

Position Paper: Response to the 2020 European Commission's White Paper on AI

General remarks on the White Paper

The European Commission's White Paper on AI, published in February of 2020, correctly emphasizes the need for Europe to become a global leader in AI. However, while the Commission pushes for a value-based and principle-driven approach, Europe has yet to fully earn its 'right to play' in the world arena. In order to have a real say in the global development of AI, Europe must have a significant impact on both AI development and application. Otherwise, European companies risk losing competitive power even as they uphold theoretical principles. The Commission's White Paper points out two ecosystems that must be developed in order for Europe to become a global AI leader, namely, an ecosystem of excellence and an ecosystem of trust. The appliedAI Initiative and its partners welcome the Commission's ambition in this regard and the consolidated approach that emphasizes the importance of a single digital market, as well as the willingness to invest significantly in AI.

The section of the White Paper entitled "Ecosystem of Excellence" (Sec. 4) highlights measures by which Europe may be enabled to take the lead in research, skills and innovation. In this response paper, we suggest several additional measures to promote greater AI innovation and application.

The Commission's White Paper also proposes measures to foster an "Ecosystem of Trust" (Sec. 5). However, if regulation is to be emphasized, the opportunity costs of not using AI and the additional benefits that might come through the application of AI technologies should be taken into full consideration. For example, in a discussion of autonomous driving, the dramatic increase in safety compared to human drivers—not only the potential risk of an AI-induced accident—must be considered. Similar benefits have already been observed with the use of AI in China to fight the spread of the coronavirus via telemedicine and robocalls on a massive scale, as well as deployment of autonomous cleaning and delivery systems in infected areas ([South China Morning Post](#)). Moreover, the White Paper promotes a technology-specific, risk-based approach to technology that is difficult to define.

There is also an imbalance in the White Paper's discussion of these ecosystems that arises from its emphasis on regulation as being the most important factor for building trust. This assumption is not supported by the evidence when international attitudes towards technology are analyzed.

In addition, the Commission's focus on developing an ecosystem of excellence and an ecosystem of trust offers scarce consideration of speed or agility, although the opening sentence of the White Paper itself acknowledges the importance of this concept

("Artificial Intelligence is developing fast"). The appliedAI Initiative and its partners would welcome this factor being given greater prominence in all proposed measures or even being added as a third pillar—the speed of AI development being, arguably, a top priority at this stage. This is particularly critical in view of the peculiarities of European decision making: It is comparatively easy to agree on an initial joint framework, since the value of commonality is so high. These dynamics are absent, however, when existing legislation must be changed, as will be required on an ongoing basis given the ever-changing nature of AI-based systems and services. Finding methods to dynamically adjust to the advancements of AI is mandatory if we follow a regulatory approach. The EC should also be attentive to regulatory developments in other regions of the world to make sure that EU companies are not slowed down and therefore disadvantaged by new EU regulations.

Overall, the Commission's White Paper adopts a reactive and conservative perspective rather than encouraging Europe to take a more proactive role in influencing and driving the future development of AI. However, the adoption of the latter approach is essential if European companies are to effectively compete on the global stage and the EU is to truly shape the development of AI. The Commission needs to promote a positive vision of the application of AI in order to promote interest and acceptance of this powerful technology and to avoid an overly risk-focused perspective. Measures should be more target-oriented and visionary. Regulations should only be put in place when needed; agility and speed need to serve as our north star. Only then can Europe fully exercise its strengths and compete globally.

1. Ecosystem of Excellence: The ambition vs. outlined measures

The Commission's broad goals for EU involvement in AI are expressed in the Introduction to the White Paper:

“The Commission is committed to enabling scientific breakthrough, to preserving the EU’s technological leadership and to ensuring that new technologies are at the service of all Europeans—improving their lives while respecting their rights” (Introduction, p. 1).

While the goal of improving lives is certainly laudable, the appliedAI Initiative and its partners maintain that this goal can only be achieved through innovation and application and not merely on a theoretical level. In our opinion, the actions outlined in the White Paper will not be sufficient to achieve for the EU the goal of global leadership. The White Paper takes a single European perspective (“The European approach for AI aims to promote ... across the EU economy,” p. 25) but this perspective must take into account the broader range of global activities, target the application of European-trustworthy AI also outside of Europe, and anticipate the actions and reactions of other AI countries. Missing from the Commission’s White Paper is a strategy which builds on Europe’s strength while offsetting its weaknesses. In particular, more emphasis should be placed on startup and innovation activities. Moreover, many actions recommended in the White Paper may even slow Europe’s progress in comparison to other parts of the world. Currently, when compared to China or the US—or even Israel and Canada—Europe does not appear to create new global leaders in AI. It should also be noted that none of these regions strive simply for a parochial ‘regional leadership’; they all compete at the global level. It may be worthwhile to outline and monitor KPIs for “global leadership” to better focus on the most effective actions.

In addition to the actions mentioned in this chapter, the sought after ecosystem of trust should also reflect the same mindset as the ecosystem of excellence. Accordingly, regulation should encourage innovation and support the goal of global EU leadership. Excellence in regulation would entail close coordination within a common European framework. In contrast, the “forum for a regular

exchange of information and best practice identifying emerging trends, advising on standardisation activity as well as on certification” with a “cooperation of national competent authorities” as outlined on p. 24 of the White Paper describes a cumbersome, slow process that will be inadequate if the EU is to attain global leadership and avoid the fragmentation of its internal market. The appliedAI Initiative and its partners would strongly welcome a consistent, excellence-driven European approach (a “European AI Core”) built on standards, norms, and certification.

1.1. Action 1: Working with member states (p. 5)

The appliedAI Initiative and its partners welcome the Commission’s plan to invest 20 bn EUR per year in order to remain competitive globally, although it remains unclear whether this amount would be in addition to existing investments or whether it is represented to some extent by already ongoing activities. In light of Covid-19 and the challenges brought by climate change, it should be emphasized that AI helps address these challenges—indeed, we will not manage without it—and thus AI has considerable benefits for the whole of society as much as for individuals.

1.2. Action 2: Focusing the efforts of the research and innovation community (p. 6)

The Commission correctly points out that “Europe cannot afford to maintain the current fragmented landscape of centres of competence with none reaching the scale necessary to compete with the leading institutes globally.” A long term, outcome-oriented commitment must be made—most importantly to the speed of AI development—and additional effort must be put into creating centers of excellence so as to avoid member states creating a chain of repetitive, subscale activities. The Commission should clearly state that it will move forward with willing member states. This will help to prevent the development of AI from becoming delayed while the EU waits for laggards to join. It is not clear in the White Paper whether a lighthouse centre of research is to be considered a single centre or a virtual centre consisting of many

existing organizations that have joined forces. Current activities (e.g. ICT-48, ICT-26 calls) point to a networked approach. The appliedAI Initiative and its partners would welcome a well-orchestrated network. The creation of networked lighthouses and centres requires clear communication and rigorous quality assessment to support dissemination of information and knowledge transfer. An approach similar to the DARPA challenges, which includes substantial funding and has been proven to be outcome-oriented and to attain a high level of quality, might be adopted for the operation of these centres in order to maintain competitiveness. Moonshot activities (e.g. the lighthouse cluster on trust, the self-driving EU car, the European AI center against climate change) and the maintenance of close ties to the Commission might also be used to attract top experts and talent.

In regards to the test centres as well as the lighthouses, it is also not clear in the White Paper whether they are research- or application-driven. We would strongly encourage an application-driven approach.

1.3. Action 3: Advanced skills (p. 6)

The appliedAI Initiative and its partners welcome the measures described by the Commission to increase capabilities and to attract talent to Europe, although it is important that these activities are aligned with the overall goal of the Commission. The measures should encompass not only application fields but also address less attractive research directions such as liability, testing methods, or transparency.

1.4. Action 4: Focus on SMEs (p. 7)

While the appliedAI Initiative and its partners are in favor of the actions proposed by the Commission to support SMEs and the startup ecosystem, the measures outlined in the White Paper do not suffice. Strategic European Champions in the startup sector must compete globally with teams that receive significant public long-term contracts (e.g. SpaceX/NASA, Sensetime/Chinese Cities). Besides a financing pillar, which needs to be well beyond 10bn EUR, tender procedures need to be adjusted to allow for the creation of new global champions. These measures might be tied to the lighthouse centers and to a setting much like the DARPA challenges. Digital Innovation Hubs require substantial funding if they are to network and share high quality knowledge and assets so that SMEs will be adequately supported. It goes without saying that innovation is not dependent on the size of the company (start-up, SME, or large corporate) and that a level playing field should be maintained.

1.5. Action 5: Partnership with the private sector (p. 7)

The appliedAI Initiative and its partners recognize and fully support the significance of a European data strategy as an essential foundation for AI, along with the important role of international cooperation, standardization, harmonization, and mutual recognition of standards and a regulatory structure.

Today's private sector acts in a globalized world. While the appliedAI Initiative's corporate partners welcome partnering on AI with the Commission, it needs to be emphasized that AI will inevitably become a competitive factor in the world economy; therefore, a global perspective is paramount for global players. It is important not to view the topic from an isolated European standpoint that doesn't take into consideration how the rest of the world approaches AI. Of central concern is the question of

how we might maintain accelerated development (as can be seen occurring in other parts of the world) within the context of a values-based approach to AI.

As a means of directly contributing to the European AI agenda, a feedback and suggestion mechanism from the private sector could be established so as to provide application-focused input to the European Research Agenda.

1.6. Action 6: Promoting the adoption of AI by the public sector (p. 8)

The "Adopt AI" program as outlined in the White Paper is a crucial component for the uptake of AI in Europe. The appliedAI Initiative and its partners would welcome an ambitious budget to support this program.

2. Ecosystem of Trust: The ambition vs. outlined measures

In Section 5, the “Ecosystem of Trust,” the Commission states, “Given how fast AI is evolving, the regulatory framework must leave room to cater for further developments” (p. 10). The Commission here rightly acknowledges the speed of AI development. Yet, the White Paper falls short of connecting the proposed measures with this most central acknowledgement. This section of the White Paper gives the impression that trust in AI technology can be achieved through regulation and certification alone:

“A clear European regulatory framework would build trust among consumers and businesses in AI, and therefore speed up the uptake of the technology” (p. 9).

Yet, trust can also be built through technology itself, standards, or market-driven approaches. Regulation should only be applied when needed and should avoid addressing factors that are software-specific and not limited to AI.

A principles-based global framework on Data Ethics and AI could be beneficial by reflecting a broader common understanding of the relevant existing legislation. Companies might either sign a public commitment to comply with the framework or become certified. In order to avoid market fragmentation and impediments to innovation, any EU framework should reflect global principles and add more particular requirements only if necessary to prevent harm. It needs to be recognized that AI is currently being developed without an established framework for certification, with constant progress on explainability, with the industry-driven development of standards to mitigate bias, and with a massive amount of research that is occurring around the globe. The ecosystem of trust and the regulations outlined by the Commission must reflect this dynamic environment in order to preclude major developments and new applications from happening only outside of Europe. Regulation must cater to rapid changes and anticipate future developments that are not based merely on the status quo. In addition, a risk-based approach needs to make room for ongoing reassessment. National bodies must be able to constantly monitor the advances made in AI, as failure to do so would mean falling behind the more innovation-friendly environments

in China, the US, Singapore, and other countries. Given this challenge, the appliedAI Initiative and its partners would welcome the formation of a competence center within the EU that would guide and support standardization efforts, consult Member States and the industry on regulatory measures, and constantly monitor technological developments and the end-to-end effects of existing legislation.

Overall, the appliedAI Initiative and its partners welcome the guidelines for a trustworthy AI and a risk-based approach, given that there is no “one-size-fits-all” solution for the multitude of applications affected by AI. However, the ecosystem of trust outlined in the Commission’s White Paper presents a relatively defensive approach to AI that may make it quite hard if not impossible to achieve the goal of global EU leadership in practice. Therefore, before regulating, the EC should assess the impact of proposed regulation on AI innovation and growth and take an active role in guiding its development. In that spirit, we propose significant adjustments to the White Paper’s described ecosystem of trust, as summarized in the following sections:

- **General remarks on trustworthy AI (p. 9)**
- **Liability (pp. 12-15)**
- **Standardization before regulation**
- **Risk-based approach (p. 17)**
- **Specific requirements for high-risk cases (pp. 18-22)**
- **Effects on sharing and open sourcing**
- **Monitoring (pp. 23-24)**

2.1. General remarks on trustworthy AI (p. 9)

The guidelines for trustworthy AI provided by the AI High Level Expert Group of the European Commission delineate principles that should be followed within the EU. However, the application of the principles to any AI use case might be highly case-specific and not limited to AI. Therefore, it seems necessary to comment on several of the principles.

Bias, discrimination and fairness: On the one hand, bias and discrimination may or may not be challenges encountered with AI technologies. Bias is a well established concept in data science education, and thus various methods for handling bias are available. If bias is unwanted and needs to be reduced in a specific use case, bias reduction is an established step in the process of building AI applications. On the other hand, humans are known to routinely make biased decisions. AI is measured much more rigorously than human activity and can be re-engineered instantly; therefore, it may actually be better at reducing bias than human deciders could ever be. Although the data on which an algorithm is based still play a significant role, potential discrimination only occurs when the trained algorithm is applied and specific criteria are imposed on the results. Any potential regulation should take this fact into account and not create major additional obstacles for the recording and quality of training data in itself. Instead, the requirements written into regulations should be worded in such a way that any potential discrimination in the selection and use of data in training algorithms is sufficiently considered. In fact, discrimination has been covered by law for a long time, and we do not find any convincing argument in the Commission's White Paper that existing laws are insufficient. Note that whereas bias and discrimination can be handled rather straightforwardly, fairness (outside a definition of "significant influence of a random irrelevant variable on the result") is a relatively difficult concept to address—a fact that should be kept in mind when deciding whether to include any related requirements in the regulation of AI applications.

Explainability and Transparency: The AI-specific challenges of explainability (and transparency) are of an inherently technical nature. A significant amount of research is being conducted to resolve these challenges in scientific terms, and the solutions that result are highly relevant to the interests of the industry. Findings and technical solutions should be translated into sector-neutral standards without significant regulation. Following each scientific advancement, a reasonable time frame for implementation and integration into existing AI applications must be provided. Instead of new regulation, the appliedAI Initiative and its partners fully support the extension of existing transparency rules to cover alternative solutions that provide equivalent benefits to the customer (e.g. a customer's right to ask for human validation or revalidation of the correctness of an algorithm-generated result).

2.2. Liability (pp. 12-15)

The appliedAI Initiative and its partners believe that the existing technology-neutral liability regime is quite comprehensive and should be applied to AI—with clarifications as needed—before new concepts are introduced. A separate civil liability regime for AI as suggested by the EU Parliament might actually hinder innovation and be counter-productive, given that it introduces strict liability for high-risk applications in the public sphere. Additionally, the proposal suggests fault liability for non-high-risk applications beyond contractual relationships alongside the existing EU civil liability regime. The existing EU Product Liability Directive (EU 85/374) should be amended to also provide guidance on the matter of liability for embedded software, including

AI-based (self-learning) algorithms and applications. Final documentation and duty to provide information under the PLD should be equally applicable to any AI applications that can have an impact on customers and citizens, irrespective of the assumed qualification of high risk. However, any new compulsory requirements such as ex-ante testing and approval by authorities should be limited to high-risk applications only. The PLD should include ongoing monitoring and updating of obligations to the developer and deployer of AI-based products.

2.3. Standardization before regulation

The Commission rightfully points out the relevance of trust. Trust is in the interest of the market and of each active participant, as trust broadens the acceptance and application of AI technologies. Therefore, we hold that standardization and certification activities for adhering to methods and procedures are of the highest priority for the market. Regulation, however, should be applied only if it is anticipated that market forces will not achieve the principles outlined in the White Paper. Playbooks for the use of AI and interpretations of existing legislation may yield faster and more targeted results than would new regulatory activities. The Commission should prioritize clarification and guidance in regards to existing legislation before creating new legislation.

2.4. Risk-based approach (p. 17)

The appliedAI Initiative and its partners welcome the risk-based approach in principle. Yet, there are several aspects of this approach as presented in the Commission's White Paper that we wish to highlight for reconsideration in order to ensure that any potential regulation is targeted at the right use cases, provides legal certitude, and does not discourage the development and diffusion of AI.

- **Risk:** The approach outlined in the White Paper seems to define risk without considering the costs of any alternative options or the potential good of the AI solution. Even though the use of AI might involve risk, there may well be greater harm if AI is not used. For example, AI-based cancer detection may be wrong in 5% of cases, but if an average physician has a 20% probability for misdiagnosis, the AI solution might be preferred. Similarly, autonomous cars are likely to cause far fewer fatal accidents than human drivers even though the risk of an AI-caused accident cannot be entirely eliminated. In addition, as demonstrated in the fight against the spread of Covid-19, AI's speed and scale can be a tremendous advantage in saving lives. The appliedAI Initiative and its partners propose a balanced risk assessment with both negative and positive effects being considered in the classification process.
- **Two classes:** In the proposal put forward by the German Data Ethics commission, there were five levels to allow a more differentiated view of risk assessment. We would similarly welcome at least three classes (high-, medium- and low-risk), as such differentiation offers several advantages. There is some risk involved in any application; hence, some applications may require regulation. However, only a few need drastic external involvement. With the two-class system proposed by the Commission, it is possible that many relatively low-risk applications might be nonetheless classified as high risk simply because the lower risk class appears too relaxed. Thus, a system with three classes of risk assessment would not only allow for greater flexibility but may prevent the high-risk class from becoming a catch-all for an increasing number of applications that pose relatively little risk.

- **Sector-based classification:** The sector classification presented in the White Paper seems to be unnecessary. In every sector, AI can be used for safety-critical applications as well as purely supportive functions. Thus, the Commission should avoid classifying entire sectors as high-risk. The appliedAI Initiative and its partners would propose a technology and system- or application-based classification model that follows existing sectoral regulation. More precisely, the decisive criterion for a classification should not be made on a component level (the AI technology) but on a system level, taking into consideration the intended use of the AI within the larger system. This approach would recognize the varying degrees of relevance of AI technology within whole systems (e.g. backup, recommendation, and autonomous decision). Any additional regulation should be applied within existing sector-specific frameworks, which in many fields seem to be sufficient already (e.g. autonomous driving and healthcare).
- **Measurement:** The risk classification should be formulated so as to prevent legal uncertainty and allow for self-assessment. Therefore, the classification process should be precisely defined with the inclusion of white lists of exemplary cases and an explicit classification rationale for each class. The goal should be for every company (especially SMEs) to be enabled to assess risks without external assistance.
- **Probability assessment:** Regulation of AI applications must take into account the evolving risk of any particular AI application along its life-cycle. That is, the intended use of an AI application as well as the probability that a particular risk will be manifested varies throughout the lifetime of an AI application and appears very difficult (without proper technical support) to predict and to control.

2.5. Requirements for high-risk cases (pp. 18-22)

In its White Paper, the Commission describes specific requirements for high-risk cases. In general, the appliedAI Initiative and its partners propose that most of these requirements require amendment in order to avoid hindering innovation by European companies. The requirements should focus only on objectives and leave to the companies the precise ways in which these objectives are operationalized. Otherwise, the requirements may quickly become outdated, inconsistent or even lead to contradictory rules and uncertainty in the application.

Training data: There is too much emphasis in the White Paper on training data quality, thus reflecting a focus on past standard modes of supervised learning from labeled data, not on future AI technologies. Data augmentation, transfer learning, generative adversarial methods or even model-based reinforcement learning approaches will prove elusive. Besides, a high quality of training data is already in the core interest of the company. Moreover, any rules for transparency beyond the existing regulations (e.g. GDPR) might affect IPR and the trade secrets of a company and, thus, should be avoided. More important from a regulator's perspective are standards for test data and testing environments in order to assess the quality of an AI application. It must be noted, however, that due to the nature of AI systems, it is not possible to test 100% of all possible scenarios.

Data and record-keeping: The appliedAI Initiative and its partners would welcome from the Commission a clarification of the documentation and retention obligation for development documentation. A general preservation of datasets should, however, not be mandatory. On the one hand, such a require-

ment is likely to conflict with GDPR provisions requiring deletion of personal data. On the other hand, a general requirement of this sort conflicts with copyrighted datasets authorized only for short-term access (e.g. a one-year license for input data allows a company to use the trained model afterwards but not to keep the data itself). In addition, any change to the training data would make reproduction impossible. Moreover, a general preservation requirement would destroy the privacy benefits of on-device processing because it would effectively force data to be collected and stored centrally. Ultimately, the requirement also conflicts with the targets of the Green Deal, as significant resources would be consumed for the ongoing storage of data sets. Therefore, we strongly recommend that any decision about storing or deleting data (except for limited cases) should be left to the companies.

Robustness and accuracy: It should be understood and accepted that AI will make mistakes and that 100% accuracy is not possible. Due to the nature of trained models being based on historic data in an ever-changing world, no developer can possibly ensure complete accuracy during all life cycle phases. The appliedAI Initiative and its partners propose scenario-based assessments for high-risk cases following best practices in the financial industry and the validation mechanisms for automated driving.

Human oversight: AI-based systems should be considered for situations in which automation offers an improvement over human performance (e.g. split-second decisions or highly-complex situations). Thus, human oversight of AI is applicable only in limited situations and when oversight is interpreted as monitoring and ex-post reaction, not prior clearance of each decision. If AI-based systems are designed to augment the human decision-making process (e.g. providing recommendations to radiologists), then human oversight is given by design.

2.6. Effects on sharing and open sourcing

AI thrives in a vivid open-source ecosystem in which training data sets, pre-trained models, or network architectures are shared within the community. The proposals in the Commission's White Paper regarding the establishment of an ecosystem of trust and, more specifically, the requirements for high-risk cases would limit or even eliminate the open source ecosystem. Strict liability rules or requirements for data storage and documentation that fall back on the developers of open-sourced data sets would make sharing impossible. In addition, no one would be able to use pre-trained models if the original training data were not published as well (e.g. a trained generic model for basic language understanding that a developer makes available to be further specified through the user's own data; i.e. transfer learning). This result contradicts the principles that the Commission itself outlined in the White Paper as representing our European values and that serve as the targets within the ecosystem of excellence. Therefore, every regulation should be assessed against unwanted effects on sharing, open sourcing and cross-company collaboration or even be evaluated as to whether the regulation leads to improvement of these factors. Checks should be applied to the results of AI-based systems, not to the input.

2.7. Monitoring (pp. 23-24)

If new conformity assessments for high-risk AI applications in non-harmonized sectors become necessary, the appliedAI Initiative and its partners propose a two-step approach: 1) perform ex-ante self-assessment against agreed international standards, coupled with ex-post market surveillance, and 2) the EC or a supervisor monitors and evaluates the application of this framework to determine the need for modifications in the light of technological or market developments. Any ex-post testing should be proportionate to the level of risk of the AI-based application.

In general, ongoing testing and alerts throughout the entire life cycle of an AI application will prove more important than upfront testing. Such testing could include “stability over time,” “scenario-based testing,” “benchmarking against a standard proprietary test set,” and explainability tests (e.g. “feature relevance”).

3. Additional Comments

Data: Data are a resource that can be used for various applications. Therefore, it is particularly difficult to define the requirements for data whose suitability depends significantly upon the data’s use in a specific product or service. Legally binding regulated guidelines on data quality requirements would need to be described, if at all, in the context of their real-world application. As data can generally be used in or to train multiple different applications, it will be difficult to stipulate absolute data requirements beyond minimum standards for data quality, completeness, and representativeness. It is important to provide clear purpose limitations and meaningful metadata to describe data sets when making and documenting a choice for product development. However, this function is already covered by the extensive regulation of product liability, quality, safety, and reliability.

Data Sharing: The appliedAI Initiative and its partners believe that a liberal data economy fostering fair access to and free flow of data while protecting investments and trade secrets enhances innovation and thus leads to better services and products for European citizens. Nonetheless, specific value-characteristics of data need to be considered:

- Time value (e.g. real-time data is almost useless if provided late)
- Context-specific value (e.g. a search item is more valuable if the location or situation is known; when provided without context, e.g. anonymized, the value of the data is marginal)
- Explicit value (e.g. a data set with raw data and a label providing information about the data)
- Knowledge value (e.g. data is the information, as with data representing a protein structure)

Any regulation attempting to enforce data sharing should be informed by these different elements and the value creation that exists in the data collection. While public (tax-payer) data should be shared (anonymized or not anonymized) for dedicated use cases, company data need to be handled sensitively due to their mostly unknown value as well as unclear access rights. An obligation to share with other industries or competitors data that have been enhanced, enriched or aggregated (i.e. business secrets) must be avoided, as such activity would disclose core business strategies. In our view, a set of minimum requirements regarding access to and portability of data within the EU (e.g. standardization and interoperability) would help enhance the free flow of data. The Commission’s aim for the EU is global leadership, but global leaders need to keep valuable data proprietary.

Authors

Spearheaded by Dr. Andreas Liebl, Managing Director at UnternehmerTUM and appliedAI, this paper is a joint collaboration between appliedAI and many individuals from our partner organizations. Co-Authors include experts from Dell, Intel, Allianz, BMW Group, EnBW, Infineon, MunichRE and Telekom.

About appliedAI

The appliedAI Initiative, Europe's largest non-profit initiative for the application of artificial intelligence technology, aims to bring Germany into the AI age and offers its wide ecosystem of established companies, researchers, and startups neutral ground on which to learn about AI, implement the technology, and connect with each other.

**BY
UNTER
NEHMER
TUM**

 **INITIATIVE FOR
APPLIED ARTIFICIAL
INTELLIGENCE**

Contributing Partners



Supporting Partners

