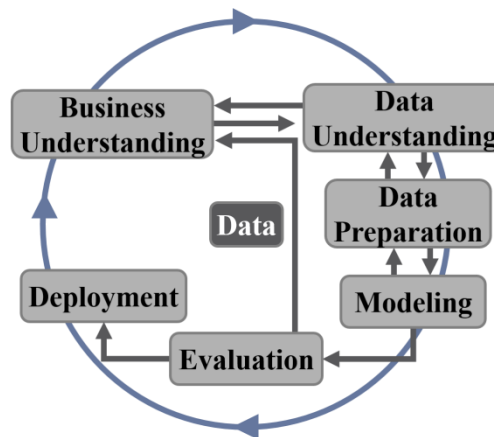


Figure 2: The CRISP-DM process cycle acc. to (Chapman et al., 1999), (IBM Corp, 2012) and (Wirth and Hipp, 2000)

The dependencies between the different phases are represented by arrows. The sequence of the phases is not predefined, but varies according to the project task. The outer circle symbolizes the general cycle of Data Mining: Once a solution has been identified, the knowledge gained during the process or the new information gained from the specific results can be used to find further questions for the original problem.

In general, the CRISP-DM model is very flexible and applicable to many problems. For further information about the different phases, more detailed explanations and precise formulations of the tasks, please refer to the sources (Chapman *et al.*, 1999) and (IBM Corp, 2012).

3.2 Data Mining methods

The determination of the Data Mining method is strongly dependent on the objective of the application. Therefore, Data Mining methods can be divided into various types of analyses, which all provide different types of results, e.g. number values, assigned classes or groups (Vajna *et al.*, 2018). A selection of process classes and their objectives is shown in Table 1. The classification and regression method are explained in more detail since they are used in the following methods.

Table 1: Data Mining methods acc. to (Cleve and Laemmel, 2016) and (Runkler, 2010)

Data Mining methods	
Method	Objective
Regression	Identification of relationships between variables
Classification	Assignment of objects to predefined classes
Anomaly Detection	Identification of unusual or erroneous data
Association rule learning	Identification of dependencies in the form of rules
Clustering	Grouping of objects into a predefined number of clusters
Summarization	Providing different, compact representation of the data

3.2.1 Classification

The aim of classification is to assign data to previously defined classes. The classification problem is related to cluster analysis. However, the aim of cluster analysis is to combine values into a predefined number of groups or classes. These classes already exist in the classification, but the objects must be assigned according to their properties. (Runkler, 2010), (Tan et al., 2006), (Han et al., 2012) and (Aggarwal, 2015)

According to (Runkler, 2010) different examples of classification methods are:

Artificial Neural Networks, k-nearest neighbor classifiers, decision trees, rule-based classifiers, Support Vector Machines or Naïve Bayes classifiers.

3.2.2 Regression

The regression analysis determines the functional dependencies between characteristics in order to derive correlations from the data. Therefore, it is often used for numerical prediction, for example to improve product quality by selectively controlling individual production parameters. The necessary values can be determined by regression in order to achieve the desired quality. Different methods of regression are:

Linear regression (Larose, 2006), multiple linear regression (Larose, 2006), robust regression (Wolf and Hennig, 2010), nonlinear regression (Runkler, 2010), multilayer perceptron (Cleve and Laemmel, 2016) or Radial Basis Function networks (Runkler, 2010).

4 ASSISTANCE SYSTEMS FOR SUPPORTING STATIC MECHANIC FINITE-ELEMENT SIMULATIONS

The setup and interpretation of finite-element simulations for static mechanic problems is both a challenging and responsible task. This is because such simulations require detailed knowledge of various setup options in commercial FE-tools and their implications on the obtained results. Moreover, careful attention should be paid to the result interpretation, since implausible simulation results may lead to wrong design decisions. Motivated by these issues, different approaches have been developed for supporting the pre-processing procedure as well as for the plausibility check of static mechanic FE simulations, which are briefly described in the following sections 4.1 and 4.2.

4.1 Supporting the pre-processing procedure

The concept of (Kestel and Wartzack, 2015) applies the two above-mentioned Data Mining methods. This procedure uses a knowledge acquisition component and a corresponding knowledge database to access and save the provided simulation data. Furthermore, this knowledge databases contains design and simulation rules and methods, to be able to create design-accompanying simulations. This tool is implemented in the CAD-environment to ensure easy operation by design engineers.

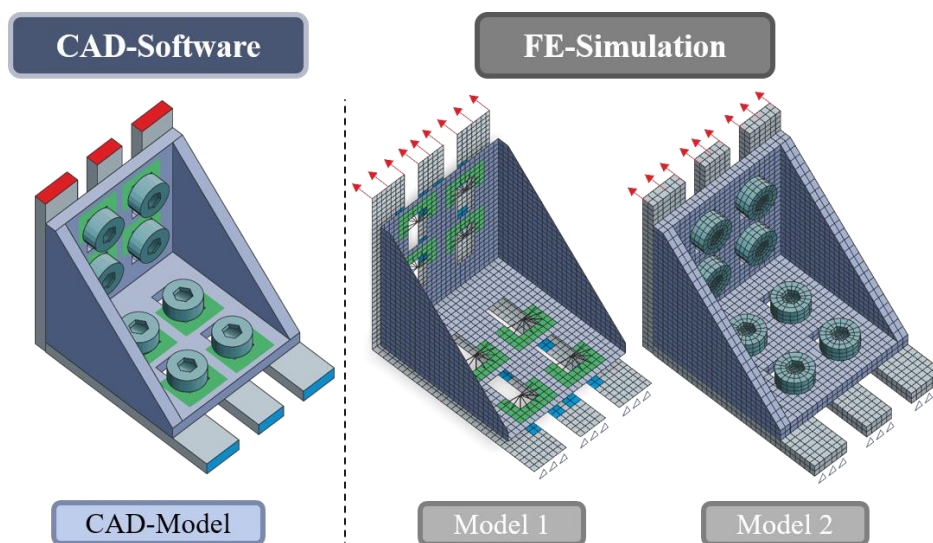


Figure 3: Example of knowledge supported generation of simulations acc. to (Kestel and Wartzack, 2015) and (Kestel and Wartzack, 2016)

Figure 3 shows the application of this tool with the example of an angle bracket. According to the simulation objective, a fitting simulation model is determined by classification. In Figure 3, two possible simulation models are shown. The first one is a simplified model with shell elements, which is used to determine deflection or contact forces. For the analysis of the force profile in the component, the second model should be applied, which uses volume elements to fulfil the defined tasks. This procedure can be carried out in even more detail so that e.g. contact settings, element settings and element types are automatically determined for the simulation. For this type of prediction, regression, classification and neural networks must be used. More information about this method can be found in (Kestel *et al.*, 2016), (Kestel and Wartzack, 2015) and (Kestel and Wartzack, 2016).

4.2 Plausibility checks for structural mechanic FE simulations with Deep Learning

During the last years, a methodology for plausibility checks of linear structural mechanic finite element simulations has been developed. This approach is able to transfer arbitrary geometry data to machine learning algorithms, such as Artificial Neural Networks. The methodology uses spherical detector surfaces to transform any FE mesh of a simulation to a matrix of fixed size with numerical values. As all necessary information in finite element simulations is node bound, also the boundary conditions and the result files can be transformed into numerical matrices of the same size with the new methodology. Figure 4 shows the different sub steps of the methodology that are explained in the following sections.

4.2.1 Spherical detector surface, Node projection, and Node matrix

The start of each investigation is a performed FE simulation with the underlying FE mesh, the known boundary conditions and the result variables. At first, a spherical detector surface is built around the FE mesh. A good number of pixels on the surface is 100x100 pixels, which results in 10,000 pixels of the same surface area. As parts can be oriented differently in 3D space, a uniform orientation on basis of principal component analysis must be performed to the point cloud. The center of the sphere is the same as the center of gravity of the point cloud of the rotated FE mesh.

In the second step, each point of the FE mesh is projected onto the surface of the detector sphere. The center of projection is always the center of gravity of the point cloud and the center of the sphere (which are the same). In each pixel, the number of projected nodes is counted.

Just as the surface of the earth can be projected onto a 2D map, the surface of the detector sphere can be converted into a 2D matrix. Each entry of the matrix corresponds to one pixel of the sphere. The numeric value in the matrix is the number of projected nodes to that specific pixel. At this point, the 3D geometry was converted to a 2D representation of the underlying FE mesh. This matrix could now already be used to perform part recognition. This methodology and first result are published in (Spruegel and Wartzack, 2016).

4.2.2 Matrices for boundary conditions and results from FE simulation

All the information in finite element simulation correspond to certain nodes of the FE mesh. Consequently, this information can also be converted to matrices of fixed size. In Figure 4 the procedure is shown for the fixed displacement support in Y-direction (each node, which has predefined 0 displacement in Y-direction, is projected onto the surface), the force in Y-direction and the equivalent stress (the values of the equivalent stress from all corner nodes of each FE element are accumulated in the detector pixel). This procedure is then performed for all the other inputs and outputs of the simulation (fixed rotations, applied forces, applied moments, deformations, stresses, etc.). When all of these matrices are combined to one large matrix with numerical values, the DNA of an FE simulation can be obtained. The advantage is that the DNA represents the FE simulation and is a numerical matrix of fixed size even if very different parts are simulated with finite element analysis.

4.2.3 Labelled data base

If a database is built up with validated simulations, the corresponding DNA must be labelled with certain additional information. If the DNA should be used to train a CNN for plausibility checks, the already available DNAs of previous simulations must be labelled with 'plausible' or 'non-plausible'. Of course, if the DNA will be used for different investigations, other labels must be attached (e.g. 'missing support in Y-direction').

4.2.4 Automatic plausibility check with Deep Learning Convolutional Neural Networks

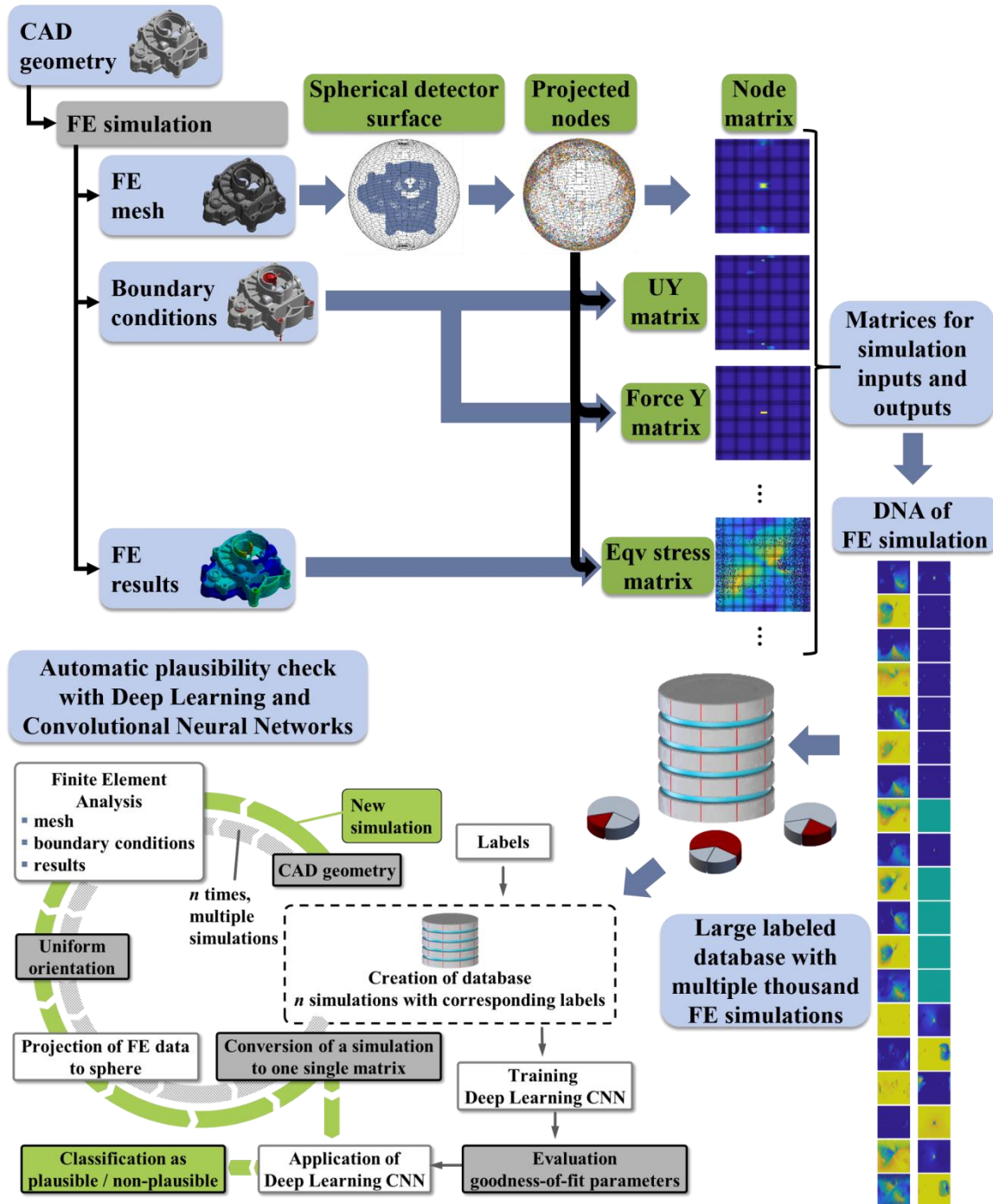


Figure 4: Methodology for plausibility checks of arbitrary FE simulation data with spherical detector surfaces and CNN Deep Learning

In the lower left corner of Figure 4, the whole process for the training and the application of CNN is shown. A detailed explanation can be found in (Spruegel *et al.*, 2018). For the training and testing of the CNN many simulations and the corresponding DNAs are needed. Afterwards the CNN can be trained with different settings of the hyper-parameters (e.g. filter size, number of filters, number of convolution layers, max pooling options, number of fully connected layers, number of neurons in these layers etc.). The evaluation of the performance with data that was not used for the training process is needed to detect overfitting of the network. After the training of a CNN with high prediction quality, a new simulation can be evaluated (e.g. classified as 'plausible'). This application process is shown in Figure 4 in the lower left corner with the outer circle for a new FE simulation. More detailed information about the methodology can be found in (Spruegel *et al.*, 2018). Another example for the application of machine learning methods in FE simulations can be found in (Bohn *et al.*, 2013).

5 COMPARISON OF PRODUCT DEVELOPMENT PROCESSES TODAY AND IN THE NEAR FUTURE

With the introduction of virtual development tools in the product development process certain shifts and improvements have been made. The transition from the established process, presented in section 2, to the virtual product development for the model- and prototype-phase is shown in Figure 5. The procedure on the left side demonstrates the product development process according to (Ehrlenspiel and Meerkamm, 2013). The illustration on the right side of Figure 5 shows the same two development phases (Model and Prototype) in the virtual product development.

This small process shows one product design iteration, using CAD- and FE-Simulation tools. For each step, at least one engineer is necessary to fulfil the given tasks. The design engineer generates the part, using 3D-CAD-Systems, and then transfers the data to the simulation engineer. With the new simulation-generated information, design enhancements are made by the design department. After these refinements, the part will be simulated again and is then ready for manufacturing or testing.

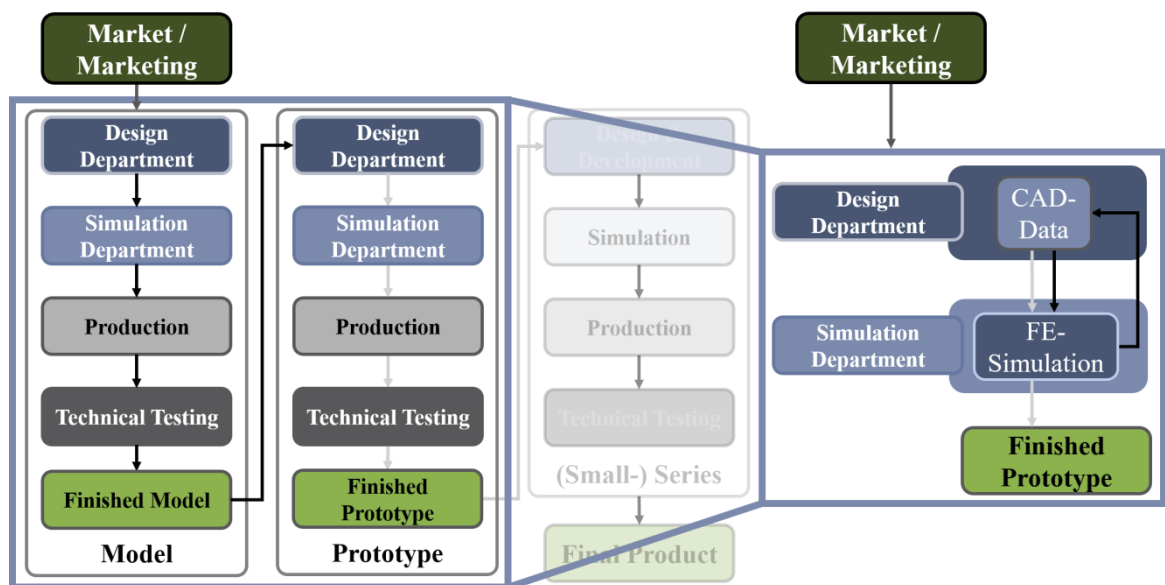


Figure 5: Comparison of the established (left) and the virtual development process (right)

5.1 Example of changing tasks for design and simulation engineers

Figure 6 shows a possible example how the general roles of design and simulation engineers could change with data-driven engineering tools. In contrast to the common process in today's product development environment (Figure 5 on the right side), data-driven product development shifts the roles of design and simulation departments. The design division is still responsible for the generation of 3D-design data, but by using the new data-driven software tools, they are able to generate and start an automatic FE-Simulation. A possible solution for this method is elaborated in chapter 4.1, with the supporting pre-processing procedure. After a successful automatic simulation run, the result can be reviewed with the plausibility check method. This method is explained in detail for the simulation of a single part in chapter 4.2.4.

If the outcome of the method is plausible, the design engineers can use the simulation result for a part revision. If the simulation is classified 'non-plausible', the design engineer should repeat the simulation or work together with the simulation department to achieve a plausible simulation result. After the reengineering of the design, a second simulation can be initiated, again with the application of the data-driven automated pre-processing tool. But this time, the simulation department is responsible for the simulation. The engineer can use the assisting tool, to reduce the time of standard and repetitive procedures, while generating a simulation. Therefore, the employees are able to concentrate more on the complex items of the simulation e.g. non-linear materials or contact problems. If the virtual testing shows a positive result, the part is ready for manufacturing or final testing. Otherwise, more iterations need to be arranged.

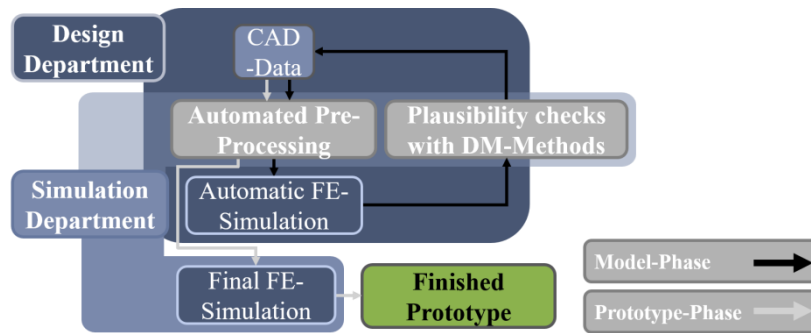


Figure 6: Example of new roles for design and simulation departments in the design process

In this small example, many role changes can be observed. An overview of these changes is shown in Figure 7. The illustration represents the new responsibilities of design and simulation engineers. In the picture, the color distinguishes the tasks between new and consisting tasks. For the design department the main business, the part design, is still an important component of their work. Additionally to that, the application of data-driven tools becomes inevitable. Furthermore, the evaluation of the results will be a new area of responsibility, which overlaps with the simulations department. Both divisions need to work together to achieve the best possible results.

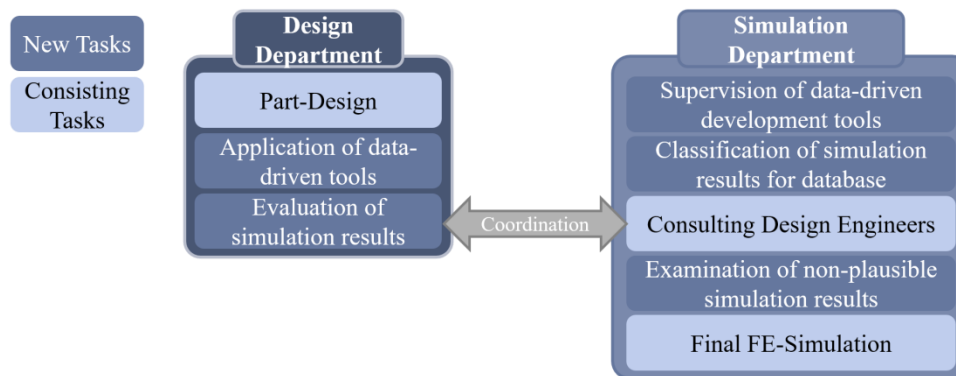


Figure 7: Comparison of the new responsibilities between design and simulation engineers

The duties of the simulation engineers are changing likewise. One important new task is the classification of simulation results for the database. The simulation engineer has to provide and maintain the database with plausible simulation results for given parts or problems. Only with this database, the data-driven tools are able to perform good results. In general, the new assistant-systems need to be supervised by the simulation engineers to check if they are still working correctly and providing good results. A consisting task is and will be the final FE-Simulation to examine the product properties and compare the result to the defined specifications. Another new work package for simulation engineers is the examination of non-plausible results. If a design engineer receives a non-plausible result and is not able to solve the problem by its own, the simulation engineer provides assistance to overcome the difficulty with the specific simulation. In order to cope with these new responsibilities, the existing expertise of the simulation department is necessary. These new changes in the distribution of roles lead to general advantages for the product development process and the according employees. Figure 8 lists all these positive influences.

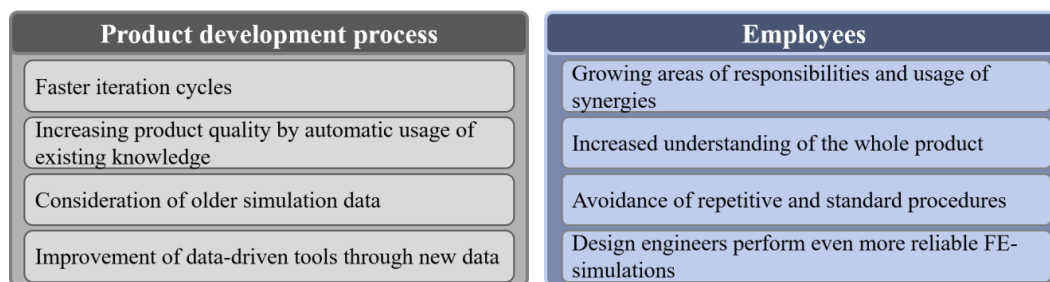


Figure 8: Advantages of the data-driven product development

5.2 Roles / Characters

With a changing product development process, also the number of different people working in these domains will change as follows (Figure 9 demonstrates the changes between the current development process on the left and the future process on the right):

- Design engineers are becoming more and more product development engineers and the total number will increase with more products coming to the market in shorter time.
- Both product and process simulation is getting more important and the number of people needed in these domains is increasing.
- The number of testing engineers will decrease as well as the number of engineers in manufacturing. Especially because the available data helps to understand, which tests are necessary and therefore increase the productivity and decrease the scrap rate in production.
- Data can help all different domains to perform their job faster and better. However, new roles are necessary to do this. At first, data scientists are needed to cope with the amount of data and to perform stable and statistically verifiable analyses. Also, expert knowledge about the underlying processes in development or production are crucial.

Therefore, data affine experts from the different fields (product development, simulation, testing, production, etc.) are very important people and the industry must acquire available experts, before the competitor acquires these data affine engineers with expert knowledge in their domain. Often these people can be found at research institutes or directly in the master programs of universities (Davenport and Patil, 2012).

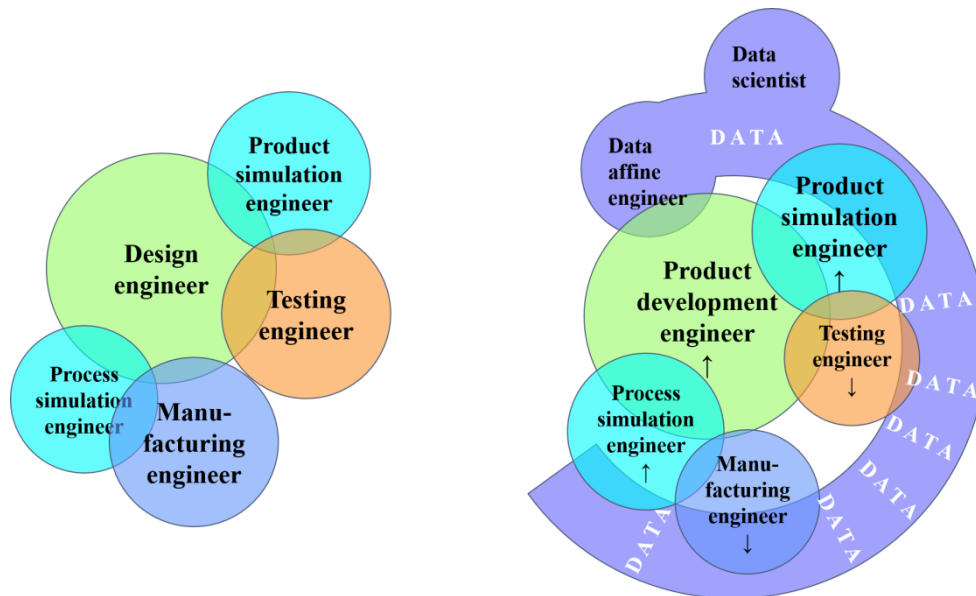


Figure 9: Qualitative shift in the number of people (size of the balloons) in certain domains and the importance of data in a future product development process (right) compared to the current people and their domains in current product development processes (left)

6 CONCLUSION AND OUTLOOK

In this paper, the transition from current product development processes to a data driven / simulation driven product development process is described. A tool for automatic plausibility checks of simulation data for the product developer and a method to facilitate the process of setting up a simulation are presented briefly, including the impact and the advantages of such tools on a future product development process. The transitions and changes of different roles and domains are shown. The adjustments in the process of product development are described and compared to the current status. Data is an essential element in the transition of currently established methods and tools towards Digital Engineering. Besides the important steps of relational storage and collection of the data, the analysis of the data with expert knowledge from the domain is and will be highly relevant.

The volume of data will increase in the coming years and therefore the processing of data will become more important. With the right and efficient use of the provided data, it is possible to decrease the development time and increase the quality of new products.

REFERENCES

- Aggarwal C. (2015), *Data Mining. The Textbook*. Springer International Publishing, New York
<https://doi.org/10.1007/978-3-319-14142-8>
- Bohn B., Garcke J., Iza-Teran R., Paprotny A., Peherstorfer B., Schepsmeier U., and Thole C.A. (2013), “Analysis of Car Crash Simulation Data with Nonlinear Machine Learning Methods”, *International Conference on Computational Science, ICCS 2013*, Procedia Computer Science, pp. 621–630.
- Chapman P., Clinton J., Kerber R., Khabaza T., Reinartz T., Shearer C., and Wirth R. (1999), *CRISP-DM 1.0. Step-by-step data mining guide*, CRISP-DM consortium.
- Cleve, J. and Laemmel, U. (2016), *Data Mining*. De Gruyter, Berlin <https://doi.org/10.1515/9783110456776>
- Davenport, T.H. and Patil, D.J. (2012), *Data scientist: The sexiest Job of the 21st century*. Harvard Business Review, pp. 70–76.
- Ehrlenspiel K. and Meerkamm H. (2013), *Integrierte Produktentwicklung. Denkabläufe, Methodeneinsatz, Zusammenarbeit*. Carl Hanser Verlag, München <https://doi.org/10.3139/9783446421578>
- Han J., Kamber M. and Pei J. (2012), *Data Mining. Concepts and Techniques*. Morgan Kaufmann, Waltham
<https://doi.org/10.1016/C2009-0-61819-5>
- IBM Corporation. (2012), “IBM SPSS Modeler CRISP-DM Handbuch”. IBM Corporation. Available at:
<ftp://public.dhe.ibm.com/software/analytics/spss/documentation/modeler/15.0/de/CRISP-DM.pdf>
 (11.06.2018 - 15:36)
- Kestel P. and Wartzack S. (2015), “Konzept für ein wissensbasiertes FEA-Assistenzsystem zur Unterstützung konstruktionsbegleitender Simulationen”, *DfX Symposium*, Herrsching, 2015, 10.07. - 10.08.2015, Tutech Verlag, Hamburg, pp. 87–98.
- Kestel, P., Schneyer, T. and Wartzack, S. (2016), “Feature-based approach for the automated setup of accurate, design-accompanying Finite Element Analyses”. *14th International Design Conference*. Dubrovnik, 05.16.2016 – 05.19.2016, pp. 697–706.
- Kestel P., and Wartzack S. (2016), “Wissensbasierter Aufbau konstruktions-begleitender Finite-Element-Analysen durch FEA-Assistenzsystem”, *Entwerfen Entwickeln Erleben*, Dresden Germany, 06.30. – 07.01.2016, TUDpress, Dresden, pp. 315–329.
- Larose D. (2006), *Data Mining Methods and Models*. John Wiley and Sons, Hoboken, New Jersey
<https://doi.org/10.1002/0471756482.index>
- Runkler T. (2010), *Data-Mining. Methoden und Algorithmen intelligenter Datenanalyse*. Vieweg + Teubner, Wiesbaden <https://doi.org/10.1007/978-3-8348-9353-6>
- Spruegel, T.C. and Wartzack, S. (2016), “Das FEA-Assistenzsystem - Analyseteil FEdeIM”, *Entwerfen Entwickeln Erleben*, Dresden Germany, 06.30. – 07.01.2016, TUD press, Dresden, pp. 463–474.
- Spruegel, T.C., Rothfelder, R., Bickel, S., Grauf, A., Sauer, C., Schleich, B. and Wartzack, S. (2018), “Methodology for plausibility checking of structural mechanics simulations using Deep Learning on existing simulation data”, *NordDesign 2018*, Linköping Sweden, 08.14. – 08.17.2018, LiU Tryck, Linköping, Session 1A Machine Learning.
- Tan P.-N., Steinbach M., and Kumar V. (2006), *Introduction to Data Mining*. Pearson/ Addison-Wesley, Boston.
- Trage, S., Saier, M., Amadori, D. and Reschke, K. (2018), “Whitepaper Innovationen wie am Fließband – Auswirkungen der Digitalisierung auf die Innovation und Entwicklung von Produkten in Fertigungsunternehmen” [online] KPMG. Available at: www.hub.kpmg.de (11.06.2018 - 16:56).
- Wirth R. and Hipp J. (2000), “CRISP-DM: Towards a standard process model for data mining”, *4th International Conference on the Practical Applications of Knowledge Discovery and Data Mining*, Manchester, pp. 29–39.
- Wolf C. and Hennig B. (2010), *Handbuch der sozialwissenschaftlichen Datenanalyse*. Springer, Wiesbaden
<https://doi.org/10.1007/978-3-531-92038-2>
- Vajna S., Weber C., Zeman K., Hehenberger, Gerhard D. and Wartzack S. (2018), *CAX für Ingenieure: Eine praxisbezogene Einführung*. Springer, Berlin <https://doi.org/10.1007/978-3-662-54624-6>

ACKNOWLEDGMENTS

The authors would like to thank the NVIDIA Corporation and the academic GPU Grant Program for the donation of a Tesla GPU.

This research work is part of the FAU “Advanced Analytics for Production Optimization” project (E|ASY-Opt) and funded by the Bavarian program for the “Investment for growth and jobs” objective finance by the European Regional Development Fund (ERDF), 2014-2020. It is managed by the Bavarian Ministry of Economic Affairs and Media, Energy and Technology. The authors are responsible for the content of this publication.