



Quality 4.0: a review of big data challenges in manufacturing

Carlos A. Escobar¹ · Megan E. McGovern¹ · Ruben Morales-Menendez²

Received: 21 November 2020 / Accepted: 19 March 2021

© The Author(s), under exclusive licence to Springer Science+Business Media, LLC, part of Springer Nature 2021

Abstract

Industrial big data and artificial intelligence are propelling a new era of manufacturing, *smart manufacturing*. Although these driving technologies have the capacity to advance the state of the art in manufacturing, it is not trivial to do so. Current benchmarks of quality, conformance, productivity, and innovation in industrial manufacturing have set a very high bar for machine learning algorithms. A new concept has recently appeared to address this challenge: *Quality 4.0*. This name was derived from the pursuit of performance excellence during these times of potentially disruptive digital transformation. The hype surrounding artificial intelligence has influenced many quality leaders take an interest in deploying a *Quality 4.0* initiative. According to recent surveys, however, 80–87% of the *big data* projects never generate a sustainable solution. Moreover, surveys have indicated that most quality leaders do not have a clear vision about how to create value of out these technologies. In this manuscript, the process monitoring for quality initiative, *Quality 4.0*, is reviewed. Then four relevant issues are identified (paradigm, project selection, process redesign and relearning problems) that must be understood and addressed for successful implementation. Based on this study, a novel 7-step problem solving strategy is introduced. The proposed strategy increases the likelihood of successfully deploying this *Quality 4.0* initiative.

Keywords Quality 4.0 · Quality control · Manufacturing systems · Artificial intelligence · Big data

Introduction

Manufacturing is a primary economic driving force. Modern manufacturing began with the introduction of the factory system in late eighteenth century Britain. This precipitated the first industrial revolution (mechanical production) and later spread worldwide. The defining feature of the first factory system was the use of machinery, initially powered foot treadles, soon after by water or steam and later by electricity. The second industrial revolution (science and mass production) was a phase of electrification, rapid standardization, and production lines, spanning from the late nineteenth to the early twentieth century. The introduction of the internet, automation, telecommunications, and renewable energies in the late twentieth and early in the twenty-first century gave rise to the third industrial revolution (globalization and digital revolution) [1].

Today, at the brink of the fourth industrial revolution (interconnectedness, data and intelligence) or *Industry 4.0* [2], technologies such as *Industrial Big Data (IBD)* [3], *Industrial Internet of Things (IIoT)* [4], acsensorization [5], *Artificial Intelligence (AI)* [6], and *Cyber-Physical Systems (CPS)* [2] are propelling the new era of manufacturing, *smart manufacturing* [7]. Mass customization, just-in-time maintenance, solutions to engineering intractable problems, disruptive innovation, and perfect quality are among the most relevant challenges posed to this revolution. Manufacturing systems are upgraded to an intellectual level that leverages advanced manufacturing and information technologies to achieve flexible, intelligent, and reconfigurable manufacturing processes to address a dynamic, global marketplace. Some technologies also have *AI*, which allows manufacturing systems to learn from experiences to realize a connected, intelligent, and ubiquitous industrial practice [8].

Although the driving technologies in the fourth industrial revolution have the capacity to move the state of the art in manufacturing forward, it is not trivial to do so. Current benchmarks of quality, productivity, and innovation in industrial manufacturing are very high. Rapid innovation shorten process lifespans, leaving little time to understand and solve engineering problems from a physics perspective.

✉ Carlos A. Escobar
carlos.1.escobar@gm.com

¹ General Motors, 30470 Harley Earl Blvd., Warren, MI 48092, USA

² Tecnológico de Monterrey, Av E Garza Sada 2501, Monterrey, NL 64849, Mexico

Moreover, application of *Statistical Quality Control (SQC)* techniques have enabled most modern processes to yield very low *Defects per Million of Opportunities (DPMO)* (6σ , 3.4 DPMO). This ultra-high conformance rate sets a very high bar for AI. These challenges will only intensify, since *smart manufacturing* is driven by processes that exhibit increasing complexity, short cycle times, transient sources of variation, hyper-dimensional feature spaces, as well as non-Gaussian pseudo-chaotic behavior.

To address this in the age of the fourth industrial revolution, a new concept has recently appeared: *Quality 4.0*. This name was derived from the pursuit of performance excellence during these times of potentially disruptive digital transformation [9].

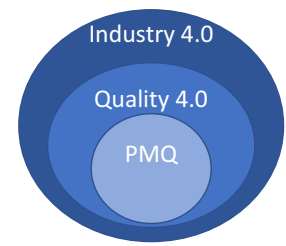
The authors define *Quality 4.0* as:

Quality 4.0 is the fourth wave in the quality movement (1. *Statistical Quality Control*, 2. *Total Quality Management*, 3. *Six sigma*, 4. *Quality 4.0*). This quality philosophy is built on the statistical and managerial foundations of the previous philosophies. It leverages industrial big data, industrial internet of things, and artificial intelligence to solve an entirely new range of intractable engineering problems. *Quality 4.0* is founded on a new paradigm based on empirical learning, empirical knowledge discovery, and real-time data generation, collection, and analysis to enable smart decisions.

In manufacturing, the main objectives of *Quality 4.0* are: (1) to develop defect-free processes, (2) to augment human intelligence, such as empirical knowledge discovery for faster root cause analyses, (3) to increase the speed and quality of decision-making, and (4) to alleviate the subjective nature of human-based inspection [9–11]. According to recent surveys, although 92% of surveyed leaders are increasing their investments in *IBD* and *AI* [12], a substantial portion of these quality leaders do not yet have a deployment strategy [11] and universally cite difficulty in harnessing these technologies [13]. This situation is reflected in the deployment success rate: 80–87% of the *big data* projects never generate a sustainable solution [14,15].

Process Monitoring for Quality (PMQ) [5] is a data-driven quality philosophy aimed at addressing the intellectual and practical challenges posed to the application of *AI* in manufacturing. It is founded on *Big Models (BM)* [16], a predictive modeling paradigm that applies machine learning, statistics, and optimization to process data to develop a classifier aimed at real-time defect detection and empirical knowledge discovery aimed at process redesign and improvement. To effectively solve the pattern classification problem, *PMQ* proposes a 4-step approach—*acsensorize*, *discover*, *learn*, *predict (ADLP)*. Figure 1 describes the *PMQ Quality 4.0* initiative in the context of *Industry 4.0*.

Fig. 1 *PMQ* in the context of *Industry 4.0*



In this review, four relevant-to-manufacturing issues are identified that quality and manufacturing engineers/researchers, managers, and directors must understand and address to successfully implement a *Quality 4.0* initiative based on *PMQ*: (1) the paradigm problem, (2) project selection problem, (3) process redesign problem, and (4) relearning problem. Based on this study, the original *PMQ* 4-step problem solving strategy is revised and enhanced to develop a comprehensive approach that increases the likelihood of success. This novel strategy is an evolution of the Shewhart learning and improvement cycle (widely used in manufacturing for quality control) to address the challenges posed to the implementation of this *Quality 4.0* initiative.

This paper does not attempt to offer new methods, algorithms or techniques. The agenda of the problem solving strategy is pragmatic: defect detection through binary classification and empirical knowledge discovery that can be used for process redesign and improvement. Its scope is limited to discrete manufacturing systems, where a deterministic solution for process monitoring and control is not available or feasible.

The rest of the paper is organized as follows: a review of the state of the art in “Similar work” section. A brief description of *PMQ* in “*PMQ* and its applications” section. The four problems are discussed in “Challenges of big data initiatives in manufacturing” section. Then, the *PMQ* 4-step problem solving strategy is updated in “Problem solving strategy for a *Quality 4.0* initiative” section. An implementation strategy is presented in “Strategy development and adoption for *Quality 4.0*” section. Finally, conclusions and future research are contained in “Conclusions” section.

Similar work

Manufacturing has been a primary economic driving force since the first industrial revolution in the nineteenth century and continues to be so today. The *Industry 4.0* paradigm has necessitated a similar paradigm shift in manufacturing to keep up with quick launch, low volume production, and mass customization of products [17], all while conforming to high quality and productivity requirements. This new manufacturing paradigm has been termed *smart manufacturing*, and it has been enabled by technologies such as *IIoT*, *CPS*, cloud computing [18], *Big Data Analytics*, Information and Com-

munications Technology, and AI [8]. *IBD* offers an excellent opportunity to enable the shift from conventional manufacturing to *smart manufacturing*. Companies now have the option of adopting data-driven strategies to provide a competitive advantage.

In 2017, Kusiak [19] advocated for manufacturing industries to embrace and employ the concepts of *Big Data Analytics* in order to increase efficiency and profits. Tao et al. [20] studied *big data*'s role in supporting this transition to *smart manufacturing*. Fisher et al. [21] investigated how models generated from the data can characterize process flows and support the implementation of circular economy principles and potential waste recovery.

Though *smart manufacturing* has taken off in more recent years, the concept has been around for decades. In 1990, Kusiak [22] authored *Intelligent Manufacturing Systems*, which was designed to provide readers with a reference text on such systems. For a more recent reviews, the reader is referred to [23,24]. The wealth of data has both contributed to and enabled the rise of *smart manufacturing* practices. In 2009, Choudhary et al. [25] provided a detailed literature review on data mining and knowledge discovery in various manufacturing domains. More recently (2020), Cheng et al. [26] proposed data mining technique to enhance training dataset construction to improve the speed and accuracy of makespan estimation. Wei et al. [27] proposed a data analysis-driven process design method based on data mining, a *data + knowledge + decision* backward design pattern, to enhance manufacturing data.

Quality is a fundamental business imperative for manufacturing industries. Historically, i.e., pre-*Industry 4.0*, manufacturing industries have produced higher volumes for products with longer intended lifecycles. Consequently, there was more time and data on which to refine processes and reduce process variance. In other words, there was a longer learning curve to enable the production of high-quality products. In the era of *smart manufacturing*, the luxury of these long learning curves is removed. This has propelled the adoption of intelligent process monitoring and product inspection techniques. Some examples of these *Quality 4.0* applications will now be presented.

In 2014, Wuest et al. [28] introduced a manufacturing quality monitoring approach in which they applied cluster analysis and supervised learning on product state data to enable monitoring of highly complex manufacturing processes. Malaca et al. [29] designed a system to enable the real-time classification of automotive fabric textures in low lighting conditions. In 2020, Xu and Zhu [30] used the concept of fog computing to design a classification system in order to identify potential flawed products within the production process. Cai et al. [31] proposed a hybrid information system based on an artificial neural network of short-term memory to predict or estimate wear and tear on manufac-

turing tools. Their results showed outstanding performance for different operating conditions. Said et al. [32] proposed a new method for machine learning flaw detection using the reduced kernel partial least squares method to handle non-linear dynamic systems. Their research resulted in a calculation time reduction and a decrease in the rate of false alarms. Hui et al. [33] designed data-based modeling to assess the quality of a linear axis assembly based on normalized information and random sampling. They used the replacement variable selection method, synthetic minority oversampling technique, and optimized multi-class *Support Vector Machine*.

The cases briefly presented include fault detection systems, classifiers development techniques, or process control systems. The common denominator is that they are all data-driven approaches based on artificial intelligence algorithms. In most cases, a particular problem is solved. Therefore, unless managers or directors have a similar problem, they cannot develop a high impact *Quality 4.0* initiative from what is reported. In this context, this article highlights several manufacturing issues that must be understood and addressed to successfully deploy a *Quality 4.0* initiative across manufacturing plants.

PMQ and its applications

Though quality inspections are widely practiced before, during, and after production, these inspections still highly rely on human capabilities. According to a recent survey, almost half of the respondents claimed that their inspections were mostly manual (i.e., less than 10% automated) [34]. Manual or visual inspections are often subject to the operators' inherent biases, and thus only about 80% accurate [35]. This generates the hidden factory effect, where unforeseen activities contribute to efficiency and quality reduction [36]. *PMQ* proposes to use real-time process data to automatically monitor and control the processes, i.e., identify and eliminate defects. Where defect detection is formulated as a binary classification problem. *PMQ* originally proposed a 4-step problem solving strategy, Fig. 2:

- *Acsensorize*, this step refers to observe (e.g., deploy sensors, cameras) the system to generate the raw empirical data to monitor the system [5,37].
- *Discover*, refers to the step of creating features from the raw empirical data [38,39].
- *Learn*, refers to the step of applying machine learning, statistics, and optimization techniques to develop a classifier. The big models learning paradigm addresses the main challenges posed by manufacturing derived data sets for binary classification of quality [16].

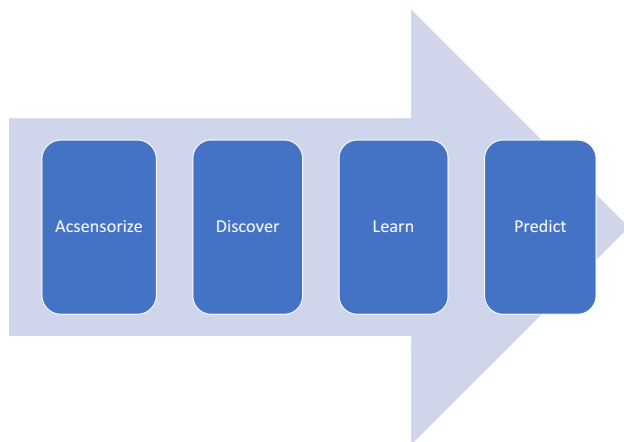


Fig. 2 Original problem solving strategy proposed by *PMQ*

Table 1 Confusion matrix

	Predicted good	Predicted bad
Good item	True negative (TN)	False positive (FP)
Defective item	False negative (FN)	True positive (TP)

- *Predict*, classifier fusion to optimize prediction is the main goal of this step. An ad-hoc multiple classifier system is presented in [40].

The confusion matrix is a table used to summarize the predictive ability of a classifier, Table 1. Where a positive result refers to a defective item, and a negative result refers to a good quality item. Since prediction is performed under uncertainty, a classifier can commit *FP* (type-I, α) and *FN* (type-II, β) errors [41].

To explain how *PMQ* is applied to advance the state of the art in manufacturing, three common quality control scenarios¹ without intelligent systems are presented, Fig. 3. Their counterparts are then presented in Fig. 4, followed by a brief description about how *PMQ* can boost them. More insights about this analysis can be found in [42].

The processes of most mature manufacturing organizations generate only a few defective items, Fig. 3. The majority of these defects are detected (*TP*) by either a manual/visual inspection, Fig. 3a or by a statistical process monitoring system (*SPC/SQC*), Fig. 3b. Detected defects are removed from the value-adding process for a second evaluation, where they are finally either reworked or scrapped. Since neither inspection approaches are 100% reliable[35,43], they can commit *FP* (i.e., call a good item defective) and *FN* (i.e., call a defective item good) errors. Whereas *FP* create the hidden factory

effect by reducing the efficiency of the process, *FN* should always be avoided as they cause warranties.

In extreme cases, Fig. 3c, time-to-market pressures may compel a new process to be developed and launched even before it is totally understood from physics perspective. Even if a new *SPC/SQC* model/system is developed or a pre-existing model or system is used, it may not be feasible to measure its quality characteristics (variables) within the time constraints of the cycle time. In these intractable or infeasible cases, the product is launched at a high risk for the manufacturing company.

Boosting the state of the art

PMQ is applied to eliminate manual or visual inspections, as well as to develop an empirical-based quality control system for the intractable and unfeasible cases, Fig. 4a. The real case of ultrasonic welding of battery tabs in the Chevrolet Volt is a good example of this approach [5]. In a process statistically under control, *PMQ* can also be applied to detect those few *DPMO* (FN) not detected by the *SPC/SQC* system to enable the creation of virtually defect-free processes through perfect detection [10], Fig. 4b. Finally, the empirical knowledge discovery component-aimed at process redesign and improvement-of *PMQ* is described in “The redesign problem” section.

From the initiative of implementing *Quality 4.0* in the plants, to actually develop and deploy a sustainable solution (i.e., predictive system) that would endure the dynamism of manufacturing systems, a plethora of managerial, intellectual, technical, and practical challenges must be effectively addressed. Manufacturing knowledge and a solid learning strategy must therefore be in place to avoid painful lessons characterized by over-promises and under-deliveries.

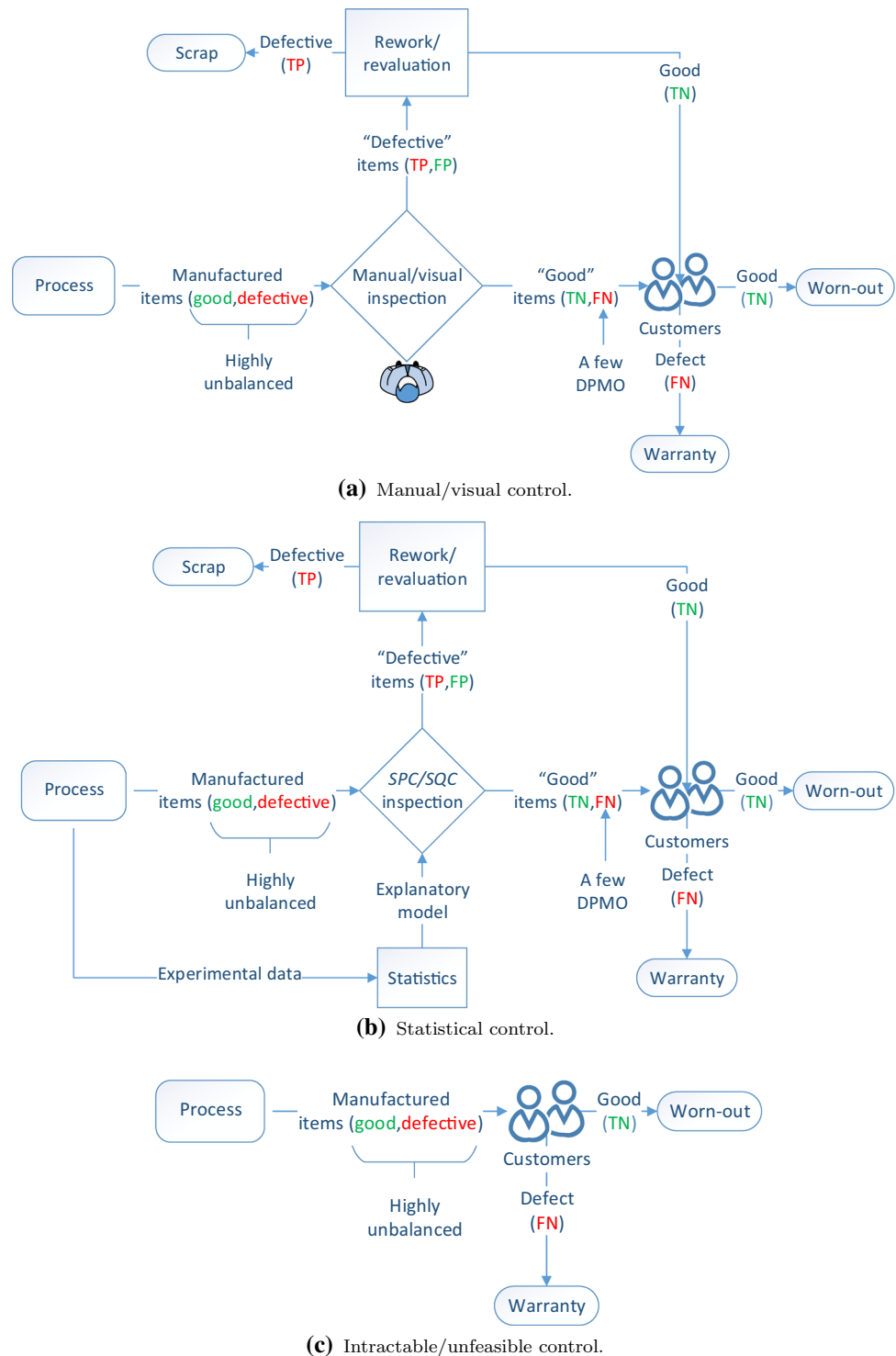
Designing the classifier

Although the intention of this paper is not to train data scientists in machine learning techniques (this task may take months or even years), this section provides an overview of the most relevant challenges posed by the development of a binary classifier using manufacturing-derived data and suggests a direction to address these challenges. This subsection aims to contextualize and provide some background for the next section.

In *Quality 4.0*, process data is generated, collected and analyzed real-time. Developing the real-time infrastructure is one of the last steps of the project and is usually addressed at the deployment stage. Proof-of-concepts are based on asynchronous analyses, as the first challenge for the data scientists is to demonstrate feasibility.

Prior to developing a model, the training data is generated or collected; each sample must have the associated

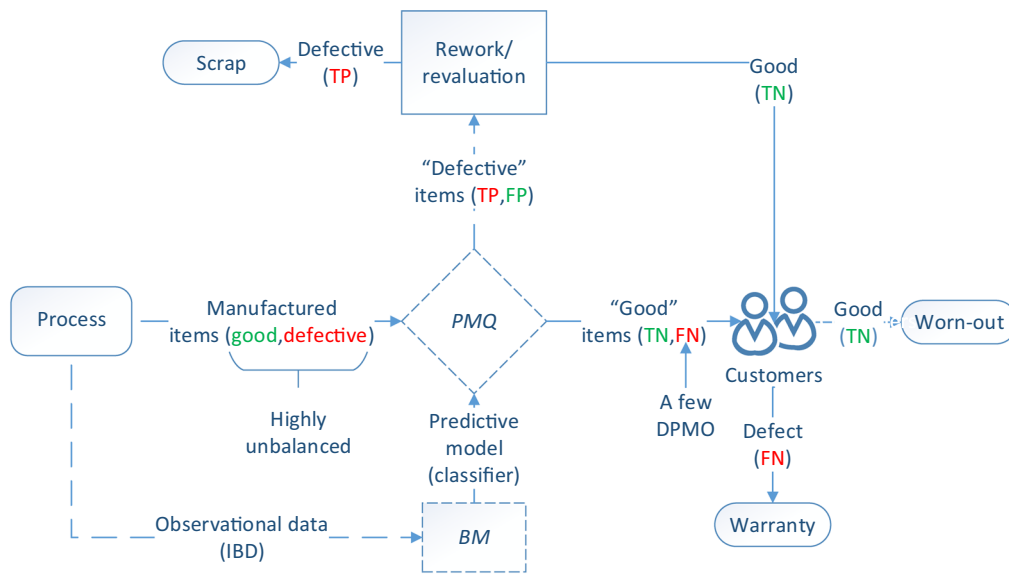
¹ Authors acknowledge that this is an oversimplified analysis, which is only presented to convey the ideas of this section.

Fig. 3 Traditional quality control scenarios

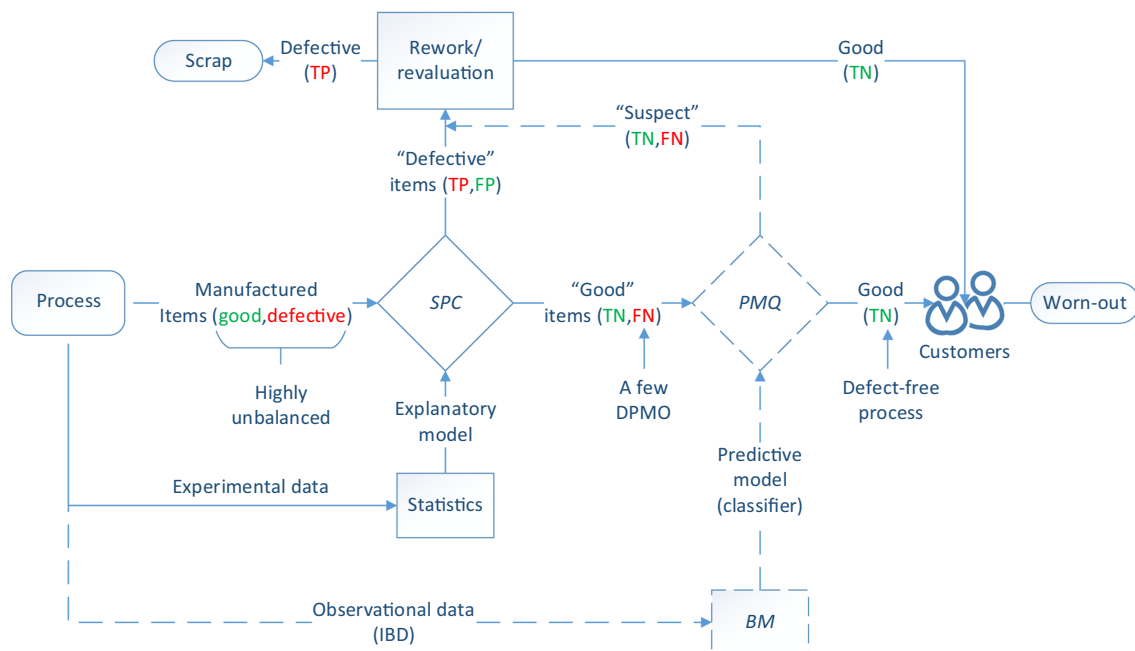
quality characteristic (good, defective) for the algorithm to learn the pattern. In general, there are two common cases: (1) plant data and (2) lab-generated data. When available, plant data is the most desirable since it is a representative, relevant environment, data-rich option in which the process is already observed (i.e., data generated and stored). This allows for feasibility assessment to be performed promptly.

Lab-generated data aims to imitate relevant environmental conditions and is usually limited to smaller data sets. The principle of *accesorization* describes good insights about exploring different options for data generation [5].

The process data is usually found in three common formats: (1) pictures, (2) signals, and (3) direct measurements of quality characteristics (features). For the first two formats,



(a) Empirical-based (data-driven) control.



(b) Boost a process statistically under control.

Fig. 4 PMQ applications

features can be created following typical construction techniques [38,39,44]. Feature creation followed by a feature relevance analysis helps engineers to determine the driving features of the system, information further used for deriving physics knowledge and augment troubleshooting. This concept is explained in more detail in the following section. Raw signals and images can also be used to train a deep neural network architecture for quality control [45,46]. Although these *black boxes* can efficiently solve many engineering

problems from a prediction perspective, they do not facilitate information extraction. More theory about deep learning architectures can be found in [47,48].

Manufacturing-derived training data for binary classification of quality poses the following challenges: (1) hyper-dimensional features spaces, including relevant, (2) highly/ultra unbalanced (minority/defective class count < 1%), irrelevant, trivial and redundant, (3) mix of numerical and categorical variables (i.e., nominal, ordinal or dichoto-

mous), (4) different engineering scales, (5) incomplete data sets and (6) time dependent.

To address these challenges, the *BM* learning paradigm has been developed [16]. This paradigm is comprised of several well known ad-hoc tools plus many particular developments aimed at addressing specific challenges. Since the data structure is not known in advance, there is no a priori distinction between machine learning algorithms [49], therefore the *BM* learning paradigm suggests to explore eight diverse options to develop a multiple classifier system [40]. Due to the time effect, models are usually validated following a time-ordered hold-out scheme. The training set is partitioned in training, validation, and test sets, where the latest set is used to emulate deployment performance and compare it to the learning targets to demonstrate feasibility. Different data partitioning strategies can also be explored based on the data characteristics (see in [50]).

To induce information extraction and model trust, *BM* is founded on the principle of parsimony [51]. Parsimony is induced through feature selection [52]² and model selection [51]. These methods help to eliminate irrelevant, redundant and trivial features. To address the highly/ultra unbalanced problem, hyperparameters are tuned with respect to the Maximum Probability of Correct Decision [53], a metric based on the confusion metric highly sensitive to *FN*. Since most machine learning algorithms work internally with numeric data, it is important to develop a strategy to deal with categorical variables [54] (i.e., encode them in a numeric form). For binary features, it is recommended to use *effect coding* (−1 and 1) instead of *dummy coding* (0 and 1) [55]. Moreover, since features tend to have different engineering scales, it is important to normalize or standardize the data before present it to the algorithm. Feature scaling generally speeds up learning and leads to faster convergence [56]. More insights about when to normalize or standardize can be found in [57,58]. To deal with missing information, deleting rows or columns is a widely used approach; however, this method is not advised unless the proportion of eliminated records is very small (< 5%) [59]. Imputation is a better approach. This technique refers to the process of replacing missing values with statistical estimations. Before selecting the imputation method, the data set should be thoroughly assessed, to determine the most appropriate statistical method (that minimizes bias) to handle missing data. A review of imputation methods can be found in [59–61]. Permuting rows and columns is a different approach that helps to minimize the bias induced by estimates, as it maximizes the number of samples that can be used for learning [62]. Finally, manufacturing systems tend to be time-dependent, the data correlation effect must be taken

into consideration when creating the features, selecting the neural network architecture and the model validation method. Case studies analyzing manufacturing-derived data sets that illustrate the application of most of the tools described in this section are presented in [5,16,52].

As reviewed in this subsection, developing a binary classifier using manufacturing-derived data is a complicated task. It is important to only select appropriate projects for the data science team. Selecting appropriate projects not only increases the likelihood of success, but also avoids unnecessary allocation of data science team resources. Moreover, developing a classifier with the capacity to satisfy the learning targets for the project (i.e., demonstrate feasibility), is just the beginning: higher order challenges must be understood and addressed to develop a sustainable solution that creates value to the company. The following section, provides deeper insights about these challenges.

Challenges of big data initiatives in manufacturing

In this section, four relevant issues posed by *big data* initiatives in manufacturing are discussed. These problems should be effectively addressed to increase the chances of success in the deployment of *Quality 4.0*. The following discussion frames the problems and takes place more from a philosophical perspective than from a technical perspective. In the context of these issues, an evolved problem solving strategy is proposed in “Problem solving strategy for a *Quality 4.0* initiative” section.

The paradigm problem

Access to large amounts of data, increased computational power, and improvement of machine learning algorithms have contributed to the sharp rise of machine learning implementation in recent years. It is important, however, to discern whether and/or which *AI* techniques are appropriate for a given situation. This section will discuss some of the pitfalls and limitations of *AI* techniques, including low “understandability” and “trustability”. Much of the following will be discussed from a manufacturing perspective.

The *big data* approach conflicts with the pre-*Industry 4.0* traditional philosophy, where models were developed and necessarily ground in physics-based theory. The inherent unknown nature of *AI*-based solutions, therefore, naturally leads to issues in trust, i.e., low “trustability.” Further contributing to the trustability issue is the fact that many machine learning techniques are *black boxes*, where the patterns used to classify a problem are not accessible or understandable. To a trained and practiced experimentalist, the inherent lack of a physics-based understanding can seem alarming. However,

² Authors do not encourage dimensionality reduction methods such as Principal Component Analysis, since they do not help to identify the driving features of the system.

from a data scientist perspective, as long as the model can accurately predict unseen data, the model is validated.

The lack of understanding also causes problems from a root-cause standpoint. If machine learning-based models are to be used from a feedback control standpoint to identify “failures”, the inability to tie the model results with the physical system render it difficult to root-cause these failures and adjust the system parameters accordingly. Moreover, correlation does not imply causation, a feature predictive power does not necessarily imply in any way that that feature is actually related to or explains the underlying physics of the system being predicted.

It is a complicated task to attempt to understand these *black box* models [63], especially if the derived solution is based on a deep learning model. Though significant efforts are taking place to elucidate *black box* models [64–66], such efforts are still in their infancy. It is therefore recommended to determine, in advance, if a *black box* solution is an acceptable approach for a given application, with the caveat that little to no time should be expended in understanding the model.

In this context, it is sensible to approach a problem by starting with simpler models (e.g. support vector machine, naive Bayes, classification trees). Most of today’s problems can be more effectively solved by simple machine learning algorithms rather than deep learning [63]. Of course, deep learning is suitable in certain applications. For example, deep learning is well suited to image classification problems, where the resulting model does not need to be well understood, but if root cause analysis (or physics knowledge creation) is part of the project goals, creating features from images [44], followed by feature ranking and interpretation is a better approach.

The existence of a large amount of available data, smart algorithms, and computational power does not mean that the creation of a machine learning model is a value-add, nor is there necessarily even a straightforward machine learning solution. This is especially true in manufacturing, where systems are complex, dynamic, with many sources of variation (e.g., suppliers, environmental conditions, etc.). It is important to discern whether an machine learning model is even required by determining if it fits in one of the following problem types: (1) regression, (2) classification, (3) clustering, or (4) time-series. If the problem does not fall into one of those four categories, it is not a straightforward machine learning problem.

Even if machine learning is an appropriate approach for a given problem, an abundance of data does not necessarily imply that such data contains discriminative patterns or contains the intrinsic information required to accurately predict. A deep understanding of the process is therefore necessary to develop a customized solution. Developing such a solution is an iterative process, and it requires collaboration between

Payoff	High	Do first	Do last
	Low	Do second	Avoid
		Low	High
		Difficulty	

Fig. 5 Lean six sigma project pick chart

the model developers and subject matter experts (engineers with domain knowledge).

In-lab solutions tend to be an overoptimistic representation of predictive system capabilities, since lab data is generated under highly controlled conditions. These conditions are likely not entirely representative of the plant environment. For this reason, lab-generated data is useful for developing proof-of-concept models; however, more work is required to increase the manufacturing readiness for plant deployment.

Many applications suffer from a scarcity of data. Lab data is often generated to develop prototype machine learning models. Been manually produced, the data is consequently often limited to hundreds (or in very fortunate cases, thousands) of samples. To perform hyperparameter tuning, these small data sets are used to create and test hundreds or even thousands of classifiers. Therefore, the chances of creating a spurious model that exhibits high prediction ability are very high [51]. But this model will never generalize beyond the training data.

As illustrated in the above text, *AI* techniques are powerful, but they are not a one-size-fits-all. Many intractable engineering problems are also *AI*-intractable. For this reason, project selection becomes an important problem and will be discussed in the following subsection.

The project selection problem

A vision can be created, resources invested, teams formed, projects selected, yet there are many situations in which little to no value is obtained. Project selection drives the success of *Quality 4.0*. Today, the hype and success stories of *AI* have influenced most organizations to have an interest in deploying an *AI* initiative. But what is the current success rate? According to recent surveys, 80–87% of projects never make it into production [14,15]. Improper project selection is the primary cause of this discouraging statistic, as many of these projects are ill-conditioned from the launch.

Whereas generic questions for choosing *AI* projects [67] and two dimensional feasibility approaches such as the pick chart (Fig. 5) from lean six sigma [68] and the machine learning feasibility assessment from Google Cloud [69] (Fig. 6) are useful and widely used [70], they do not to provide enough arguments to find the winners [71] in manufacturing applications.

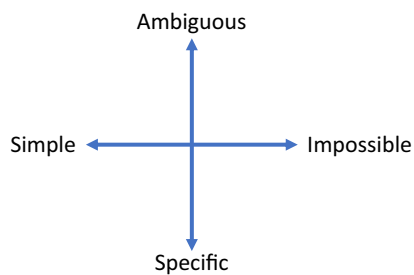


Fig. 6 Google's project pick chart

A strong project selection approach is founded on deep technical discussions, and business value analysis is required to develop a prioritized portfolio of projects. In this section, a list of 18 basic questions³ aimed at evaluating each candidate project is presented.

1. Can the problem be formulated as binary classification problem?
2. What is the profitability?
3. Are predictor features available?
4. Are all sources of variation captured by the predictor features?
5. Are the predictor features strongly related to the underlying physics of the project?
6. Is the training data (i.e., sample size) big enough?
7. Does the system exhibit chaotic behavior?
8. Are the quality labels (i.e., good, defective) available?
9. If there are no quality labels, how long it would take to generate labels?
10. Can warranty data be used?
11. Does the project align with the overall team strategy and expertise?
12. In terms of plant dynamics and restrictions (e.g, cycle times), how difficult it would be to implement the solution?
13. In terms of the plant requirements (i.e., α , β errors), what is the probability to solve the problem?
14. How long it would take?
15. Are machine learning methods more suitable than other conventional *SQC* methods?
16. Since prediction may not provide causation insights, is there any value in predicting without causation?
17. Can this project leverage/enable more projects?
18. Considering the nature of *Information Technology (IT)*, business intelligence and descriptive statistics projects, does this project promote learning, support automatic decisions and control?

³ Questions are not comprehensive and/or ordered by relevance.

Question	q ₁	q ₂	q ₃	.	.	.	q ₁₈	Total
Weights	w ₁ ∈ {0, 1, 3, 9}	w ₂	w ₃	.	.	.	w ₁₈	
Project 1	p ₁ ∈ {0, 1, 3, 9}	p ₂	p ₃	.	.	.	p ₁₈	
Project 2								
Project 3								
.								
.								
Project n								

Fig. 7 Weighted project decision matrix for *Quality 4.0*

These questions can be aggregated in a weighted decision matrix using the 0 (none), 1 (low), 3 (mid), 9 (most desirable condition) scoring system [72] to develop the data science portfolio of projects, Figs. 5, 7.

First projects should deliver value within a year. Therefore, even low rewarding “low hanging fruit” projects that may create business synergy and trust in these technologies are better starting points than highly rewarding “moon shots”. Projects with data already available can be a good place to start, since data generation is manual and time consuming.

If there is no in-house *AI* expertise, working with external partners is a good idea-while developing in-house expertise-to catch up with the pace of *AI*'s rise. However, the overoptimistic bias induced by vendors and the fallacy of planning should be avoided when selecting and defining a project.

It is important to keep in mind that data science projects are not independent from one other. To increase success likelihood, data science teams should be working on several projects concurrently. Once the first successful project is implemented, it will inspire new projects and ignite new ideas towards the development of the mid/long term visions. Companies with a well defined project selection strategy will be more likely to succeed in the era of *Quality 4.0*.

The redesign problem

Machine learning algorithm are trained using observational data [73]. In contrast with experimental data that can be used for causal analysis, machine learning-derived information is ill conditioned for this task [63,64,74,75], as discovered patterns may be spurious representations. To address this situation, the relevant features identified through data-driven methods (e.g., feature ranking/selection) should be used to extract information (e.g., uncover hidden patterns, associations, and correlations). This empirically-derived information, along with a physics analysis, can be used to generate useful hypotheses about possible connections between the features and the quality of the product. Then statistical analyses (e.g., randomized experiments) can be devised to establish causality to augment root-cause analyses and to

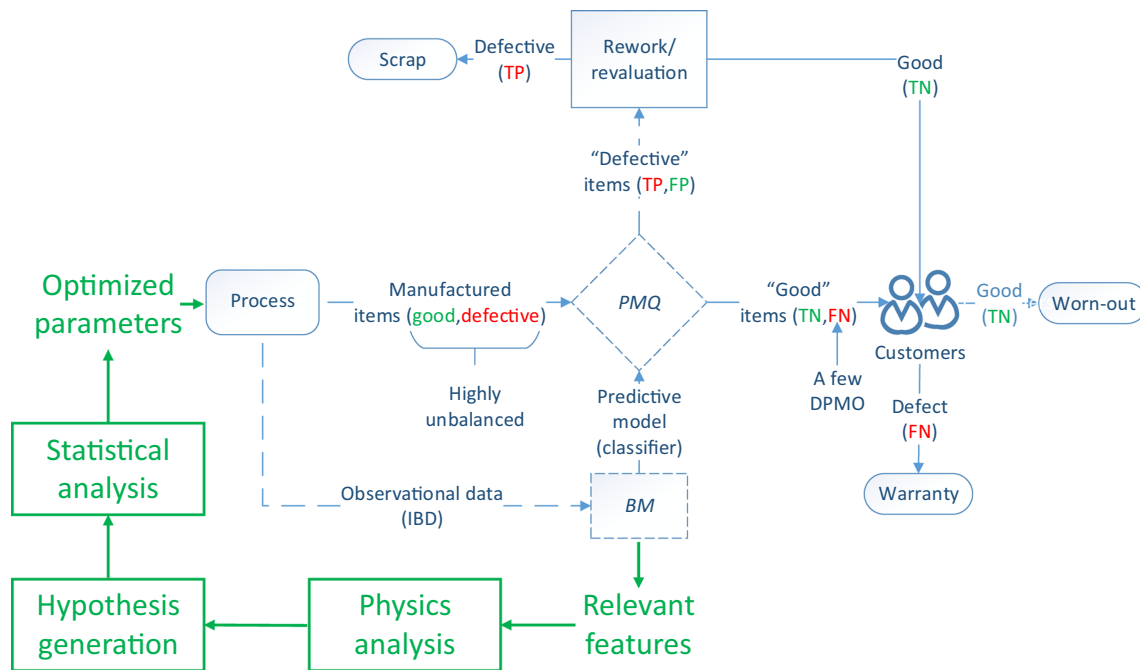


Fig. 8 Causal analysis and a road-map for process parameter optimization

find optimal parameters to redesign the process, as shown in Fig. 8. A real case study of this concept is presented in [10].

The relearning problem

In machine learning, the concept of drift embodies the fact that the statistical distributions of the classes of which the model is trying to predict, change over time in unforeseen ways. This poses a serious technical challenges, as the classifiers assume a stationary relationship between features and classes. This static assumption is rarely satisfied in manufacturing. Consequently, the prediction capability of a trained model tends to significantly degrade overtime, resulting in friction, distrust, and a bad reputation for the data science teams. Although this concept has been widely studied [76,77], in this section a few insights from manufacturing perspective are provided.

The transient and novel sources of variations cause manufacturing systems to exhibit non-stationary data distributions. Therefore, developing an in-lab model with the capacity to predict well the items in the test set and satisfy the learning targets for the project, is only the beginning since the model will usually degrade after deployment. In Fig. 9a, a real situation is presented, in which the model exhibited less than 1.5% of α error (target set by the plant) in the test set (lab environment) to satisfy the defect detection goals ($\beta < 5\%$) of the project. Immediately after deployment, the α error increased to 1.78%. A few days later, the α increased to 4%,

an unacceptable FP rate for the plant. Although this model exhibited good prediction ability and made it into production, it was not a sustainable solution, and therefore never created value.

To put prediction decay into a binary classification context, it means that either or both α or β errors increase to a problematic point. In the case of α error, excessive FP generate the hidden factory effect (e.g., reducing process efficiency), whereas an increase of the β error, would generate more warranties, a less desirable situation.

The main goal of relearning, is to keep the predictive system in compliance with the restrictions set by the plant (α error), Fig. 9b, and the detection goals (β error). This is accomplished by ensuring that the algorithm is learning the new statistical properties of both classes (good, defective). Continual learning or auto-adaptive learning is a fundamental concept in artificial intelligence that describes how the algorithms should autonomously learn and adapt in production as new data with new patterns comes in. Broadly speaking, a relearning scheme should include the following four components:

1. **Learning strategy** A learning scheme must be defined in advance. This is basically a research challenge where the following topics are addressed: (1) data generation and pre-processing, (2) machine learning algorithm(s), (3) hyperparameter tuning strategy, and (4) model validation. The outcome of this component is the *final model* along

with the estimated prediction performance on unseen data. This information is compared with the learning goals before deployment.

2. **Relearning data set** Since the rework/revaluation station contain *TP* and *FP* cases that are usually further inspected by a human, the labels of the latter can be updated to *TN*. Then, this information along with warranty data should be used to generate the retraining data set (note that the *FN* labels of the warranties must be updated too). Thus, the algorithm will be able to constantly learn the statistical properties of the good and defective items, Fig. 10. Some algorithms use a window of a fixed size (e.g., time windowed forgetting), while others use heuristics to adjust the window size to the current extend of the concept drift (e.g., adaptive size) [78].
3. **Relearning schedule** Algorithms should be retrained as often as possible to promptly adapt the classifier to the new sources of variations. Therefore, a relearning schedule based on the plant dynamics should be defined in advance.
4. **Monitoring system** Even if a relearning scheme has been defined, there is no guarantee that the classifier would perform within tolerances all the time. Therefore, it is important to make sure that if the model is losing its prediction ability (e.g., corrupted data), there is an alerting mechanism in place to implement temporary solutions (e.g., 100% manual inspections or random sampling). Basically a monitoring system it is necessary to keep track of both errors, Fig. 11. When the statistical properties of the good classes change, usually the number of *FP* significantly increase, overloading the inspection/revaluation station. Whereas if novel defects are created or the statistical properties of the defective class change, the number of *FN* increases, generating more warranties than expected. In general, the main goal here is to develop an algorithm with the minimal number of false alarm rates and the maximal number of early drift detections [77].

Other applications of machine learning tend to be more stable. For example, a neural network is trained to identify a stop sign. The stop sign will not change over time, and the algorithm can only improve with more data. This is not true for many manufacturing scenarios. To develop a sustainable solution, a relearning scheme should be developed before deployment that considers the type of drift the manufacturing system exhibits, e.g., temporal, gradual, recurring or sudden [78]. Although continual learning will ultimately optimize models for accuracy and save manual retraining time by making models auto-adaptive, even the more advance solutions would need the human supervision to ensure the predictive system is behaving as expected.

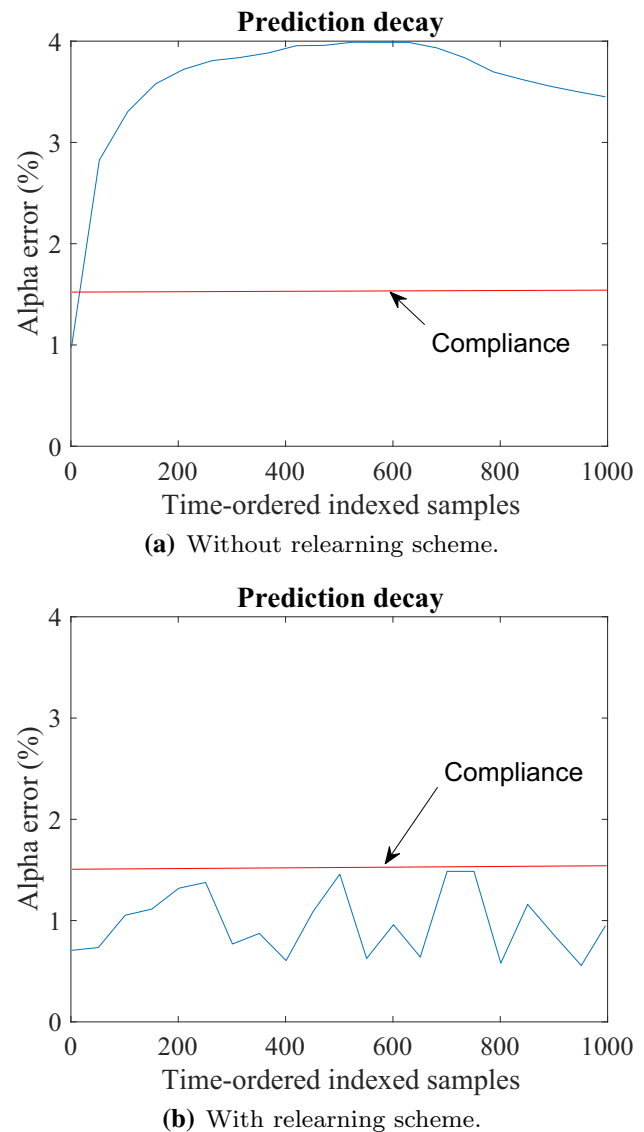


Fig. 9 Error analysis after deployment

Problem solving strategy for a Quality 4.0 initiative

In the context of the issues presented in “Challenges of big data initiatives in manufacturing” section, an evolved problem solving strategy is presented to effectively implement AI in manufacturing. This evolved strategy is an update to the initial problem solving strategy—*acsensorize, discover, learn, predict*—proposed by *PMQ* [5]. Three more issues—*identify, redesign, relearn*—relevant to *big data* in manufacturing are included as extra steps to develop a comprehensive problem solving strategy for a successful implementation of a *Quality 4.0* initiative, Fig. 12.

Since its introduction into the manufacturing world, the quality movement has continuously included new tech-

Fig. 10 Relearning data

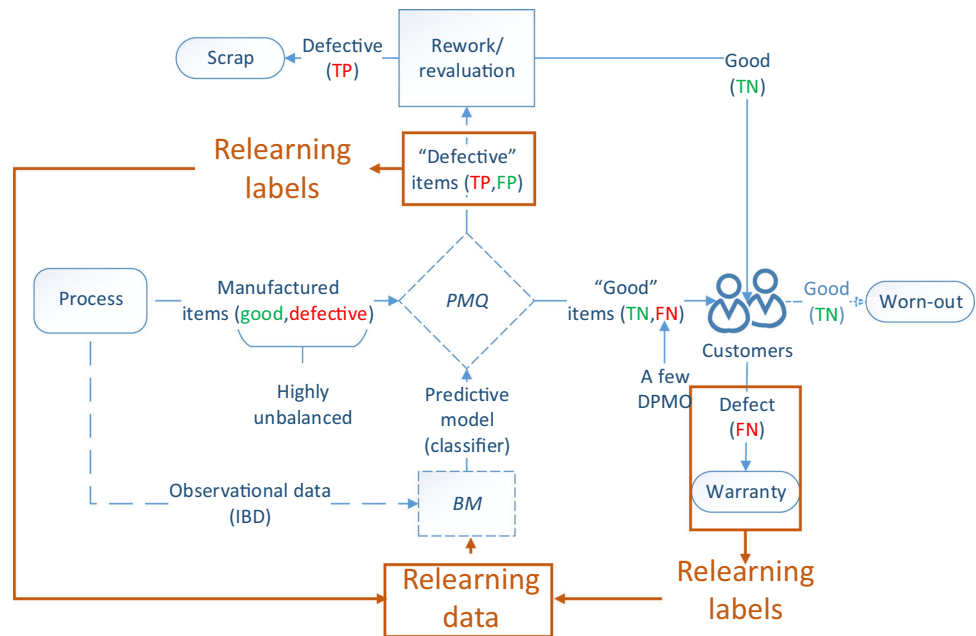
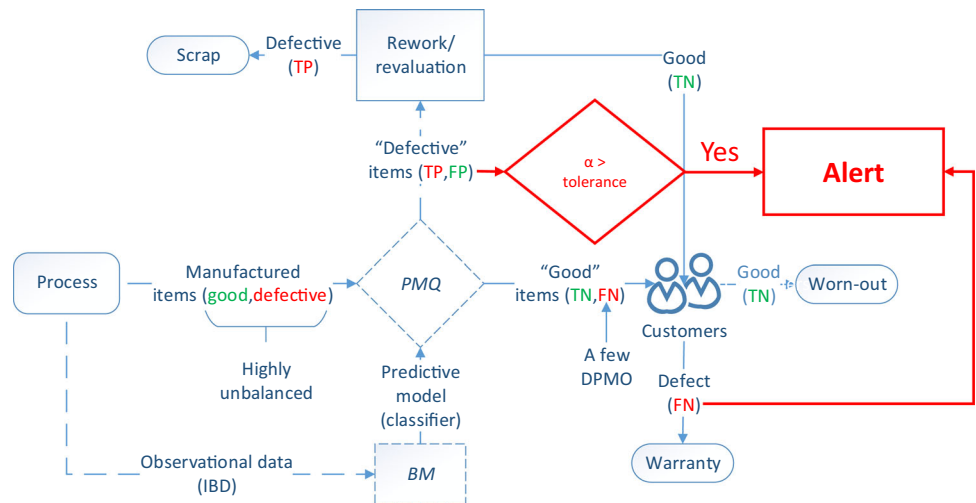


Fig. 11 Alerting system



niques, ideas, and philosophies from the various fields of engineering and science. Modern quality control can be traced back to the 1930's when Dr. Walter Shewhart developed a new industrial *SQC* theory [79]. He proposed a 3-step problem solving strategy (*specification, production, inspection—SPI*) based on the scientific method. This problem solving method is known as the Shewhart learning and improvement cycle [80], Fig. 13a.

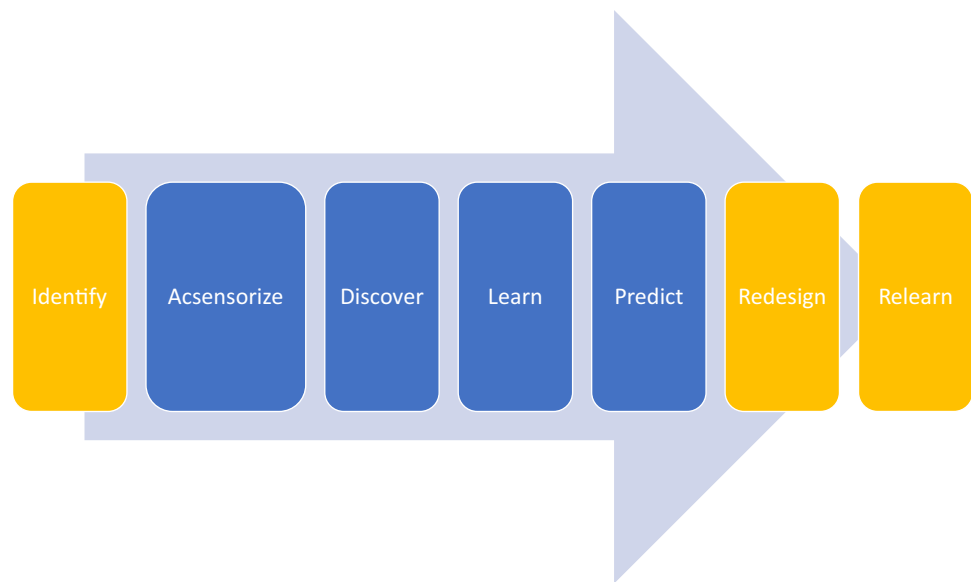
The *TQM* philosophy emerged in the 1980's [81], where the Shewhart's problem solving strategy was refined by Deming, also known as the Deming cycle, Fig. 13b, to develop a more comprehensive problem solving approach (*plan, do, check/study, act,—PDCA or PDSA*).

A few years later, in 1986, six sigma was introduced by Bill Smith at Motorola. Six sigma is a reactive approach, based on a 5-step problem solving strategy (*define, measure, ana-*

lyze, improve, control—DMAIC) aimed at identifying and eliminating the causes of defects and/or sources of variation from processes. Its complement *Design for Six Sigma (DFSS)* appeared shortly after, a proactive approach aimed at designing robust products and processes so that defects are never generated, Fig. 13d, [82,83]. It is founded on optimization techniques which are usually applied through a 6-step problem solving strategy (*define, measure, analyze, design, optimize, verify—DMADOV*).

To cope with the ever increasing complexity of manufacturing systems, each new quality philosophy has redefined the problem solving strategy. It started with a basic 3-step approach and has gradually evolved and increased its complexity up-to the 7-step approach (*identify, acsensorize, discover, learn, predict, redesign, relearn—IADLPR²*) herein proposed, Fig. 13.

Fig. 12 Comprehensive problem solving strategy for *Quality 4.0* (new steps in yellow)



Quality philosophy				
(a) SQC	(b) TQM	(c) Six sigma	(d) DFSS	(e) Quality 4.0
Controlling	Managing	Reactive	Proactive	Predictive
Specification	Plan	Define	Define	Identify
Production	Do	Measure	Measure	Acensorize
Inspection	Check	Analyze	Analyze	Discover
	Act	Improve	Design	Learn
		Control	Optimize	Predict
			Verify	Redesign
				Relearn

Fig. 13 Evolution of quality, problem solving strategy by quality philosophy

Manufacturing systems pose particular intellectual and practical challenges (e.g., cycle times, transient sources of variation, reduced lifetime, high conformance rates, physics-based) that must be effectively addressed to create value out of *IBD*. The *IADLPR*² problem solving strategy is founded on theory and our knowledge studying complex manufacturing systems. According to empirical results, this new problem solving strategy should increase the chances of successfully deploying a *Quality 4.0* initiative.

Strategy development and adoption for *Quality 4.0*

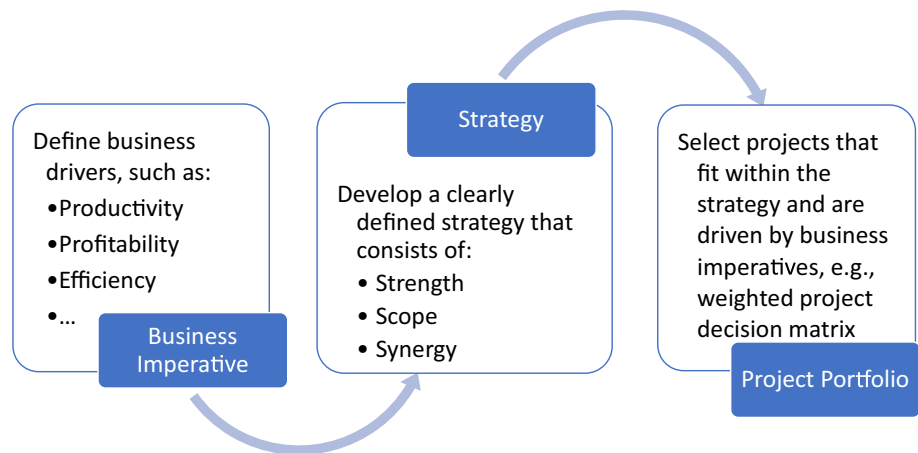
Adopting a new philosophy will inevitably suffer from resistance. Broadly, this resistance to adoption falls into three categories: psychological reservations, infrastructural limitations, and business impediments. A strategic vision for adopting the *Quality 4.0* philosophy can help mitigate against this. The strategy should be developed and driven by business

imperatives, and should be used to create a project portfolio. Any successful strategy should comprise strength, scope, and synergy. Strength defines the competitive advantage, i.e., what makes the company “stand out” from the rest. Scope encompasses not only what the company should do, but what they choose not to do. Finally, synergy ensures that all parts work together in support of the overall vision, i.e., there are no contradictory practices. Figure 14 contains a map for strategy development and adoption of *Quality 4.0*. In this section, some of the considerations that should be taken into account when developing/adopting such a strategic vision will be discussed. In the interest of brevity, this discussion will be kept short. The reader is referred to Porter [84], Cheatham et. al [85], Kotter [86], Bughin [87], and Davenport and Ronanki [88] for more detailed discussions on *AI* and strategy.

Having the appropriate infrastructure is crucial to the ability to adopt *Quality 4.0* practices. *Quality 4.0* practices require a fully connected lab/production environment to enable data acquisition and analysis. *IT* should support and enable any updates required while ensuring that adequate security and access policies are adopted. Adequate infrastructure should exist for data storage, computation, and deployment (to enable scaled-up solutions). Ultimately, for a new vision to be successfully adopted, the company culture must accept it. For this reason, it is important that the philosophy is advocated for and supported by management. It can help the overall attitude to locally implement relevant aspects of *Quality 4.0* practices, i.e., not all at once. Local successes can promote positive attitudes towards adoption.

Companies should perform a risk-benefit analysis, such as projections on investments versus the benefits of developing *Quality 4.0* capabilities. Since *AI* in manufacturing is in an early stage, there is no guarantee of return-on-investment (at

Fig. 14 Map of strategy development and adoption of *Quality 4.0*



least in the short term). The company therefore may need to incorporate this analysis into budget planning to determine if deployment of *Quality 4.0* resources (e.g., staff, hardware and software) is possible. Once the company culture, infrastructure, and business conditions allow for adopting a *Quality 4.0* strategy, successful deployment should consist of cross-functional teams. Cross-functional teams enable the successful execution of *Quality 4.0* projects. For example, management should select impactful projects. In house data science expertise should exist (although this can be complemented with some consulting work too). Data science teams should confirm that the projects that are actually machine learning projects (i.e., fit within the scope of *Quality 4.0* paradigm). Management and *IT* should coordinate with plant teams to ensure proper collection of data. *IT* should facilitate projects by ensuring collected data is available for analysis and the solution scaled and deployed. Finally, engineers and subject matter experts should confirm that the solution satisfies the engineering requirements and that the model is trustable (i.e., it captures the underlying physics of the process).

Conclusions

The hype surrounding artificial intelligence and *big data* have propelled an interest by many quality leaders to deploy this technology in their companies. According to recent surveys, however, 80–87% of these projects never develop a production-ready sustainable solution and many of these leaders do not have a clear vision of how to implement *Quality 4.0* practices. In this article, four *big data* challenges in manufacturing are discussed: the change in paradigm, the project selection problem, the process redesign problem, and the relearning problem. These challenges that must be fully understood and addressed to enable the successful deployment of *Quality 4.0* initiatives in the context of *Pro-*

cess Monitoring for Quality. Based on this study, a novel 7-step problem solving strategy (*identify, acsensorize, discover, learn, predict, redesign, relearn*) is proposed. This comprehensive approach increases the likelihood of success.

The importance of developing and adopting a *Quality 4.0* strategy is also discussed. A well defined strategy ensures that business imperatives are met, mitigates against resistance to adoption, and enables the appropriate selection of projects.

Future research should focus on evaluating project potential via the 18 questions herein presented in order to describe the characteristics associated at each level, i.e., none, low, mid and desirable. This would increase the effectiveness of using the weighted project selection matrix to rank (prioritize) projects. This work can also be extended by redefining the problem solving strategy based on deep learning, as the discovered step (i.e. future creation/engineering) is not required.

References

1. Kapás, J. (2008). Industrial revolutions and the evolution of the firm's organization: An historical perspective. *Journal of Innovation Economics Management*, 2(2), 15–33.
2. Kagermann, H., Helbig, J., Hellinger, A., & Wahlster, W. (2013). *Recommendations for implementing the strategic initiative Industrie 4.0: Securing the future of German manufacturing industry; final report of the Industrie 4.0 working group*. Forschungsunion.
3. G.I.P.C. Information, "The rise of industrial big data," (2012). (Online). <http://www.geautomation.com/download/rise-industrial-big-data>
4. Sadiku, M. N., Wang, Y., Cui, S., & Musa, S. M. (2017). Industrial internet of things. *International Journal of Advanced Research in Science, Engineering*, 3, 4724–4734.
5. Abell, J. A., Chakraborty, D., Escobar, C. A., Im, K. H., Wegner, D. M., & Wincek, M. A. (2017). Big data driven manufacturing—Process-monitoring-for-quality philosophy. *ASME Journal of Manufacturing Science and Engineering on Data Science-Enhanced Manufacturing*, 139(10), 101009.
6. Schwab, K. (2016). *The fourth industrial revolution: what it means, how to respond*. World economic forum.

7. Yin, S., & Kaynak, O. (2015). Big data for modern industry: Challenges and trends [point of view]. *Proceedings of the IEEE*, 103(2), 143–146.
8. Zhong, R. Y., Xu, X., Klotz, E., & Newman, S. T. (2017). Intelligent manufacturing in the context of industry 4.0: a review. *Engineering*, 3(5), 616–630.
9. Radziwill, N. M. (2018). *Quality 4.0: Let's get digital-the many ways the fourth industrial revolution is reshaping the way we think about quality*. Preprint [arXiv:1810.07829](https://arxiv.org/abs/1810.07829)
10. Escobar, C. A., Arinez, J., & Morales-Menendez, R. (2020). *Process-monitoring-for-quality—A step forward in the zero defects vision*. SAE technical paper, no. 2020-01-1302, 4 (Online). <https://doi.org/10.4271/2020-01-1302>
11. Dan, J. (2017). *Quality 4.0 fresh thinking for quality in the digital era*. Quality digest.
12. NewVantage Partners, L. (2019). *Big data and AI executive survey 2019: Data and innovation how big data and AI are accelerating business transformation*.
13. Ortega, M. (2018). *Cio survey: Top 3 challenges adopting AI and how to overcome them*. Databricks (Online). <https://databricks.com/blog/2018/12/06/cio-survey-top-3-challenges-adopting-ai-and-how-to-overcome-them.html>
14. STAFF, V. O. “Why do 87 production?” (Online). <https://venturebeat.com/2019/07/19/why-do-87-of-data-science-projects-never-make-it-into-production/>
15. Research, G. (2018). Predicts 2019: Data and analytics strategy. Nov. (Online). <https://emtemp.gcom.cloud/ngw/globalassets/en/doc/documents/374107-predicts-2019-data-and-analytics-strategy.pdf>
16. Escobar, C. A., Abell, J. A., Hernández-de Menéndez, M., & Morales-Menendez, R. (2018). Process-monitoring-for-quality—Big models. *Procedia Manufacturing*, 26, 1167–1179.
17. Zhong, R. Y., Dai, Q., Qu, T., Hu, G., & Huang, G. Q. (2013). RFID-enabled real-time manufacturing execution system for mass-customization production. *Robotics and Computer—Integrated Manufacturing*, 29(2), 283–292.
18. Xu, X. (2012). From cloud computing to cloud manufacturing. *Robotics and Computer—Integrated Manufacturing*, 28(1), 75–86.
19. Kusiak, A. (2017). Smart manufacturing must embrace big data. *Nature*, 544(7648), 23–25.
20. Tao, F., Qi, Q., Liu, A., & Kusiak, A. (2018). Data-driven smart manufacturing. *Journal of Manufacturing Systems*, 48, 157–169.
21. Fisher, O. J., Watson, N. J., Escrig, J. E., Witt, R., Porcu, L., Bacon, D., et al. (2020). Considerations, challenges and opportunities when developing data-driven models for process manufacturing systems. *Computers and Chemical Engineering*, 140, 106881.
22. Kusiak, A. (1990). *Intelligent manufacturing systems*. London: Prentice Hall Press.
23. Cadavid, J. P. U., Lamouri, S., Grabot, B., Pellerin, R., & Fortin, A. (2020). Machine learning applied in production planning and control: A state-of-the-art in the era of industry 4.0. *Journal of Intelligent Manufacturing*, 2020, 1–28.
24. Feng, Y., Zhao, Y., Zheng, H., Li, Z., & Tan, J. (2020). Data-driven product design toward intelligent manufacturing: A review. *International Journal of Advanced Robotic Systems*, 17(2), 1729881420911257.
25. Choudhary, A. K., Harding, J. A., & Tiwari, M. K. (2009). Data mining in manufacturing: A review based on the kind of knowledge. *Journal of Intelligent Manufacturing*, 20(5), 501.
26. Cheng, L., Tang, Q., Zhang, Z., Wu, S., et al. (2020). Data mining for fast and accurate makespan estimation in machining workshops. *Journal of Intelligent Manufacturing*, 32, 1–18.
27. Wei, W., Yuan, J., & Liu, A. (2020). Manufacturing data-driven process adaptive design method. *Procedia CIRP*, 91, 728–734.
28. Wuest, T., Irgens, C., & Thoben, K.-D. (2014). An approach to monitoring quality in manufacturing using supervised machine learning on product state data. *Journal of Intelligent Manufacturing*, 25(5), 1167–1180.
29. Malaca, P., Rocha, L. F., Gomes, D., Silva, J., & Veiga, G. (2019). Online inspection system based on machine learning techniques: Real case study of fabric textures classification for the automotive industry. *Journal of Intelligent Manufacturing*, 30(1), 351–361.
30. Xu, C., Zhu, G., et al. (2020). Intelligent manufacturing lie group machine learning: Real-time and efficient inspection system based on fog computing. *Journal of Intelligent Manufacturing*, 32, 1–13.
31. Cai, W., Zhang, W., Hu, X., & Liu, Y. (2020). A hybrid information model based on long short-term memory network for tool condition monitoring. *Journal of Intelligent Manufacturing*, 31, 1–14.
32. Said, M., Ben-Abdellafou, K., & Taouali, O. (2019). Machine learning technique for data-driven fault detection of nonlinear processes. *Journal of Intelligent Manufacturing*, 2019, 1–20.
33. Hui, Y., Mei, X., Jiang, G., Zhao, F., Ma, Z., & Tao, T. (2020). Assembly quality evaluation for linear axis of machine tool using data-driven modeling approach. *Journal of Intelligent Manufacturing*, 2020, 1–17.
34. Belfiore, M. (2016). *Automation opportunities abound for quality inspections*. Automation World (Online). <https://www.automationworld.com/products/software/article/13315584/automation-opportunities-abound-for-quality-inspections>
35. See, J. E. (2015). Visual inspection reliability for precision manufactured parts. *Human Factors*, 57(8), 1427–1442.
36. Miller, J. G., & Vollmann, T. E. (1985). *The hidden factory*. Harvard Business Review (Online). <https://hbr.org/1985/09/the-hidden-factory>
37. EY, A. (2019). Oxford, “Sensors as drivers of industry 4.0—A study on Germany, Switzerland and Austria (Online). <https://www.oxan.com/insights/client-thought-leadership/ey-sensors-as-drivers-of-industry-40/>
38. Huan, L., & Motoda, H. (eds) (1998). *Feature extraction, construction and selection: A data mining perspective*. Vol. 453. Springer Science & Business Media
39. Boubchir, L., Daachi, B., & Pangracious, V. (2017). *A review of feature extraction for EEG epileptic seizure detection and classification*. In: 2017 40th International Conference on Telecommunications and Signal Processing (TSP). IEEE
40. Escobar, C. A., Macias, D., Hernández-de Menéndez M., & Morales-Menendez, R. (2021). Process-monitoring-for-quality—Multiple classifier system for highly unbalanced data. *Studies in Big Data*.
41. Devore, J. (2015). Probability and statistics for engineering and the sciences. *Cengage Learning*, 2015, 768.
42. Escobar, C. A., Wincek, M. A., Chakraborty, D., & Morales-Menendez, R. (2018). Process-monitoring-for-quality—Applications. *Manufacturing Letters*, 16, 14–17.
43. Wuest, T., Irgens, C., & Thoben, K. (2013). An approach to quality monitoring in manufacturing using supervised machine learning on product state based data. *Journal of Intelligent Manufacturing*, 20, 1.
44. Christ, M., Kempa-Liehr, A. W., & Feindt, M. (2016). Distributed and parallel time series feature extraction for industrial big data applications. Preprint [arXiv:1610.07717](https://arxiv.org/abs/1610.07717)
45. Granstedt Möller, E. (2017). *The use of machine Learning in industrial quality control*. Thesis by Erik Granstedt Möller for the degree of Master of Science in Engineering. KTH Royal Institute of Technology. Stockholm, Sweden. <https://www.diva-portal.org/smash/get/diva2:1150596/FULLTEXT01.pdf>
46. Li, Y., Zhao, W., & Pan, J. (2016). Deformable patterned fabric defect detection with fisher criterion-based deep learning. *IEEE Transactions on Automation Science and Engineering*, 14(2), 1256–1264.
47. Charniak, E. (2019). *Introduction to deep learning*. London: The MIT Press.

48. Khan, A., Sohail, A., Zahoor, U., & Qureshi, A. S. (2020). A survey of the recent architectures of deep convolutional neural networks. *Artificial Intelligence Review*, 53, 1–62.
49. Wolpert, D. H. (1996). The lack of a priori distinctions between learning algorithms. *Neural Computation*, 8(7), 1341–1390.
50. Friedman, J., Hastie, T., & Tibshirani, R. (2001). *The elements of statistical learning*. Statistics (p. 1). Berlin: Springer.
51. Burnham, K. P., & Anderson, D. R. (2003). *Model selection and multimodel inference: A practical information-theoretic approach*. Berlin: Springer.
52. Escobar, C. A., Arinez, J., Macias, D., & Morales-Menendez, R. (2020). A separability-based feature selection method for highly unbalanced binary data. *International Journal on Interactive Design and Manufacturing*, 2020, 1.
53. Escobar, C. A., & Morales-Menendez, R. (2019). Process-monitoring-for-quality—A model selection criterion for support vector machine. *Procedia Manufacturing*, 34, 1010–1017.
54. Wright, M. N., & König, I. R. (2019). Splitting on categorical predictors in random forests. *PeerJ*, 7, e6339.
55. Sarle, W. S. (2002). comp.ai.neural-nets faq, part 2 of 7: Learning. faqs.org (Online). <http://www.faqs.org/faqs/ai-faq/neural-nets/part2/>
56. Ioffe, S., & Szegedy, C. (2015). Batch normalization: Accelerating deep network training by reducing internal covariate shift. *International Conference on Machine Learning*, 2015, 448–456.
57. Mohamad, I. B., & Usman, D. (2013). Standardization and its effects on k-means clustering algorithm. *Research Journal of Applied Sciences, Engineering and Technology*, 6(17), 3299–3303.
58. Juszczak, P., Tax, D., & Duin, R. P. (2002). Feature scaling in support vector data description. In *Proceedings of the asci* (pp. 95–102). Citeseer.
59. Jakobsen, J. C., Gluud, C., Wetterslev, J., & Winkel, P. (2017). When and how should multiple imputation be ESED for handling missing data in randomised clinical trials—A practical guide with flowcharts. *BMC Medical Research Methodology*, 17(1), 162.
60. Barnard, J., & Meng, X.-L. (1999). Applications of multiple imputation in medical studies: From AIDS to NHANES. *Statistical Methods in Medical Research*, 8(1), 17–36.
61. Rahman, M. M., & Davis, D. N. (2013). Machine learning-based missing value imputation method for clinical datasets. *IAENG Transactions on Engineering Technologies*, 2013, 245–257.
62. Escobar, C. A., Arinez, J., Macias, D., & Morales-Menendez, R. (2020). Learning with missing data. In *2020 IEEE international conference on big data* (pp. 5037–5045).
63. Bathaee, Y. (2017). The artificial intelligence black box and the failure of intent and causation. *Harvard Journal of Law & Technology*, 31, 889.
64. Shmueli, G. (2010). To explain or to predict? *Statistical Science*, 25, 289–310.
65. Garreau, D., & von Luxburg, U., (2020). Explaining the explainer: A first theoretical analysis of lime. Preprint [arXiv:2001.03447](https://arxiv.org/abs/2001.03447)
66. Ribeiro, M. T., Singh, S., & Guestrin, C. (2016). Why should I trust you? Explaining the predictions of any classifier. In *Proceedings of the 22nd ACM international conference on knowledge discovery and data mining* (pp. 1135–1144).
67. Ng, A. (2019). How to choose your first AI project. In *Artificial intelligence: The insights you need from Harvard Business Review* (pp. 79–88).
68. George, M. (2013). *Lean six sigma for service*. Maxima.
69. Cloud, G. (2021). Managing machine learning projects with google cloud. Coursera [Online]. <https://www.coursera.org/lecture/machine-learning-business-professionals/practice-assessing-the-feasibility-of-ml-use-cases-Lhquq>
70. Pestana, D. (2020). *How to say no to useless data science projects and start working on what you want*. Towards Data Science (Online). <https://towardsdatascience.com/how-to-say-no-to-projects-d6f88641ab3c>
71. Mason, H. (2018). *How to decide which data science projects to pursue*. Harvard Business Review. (Online). <https://hbr.org/2018/10/how-to-decide-which-data-science-projects-to-pursue>
72. Cohen, L., (1995). *Quality function deployment: how to make QFD work for you*. Prentice Hall
73. Montgomery, D. (2017). Exploring observational data. *Quality and Reliability Engineering International*, 33(8), 1639–1640.
74. Wasserman, L. (2013). *All of statistics: A concise course in statistical inference*. Berlin: Springer.
75. Leetaru, K. (2019). *A reminder that machine learning is about correlations not causation*. Forbes (Online). <https://www.forbes.com/sites/kalevleetaru/2019/01/15/a-reminder-that-machine-learning-is-about-correlations-not-causation/#7d6268d46161>
76. Webb, G. I., Lee, L. K., Petitjean, F., & Goethals, B. (2017). *Understanding concept drift*. Preprint [arXiv:1704.00362](https://arxiv.org/abs/1704.00362)
77. Wang, H., & Abraham, Z. (2015). Concept drift detection for streaming data. In *2015 international joint conference on neural networks (IJCNN)* (pp. 1–9). IEEE.
78. Tsybmal, A. (2004). The problem of concept drift: Definitions and related work. *Computer Science Department, Trinity College Dublin*, 106(2), 58.
79. Juran, J. M. (1997). Early SQC: A historical supplement. *Quality Progress*, 30(9), 73.
80. Moen, R. D., & Norman, C. L. (2010). Circling back. *Quality Progress*, 43(11), 22.
81. Mandru, L., Patrascu, L., Carstea, C.-G., Popescu, A., & Birsan, O. (2011). Paradigms of total quality management. In *Recent research in manufacturing engineering*, pp. 121–126, Transilvania University of Braşov, Romania.
82. Chowdhury, S. (2002). *Design for six sigma*. Upper Saddle River, New Jersey 07458: Financial Times Prentice Hall.
83. Basem, E.-H. (2008). *Design for six sigma: A roadmap for product development*. New York, NY: McGraw-Hill Publishing.
84. Porter, M. E., et al. (1996). What is strategy? *Harvard Business Review*, 74(6), 61–78.
85. Cheatham, B., Mehta, C. A. N., & Shah, D. (2020). *How to build AI with (and for) everyone in your organization*. Technical report, McKinsey Analytics.
86. Kotter, J. P. (1995). Leading change: Why transformation efforts fail. *Harvard Business Review*.
87. Bughin, J., & Hazan, E. (2017). Five management strategies for getting the most from AI. *MIT Sloan Management Review*.
88. Davenport, T. H., & Ronanki, R. (2018). Artificial intelligence for the real world. *Harvard Business Review*, 96(1), 108–116.

Publisher's Note Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.