

Position Paper: Minimizing Risks in Artificial Intelligence Projects

Fernando Pereira dos Santos, Bruno Miguel de Souza, and Maurício Schiezero

Venturus - Innovation & Technology

Campinas, SP, Brazil

{fernando.santos, bruno.souza, mauricio.schiezero}@venturus.org.br

Abstract—To successfully develop an Artificial Intelligence (AI) project, a deep knowledge of learning-based algorithms is essential. However, understanding the concepts of innovation, scientific methodology, and project management can be key to a smooth design. The risks in AI projects are numerous, ranging from data availability, tight budgets, and deadlines to a lack of awareness of how AI actually works among stakeholders. These setbacks are likely to jeopardize progress and may even lead to project dissolution. Therefore, accurate and early management of these elements can increase the likelihood of success. In accordance, this position paper lists the main liabilities in AI development and discusses ways to anticipate and minimize them. Our goal is to enhance the discussion on how we can effectively manage AI projects, highlighting issues that might go unnoticed during the entire project. By planning for these risks, AI projects are more likely to be successfully deployed. Our paper presents vital concepts of learning-based algorithms, as well as project management, innovation, and scientific methodology. Finally, we discuss the major risks found in the literature and the possible solutions for fluid development.

Index Terms—artificial intelligence, risk management, innovation, scientific methodology.

I. INTRODUCTION

Artificial Intelligence (AI) is currently a generic nomenclature for any application that simulates human intelligence. Terms such as Machine Learning, Deep Learning, Computer Vision, Time Series, and Natural Language Processing may come to mind for those who delve deeper into the topic. Various definitions of AI have been proposed over the years, making it difficult to establish a consensus. A more rational statement defines AI as “the study of mental faculties through the use of computational models” [1]. If we consider AI to act similarly to humans, the definition “the art of creating machines that perform functions that require intelligence when performed by people” fits better [2]. Regardless of the precise theoretical conceptualization, AI has permeated people’s daily lives in all aspects of society, from posting a video on the Internet [3] to self-driving cars [4]. However, AI goes far beyond leisure and comfort. AI is present in the quality of industrial production [5], vaccine manufacturing [6], genetic sequence matching [7], and other highly complex functionalities [8]. If well designed, AI can be a powerful resource for reducing costs and enabling scientific advances in all fields.

The PMBOK Guide [9] describes a project as a temporary effort to create a product to be launched, a service to be provided, or a unique result, serving as a strategic position.

Accordingly, **projects** are governed to achieve objectives that guide the work to be carried out. Like any segment, developing AI initiatives requires preparation to formalize the project. Each year, the development of AI applications has grown considerably [8]. This advance comes from the large number of scientific papers published in the field recently. However, do these academic studies and projects follow **Project Management** principles? It is hard to say! Most scientific papers seek to publicize successes rather than failed experiments. The same occurs in commercial and industrial applications.

The **concept of risk**, according to the Project Management panorama, can encompass everything that causes losses or damage to a project, and this risk may be present from its conception but may never occur [10]. A crucial factor in this position paper is that we do not refer to risks when intelligence systems can harm people or the environment. Risk, in our paradigm, is the likelihood of negative consequences from events that could hinder the effective development of AI solutions. Therefore, having a precise mapping of what could harm the execution of a project is an essential planning step [11], especially in scenarios where progress is entirely dependent on iterations with stakeholders and in which the development process is guided by research. AI is a very dynamic area in which the development of a new application involves a large, adaptive, and innovative process [12]. As a result, the techniques inherent to the Predictive Project Management approach are not fully suitable for this scenario. Unlike this more rigid management, the **Adaptive Project Management** approach makes it very difficult to precisely define the deadlines, costs, and metrics to be achieved, as these factors are complex in each project. Generally, AI development is governed by the performance achieved by computational models structured with specific objectives, which may create unrealistic expectations due to the non-deterministic structure common to other software. Thence, risks such as a lack of knowledge about how AI works, the absence of transparency in results, high complexity, low error tolerance, data availability, and a limited budget are part of the development cycle of this category of project [13].

Based on this challenging context, this **position paper** provides fundamental guidelines for the development of AI applications conditioned on Project Management. Our goal is not to outline or determine how teams should conduct their day-to-day work. Our goal is to elucidate the aspects that must

be observed to increase the probability of success. Considering the fundamentals of Project Management combined with AI principles, developers can pay attention to details that might otherwise go unnoticed. We also do not intend to address all aspects of Project Management, but rather to **focus on minimizing risks**. Learning-based algorithms depend on the quality of the data and their computational complexity to obtain adequate inference [14], but the guaranties that a methodology will be efficient are not anticipated. Developing AI requires technological advancement and innovation, which increases the risk of failure.

II. LITERATURE REVIEW

A. Learning-based Algorithms

Machine Learning and **Deep Learning** algorithms are distinguished by their “intelligence” derived from experience [15]. Based on this foundation, data are supplied to a mathematical formulation that, through training routines, adjusts its parameters, learning the intrinsic data distribution [16]. As a consequence, if the training set is consistent with the expected future data, the model will perform well [17]. The difficulty in developing these algorithms lies in finding a match between the model (a mathematical function adjusted by the data [18]) and the provided data itself. Theorems such as “No Free Lunch” demonstrate that there is no superior model for all scenarios [19]. Experiments and new modeling become recurrent until the best function is found – one that translates learning into generalization. The main concern at the beginning of any modeling process is the **quantity and quality of the data**. More complex models require a larger amount of data to avoid overfitting during training [20]. Conversely, simplified models may be underfitted due to a lack of parameterization necessary to learn all the nuances of the data distribution. With respect to data quality, it originates from creating a distribution that covers all possible scenarios. This consistency comes at a significant financial cost in annotations for supervised models, mainly in scenarios that require a high specialization of knowledge [21] [22]. Furthermore, highly parameterized models have longer **response-time** and require more processing [18], which may be unfeasible in scenarios with limited budgets or embedded devices. Therefore, alternatives that compress models are attractive, such as Knowledge Distillation, Network Pruning, and Network Quantization [23]. All these principles must be taken into account when defining the project objective.

B. Project Management

Figure 1 illustrates **Organizational Project Management** (OPM), which aims to ensure that the organization selects projects and / or programs that are more in line with the objectives of its strategic plan. The portfolio encompasses the organization’s value decisions and aligns strategies with the prioritization of programs and projects for value delivery. Programs control the interdependence of components, and projects enable the organization to achieve its objectives. When a program or project reaches operation, the organization

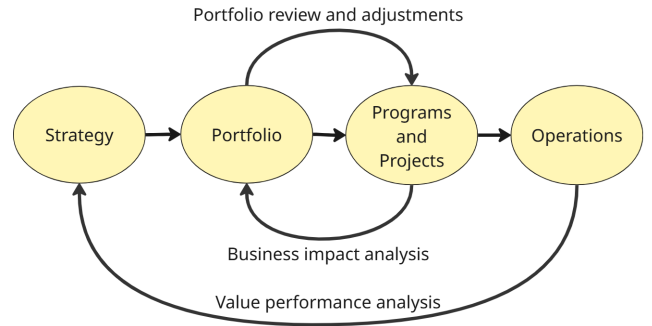


Fig. 1. Organizational Project Management. Adapted from [9].

accomplishes its business value [9]. The OPM precedes Project Management and is permanently active, being headed by directors of the organization or collaborators who have a global vision of the different sectors and how they integrate to coordinate all human resources with the highest-priority initiatives.

More specifically, in our focus on Project Management, each project has a life cycle to achieve its targets: initiation; planning; execution; monitoring and control; and closing. Following recommendations from the PMBOK Guide [9], each of these stages focuses on a specific group of processes in different areas of knowledge: **Integration Management** focuses on the processes and activities required to execute the project; **Scope Management** ensures that the project completes its functionality successfully; **Schedule Management** assures on-time completion; **Cost Management** involves estimating, financing, and cost control; **Quality Management** includes processes to control stakeholder expectations and requirements; **Resource Management** ensures the identification of what is needed to execute the project; **Communication Management** establishes that project information is properly organized; **Risk Management** includes response and implementation planning, as well as monitoring potential project losses; **Acquisition Management** encompasses processes for purchasing products or ordering services; and **Stakeholder Management** identifies people who may be impacted by the project and develops strategies for their engagement. The grid among stages and areas of knowledge is presented in Figure 2 (for in-depth knowledge of Project Management, we recommend carefully reading the PMBOK Guide [9]). For a project to be successful, all these areas must be executed correctly at every stage.

The more complex the project, the tighter the controls must be. However, AI projects do not fit adequately into the rigidity of Predictive Project Management (also called Traditional or Waterfall Management [24]), requiring the adaptability of processes due to their innovative nature and experimental orientation. In adaptive environments, also called Agile Methodology [25], with the notable Scrum [26], team engagement is crucial for determining how **integration** will be carried out. Due to uncertainty, the **scope** is often not fully established at the beginning, particularly in discovery projects. The format and the ultimate goal are then constantly refined according

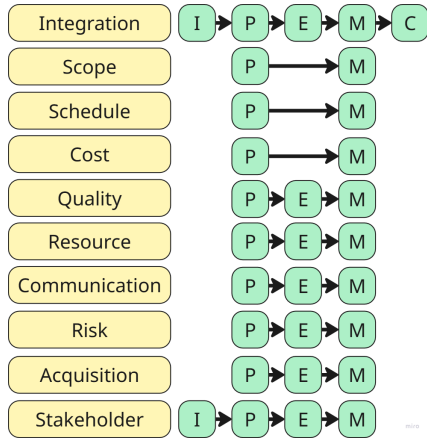


Fig. 2. Areas of knowledge for existing stages: (I) initiation; (P) planning; (E) execution; (M) monitoring and control; and (C) closing. Adapted from [9].

to current requirements [27], where tasks are prioritized and backlogs are formed [9]. Another significant change concerns the **schedule**, which differs from traditional management, where it is possible to certify the entire deadline. The Agile Methodology allows for short cycles, the analysis of results, and the adaptation of strategies. Similarly, the total **costs** of the project cannot be precisely defined at the outset and are constantly adjusted. The project continued until the desired outcome was achieved or until resources became available, especially if the process included discoveries and potential avenues for innovation. Cost is a high-impact risk, especially when developing projects for customers, as estimates can even render the initiative unfeasible [9]. For **quality** control, requirements should be constantly evaluated throughout each cycle rather than being concentrated at the end of execution. Retrospectives help validate open issues and suggest new strategies, providing insights. Human **resources** must collaborate even more to maximize focus, increase productivity, and lead to innovative solutions [9]. In addition, evolving environments require more assertive **communication** due to frequent, rapid changes. As a result, transparency improves the dynamics of information access. In our focus, adaptive projects naturally have more uncertainty [27]. Hence, **risks** must be assessed for each cycle, following essentially the same processes as traditional management, but on a smaller scale. **Acquisitions** must be conducted cautiously to avoid unnecessary purchases for future cycles. Lastly, **stakeholder** engagement must involve a high degree of interaction due to recurring changes [27]. Discussions, brainstorming, and decision-making should occur frequently, fostering a dynamic and collaborative environment [9]. Because of these characteristics, the Agile Methodology has been preferred in software development due to its improved teamwork, collaboration with stakeholders, and increased process efficiency [24] (we also endorse reading the Agile Manifesto [28]). Research shows that adaptive management offers a greater chance of success than the traditional approach [29]. Essential elements in designing innovative projects to minimize risks.

C. Innovation and Scientific Methodology

For effective AI development, in addition to Project Management, we need to pay attention to the concepts of Innovation and Scientific Methodology. **Innovation** encompasses human, technological, and economic capital, as well as competitive advantages [30] through activities that transform creative thinking into products or services beneficial to a specific audience [31]. Thus, innovation is the process of creating something unique and useful [32]. An organization that innovates through AI has a greater capacity to add value to its business; e.g., a process that reduces production costs or improves quality analysis. The significant advantage of AI transformation is its potential to radically change the innovation process, significantly improving quality and productivity in all sectors [33], including assisting Project Management [34]. Innovative projects run the risk of not being completed or even being discarded after approval [35] due to rapid changes in public demand. Accordingly, appropriate Risk Management must be maintained throughout the project.

Competitive advantages through AI innovation must be built on **Scientific Methodology**. As we saw previously, learning-based algorithms have theoretical principles that govern experimentation and the development of intelligent systems. Scientific rigor tends to minimize technological risks [8]. Considering scientific methodology, two rules stand out. First, “the game of science is never-ending” [36]. Applying this rule to AI development, we can say that there is always room for improvement, as it can lead to the project lacking a well-defined end or scope. The second rule is “once a hypothesis has been proposed and tested, and its merits have been proven, it cannot be discarded without good reason” [36]. Thus, once the algorithm is operational, it becomes obsolete only if it requires innovation (e.g., data shift) or is no longer needed. These rules fit perfectly with Occam’s Razor for algorithms: a more complex solution should only be developed if the current solution is insufficient to solve the problem [37].

III. RISK MANAGEMENT IN AI PROJECTS

Since our focus is on minimizing risks, a more detailed review of **Risk Management** is provided so that readers are aware of some strategies that can be used in their projects. Due to its innovative nature, the agile approach easily adapts to drastic changes in requirements. Aven and Krohn [38] state that **risk assessment** R is expressed by a function $R = f(s, p, c, k)$ of four variables: what can go wrong s , the probability that the risk can occur p , the degree of its severity c , and the knowledge needed to solve the problem when it arises k . Contextually, the risk assessment function is described by **several processes** in the PMBOK Guide [9] to ensure adequate planning, execution, monitoring, and controlling for when the risk actually becomes a reality.

The activities that comprise Risk Management include [9]: i) **planning and management** of how to conduct activities, considering expert opinions, data analysis, work methodologies, responsibilities, financing, deadlines, and stakeholders, among others; ii) **identification** of uncertainties found in other

TABLE I
MOST COMMON AI RISKS [13] [35] [39]–[41].

Areas	Risk
Stakeholder	understanding of AI technical dialogue little acceptance
Scope	expansion
Integration	regulations, ethics, and social responsibility fear of innovation
Quality	model instability transparency in predictions manipulation of results confidence level of the predictions in-depth knowledge of the team
Resource	data availability (quantity and quality) data governance infrastructure
Cost	suitable professional profiles innovation is expensive returns on investment value to the business

areas of knowledge and documentation of their characteristics (this activity can be carried out through brainstorming to map adverse situations along with causes, premises, and constraints); iii) **qualitative analysis** of their severity levels and prioritization, assessing probability and impact, as well as the abilities of the team to resolve the issues; iv) **quantitative analysis** to numerically calculate the effects of occurrences through simulations and other mathematical and graphical resources; v) **response planning** to develop alternatives and action strategies that involve previous experience, assignments within the team, and strategies to mitigate or prevent issues; vi) **response implementation** of the agreed plans, recording successes and failures for future projects; and vii) **monitoring** to evaluate the effectiveness of the established process. The premise is to take controlled risks to create value. These activities can be applied using the Agile Methodology in each cycle, taking advantage of previously used materials as long as the mapped risk remains in the current stage, discarding those with no probability of occurrence and identifying new ones. Knowing how to resolve a risk should it arise is crucial. Even if the project is completed successfully, research shows that many applications do not have the expected impact on their purpose or are nothing more than approved pilots [8] [35].

Developing AI requires different skills and is more complex than typical software projects [42]. Different studies report several causes of failures and challenges in the successful completion of AI applications. These studies point to a common convergence of risks that delay AI development, in which we group and allocate areas of knowledge according to our understanding of these risks (see Table I). Based on Westenberger et al. [13] and Cooper [35], stakeholders lack **an understanding of how AI actually works** and usually end up idealizing unrealistic features that cannot currently be integrated for any reason. Also, the feeling that the application does not add **value to the business** is a liability, even when starting the project [39]–[41]. Therefore, the **scope might constantly expand**, creating a never-ending project to accomplish more functionalities [13]. Additionally, **technical dialogue**, which

is far from a business approach, hinders communication between stakeholders and the development team [35] [39]. This happens because stakeholders have difficulty perceiving the technical results and require a user-friendly interface with multiple levels of experience [39]. Due to these factors, AI projects tend to have **little acceptance and a fear of innovation** at the management and user levels [40] [41]. Recurrently overlooked by developers, **regulations, ethics, and social responsibility** pose risks primarily encountered when a project is completed successfully but cannot be approved due to constraints not initially mapped [13] [41]. In addition, data privacy and security can go unnoticed [39] [40]. Quality is closely related to the theoretical knowledge of AI. Thus, **model instability**, inadequate **transparency in predictions**, and **manipulation of results** are also risks cataloged [13] [39] [41] that are better mitigated with an in-depth knowledge of the technology employed, theoretical concepts, and professional experience. Regardless, many stakeholders believe that AI applications are deterministic, as is standard software. As a consequence, the **confidence level of the predictions** must be extremely high to ensure satisfaction [13] [35] [39]. Of course, this is a valid premise in any application; however, small margins of error may be unacceptable to some people. Key resources are also a major cause of failure: the absence of **suitable professional profiles** for specialized tasks [13] [35] [40] [41] may be due to poor planning, high market demand for these professionals, or the significant financial costs of hiring them. In addition, for human resources, the scarcity of **in-depth knowledge within the team** to solve the complexity of the application is another factor that constitutes a high risk for the project [13] [39] [41]. Technically, learning-based algorithms require a lot of data, which, in many cases, are available only in limited quantities and quality. Furthermore, data may be noisy, incomplete, and contain outliers [39]. The lack of **data availability** [13] [35] [41] and its proper governance [39] are perhaps among the greatest technological risks to AI projects. Insufficient infrastructure to support big data causes drastic impacts on development and deployment [39] [41]. Considering costs, many projects often begin as speculations, attempts to improve the current process, or to achieve numerous benefits with limited financial resources. Nevertheless, as the project progresses, it becomes clear that **innovation is expensive** due to the need to hire experts and purchase equipment [13] [40] [41]. Finally, **returns on investment** [13] are regularly calculated poorly. This mistake results in costs that exceed potential future profits or savings. Then, **how can we proceed?**

IV. WHAT CAN BE DONE TO MINIMIZE AI RISKS?

A. People and Dialog

One of the biggest challenges faced by AI developers is that stakeholders lack a complete understanding of how AI works. Many of these people believe that development is similar to standard software, with deadlines, costs, and accurate results (the models are probabilistic); or that it will even replace them [8]. For that reason, disappointments and misunderstandings are common. AI must be understood as a tool that will

improve employee performance [8] [41]. To avoid this significant risk, all meetings should **clarify the entire development process** [35], describing the risks and materials (data availability, infrastructure, and expert personnel, among others). Stakeholders should **be aware that the innovative process and scientific methodology** are integral to the project, and that complex problems are naturally more difficult to solve, requiring greater participation from all interested parties. As a result, expectations are aligned, and unrealistic ideals are crossed out. An interesting approach is to present the project using the Cross Industry Standard Process (CRISP), which is based on well-defined steps and development cycles that help stakeholders visualize how the team will act [43]. Another contribution that CRISP offers is clearer dialog. Thus, **a process-based approach reduces the distance between stakeholders and developers**. Stakeholders seek real business impacts [39], and developers possess more technological than managerial knowledge. **Translating assertiveness metrics** into profits, cost reductions, or productivity is essential to the success of the application (we will discuss costs in Subsection IV-F). For users, communication must be ongoing, **focusing on the value that the tool adds** to their daily work [8]. Presenting technical visualizations, results that are difficult to interpret, or measurements that are not part of the user's work is pointless, as it only adds more difficulties [39]. It is important to **emphasize that the tool is reliable** due to its validations and a user-friendly interface (more details on transparency can be found in Subsection IV-D). A lack of understanding and difficulties in dialog lead to low acceptance of the AI application by users. Low acceptance is also caused by fear, poor transparency, high complexity, and low usefulness. Fear stems from the possibility of being replaced by tools and process changes, while poor transparency makes users distrust the results. High complexity prevents users from utilizing the tool. Low usefulness arises from the tool being built without the input of key users, who do not perceive it as beneficial when integrated into their daily lives. To minimize these risks, organizations **must prepare their employees or customers** with the appropriate expectations [44], being clear about the objectives and scope of the application; e.g., the solution will increase productivity in task *A* or AI will reduce the time spent on task *B*. Managers must support its use **by changing processes to include the tool and by providing training**.

B. Scope and Objective

Defining the scope and objectives of a project is arduous. Stakeholders always want to include additional features, while vendors promise delivery. Thence, the higher the expectations, the greater the risk of failure [35]. One of the main causes of disappointment is a shortage of understanding of user needs, which results in an unclear, ambiguous, and imprecise scope. It is also necessary to determine the technology to be used, as well as the deployment process. Lean methodologies can be used to guide the process of defining cycles, tasks, and intermediate deliverables [35]. **Detailed meetings with stakeholders and users** should take place at the beginning of

the project, at different levels of usability, to **identify current needs, desires, and problems**. Only then should the project actually begin. Additionally, **constant scope validation and objective refinement** should be discussed during execution, **presenting expected outcomes and project approval criteria**. A coherent approach to the innovative process and adaptive management is to **create a Proof of Concept (PoC)**. Conceptually, the objective of a PoC is to identify and demonstrate the feasibility of an approach [45]. Thus, we can **narrow the scope and focus on the larger problem** to assess the difficulties and requirements. Afterward, **building a Minimum Viable Product (MVP)** [46] can validate the overall scope. Smaller features kept in the backlog can be added over time, thereby strengthening the main objective.

C. Regulations and Culture

The scope and objectives must be in accordance with current AI regulations **while observing ethical standards and social responsibility**. Several organizations [47] [48] have published reports advising that AI applications should provide clarification about their predictions and internal functionality. A major social concern is biased datasets that may contain ethnic, religious, and gender issues [49]. Teams must **build a dataset that prevents social exclusion** (see Subsection IV-E), even unintentionally, and avoids legally discriminatory situations [50]. **Data privacy and security must also be observed** [39]. Breaching confidentiality leads to a lack of trust and user churn. A project may have been successfully developed, with excellent predictive quality and low computational cost; however, if social responsibility is not observed, it will be rejected.

Regarding the cultural aspect, we can distinguish two views: i) a macro-perspective on the culture of a people, which must be considered when designing the project, where social factors, such as customs and gestures, may differ from one location to another (unfamiliarity can lead to offenses and even illegalities); and ii) a micro-perspective on the organizational perception of innovation. **Adopting an innovation** comprises three stages [51]: i) initiation, in which the organization needs to recognize the necessity and be aware of the process; ii) adoption decision, which involves evaluating the innovation to approve or reject it; and iii) implementation, to actually acquire and use it. Adopting an AI tool is highly complex, requires process changes, and the solution must be customized for users [41]. Organizations should **examine AI as part of solving internal problems and seizing opportunities**, thereby adding competitive advantages [41]. An innovative culture must be integrated into people and processes. Innovative individuals are prone to experimentation, risk taking, and problem solving situations [52]. With **innovative thinking**, employees become aware of expectations and changes, and the AI adoption process becomes smoother. The development team can **encourage stakeholders and users**, especially when they are closely involved, to cultivate this skill through active participation in the project. Managers also have an important role in **clarifying the benefits of an innovative culture** and how AI can help in the execution of repetitive tasks [53].

D. Quality and Technological Knowledge

The **level of expertise of the team** is extremely critical, depending on the complexity of the application. Naturally, the more complex the project, the greater the expertise required. Therefore, theoretical knowledge and previous experience in similar projects tend to minimize predictive quality risks. Of course, each project has its own particularities, and quality assurance is not guaranteed in terms of prediction. Nonetheless, knowledge of mathematics, statistics, and technology drastically reduces the chances of failure. An extremely important step in AI projects is an **effective literature review** guided by the project's objectives. Surveys are excellent examples of gaining a broad overview of the subject by delving into the cited references to specify knowledge. This preparation anticipates the potential risk of lacking know-how. Based on this foundation, **initial results (baseline) ensure the minimum performance** that must be achieved. Hence, through scientific methodology and the innovative process, the team can find more complex solutions to improve the current model. In addition, development for commercial purposes should differ from academic research. Many academics seek small gains in assertiveness, often in exchange for high computational costs [39]. This approach is not commercially profitable. Computational costs hinder projects from their inception. Thus, small **assertiveness drawbacks can be beneficial** for an application as long as the quality remains acceptable. Assuming that the team has sufficient technical knowledge, the primary quality risk is that the model will not work when it becomes available (model generalization). As mentioned, the review of the literature contributes to minimizing this risk, as development will be based on previously established studies. **Following the principles of learning-based algorithms** also guarantees risk minimization. An alternative that should only be pursued if necessary is to **increase the complexity of the predictive model** because the computational cost and the amount of data provided for learning will increase. Furthermore, **extensive validation** sessions with diverse metrics and user trials [35] are consistent methods for enhancing the model's generalization before deployment.

One success factor that can be overlooked during scope definition is **the agreement on the approval criteria**. Accurately defining this requirement ensures that development is directed toward this milestone. However, deployment alone may not be enough to ensure success. A key to delivering quality to users is the transparency of results and usability. Critical domains, such as diagnostics, must have interpretable or explainable results [22]. Complex models, such as deep neural networks, operate as closed boxes, making it practically impossible to visualize how the output is generated. Thence, **explainability methods must be coupled** to ensure that transparency reaches the user. The absence of trust in the model drastically reduces the usability of the application [39]. One important point is how the user will access the results. Having **practical interfaces with objective and direct results** increases the user's chances of successfully using the application.

Another essential quality requirement is **the continuous monitoring of the model** [35]. Model instability can occur due to changing technologies, new data distributions, or erroneous validations during development (data quality or algorithm limitations [13]). The measurement of the quality of a deployed model is a costly endeavor. Yet, without it, the model is doomed to fail. Fortunately, teams have automated tools for this process, validating the data and its respective predictions, comparing current results with previous ones, and verifying changes in assertiveness [39]. The robustness of the model must persist after the project is completed.

E. Data Governance and Infrastructure

Having high-quality and large-scale data is a fundamental requirement for learning-based algorithms. The more complex the model, the larger the amount of data [50]. Accordingly, data availability is a significant risk in all AI projects [13] [35] [41]. Data can be structured (in table format) or unstructured, consisting of images, videos, texts, signals, or audios. Sometimes, multimodal data are used for even more complex systems. Structured data are easily manipulated and provide direct meaning to users [54]; however, these data require careful processing for completeness and correctness [55]. In contrast, unstructured data require complex models and infrastructure for storage, training, and processing. Regardless of the type of data, noise, imbalance, and bias during data collection reduce the assertiveness of the model [56]. Moreover, **continuously evaluating the quality of the dataset** is essential for success. Quality also derives from variability, especially regarding unstructured data. Obtaining images with varying angles and lighting, for instance, improves generalization. **Applying data augmentation** is feasible as long as there is no saturation in variability [57]. Another common bottleneck is data annotation, which is a tedious, repetitive, and time-consuming task. Supervised learning outperforms unsupervised learning in most cases; however, **zero-shot, few-shot, semi-supervised, and active learning techniques may be interesting alternatives** for consistent performance [58].

Assuming that the quality and size of the dataset are sufficient, handling the entire database can be problematic. Large databases require robust feature engineering, including high levels of security and privacy, appropriate storage with backup, and processing with adequate response time [39]. Therefore, **data governance must be a priority**, with clear and well-defined processes, data checks, and constant monitoring [35], especially for handling sensitive data. A well-established workflow allows developers to access and move large amounts of data efficiently, without the risk of loss, duplication, or damage. **Tools that automate pipelines are excellent options for proper data management** [41]. These actions are important during project development; besides, we cannot forget that data continues to be the target after deployment. New data will arrive for processing, possibly in large quantities and at a high frequency. In consequence, **mechanisms that frequently validate assertiveness must be implemented** to certify that the data distribution has not

changed [39]. To efficiently manage big data initiatives, the **development infrastructure must be modular** to facilitate the integration of new projects and **must have high-capacity hardware** to allow the training of complex models [41]. Routines with highly parameterized models require hardware with a large memory capacity and Graphics Processing Units (GPUs) with fast access [39]. This demand can generate competition among teams for resources, which **compels co-operation and the definition of clear priorities**. Related issues involve developing systems in one environment and then having to port them to another. This can happen due to the wide variety of devices that will use the application (embedded and mobile, for instance) and the customer's internal operating infrastructure. Often, the customer infrastructure is not prepared for the necessary computational demands, and the cost of preparing it can be prohibitive [39], and adjustments must be made to the current infrastructure. Then, to avoid inconsistent development, it is essential to **comprehend the deployment environment at the beginning** of the project so that stakeholders have estimates of costs and necessary changes. Consequently, the development team can work with models supported by the final infrastructure, and an approach that allows extensive standardization of the infrastructure is the **use of cloud environments** [39].

F. Cost and Business Value

Adopting and developing AI are expensive and require investment. For customers, a high-performance application will only have productive effects if **users are encouraged to change their behavior** [59], which requires time and the support of a culture that embraces the new. As a result, a strategic view of the benefits of AI can shape users to provide higher quality data [50] and balance expectations [41]. For developers, personnel need to constantly **study and stay up-to-date** with new techniques, architectures, and frameworks. **Providing time allocation for personal improvement of knowledge** should be on the agenda of project managers. Additionally, developers should be encouraged to **propose innovative solutions**. The absence of qualified professionals to develop the project is one of the main causes of failure [35]. Allocating a suitable professional late causes low stakeholder confidence, additional costs, and missed deadlines [13]. Thence, hiring people with this know-how is expensive, as this role requires high levels of specialization. In contrast, having a highly capable technical team may not be enough. **Teams should be composed of profiles from interdisciplinary areas**, such as business and security, working in harmony while maintaining focus. Therefore, the project manager **must foster a collaborative environment**.

Considering business aspects, managers want to know the costs and deadlines in advance (although these points in an adaptive and innovative process are merely estimates) to calculate the return on investment. However, development is often the only consideration [13], disregarding process adaptation, user training, and other factors. This leads to disappointment, lower-than-expected productivity and cost savings,

and limited profits. Accordingly, the return on investment **must be calculated by a person who understands the environment**, considering all variables, even if they are difficult to quantify. Another fundamental point is to **validate the application performance** using Key Performance Indicators (KPIs). The relevant metrics for the model must satisfy the requirements of business performance analysis. Thus, **setting management metrics**, such as productivity, the quality of activities performed, and time savings, which can be translated into monetary values or competitive advantages, is a manner to enhance the value added to the business.

V. CONCLUSIONS

This position paper offers indispensable recommendations for minimizing risks and facilitates the recognition of the difficulties encountered in AI projects grounded in the principles of learning-based algorithms, project management, innovation, and scientific methodology. Many projects fail or are canceled due to their high complexity and poor results, as well because of details that may go unnoticed in relation to people and business value. Technical knowledge is extremely important in development; however, knowing how to deal with stakeholders, precisely defining the scope and objectives, and calculating the return on investment are intrinsic factors for success. Thence, we address essential topics present in commercial development and promote perspectives that can be adopted or improved to increase triumph. We hope that this paper serves as a reminder of the factors and actions that can be used to leverage the chances of success, and as an incentive for more in-depth discussions on how we can measure the approval of an AI application beyond the accuracy of the model.

REFERENCES

- [1] E. Charniak, *Introduction to artificial intelligence*. Pearson Education India, 1985.
- [2] R. Kurzweil, R. Richter, R. Kurzweil, and M. L. Schneider, *The age of intelligent machines*. MIT press Cambridge, 1990, vol. 579.
- [3] J. Davidson, B. Liebald, J. Liu, P. Nandy, T. Van Vleet, U. Gargi, S. Gupta, Y. He, M. Lambert, B. Livingston *et al.*, "The youtube video recommendation system," in *Proceedings of the fourth ACM conference on Recommender systems*, 2010, pp. 293–296.
- [4] C. Badue, R. Guidolini, R. V. Carneiro, P. Azevedo, V. B. Cardoso, A. Forechi, L. Jesus, R. Berriel, T. M. Paixao, F. Mutz *et al.*, "Self-driving cars: A survey," *Expert systems with applications*, vol. 165, p. 113816, 2021.
- [5] E. K. Nouchi and F. P. dos Santos, "Neural network compression for visual quality control in industrial edge computing," in *Conference on Graphics, Patterns and Images (SIBGRAPI)*. SBC, 2025, pp. 309–316.
- [6] D. B. Olawade, J. Teke, O. Fapohunda, K. Weerasinghe, S. O. Usman, A. O. Ige, and A. C. David-Olawade, "Leveraging artificial intelligence in vaccine development: A narrative review," *Journal of microbiological methods*, vol. 224, p. 106998, 2024.
- [7] J. Wei, Y. Lin, X. Yao, J. Zhang, and X. Liu, "Differential privacy-based genetic matching in personalized medicine," *IEEE Transactions on Emerging Topics in Computing*, vol. 9, no. 3, pp. 1109–1125, 2020.
- [8] N. Maslej, L. Fattorini, R. Perrault, Y. Gil, V. Parli, N. Kariuki, E. Capstick, A. Reuel, E. Brynjolfsson, J. Etchemendy *et al.*, "Artificial intelligence index report 2025," *arXiv preprint arXiv:2504.07139*, 2025.
- [9] P. M. Institute, "A guide to the project management body of knowledge (pmbok guide)." Project Management Institute, 2000.
- [10] E. Rios, "Análise de riscos em projetos de teste de software," *Rio de Janeiro, RJ: ALTA BOOKS*, 2005.

- [11] M. Schieg, "Risk management in construction project management," *Journal of Business Economics and Management*, vol. 7, no. 2, pp. 77–83, 2006.
- [12] G. Chin, "Agile project management," *AMACOM, New York*, pp. 4–5, 2004.
- [13] J. Westenberger, K. Schuler, and D. Schlegel, "Failure of ai projects: understanding the critical factors," *Procedia computer science*, vol. 196, pp. 69–76, 2022.
- [14] S. Salman and X. Liu, "Overfitting mechanism and avoidance in deep neural networks," *arXiv preprint arXiv:1901.06566*, 2019.
- [15] M. Learning, "Tom mitchell," *Publisher: McGraw Hill*, p. 31, 1997.
- [16] T. F. Sterkenburg, "Statistical learning theory and occam's razor: the core argument," *Minds and Machines*, vol. 35, no. 1, p. 3, 2024.
- [17] F. J. Bargagli Stoffi, G. Cevolani, and G. Gnecco, "Simple models in complex worlds: occam's razor and statistical learning theory," *Minds and Machines*, vol. 32, no. 1, pp. 13–42, 2022.
- [18] E. S. Ortigossa, T. Gonçalves, and L. G. Nonato, "Explainable artificial intelligence (xai)—from theory to methods and applications," *IEEE Access*, vol. 12, pp. 80 799–80 846, 2024.
- [19] P. E. Hart, D. G. Stork, and R. Duda, *Pattern classification*. Wiley Hoboken, 2001.
- [20] E. Briscoe and J. Feldman, "Conceptual complexity and the bias/variance tradeoff," *Cognition*, vol. 118, no. 1, pp. 2–16, 2011.
- [21] F. P. dos Santos and M. A. Ponti, "Robust feature spaces from pre-trained deep network layers for skin lesion classification," in *2018 31st SIBGRAPI Conference on Graphics, Patterns and Images (SIBGRAPI)*. IEEE, 2018, pp. 189–196.
- [22] J. V. R. Santillo, F. Pereira dos Santos, and R. A. F. Romero, "Saliency map explainability for breast ultrasounds," in *2025 International Joint Conference on Neural Networks (IJCNN)*, 2025, pp. 1–8.
- [23] A. Lopes, F. P. dos Santos, D. de Oliveira, M. Schiezo, and H. Pedrini, "Computer vision model compression techniques for embedded systems: A survey," *Computers & Graphics*, vol. 123, p. 104015, 2024.
- [24] H. Dong, N. Dacre, D. Baxter, and S. Ceylan, "What is agile project management? developing a new definition following a systematic literature review," *Project Management Journal*, vol. 55, no. 6, pp. 668–688, 2024.
- [25] S. Ilieva, P. Ivanov, and E. Stefanova, "Analyses of an agile methodology implementation," in *Proceedings. 30th Euromicro Conference, 2004*. IEEE, 2004, pp. 326–333.
- [26] A. Srivastava, S. Bhardwaj, and S. Saraswat, "Scrum model for agile methodology," in *2017 International Conference on Computing, Communication and Automation (ICCCA)*. IEEE, 2017, pp. 864–869.
- [27] P. Marnada, T. Raharjo, B. Hardian, and A. Prasetyo, "Agile project management challenge in handling scope and change: A systematic literature review," *Procedia Computer Science*, vol. 197, pp. 290–300, 2022.
- [28] K. Beck, M. Beedle, A. Bennekum, A. Cockburn, W. Cunningham, M. Fowler, J. Grenning, J. Highsmith, A. Hunt, R. Jeffries, J. Kern, B. Marick, R. Martin, S. Mellor, K. Schwaber, J. Sutherland, and D. Thomas, "Manifesto for agile software development," 2001. [Online]. Available: <https://agilemanifesto.org/>
- [29] I. Mergel, S. Ganapati, and A. B. Whitford, "Agile: A new way of governing," *Public administration review*, vol. 81, no. 1, pp. 161–165, 2021.
- [30] N. Haefner, J. Wincent, V. Parida, and O. Gassmann, "Artificial intelligence and innovation management: A review, framework, and research agenda," *Technological Forecasting and Social Change*, vol. 162, p. 120392, 2021.
- [31] S. Bahoo, M. Cucculelli, and D. Qamar, "Artificial intelligence and corporate innovation: A review and research agenda," *Technological Forecasting and Social Change*, vol. 188, p. 122264, 2023.
- [32] R. J. Sternberg and L. A. O'Hara, "13 creativity and intelligence," *Handbook of creativity*, vol. 251, 1999.
- [33] T. F. Bresnahan and M. Trajtenberg, "General purpose technologies 'engines of growth'?" *Journal of econometrics*, vol. 65, no. 1, pp. 83–108, 1995.
- [34] I. Taboada, A. Daneshpajouh, N. Toledo, and T. De Vass, "Artificial intelligence enabled project management: a systematic literature review," *Applied Sciences*, vol. 13, no. 8, p. 5014, 2023.
- [35] R. G. Cooper, "Why ai projects fail: Lessons from new product development," *IEEE Engineering Management Review*, vol. 52, no. 4, pp. 15–21, 2024.
- [36] K. R. Popper, *A lógica da pesquisa científica*. Editora Cultrix, 2004.
- [37] P. Domingos, "The role of occam's razor in knowledge discovery," *Data mining and knowledge discovery*, vol. 3, no. 4, pp. 409–425, 1999.
- [38] T. Aven and B. S. Krohn, "A new perspective on how to understand, assess and manage risk and the unforeseen," *Reliability Engineering & System Safety*, vol. 121, pp. 1–10, 2014.
- [39] L. Baier, F. Jöhren, and S. Seebacher, "Challenges in the deployment and operation of machine learning in practice," in *ECIS*, vol. 1, 2019, pp. 1–15.
- [40] M. Bauer, C. van Dinther, and D. Kiefer, "Machine learning in sme: an empirical study on enablers and success factors," 2020.
- [41] J. Jöhnk, M. Weißert, and K. Wyrski, "Ready or not, ai comes—an interview study of organizational ai readiness factors," *Business & information systems engineering*, vol. 63, no. 1, pp. 5–20, 2021.
- [42] D. Schlegel, K. Schuler, and J. Westenberger, "Failure factors of ai projects: results from expert interviews," *International journal of information systems and project management: IJISPM*, vol. 11, no. 3, pp. 25–40, 2023.
- [43] F. Martínez-Plumed, L. Contreras-Ochando, C. Ferri, J. Hernández-Orallo, M. Kull, N. Lachiche, M. J. Ramírez-Quintana, and P. Flach, "Crisp-dm twenty years later: From data mining processes to data science trajectories," *IEEE transactions on knowledge and data engineering*, vol. 33, no. 8, pp. 3048–3061, 2019.
- [44] F. D. Davis, "A technology acceptance model for empirically testing new end-user information systems: Theory and results," Ph.D. dissertation, Massachusetts Institute of Technology, 1985.
- [45] G. R. Waissi, M. Demir, J. E. Humble, and B. Lev, "Automation of strategy using idef0—a proof of concept," *Operations Research Perspectives*, vol. 2, pp. 106–113, 2015.
- [46] V. Lenarduzzi and D. Taibi, "Mvp explained: A systematic mapping study on the definitions of minimal viable product," in *2016 42th Euromicro Conference on Software Engineering and Advanced Applications (SEAA)*. IEEE, 2016, pp. 112–119.
- [47] E. Regulation, "2016/679 of the european parliament and of the council of 27 april 2016 on the general data protection regulation," 2025. [Online]. Available: <http://data.europa.eu/eli/reg/2016/679/oj>
- [48] W. H. Organization, "Ethics and governance of artificial intelligence for health: Who guidance," 2025. [Online]. Available: <https://www.who.int/publications/i/item/9789240029200>
- [49] T. Davidson, D. Bhattacharya, and I. Weber, "Racial bias in hate speech and abusive language detection datasets," *arXiv preprint arXiv:1905.12516*, 2019.
- [50] A. Agrawal, J. Gans, and A. Goldfarb, *Prediction machines, updated and expanded: The simple economics of artificial intelligence*. Harvard Business Press, 2022.
- [51] E. M. Rogers, A. Singhal, and M. M. Quinlan, "Diffusion of innovations," in *An integrated approach to communication theory and research*. Routledge, 2014, pp. 432–448.
- [52] F. Yuan and R. W. Woodman, "Innovative behavior in the workplace: The role of performance and image outcome expectations," *Academy of management journal*, vol. 53, no. 2, pp. 323–342, 2010.
- [53] E. Brynjolfsson and A. McAfee, "The business of artificial intelligence," *Harvard business review*, vol. 7, no. 1, pp. 1–2, 2017.
- [54] Y. Lin, Z. Jun, M. Hongyan, Z. Zhongwei, and F. Zhanfang, "A method of extracting the semi-structured data implication rules," *Procedia computer science*, vol. 131, pp. 706–716, 2018.
- [55] F. Sidi, P. H. S. Panahy, L. S. Affendey, M. A. Jabar, H. Ibrahim, and A. Mustapha, "Data quality: A survey of data quality dimensions," in *2012 International Conference on Information Retrieval & Knowledge Management*. IEEE, 2012, pp. 300–304.
- [56] N. Werts and M. Adya, "Data mining in healthcare: issues and a research agenda," *AMCIS 2000 Proceedings*, p. 98, 2000.
- [57] M. A. Ponti, F. P. dos Santos, L. S. Ribeiro, and G. B. Cavallari, "Training deep networks from zero to hero: avoiding pitfalls and going beyond," in *2021 34th SIBGRAPI Conference on Graphics, Patterns and Images (SIBGRAPI)*. IEEE, 2021, pp. 9–16.
- [58] Y. Wang, Q. Yao, J. T. Kwok, and L. M. Ni, "Generalizing from a few examples: A survey on few-shot learning," *ACM computing surveys (csur)*, vol. 53, no. 3, pp. 1–34, 2020.
- [59] W. Hummer, V. Muthusamy, T. Rausch, P. Dube, K. El Maghraoui, A. Murthi, and P. Oum, "Modelops: Cloud-based lifecycle management for reliable and trusted ai," in *2019 IEEE International Conference on Cloud Engineering (IC2E)*. IEEE, 2019, pp. 113–120.