

Article

Data-Centric and Model-Centric AI: Twin Drivers of Compact and Robust Industry 4.0 Solutions

Oussama H. Hamid 

Faculty of Computer Information Sys, Higher Colleges of Technology,
Abu Dhabi P.O. Box 41012, United Arab Emirates; ohamid@hct.ac.ae

Abstract: Despite its dominance over the past three decades, model-centric AI has recently come under heavy criticism in favor of data-centric AI. Indeed, both promise to improve the performance of AI systems, yet with converse points of focus. While the former successively upgrades a devised model (algorithm/code), holding the amount and type of data used in model training fixed, the latter enhances the quality of deployed data continuously, paying less attention to further model upgrades. Rather than favoring either of the two approaches, this paper reconciles data-centric AI with model-centric AI. In so doing, we connect current AI to the field of cybersecurity and natural language inference, and through the phenomena of ‘adversarial samples’ and ‘hypothesis-only biases’, respectively, showcase the limitations of model-centric AI in terms of algorithmic stability and robustness. Further, we argue that overcoming the alleged limitations of model-centric AI may well require paying extra attention to the alternative data-centric approach. However, this should not result in reducing interest in model-centric AI. Our position is supported by the notion that successful ‘problem solving’ requires considering both the way we act upon things (algorithm) as well as harnessing the knowledge derived from data of their states and properties.

Keywords: artificial intelligence; adversarial samples; current AI; data-centric AI; deep learning; Industry 4.0; machine learning; model-centric AI



Citation: Hamid, O.H. Data-Centric and Model-Centric AI: Twin Drivers of Compact and Robust Industry 4.0 Solutions. *Appl. Sci.* **2023**, *13*, 2753. <https://doi.org/10.3390/app13052753>

Academic Editors: Qiang Zhang and Yifeng Zeng

Received: 3 January 2023

Revised: 13 February 2023

Accepted: 16 February 2023

Published: 21 February 2023



Copyright: © 2023 by the author. Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<https://creativecommons.org/licenses/by/4.0/>).

1. Introduction

With the fourth industrial revolution, also referred to as Industry 4.0, a new wave of automation set off carrying novel industrial applications that promise to redefine the interaction between man and machine [1]. Businesses are increasingly shifting their service models from traditional automation, in which the integration of software systems on service platforms mainly depends on the outstanding programming skills of human software developers and the deployment of appropriate application programming interfaces (APIs) [2,3], to engineered systems that integrate sensing abilities, computational prowess, control, and network utilities into *cyber*-physical objects that can connect to each other over the ‘Internet of Things’ (IoT) [4].

Data are a conspicuous hallmark of Industry 4.0 technologies [5]. Examples range from cloud computing [6], where large volumes of data can be stored, analysed, and processed more efficiently and cost-effectively with the cloud, over edge computing [7], which provides low end-to-end latency in real-time production operations [8] by allowing efficient near-sensor data analysis [9], to digital twins [10], where simulations of multiple systems’ processes can be run on virtual environments, pulling data from IoT sensors and whatever objects connected to the Internet [11].

Of particular importance in moving data around the grid are the fields of artificial intelligence (AI) and machine learning (ML) [12]. Together, they can facilitate the creation of sustainable solutions and scalable business models by taking full advantage of available data volumes in a way that exceeds the use of the data generated within the concerned business field [13,14] to include the data obtained from other fields and businesses [15].

At its core, an AI system is mainly an algorithm (code) that solves a problem by learning prototypical features (patterns) from large volumes of data. Data could be in any of different forms (for example, text, audio, image, and/or video). ML is a subfield of AI that characterises the system's capacity to spot patterns that are otherwise undetectable by humans [16]. It achieves this by using general-purpose procedures, which enable the AI system to solve the problem at hand without being explicitly programmed to do so [17]. Yet, for this to occur, the system needs to process data in a way that facilitates interpretation and provides context [16]. Typically, this is made possible by letting human experts label the data as part of the data preparation process, which includes collecting, curating, cleaning, and transforming raw data before processing and analysing it. This subtask is known as 'data annotation'.

Currently, there are two different views within the ML community on how to improve the performance of AI systems: model-centric AI and data-centric AI. In model-centric AI, developers of an AI system successively upgrade a devised model (algorithm/code), while holding the volume and type of data collected fixed. Conversely, one can hold the model fixed and continuously improve the quality of data until reaching a high level of overall performance in terms of dealing with noise in the data. This is the data-centric AI approach (Figure 1). Due to the fact that the training processes in both approaches run in a continuous, iterative fashion, some refer to them as 'model-cycle' and 'data-cycle' approaches, respectively, [18].

Over the past three decades, model-centric AI was dominantly used in both research and industry. In research, for example, more than 90% of the published AI research projects utilised a model-centric approach [19]. Recently, however, voices are getting louder in the promotion of data-centric AI. Reasons for this change of heart range from the lack of sufficiently large datasets [20] to the need for highly customised solutions [21].

The concept was introduced by Andrew Ng, founder of 'DeepLearning.AI' (an educational platform) and Adjunct Professor at Stanford University. In a live stream presentation on 24 March 2021, Ng showcased the use of a data-centric approach in the detection of foreign-particle defects in steel sheets, which resulted in a considerably higher performance of the baseline model of an AI system devised to detect steel defects, compared with the almost no-improvement outcome when a model-centric approach was applied to the baseline model [19]. Furthermore, Ng showed that the application of the data-centric approach to other tasks such as solar defect detection and surface inspection confirmed its superiority over the model-centric approach in both model accuracy and speed of learning [19]. Strikingly, when asked whether to improve the code or the data in order to enhance the accuracy of the devised model in the steel inspection problem, the majority ($\approx 80\%$) of approximately 2000 voters in the live stream presentation chose improving the data.

Rather than the typical 'either/or' approach, this work supports a 'both/and' one. Though a study aiming at illustrating the relative potency of data-centric AI would benefit from experimental procedures, in which data-centric AI is tested against its model-centric counterpart. Doing this, however, is beyond the scope of our goal planning for the current study. Specifically, within the AI/ML community, we observe a kind of 'paradigm shift' in the way AI/ML models are devised. Within this constellation, more and more researchers are calling for replacing the model-centric approach with a data-centric alternative. They support their position with experimental findings from the applications they work on. That goes well when the aim is to show the potency of the new approach (i.e., the data-centric AI). Yet, in times of paradigm shifts, success stories from isolated experimental research works provide no proof supporting the supremacy of one approach over the other alternatives. Rather, one should show the supremacy of the supported approach (data-centric AI) over other alternatives (model-centric AI) through experimental results in every possible application. However, no research study can afford this.

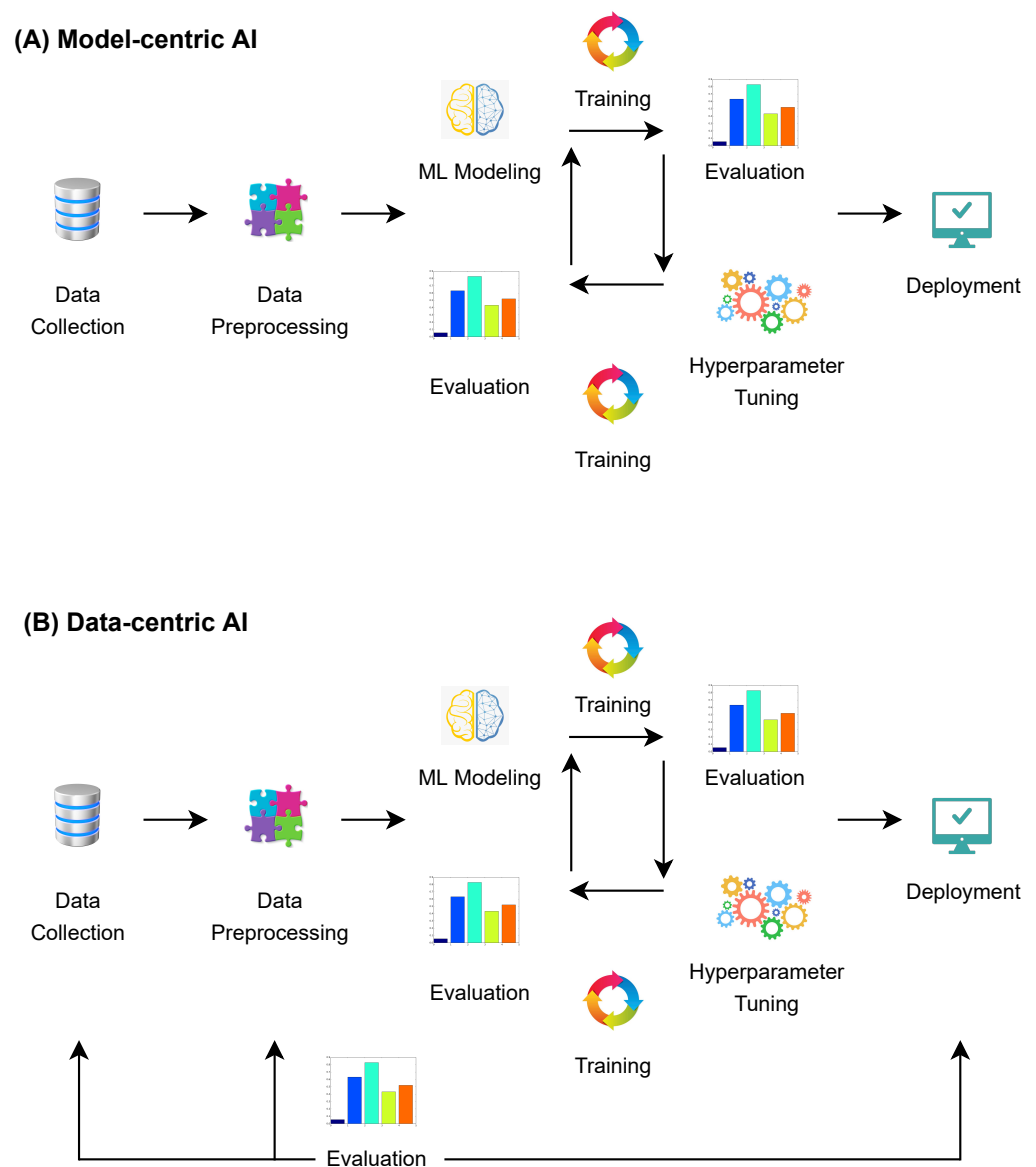


Figure 1. Cycles of model-centric AI and data-centric AI (source: [17]).

Indeed, what advocates of data-centric AI are demonstrating through their experimental findings is merely the limitation of model-centric AI. Though we share this view with them, our central argument in the current work remains that rather than championing any of the two alternatives, we opt for a ‘both/and’ approach. We argue that overcoming the limitation of model-centric AI does, indeed, require paying extra attention to the alternative data-centric AI approach. However, this should not result in reduced interest in the model-centric approach, for it provides us with the anchor necessary to assess the quality of an AI system over the course of upgrading the dataset.

Our Contribution

This paper has been significantly extended from previous work [17]. In the current work, we use a comparative analysis methodology, where data-centric AI and model-centric AI are compared and contrasted (Table 1). As elements of the compare-and-contrast approach, the paper takes Andrew Ng’s live stream presentation on 24 March 2021 [19] as frame of reference. Moreover, grounds for comparison are given through the increasing number of works by other researchers, who suggest to shift from model-centric AI to data-centric AI. The main thesis in this work is: “whereas model-centric AI suffers from

performance limitations, developing both approaches in a complimentary interplay would benefit current AI more than focusing on improving datasets only (albeit important). The main contributions of the present paper are as follows:

1. We compactly review and discuss the deep learning technique, highlighting its role in driving current AI hype (Section 2.1);
2. We connect current AI to the fields of cybersecurity and natural language inference, and through the phenomena of ‘adversarial samples’ and ‘hypothesis-only biases,’ respectively, showcase the limitations of model-centric AI in terms of algorithmic stability and robustness (Sections 3.2 and 3.3);
3. We further motivate a data-centric AI approach by elucidating the effect of the ongoing growth of the IoT, supporting our approach with the latest relevant data (Section 4);
4. We reconcile data-centric AI with model-centric AI, providing further arguments for the ‘both/and’ view instead of the evidently suboptimal alternative of the ‘either/or’ perspective (Section 6).

2. Related Work

2.1. Deep Learning: A Model-Centric Key Driver of Current AI

Historically, John McCarthy (1927–2011) was the first to coin the term ‘Artificial Intelligence’ in 1956. He defined it as the “science and engineering of making intelligent machines, especially intelligent computer programmes” [22]. Today, much of the hype around AI is attributed to a particular technique called deep learning (DL) [23] which is a modified version of artificial neural networks (ANNs) (Figure 2).

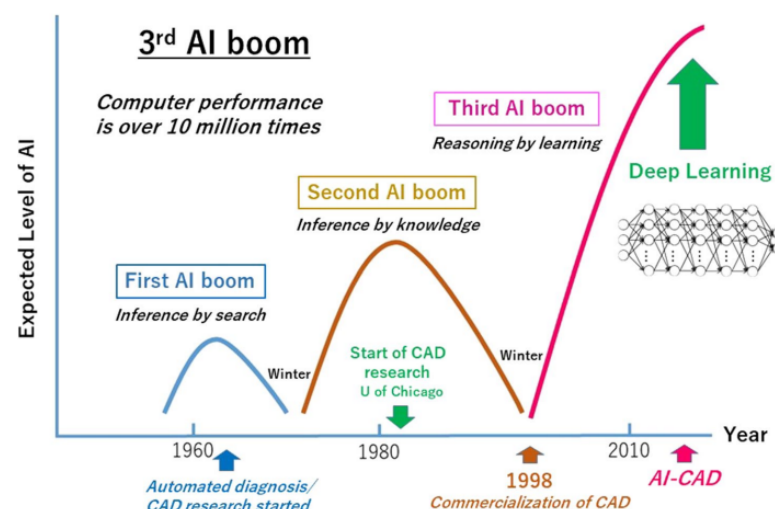


Figure 2. Current AI. The 3rd AI hype, in which previous computer-aided detection/diagnosis for medical images is augmented by deep learning, is preceded by two cycles of AI hype; source: [24].

In an accelerated pace of enhancing AI/ML prowess, DL was repeatedly reported to have accomplished human-level performance in a variety of applications. Examples range from object classification [25], over beating world-class players in Go and Poker [26,27], to detecting cancer from X-ray scans [28], and translating text across languages [29]. Although the theoretical foundations of DL were laid in the early 1940s, it was not until 2012 that researchers from AI/ML community and other related disciplines celebrated it as a technique most mimicking the human brain. This occurred when Geoffery Hinton, a British Canadian cognitive psychologist and computer scientist, along with two of his students, won the annual ImageNet contest, in its third version, with remarkably higher performance and accuracy than previous state-of-the-art algorithms [30].

2.1.1. ANN: Learning by Adjusting Weights

In a typical artificial neural network (ANN), learning occurs by adjusting the ‘weights’, w , that amplify or damp signals, x , carried by each connection between any two nodes in the ANN (Figure 3). A number of hidden layers, each constituting a module that is roughly analogous to some processing centres in the human brain, are used to backpropagate gradients. At each layer, we first compute the total input $z = \sum wx$ to each unit, which is a weighted sum of the outputs of the units in the layer before. Then, a nonlinear function $f(\cdot)$ is applied to z to obtain the output of the unit. For simplicity, bias terms were omitted. The nonlinear functions used in artificial neural networks include the rectified linear unit (ReLU) $f(z) = \max(0, z)$ along with the conventional sigmoids, such as the hyperbolic tangent, $f(z) = (\exp(z) - \exp(-z)) / (\exp(z) + \exp(-z))$ and the logistic function, $f(z) = 1 / (1 + \exp(-z))$ [31].

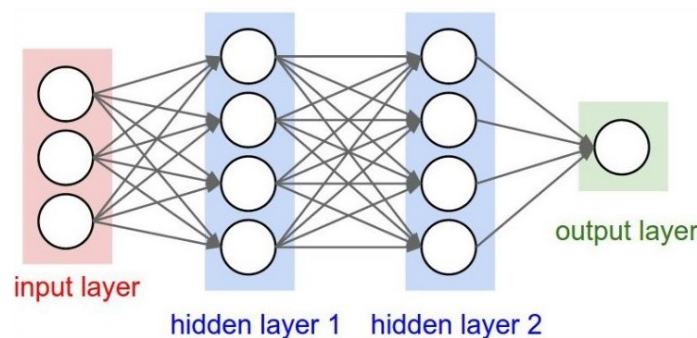


Figure 3. Artificial neural network with two hidden layers (schematic).

2.1.2. Learning over Multiple Levels of Abstraction

Compared with traditional ANNs, a DL network learns data representations over multiple levels of abstraction [31]. This is achieved by deploying significantly more layers of simple, nonlinear modules (neurons) that transform the internal representation of certain aspects of input data in one layer into another internal representation at a higher level of abstraction. A backpropagation algorithm works its way back down through the layers, fine-tuning the weights of the neural network in proportion to how much each individual weight contributes to the overall error.

Usually, the number of layers in a DL network ranges from 5 to 20 layers, hence the term ‘deep’. Today, however, neural network models in commercial applications often use over 100 layers.

Different from classical ML techniques, which require careful engineering and considerable domain expertise (by human experts) in order to extract relevant aspects of input data into a suitable internal representation (also known as feature vector), DL models can operate without the need for any explicit human intervention in the feature extraction step. Specifically, a DL algorithm learns features from data implicitly by using general-purpose procedures. Indeed, it is due to this capacity that DL was linked to the way the human brain learns to solve problems. Importantly, it has been proved that with a non-polynomial activation function, any continuous function can be approximated to any degree of accuracy [32], which makes ANN mathematically equivalent to universal computers, a result known as the universal approximation theorem.

2.2. Data: Sustainable Fuel of Current AI

Several factors explain why just in the past decade DL could evolve into what it practically became and achieved a lot of what AI researchers hoped to see in the early days of developing the technology. Most importantly, DL works well with sufficiently large volumes of data. Such data stem from various sources in the ‘Global Datasphere’ and is continuously alternating in textual, visual, and acoustic forms. The Global Datasphere consists of cloud data centres,

enterprise infrastructure such as cell towers and branch offices, and endpoints such as PCs, smartphones, sensors, social media, wearable devices, etc.

DL algorithms need data to train a model that can check statistical similarities between the currently considered instances and previously probed ones, so as to uncover hidden patterns (unsupervised learning) and make future predictions (supervised learning) about unseen instances [33]. The more granular, voluminous, and diverse the used data, the higher the performance and accuracy of the underlying algorithm. In the past, neither data nor the infrastructure required for storage and data transfer were available and mature as they are today.

According to International Data Corporation (IDC), a market intelligence company specialised in measuring created, consumed, and stored data each year, Global Datasphere would continue to expand almost exponentially (Figure 4). For example, while the volume of data created and replicated in 2010 did not exceed a lean value of 1 zettabyte (ZB), it reached a value of 64.2 ZB in 2020 [34]. It is expected that more than 180 ZB will be created in 2025 (1 ZB = 10^{12} gigabytes).

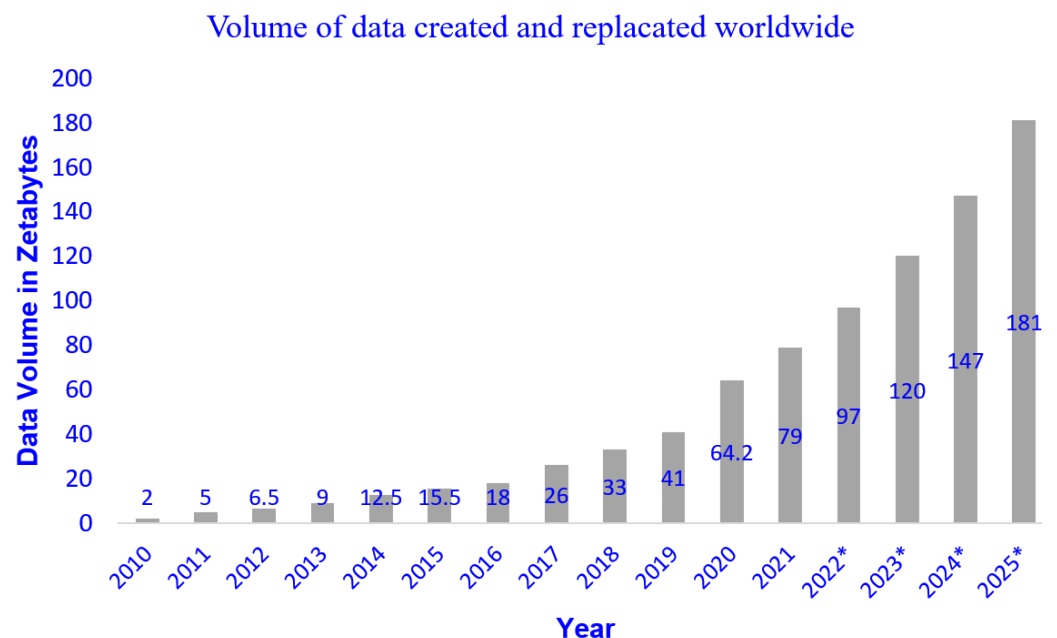


Figure 4. Annual size of data creation, consumption, and storage according to IDC. The stars denote forecast values. Data adapted from [35] with reference to [36].

3. Limitations of Model-Centric AI

Model-centric AI assumes ML solutions that mainly focus on optimising model architectures (algorithm/code) along with the underlying hyperparameters. Data within this approach are created almost only once and are kept the same over the life cycle of the AI system's development. Despite numerous success stories, model-centric AI has been placed under considerable strain. On top of it comes a narrow scope of business applicability along with the inherent vulnerability of DL/ML techniques to adversarial samples (also known as adversarial examples). The following subsections elaborate on this.

3.1. Narrow Business Applicability

The fact that training a DL algorithm requires utilising huge volumes of data causes the underlying model to suffer from an irreparable limitation [37–39]. Specifically, it works well in businesses and industries where consumer platforms with hundreds of millions of users can freely rely on generalised solutions. In such settings, a single AI system would satisfy the majority of users, while outliers would be negligibly ineffective. Examples for such businesses and industries include the advertising industry, where companies

such as Google, Baidu, Amazon, and Facebook possess mounds of data (not rarely in a standardised format), which they can deploy in creating their model-centric AI systems.

Industries such as manufacturing, agriculture, and healthcare, which require tailored solutions rather than one-size-fits-all recipes, cannot be served by standardised solutions similar to the ones provided by a single AI system. Instead, they should conceptualise their approach so as to ensure their model (algorithm) learns what it should learn by having comprehensive data that cover all important cases and are labeled consistently.

3.2. Vulnerability to Adversarial Samples

Adversarial samples are instances with small, deliberate feature perturbations that cause a DL/ML model (algorithm) to make false predictions [40]. Experts are particularly concerned about this phenomenon, for it questions the stability and robustness of DL networks with respect to noisy inputs [41]. To explain, consider a deep neural network algorithm (model) designed to perform object recognition tasks. If the algorithm can generalise well, then the network is supposed to be stable and/or robust to small perturbations to its input. This is because small disruptions in an input image would not change the categories of the objects appearing in that image. Surprisingly, though, when small, indiscernible perturbations are *deliberately* added to the test images of some objects, the network tends to change its predictions arbitrarily [41].

Attackers exploit this security gap by adding small, nonrandom, and imperceptible perturbations to benign samples—in some benchmark datasets, such as ImageNet [42], the differences between the benign samples and the adversarial ones were indistinguishable by the human eye [43]. Through this procedure, attackers aim to manipulate the DL/ML and cause it to perform false predictions [44]. Worryingly, this has severe implications with respect to adopting AI technologies in daily life applications. For instance, in the field of autonomous cars, most attacks target image-based classifiers, among which DL networks represent the most prominent approach since 2015 [45]. Here, adversaries may attack a self-driving car's system by prompting the underlying DL/ML model to ignore a 'stop' sign (e.g., due to the noise caused by placing pictures over the original stop sign or varying viewing conditions), consequently leading to a fatal crash [46]. Another critical field, in which adversarial samples may affect the comprehensive deployment of AI technologies in real-life applications is that of healthcare. Here, medical doctors and other related actors (such as insurance companies) rely on medical imaging systems to determine whether patients must undergo certain medical procedures. In order not to jeopardise patients' health and at the same time reduce the costs of treatment, concerned decision makers (medical doctors and insurance companies) try to minimise the number of unnecessary surgical procedures by deploying state-of-the-art DL systems, so as to analyse the relevant input data of patients, including, for example, dermoscopy images in dermatology, X-rays in radiology, and/or ophthalmic images, and make corresponding decisions [47,48]. Injecting adversarial samples into healthcare systems that use DL would result in an outcome that does not truly describe the underlying case and, hence, would have serious implications both for health and in terms of public and private costs.

3.3. Low Generalisation Capacity

Quite often, model-centric AI algorithms achieve high performance by utilising annotation artefacts. These are unnoticed patterns of context-free associations that come about through applying certain cognitive heuristics by human annotators [49,50]. Their effect, however, becomes obvious when the model is trained on a dataset that contains instances of such artefacts.

In natural language inference (NLI), for instance, large-scale datasets emerge from pairing a given sentence (premise) with three new ones (hypotheses), such that the given premise either entails, contradicts, or is logically neutral with respect to the generated hypothesis [49]. The hypotheses are generated by human subjects in a crowd-sourcing process. It has been observed that human annotators usually adopt cognitive strategies and

heuristics when authoring hypotheses, so as to save mental and time resources [51]. As a result, many of the NLI datasets contain annotation artefacts or biases; therefore, when used in training AI/ML algorithms, the underlying model performs surprisingly well in a dataset that contains these artefacts but fails to generalise to datasets that do not have similar artefacts. This low generalisation capacity of model-centric AI is due to the fact that the model has learned only the hypotheses (hypothesis-only bias [49,50]) but not the relationship between these hypotheses and their premises [52].

4. From Model-Centric AI to Data-Centric AI

Before AI researchers and ML practitioners began to deploy the term ‘data-centric AI’ in their scientific communications, they had years of practice in preparing datasets for training and testing their ML models. It can easily be assumed that substantial efforts were made at this stage of data preparation, in order to obtain high-quality data instances (data points). Yet, it frequently occurred that some instances were mislabelled, ambiguous, or fully irrelevant. Such data instances were considered invalid and hence had to be excluded from the datasets. However, in case the number of invalid data instances in a dataset is relatively small (in comparison with the dataset’s size), keeping these points within the dataset would have a negligible impact on the performance of an ML model, and researchers could deal with this situation without fearing further consequences [53].

There are the reasons why businesses and industries may need to focus on ensuring high-quality datasets when building their AI systems, rather than dealing with this part of the development life cycle of the system as a one-time activity.

4.1. Limited Data

Above all, the lack of sufficiently large datasets that are comprehensive both in type and volume is especially paramount. Unlike Internet companies (such as Google and Baidu), manufacturing industries are often limited in the quantity of data they possess. Usually, they run their training models on datasets with 10^2 – 10^3 relevant data points [54]. Thus, a model trained on no more than 10^3 relevant instances to detect some fault (or a rare disease) would struggle when deploying techniques that were built for hundreds of millions of data points.

4.2. Solution Customisation

There is also a need for highly customised solutions. Consider, for example, a manufacturing business with more than one product. A single one-size-fits-all AI system for fault detection across all products would not work well, as each manufactured product would require a distinctly trained ML system.

4.3. Characteristics of Data-Centric AI

Data-centric AI embraces a continuous evaluation of the devised AI model in combination with data updates (Table 1). Typically, during the production stage, a devised AI model will be trained on a dataset only once before the process of software development can be ended with the deployment of the desired functionality (Figure 1A). However, because data-centric AI assumes successive improvements in data [17], in particular, in businesses and industries that cannot afford to have millions of data points (e.g., manufacturing, agriculture, and healthcare [20]), the underlying model (represented through the implemented algorithm) will inevitably access novel instances of data points that are completely different from those encountered during the training phase. Consequently, assessing the quality of the model, too, would recur more frequently rather than occurring only once. Furthermore, with the capacity of production systems to provide timely feedback, this would result in a model (and, as a result, an AI system) being capable of recognising and, hence, reacting properly to distributional data drifts (if desired, this could serve as a prerequisite for online learning too [54]). Indeed, it is due to this faculty that the data-centric approach has an

edge over its model-centric counterpart when applied to models such as the one adopted in the steel inspection problem in [19] (Section 1).

4.4. Ongoing Growth of the IoT

According to *Statista*, a German research company specialised in market and consumer data, the number of worldwide IoT-connected devices would triple from 9.7 billion in 2020 to more than 29 billion in 2030.

A previous forecast, published in May 2020 by *Transforma Insights*, another leading research firm focused on the world of ‘digital transformation’, estimated the same figure to grow from 7.6 billion in 2019 to 24.1 billion in 2030 [55]. Later [56], however, *Transforma Insights* updated its estimate, confirming that of *Statista* (ostensibly due to taking into account the role of COVID-19 in accelerating the digital transformation, as the initial findings were published before enfolded the whole scope of the pandemic).

Not only do the analysts of both firms expect similar growth of IoT-connected devices over the next decade but they also agree on a compound annual growth rate (CAGR) of near to 10% of the IoT market, forecasting the global annual revenue of IoT to grow from USD 388 billion in 2019 to USD 1058.1 billion in 2030.

Table 1. Characteristics of model-centric AI and data-centric AI.

Category	Model-Centric AI	Data-Centric AI	References
System development lifecycle	Successive upgrade of a model (algorithm/code) with fixed volume and type of data	Continuous improvement in the quality of data with fixed model hyperparameters	[19,20]
Performance	Performs well only with large datasets	Performs well also with smaller datasets	[37–39,54]
Robustness	Susceptible to adversarial samples	Higher adversarial robustness	[50,57]
Applicability	Appropriate for testing algorithmic solutions in applications with narrow tasks	Particularly suitable for real-world scenarios	[17,21]
Generalisation	Limited capacity to generalise across datasets (due to lack of context)	May generalise well to datasets other than the one tested on	[49,51,52,58]

5. Rules and Criteria for Achieving Data-Centric AI

Adopting a data-centric approach can be achieved by combining several steps. Indeed, it is an activity for which empowering an AI system to accomplish the highest levels of performance requires incorporating interdisciplinary expertise. In the following, we consider some criteria and practical rules that help implement a data-centric approach leading to more effective AI systems.

5.1. Sufficient and Representative Data Inputs

First, datasets should be curated so as to reflect both sufficiency and representativeness of data inputs. One way to fulfil these requirements is by allowing the devised model to access as many data inputs with as much task-relevant information as required for the model to solve the task at hand. This includes cancelling out inevitable noise that is present in real life. Indeed, this is how the human brain applies ‘selective attention’ when processing visual information. Specifically, it focuses on the goal-relevant aspects of the environment while inhibiting distracting information that might be otherwise noisy for the task at hand [59].

For research teams, this can practically mean, for example, reducing the spatial complexity in an image segmentation task to the relevant image regions rather than keeping on tuning the devised model architecture, model complexity, data augmentation strategies, or related training strategies [60].

5.2. Unveiling Inherent Contexts through Textual Descriptions

Second, research teams should work towards ensuring high-quality data by revealing the inherent context within data inputs during data preparation. A vital *human* activity during the process of data preparation is that of ‘data labelling’ (also known as ‘data annotation’). Research teams usually overlook potential biases, which may have a detrimental effect on the quality of data.

Specifically, data annotators are humans with different cultural and individual backgrounds. Thus, they are susceptible to subjective judgements that could result in erroneous labels. To avert falling into the trap of biases, however, research teams may require including ‘textual descriptions’ as an intermediate step between accessing data inputs (images, videos, audio recordings, etc.) and assigning labels to them. Such textual descriptions would be some 3- to 10-word sentences in which human subjects (e.g., annotators) reflect on whatever contextual information is included in the data inputs. Although this admittedly increases the overall time required to create datasets, it helps AI engineers ensure that the collected data clearly illustrate the concepts that they need the AI to learn, resulting in AI systems that are able to learn from smaller datasets available in most industries.

5.3. Continuous Involvement of AI and Business-Domain Experts

Third, and closely related to the step above, data engineering should be performed by business-domain experts instead of AI experts. This is because AI experts have their competency mainly in representing the world in a format that enables the machine (algorithm) to learn patterns, while domain experts have comprehensive knowledge about the intricacies of a specific business use case and can hence provide a domain-relevant representation of the world [61]. Moreover, domain experts can play a role in enhancing the evaluation process by developing specific use cases that put the model into more domain-sensitive tests. This way, the ability to use AI becomes easier, and hence utilised AI system becomes more accessible to a wide range of industries.

5.4. Use of MLOps

Fourth, instead of spending time and effort on developing software, research teams and experts could reduce the maintenance cost of AI applications by using machine learning operation (MLOp) platforms. These provide much of the scaffolding software needed to facilitate the production of an AI system. As a result, the gap between proof of concept and production will massively shrink to weeks, instead of otherwise years.

There are several MLOp systems that could be applied for both data-centric and model-centric AI. For instance, one may use tools such as labelme [62] and ActiveClean [63] for data-labelling and data-cleaning activities, respectively.

In the case of model-centric AI, examples of available MLOp tools include model store systems [64], model continuous integration tools [65,66], training platforms [67], and deployment platforms [68].

6. Reconciling Data-Centric AI with Model-Centric AI

6.1. Data-Centric and Model-Centric AI Can Only Be Two Sides of One Coin

At first glance, it seems quite natural to think of tweaking the model (algorithm/code) rather than improving the data, if the goal is to enhance the system’s performance and robustness. However, with the limitations of model-centric AI solutions and the serious problems DL-based models suffer from, as described in Section 3 (or Table 1 for summary), the need to turn to data-centric approaches becomes quite understandable. We believe that this should occur while paying similar attention to the underlying model; that is, both

model-centric and data-centric approaches should be handled as two sides of the same coin, and recognising the limitations of the model-centric approach in certain businesses and industries should not lead to abandoning it altogether. This is both practical and self-evident since both the model and data affect each other interdependently.

Moreover, though shifting to data-centric AI approaches in response to the limitations of model-centric AI might seem attractive, constructing new datasets usually comes with high costs, and there is no guarantee that the new dataset would be completely free from new types of artefacts [52,69]. Indeed, due to the brain's flexibility, there is always a possibility for human annotators to apply subtle cognitive strategies during the annotation task that would limit the model's ability to generalise. So, it would be better to treat the model and data as two sides of one coin and improve them simultaneously.

6.2. Problem-Solving Requires Considering Both the How-To (Model) and the What-Is (Data)

Our intuitive understanding of problem solving puts considerable weight on the way we act upon things besides knowing their properties and facts about them. In the case of the development of AI systems, this would imply a preference towards exerting more efforts in the design and optimisation of intelligent algorithms and more vividly tweaking the underlying model hyperparameters along with the code used to implement the designed algorithms. Keeping the dataset fixed at this stage of work is even necessary since only through a fixed dataset can one have the possibility to compare models and classify them according to their performance.

On the other hand, starting the design of AI systems with a model-centric approach provides ample opportunity to accumulate experience in understanding real-world problems and potential computational solutions compared with the gains in the experience that could have resulted, had research and industry followed the track of data-centric approach. In fact, it is in the nature of things that when solving a problem, we first take action in order to produce an effect, upon which we then proceed to dig deeper to understand the properties of things around us. This is exactly what intelligent agents do to acquire experience while interacting with their environments, as postulated by reinforcement learning [70], a recently much-celebrated ML technique [26,58,71,72].

6.3. Models' Limitations Do Not Necessarily Imply the Limitation of Modelling

According to Aristotle's first-principles thinking, understanding the fundamental aspects of a problem can help provide good solutions to it [73]. In the NLI task, which was introduced in Section 3.3, the model learns, among other things, annotation artefacts. This results in high performance when tested on datasets containing such artefacts but not on ones with no such biases.

While this shows the limitations of the devised model, it does not necessarily mean that the model-centric approach itself is limited (at least in the field of NLI). Specifically, by understanding that the annotation artefacts result from ignoring the premise–hypothesis relationship, one can design models that suppress the learning of such artefacts. For instance, Belinkov et al. ([52]) proposed to input into the model both the hypothesis and entailment label so as to predict the premise instead of the typical NLI models, which learn to predict an entailment label, given a premise–hypothesis pair.

6.4. Nature Supports Learning of Models

Several strategies contribute to enhancing the robustness of DL/ML models and paving the way for a complementary interplay between data-centric and model-centric AI. First and foremost is aligning algorithms to nature with its intuitive and robust mechanisms, in particular, learning and 'problem-solving' mechanisms as performed by the human brain. One such mechanism is *common sense*. It describes the ability of humans (as well as animals) to learn *world models* by accumulating enormous amounts of *contextual* knowledge about how the world works and what causal relationships regulate the co-occurrence and succession of events, and use this knowledge in predicting future outcomes [74–76].

Typically, common-sense knowledge is acquired through unsupervised observation in a very limited number of task-irrelevant interactions with the learning environment.

Current AI and ML systems do not exhibit such a faculty. For instance, while a human car driver can draw on her/his intuitive knowledge of physics to predict the bad consequence of driving too fast, an AI system of an autonomous vehicle requires thousands of reinforcement learning trials to acquire such knowledge. A typical observation in this regard is that learning occurs at a meta-level; that is, when humans solve problems that they face for the first time, they typically rely on the affordances the brain provides, sometimes fleetingly, through interaction with the environment (see [74,77]).

6.5. Together, Data-Centric AI and Model-Centric AI Provide More Robustness and Security

The fact that data could be freely submitted into a running DL/ML algorithm (as can be experienced in cyber attacks with adversarial samples) without directly attacking the model itself or the IT infrastructure that hosts it provides a supportive argument for the significance of developing data-centric and model-centric AI in a complementary approach. Initially, it was believed that the tendency of a DL network to make false predictions in response to the inputs formed by applying small but intentionally worst-case perturbations to samples from the dataset was due to the extreme nonlinearity of DL networks, combined with insufficient model averaging and the insufficient regularisation of the purely supervised learning problem. Later, however, it was understood that adversarial samples are manifestations of the rather linear nature of neural networks in general. They are neither random artefacts of the normal variability that accumulates through different runs of propagation learning nor do they arise from overfitting or due to incomplete model training [48]. Rather, adversarial samples are robust to random noise and as such could transfer from one neural network model to another, despite differences in the number of layers with different model hyperparameters and most importantly, when trained on different sets of samples [41].

This implies that rather than being solely a matter of the datasets used to train a deep neural network model based on backpropagation, it is the way how the network is structurally connected to data distribution that matters, suggesting combining both data-centric and model-centric approaches when aiming at enhancing network robustness.

7. Conclusions

In response to the increasing interest among the AI/ML community in favoring data-centric AI over model-centric AI, which suggests a paradigm shift in the way AI/ML models are devised, we analysed the pros and cons of both approaches, only to end up reconciling data-centric AI with model-centric AI. On the one hand, we share other researchers' view that more attention should be paid to the data-centric AI approach. However, we believe that this should not result in reduced interest in model-centric AI, as promoting current AI requires considering both approaches in a complementary interplay. In particular, Industry 4.0, which promises to automate the interaction of cyber-physical objects over the IoT, would benefit from such an interplay, as both data and AI are hallmarks of Industry 4.0 technologies.

Funding: This research received no external funding.

Institutional Review Board Statement: Not applicable

Informed Consent Statement: Not applicable

Data Availability Statement: Not applicable

Conflicts of Interest: The author declares no conflict of interest.

References

1. Hamid, O.H.; Smith, N.L.; Barzanji, A. Automation, per se, is not job elimination: How artificial intelligence forwards cooperative human–machine coexistence. In Proceedings of the 2017 IEEE 15th International Conference on Industrial Informatics (INDIN), Emden, Germany, 24–26 July 2017; IEEE: Piscataway, NJ, USA, 2017; pp. 899–904.
2. Bhatt, S. The Big Fight: RPA vs. Traditional Automation. 2018. Available online: <https://www.botreetechnologies.com/blog/the-big-fight-robotic-process-automation-vs-traditional-automation> (accessed on 3 January 2023).
3. Zhang, X.; Wen, Z. Thoughts on the development of artificial intelligence combined with RPA. In *Journal of Physics: Conference Series*; IOP Publishing: Bristol, UK, 2021; Volume 1883, p. 012151.
4. Khan, Z.A.; Imran, S.A.; Akre, V.; Shahzad, M.; Ahmed, S.; Khan, A.; Rajan, A. Contemporary cutting edge applications of IoT (Internet of Things) in industries. In Proceedings of the 2020 Seventh International Conference on Information Technology Trends (ITT), Abu Dhabi, United Arab Emirates, 25–26 November 2020; IEEE: Piscataway, NJ, USA, 2020; pp. 30–35.
5. Thames, L.; Schaefer, D. Industry 4.0: An overview of key benefits, technologies, and challenges. In *Cybersecurity for Industry 4.0*; Springer: Berlin/Heidelberg, Germany, 2017; pp. 1–33.
6. Sadiku, M.N.; Musa, S.M.; Momoh, O.D. Cloud computing: Opportunities and challenges. *IEEE Potentials* **2014**, *33*, 34–36. [CrossRef]
7. Yu, W.; Liang, F.; He, X.; Hatcher, W.G.; Lu, C.; Lin, J.; Yang, X. A survey on the edge computing for the Internet of Things. *IEEE Access* **2017**, *6*, 6900–6919.
8. Yuan, L.; He, Q.; Tan, S.; Li, B.; Yu, J.; Chen, F.; Jin, H.; Yang, Y. Coopedge: A decentralized blockchain-based platform for cooperative edge computing. In Proceedings of the Web Conference 2021, Ljubljana, Slovenia, 19–23 April 2021; pp. 2245–2257.
9. Boubin, J.; Banerjee, A.; Yun, J.; Qi, H.; Fang, Y.; Chang, S.; Srinivasan, K.; Ramnath, R.; Arora, A. *PROWESS: An Open Testbed for Programmable Wireless Edge Systems*; Association for Computing Machinery: New York, NY, USA, 2022. [CrossRef]
10. Durão, L.F.; Haag, S.; Anderl, R.; Schützer, K.; Zancul, E. Digital twin requirements in the context of industry 4.0. In Proceedings of the IFIP International Conference on Product Lifecycle Management, Turin, Italy, 2–4 July 2018; Springer: New York, NY, USA, 2018; pp. 204–214.
11. Mateev, M. Industry 4.0 and the digital twin for building industry. *Industry 4.0* **2020**, *5*, 29–32.
12. Kotsiopoulos, T.; Sarigiannidis, P.; Ioannidis, D.; Tzovaras, D. Machine learning and deep learning in smart manufacturing: The smart grid paradigm. *Comput. Sci. Rev.* **2021**, *40*, 100341. [CrossRef]
13. Pareek, P.; Shankar, P.; Pathak, M.P.; Sakariya, M.N. Predicting Music Popularity Using Machine Learning Algorithm and Music Metrics Available in Spotify. *Cent. Dev. Econ. Stud.* **2022**, *9*, 10–19.
14. Murschetz, P.C.; Prandner, D. ‘Datafying’broadcasting: Exploring the role of big data and its implications for competing in a big data-driven tv ecosystem. In *Competitiveness in Emerging Markets*; Springer: New York, NY, USA, 2018; pp. 55–71.
15. Moriuchi, E. The Role of Technology in Social Media. In *Cross-Cultural Social Media Marketing: Bridging across Cultural Differences*; Emerald Publishing Limited: Bingley, West Yorkshire, UK, 2021.
16. Smith, N.; Teerawanit, J.; Hamid, O.H. AI-Driven Automation in a Human-Centered Cyber World. In Proceedings of the 2018 IEEE International Conference on Systems, Man, and Cybernetics (SMC), Miyazaki, Japan, 7–10 October 2018; IEEE: Piscataway, NJ, USA, 2018; pp. 3255–3260.
17. Hamid, O.H. From Model-Centric to Data-Centric AI: A Paradigm Shift or Rather a Complementary Approach? In Proceedings of the 2022 8th International Conference on Information Technology Trends (ITT), Dubai, United Arab Emirates, 25–26 May 2022; IEEE: Piscataway, NJ, USA, 2022; pp. 196–199.
18. Eyuboglu, S.; Karlaš, B.; Ré, C.; Zhang, C.; Zou, J. dcbench: A benchmark for data-centric AI systems. In Proceedings of the Sixth Workshop on Data Management for End-To-End Machine Learning, Philadelphia, PA, USA, 12 June 2022; pp. 1–4.
19. Ng, A. A Chat with Andrew on MLOps: From Model-centric to Data-centric AI. 2021. Available online: <https://www.youtube.com/watch?v=06-AZXmwHjo> (accessed on 3 January 2023).
20. Ng, A. AI Doesn’t Have to Be Too Complicated or Expensive for Your Business. 2021. Available online: <https://hbr.org/2021/07/ai-doesnt-have-to-be-too-complicated-or-expensive-for-your-business> (accessed on 3 January 2023)
21. Mazumder, M.; Banbury, C.; Yao, X.; Karlaš, B.; Rojas, W.G.; Damos, S.; Damos, G.; He, L.; Kiela, D.; Jurado, D.; et al. DataPerf: Benchmarks for Data-Centric AI Development. *arXiv* **2022**, arXiv:2207.10062.
22. McCarthy, J. What is artificial intelligence? 1998. Available online: <https://cse.unl.edu/~choueiry/S09-476-876/Documents/whatisai.pdf> (accessed on 3 January 2023).
23. Horvatić, D.; Lipic, T. Human-Centric AI: The Symbiosis of Human and Artificial Intelligence. *Entropy* **2021**, *23*, 332. [CrossRef] [PubMed]
24. Fujita, H. AI-based computer-aided diagnosis (AI-CAD): The latest review to read first. *Radiol. Phys. Technol.* **2020**, *13*, 6–19. [CrossRef] [PubMed]
25. He, K.; Zhang, X.; Ren, S.; Sun, J. Delving deep into rectifiers: Surpassing human-level performance on imagenet classification. In Proceedings of the IEEE International Conference on Computer Vision, Santiago, Chile, 7–13 December 2015; pp. 1026–1034.
26. Silver, D.; Schrittwieser, J.; Simonyan, K.; Antonoglou, I.; Huang, A.; Guez, A.; Hubert, T.; Baker, L.; Lai, M.; Bolton, A.; et al. Mastering the game of go without human knowledge. *Nature* **2017**, *550*, 354–359. [CrossRef]
27. Moravčík, M.; Schmid, M.; Burch, N.; Lisý, V.; Morrill, D.; Bard, N.; Davis, T.; Waugh, K.; Johanson, M.; Bowling, M. Deepstack: Expert-level artificial intelligence in heads-up no-limit poker. *Science* **2017**, *356*, 508–513. [CrossRef] [PubMed]

28. Rajpurkar, P.; Irvin, J.; Zhu, K.; Yang, B.; Mehta, H.; Duan, T.; Ding, D.; Bagul, A.; Langlotz, C.; Shpanskaya, K.; et al. Chexnet: Radiologist-level pneumonia detection on chest X-rays with deep learning. *arXiv* **2017**, arXiv:1711.05225.
29. Devlin, J.; Chang, M.W.; Lee, K.; Toutanova, K. Bert: Pre-training of deep bidirectional transformers for language understanding. *arXiv* **2018**, arXiv:1810.04805.
30. Krizhevsky, A.; Sutskever, I.; Hinton, G. ImageNet classification with deep convolutional neural networks. *Adv. Neural Inf. Process. Syst.* **2012**, *25*, 1097–1105. [\[CrossRef\]](#)
31. LeCun, Y.; Bengio, Y.; Hinton, G. Deep learning. *Nature* **2015**, *521*, 436–444. [\[CrossRef\]](#) [\[PubMed\]](#)
32. Leshno, M.; Lin, V.Y.; Pinkus, A.; Schocken, S. Multilayer feedforward networks with a nonpolynomial activation function can approximate any function. *Neural Netw.* **1993**, *6*, 861–867. [\[CrossRef\]](#)
33. Bender, E.M.; Gebru, T.; McMillan-Major, A.; Shmitchell, S. On the Dangers of Stochastic Parrots: Can Language Models Be Too Big? In Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency, Toronto, ON, Canada, 3–10 March 2021; pp. 610–623.
34. BusinessWire. Data Creation and Replication Will Grow at a Faster Rate Than Installed Storage Capacity, According to the IDC Global DataSphere and StorageSphere Forecasts. 2021. Available online: <https://www.businesswire.com/news/home/20210324005175/en/Data-Creation-and-Replication-Will-Grow-at-a-Faster-Rate-Than-Installed-Storage-Capacity-According-to-the-IDC-Global-DataSphere-and-StorageSphere-Forecasts> (accessed on 3 January 2023)
35. Hack, U. What Is the Real Story behind the Explosive Growth of Data? 2021. Available online: <https://www.red-gate.com/blog/database-development/whats-the-real-story-behind-the-explosive-growth-of-data> (accessed on 3 January 2023).
36. Reinsel, D.; Rydning, J.; Gantz, J.F. Worldwide Global DataSphere Forecast, 2021–2025: The World Keeps Creating More Data—Now, What Do We Do with It All? 2021. Available online: <https://www.marketresearch.com/IDC-v2477/Worldwide-Global-DataSphere-Forecast-Keeps-14315439/> (accessed on 3 January 2023).
37. Lowe, D. Machine Learning Deserves Better than This. 2021. Available online: <https://www.science.org/content/blog-post/machine-learning-deserves-better> (accessed on 3 January 2023)
38. Navarro, C.L.A.; Damen, J.A.; Takada, T.; Nijman, S.W.; Dhiman, P.; Ma, J.; Collins, G.S.; Bajpai, R.; Riley, R.D.; Moons, K.G.; et al. Risk of bias in studies on prediction models developed using supervised machine learning techniques: Systematic review. *BMJ* **2021**, *375*, n2281. [\[CrossRef\]](#) [\[PubMed\]](#)
39. Roberts, M.; Driggs, D.; Thorpe, M.; Gilbey, J.; Yeung, M.; Ursprung, S.; Aviles-Rivero, A.I.; Etmann, C.; McCague, C.; Beer, L.; et al. Common pitfalls and recommendations for using machine learning to detect and prognosticate for COVID-19 using chest radiographs and CT scans. *Nat. Mach. Intell.* **2021**, *3*, 199–217. [\[CrossRef\]](#)
40. Molnar, C. *Interpretable Machine Learning*; 2022. Available online: <https://christophm.github.io/interpretable-ml-book> (accessed on 3 January 2023).
41. Szegedy, C.; Zaremba, W.; Sutskever, I.; Bruna, J.; Erhan, D.; Goodfellow, I.; Fergus, R. Intriguing properties of neural networks. *arXiv* **2013**, arXiv:1312.6199.
42. Deng, J.; Dong, W.; Socher, R.; Li, L.J.; Li, K.; Fei-Fei, L. Imagenet: A large-scale hierarchical image database. In Proceedings of the 2009 IEEE Conference on Computer Vision and Pattern Recognition, Miami, FL, USA, 20–25 June 2009; IEEE: Piscataway, NJ, USA, 2009; pp. 248–255.
43. Goodfellow, I.J.; Shlens, J.; Szegedy, C. Explaining and harnessing adversarial examples. *arXiv* **2014**, arXiv:1412.6572.
44. Ren, K.; Zheng, T.; Qin, Z.; Liu, X. Adversarial attacks and defenses in deep learning. *Engineering* **2020**, *6*, 346–360. [\[CrossRef\]](#)
45. Fujiyoshi, H.; Hirakawa, T.; Yamashita, T. Deep learning-based image recognition for autonomous driving. *IATSS Res.* **2019**, *43*, 244–252. [\[CrossRef\]](#)
46. Sharma, P.; Austin, D.; Liu, H. Attacks on machine learning: Adversarial examples in connected and autonomous vehicles. In Proceedings of the 2019 IEEE International Symposium on Technologies for Homeland Security (HST), Boston, MA USA, 5–6 November 2019; IEEE: Piscataway, NJ, USA, 2019; pp. 1–7.
47. Geirhos, R.; Jacobsen, J.H.; Michaelis, C.; Zemel, R.; Brendel, W.; Bethge, M.; Wichmann, F.A. Shortcut learning in deep neural networks. *Nat. Mach. Intell.* **2020**, *2*, 665–673. [\[CrossRef\]](#)
48. Finlayson, S.G.; Chung, H.W.; Kohane, I.S.; Beam, A.L. Adversarial attacks against medical deep learning systems. *arXiv* **2018**, arXiv:1804.05296.
49. Gururangan, S.; Swayamdipta, S.; Levy, O.; Schwartz, R.; Bowman, S.R.; Smith, N.A. Annotation artifacts in natural language inference data. *arXiv* **2018**, arXiv:1803.02324.
50. Poliak, A.; Naradowsky, J.; Haldar, A.; Rudinger, R.; Van Durme, B. Hypothesis only baselines in natural language inference. *arXiv* **2018**, arXiv:1805.01042.
51. Zhang, G.; Bai, B.; Zhang, J.; Bai, K.; Zhu, C.; Zhao, T. Mitigating Annotation Artifacts in Natural Language Inference Datasets to Improve Cross-dataset Generalization Ability. *arXiv* **2019**, arXiv:1909.04242.
52. Belinkov, Y.; Poliak, A.; Shieber, S.M.; Van Durme, B.; Rush, A.M. Do not take the premise for granted: Mitigating artifacts in natural language inference. *arXiv* **2019**, arXiv:1907.04380.
53. Motamedi, M.; Sakharlykh, N.; Kaldewey, T. A data-centric approach for training deep neural networks with less data. *arXiv* **2021**, arXiv:2110.03613.
54. Berscheid, D. Data-Centric Machine Learning: Making Customized ML Solutions Production-Ready. 2021. Available online: <https://dida.do/blog/data-centric-machine-learning> (accessed on 3 January 2023)

55. Morrish, J.; Hatton, M. Global IoT Market to Grow to 24.1 Billion Devices in 2030, Generating \$1.5 Trillion Annual Revenue. 2020. Available online: <https://transformainsights.com/news/iot-market-24-billion-usd15-trillion-revenue-2030> (accessed on 3 January 2023).
56. IoT Business news. Transforma Insights Makes Powerful New IoT Forecast Resource Available for All. 2021. Available online: <https://transformainsights.com/news/powerful-new-iot-forecast-tool> (accessed on 3 January 2023).
57. Ji, X.; Tian, Q.; Yang, Y.; Lin, C.; Li, Q.; Shen, C. Improving Adversarial Robustness with Data-Centric Learning. 2022. Available online: http://alisc-competition.oss-cn-shanghai.aliyuncs.com/competition_papers/20211201/rank5.pdf (accessed on 3 January 2023).
58. Hamid, O.H.; Braun, J. Reinforcement learning and attractor neural network models of associative learning. In *Computational Intelligence: Proceedings of the 9th International Joint Conference, IJCCI 2017, Funchal-Madeira, Portugal, 1–3 November 2017*; Revised Selected Papers; Springer: New York, NY, USA, 2019; pp. 327–349.
59. van Moorselaar, D.; Slagter, H.A. Inhibition in selective attention. *Ann. N. Y. Acad. Sci.* **2020**, *1464*, 204–221. [\[CrossRef\]](#)
60. Schlegl, T.; Stino, H.; Niederleithner, M.; Pollreis, A.; Schmidt-Erfurth, U.; Drexler, W.; Leitgeb, R.A.; Schmoll, T. Data-centric AI approach to improve optic nerve head segmentation and localization in OCT en face images. *arXiv* **2022**, arXiv:2208.03868.
61. Miranda, L.J. Towards data-centric machine learning: A short review. Available one line: <https://lvmiranda921.github.io/notebook/2021/07/30/data-centric-ml> (accessed on 3 January 2023).
62. Russell, B.C.; Torralba, A.; Murphy, K.P.; Freeman, W.T. LabelMe: A database and web-based tool for image annotation. *Int. J. Comput. Vis.* **2008**, *77*, 157–173. [\[CrossRef\]](#)
63. Krishnan, S.; Wang, J.; Wu, E.; Franklin, M.J.; Goldberg, K. Activeclean: Interactive data cleaning for statistical modeling. *Proc. VLDB Endow.* **2016**, *9*, 948–959. [\[CrossRef\]](#)
64. Vartak, M.; Subramanyam, H.; Lee, W.E.; Viswanathan, S.; Husnoo, S.; Madden, S.; Zaharia, M. ModelDB: A system for machine learning model management. In *Proceedings of the Workshop on Human-In-the-Loop Data Analytics*, San Francisco, CA, USA, 26 June–1 July 2016; pp. 1–3.
65. Renggli, C.; Karlaš, B.; Ding, B.; Liu, F.; Schawinski, K.; Wu, W.; Zhang, C. Continuous integration of machine learning models with ease. *ml/ci: Towards a rigorous yet practical treatment. Proc. Mach. Learn. Syst.* **2019**, *1*, 322–333.
66. Zhang, H.; Li, Y.; Huang, Y.; Wen, Y.; Yin, J.; Guan, K. Mlmodelci: An automatic cloud platform for efficient mlaas. In *Proceedings of the 28th ACM International Conference on Multimedia*, Seattle, WA, USA, 12–16 October 2020; pp. 4453–4456.
67. Jiang, Y.; Zhu, Y.; Lan, C.; Yi, B.; Cui, Y.; Guo, C. A unified architecture for accelerating distributed {DNN} training in heterogeneous {GPU/CPU} clusters. In *Proceedings of the 14th USENIX Symposium on Operating Systems Design and Implementation (OSDI 20)*, Online, 4–6 November 2020; pp. 463–479.
68. Chen, T.; Moreau, T.; Jiang, Z.; Zheng, L.; Yan, E.; Shen, H.; Cowan, M.; Wang, L.; Hu, Y.; Ceze, L.; et al. {TVM}: An automated {End-to-End} optimizing compiler for deep learning. In *Proceedings of the 13th USENIX Symposium on Operating Systems Design and Implementation (OSDI 18)*, Carlsbad, CA, USA, 8–10 October 2018; pp. 578–594.
69. Sharma, R.; Allen, J.; Bakhshandeh, O.; Mostafazadeh, N. Tackling the story ending biases in the story cloze test. In *Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics (Volume 2: Short Papers)*, Melbourne, Australia, 15–20 July 2018; pp. 752–757.
70. Sutton, R.S.; Barto, A.G. *Reinforcement Learning: An Introduction*; MIT Press: Cambridge, MA, USA, 2018.
71. Nair, A.; McGrew, B.; Andrychowicz, M.; Zaremba, W.; Abbeel, P. Overcoming exploration in reinforcement learning with demonstrations. In *Proceedings of the 2018 IEEE International Conference on Robotics and Automation (ICRA)*, Brisbane, Australia, 21–25 May 2018; IEEE: Piscataway, NJ, USA, 2018; pp. 6292–6299.
72. Moerland, T.M.; Broekens, J.; Jonker, C.M. Emotion in reinforcement learning agents and robots: A survey. *Mach. Learn.* **2018**, *107*, 443–480. [\[CrossRef\]](#)
73. Irwin, T. *Aristotle's First Principles*; Clarendon Press: Oxford, UK, 1989.
74. LeCun, Y. *A Path Towards Autonomous Machine Intelligence*, Version 0.9. 2; Available online: <http://openreview.net> (accessed on 27 June 2022).
75. Pearl, J.; Mackenzie, D. *The Book of Why: The New Science of Cause and Effect*; Penguin Random House: London, UK, 2018.
76. Schölkopf, B. Causality for machine learning. In *Probabilistic and Causal Inference: The Works of Judea Pearl*; ACM Books: New York, NY, USA, 2022; pp. 765–804.
77. Wang, J.X.; Kurth-Nelson, Z.; Kumaran, D.; Tirumala, D.; Soyer, H.; Leibo, J.Z.; Hassabis, D.; Botvinick, M. Prefrontal cortex as a meta-reinforcement learning system. *Nat. Neurosci.* **2018**, *21*, 860–868. [\[CrossRef\]](#) [\[PubMed\]](#)

Disclaimer/Publisher's Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.