

‘Human in’ or ‘Human out’ of the Loop: Impact and Things To Consider While Deploying AI

Olivier Bouchard, SAS Institute, Paris, France

ABSTRACT

AI & GenAI are now everywhere, helping us on extracting content from long documents, summarizing notes & actions from meetings, even suggesting way to code and delivering insights to us. Now the technology can even go a step further by suggesting way to code, generating outputs for scientific or clinical ops team (Protocol Generation).

AgenticAI being the NextGen, taking action on detailed tasks and interconnected in a business processes that can be complex. In this thought-leadership talk, we wanted to have an open discussion on "where the human should stand" in such AI-driven process in our so regulated industry.

From "Human-out of the loop" to "Human on the loop" what are the things to consider to ensure people adoption of AI at scale. "Human on the loop" : where should we put the cursor to improve processes.

INTRODUCTION

The book "Nexus" written by Yuval Harari explores a future where humans and AI are interconnected, raising profound questions about the necessity of human involvement in decision-making processes. Harari's vision of a symbiotic relationship between humans and AI challenges the traditional boundaries of human autonomy and machine intelligence.

On the other hand, Thomas Senderovitz, Senior VP Datascience at NovoNordisk, argued in an interview dated January 2025, that in some cases, human involvement may not be necessary. For Senderovitz, AI can make more accurate and unbiased decisions.

This dichotomy between the philosopher/writer Harari and the business leader Senderovitz sets the stage for a deeper exploration of how engaged humans should be in the AI loop, particularly in critical sectors such as pharmaceuticals & Clinical Trial Submissions. This is a 2025 vision led by current understanding on Technology and how regulators aim to set some guardrails to have successful AI Adoption.

THE AI DREAM

Artificial Intelligence (AI) has been announced as a transformative technology with the potential to revolutionize various industries. From healthcare to finance, AI promises to enhance efficiency, improve decision-making, and create new opportunities. However, the journey from the initial promise to practical implementation is fraught with challenges and potential pitfalls.

THE UNICORN: AI'S POTENTIAL

AI's potential is often likened to a "unicorn" – a rare and valuable entity that can bring about significant positive change. With little investment it is likely expected that the "unicorn" will transform any process.

In the healthcare sector, AI can assist in diagnosing diseases, personalizing treatment plans, and predicting patient outcomes. For instance, AI models can analyze vast amounts of medical data to identify patterns and trends that may not be apparent to human clinicians. In the financial industry, AI can enhance fraud detection, optimize trading strategies, and improve customer service. AI algorithms can analyze transaction data in real-time to detect suspicious activities, thereby preventing fraud before it occurs ¹. Additionally, AI-powered chatbots can provide personalized customer support, improving the overall customer experience.



Generative AI, a subset of AI (through Large Language models – LLMs), has shown remarkable capabilities in creating new content, such as text, images, and music. This technology has the potential to drive innovation across various fields. For example, in drug discovery, generative AI can design new molecules with desired properties, accelerating the development of new medications.

In Life science Industry, despite many attempts, there is not yet one single blockbuster that has been discovered by AI Technologies. But LLMs have shown real value to automate some tasks such as summarizing texts (protocols, scientific publication...) or writing documents (CSR for Medical Teams)

THE REALITY: CHALLENGES AND RISKS

Despite its promise, AI is not without its challenges and risks. One significant concern is the potential for bias in AI models. Studies have shown that large language models (LLMs) can exhibit biases, such as recommending higher interest rates for certain demographic groups. This bias can lead to unfair and discriminatory outcomes, undermining the trust in AI systems. This is the known example of some HR bots or AI that were selecting applicants depending on race or gender.

Another challenge is the regulatory landscape for AI. Different regions have varying regulations and standards for AI,

making it difficult for organizations to navigate the compliance requirements. For instance, the EU AI Act imposes strict requirements on high-risk AI systems, such as those used in healthcare and finance. Organizations must ensure that their AI systems comply with these regulations to avoid legal and financial repercussions.

THE FAILURE: WHEN AI FALLS SHORT

AI can also fail to deliver on its promises due to technical limitations and implementation challenges. For example, AI models require large amounts of high-quality data to function effectively. In some cases, the available data may be incomplete, biased, or of poor quality, leading to inaccurate or unreliable AI predictions.

Additionally, the complexity of AI systems can make them difficult to understand and interpret. This lack of transparency, often referred to as the "black box" problem, can limit the adoption of AI in critical applications where explainability is essential. For instance, in healthcare, clinicians need to understand the reasoning behind AI-generated diagnoses to trust and act on them.

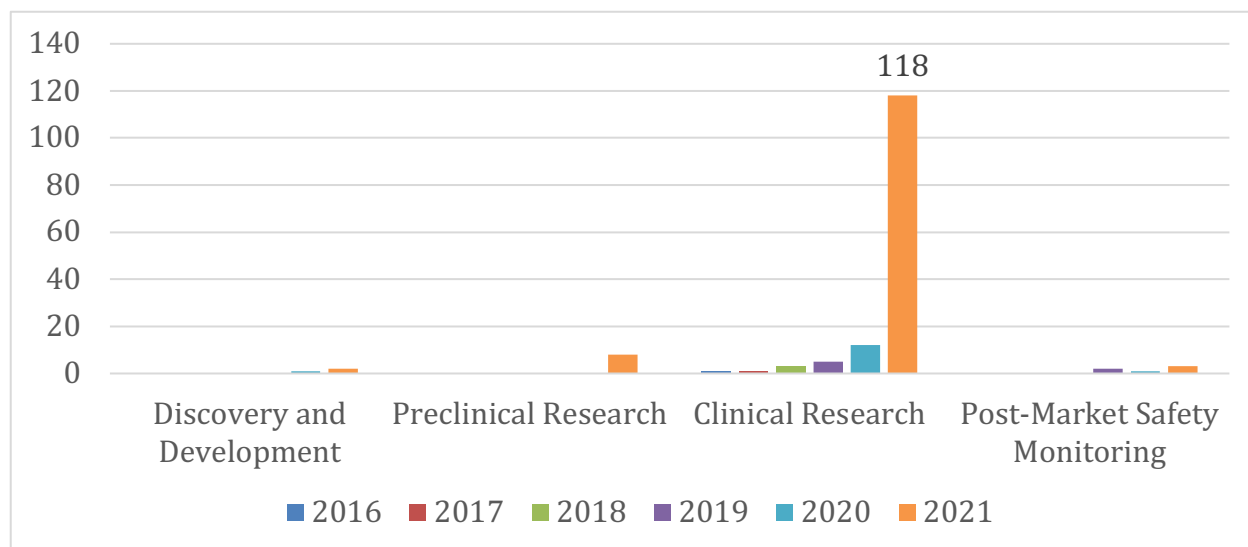
So likewise a "Unicorn", the promise of AI is immense, with the potential to drive innovation and improve outcomes across various industries. However, realizing this potential requires addressing the challenges and risks associated with AI. Organizations must invest in responsible AI practices, including bias mitigation, regulatory compliance, and transparency, to harness the full benefits of AI while minimizing its drawbacks. By doing so, AI can transition from a "unicorn" to a practical and reliable tool that enhances human capabilities and drives progress.

The Promise of AI Adoption in Pharma

Artificial Intelligence (AI) has the potential to revolutionize the pharmaceutical industry by enhancing drug discovery, development, and regulatory processes. The adoption of AI in pharma promises to improve efficiency, reduce costs, and accelerate the delivery of new treatments to patients. However, this journey is not without its challenges and risks.

IMPACT ON DRUG DEVELOPMENT

AI has already made significant strides in the pharmaceutical industry, impacting over 100 million people. One of the most notable areas of AI adoption is in regulatory submissions to the FDA. From 2016 to 2021, there has been a quite relative increase in the number of submissions that included AI/ML components, rising from just one submission in 2016 to over 200 in 2023. The most common use of AI in these submissions is in clinical research, with 118 submissions in 2021 alone. This demonstrates the growing reliance on AI to streamline and enhance various stages of drug development.



Source: [Landscape Analysis of the Application of Artificial Intelligence and Machine Learning in Regulatory Submissions for Drug Development from 2016 to 2021](#)

AI IN LIFESCIENCES VALUE CHAIN

AI is being utilized across the entire life sciences value chain, from discovery and development to post-market safety monitoring. For example, Novo Nordisk has integrated AI into their processes to optimize drug development and improve patient outcomes. AI algorithms can analyze vast amounts of data to identify potential drug candidates, predict their efficacy, and monitor their safety throughout clinical trials. This not only accelerates the drug development process but also increases the likelihood of success by identifying promising candidates early on.

Drug discovery being one of the main costs factor on delivering new drugs on the market, Pharma companies are definitely at AI to shorten drug discovery timelines. Patent duration being still the same, any help to shorten time will have financial benefits to Pharma. Getting new innovative drugs on the market is definitely one expectation from Pharma; but highlighting efficiency gains, decreasing errors, manual tasks by GenAI is looked with high scrutiny by Pharma;

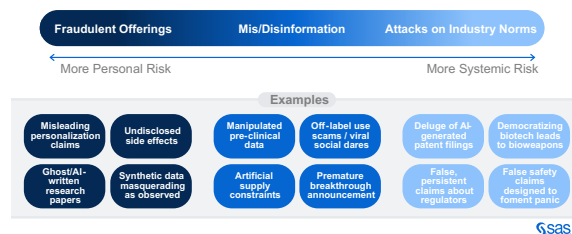
For respect of the PHUSE audience our paper and presentation will stick to Drug Development process, even if AI will definitely bring value on past market activities (better and more precise Healthcare professionals reach and information), optimized supply chain avoiding shortage and identifying fraud, and manufacturing optimizing batch and decreasing energy.

AI GOVERNANCE AND RISKS

As AI adoption in pharma continues to grow, so does the need for robust AI governance. AI governance involves establishing frameworks and guidelines to ensure that AI systems are used responsibly and ethically. This includes addressing issues such as bias, transparency, and accountability. For instance, the presentation highlights various AI risks threatening life sciences, including fraudulent offerings, misinformation, and attacks on industry norms. These risks underscore the importance of implementing strong governance practices to mitigate potential harm and build trust in AI systems.

AI Risks Threatening Life Sciences

Strategic considerations for proactive planning



STRATEGIC CONSIDERATIONS FOR PROACTIVE PLANNING

To successfully adopt AI in the pharmaceutical industry, organizations must engage in proactive planning. This involves identifying and addressing potential risks, establishing clear governance frameworks, and continuously monitoring and evaluating AI systems. By doing so, organizations can ensure that AI is used effectively and responsibly, maximizing its benefits while minimizing its drawbacks. AI governance is much more than model development eliminating bias, but need strategic considerations in areas such as data quality, regulatory compliance, and stakeholder engagement.

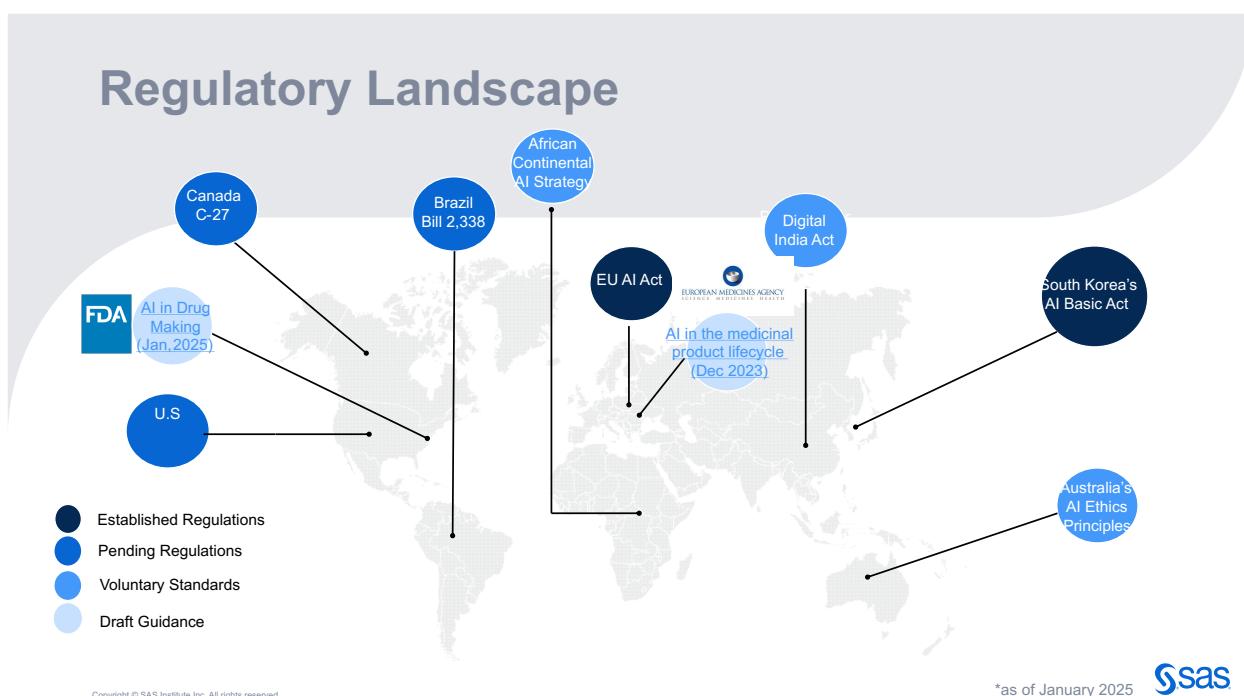
The promise of AI adoption in the pharmaceutical industry is immense, with the potential to transform drug development and improve patient outcomes. However, realizing this promise requires careful planning, robust governance, and a commitment to ethical and responsible AI practices. By addressing these challenges, the pharmaceutical industry can harness the full potential of AI to drive innovation and deliver life-saving treatments to patients more efficiently and effectively.

AI Regulation – Now and Impacts

The regulatory landscape for Artificial Intelligence (AI) is evolving rapidly as governments and organizations recognize the need to ensure the responsible and ethical use of AI technologies. The current AI regulation varies across regions, with different countries implementing their own frameworks and guidelines to address the unique challenges posed by AI.

So Governments are stepping up on AI—with regulation and investment. In 2024, U.S. federal agencies introduced 59 AI-related regulations—more than double the number in 2023—and issued by twice as many agencies. Globally, legislative mentions of AI rose 21.3% across 75 countries since 2023, marking a ninefold increase since 2016. Alongside growing attention, governments are investing at scale: Canada pledged \$2.4 billion, China launched a \$47.5 billion semiconductor fund, France committed €109 billion, India pledged \$1.25 billion, and Saudi Arabia's Project Transcendence represents a \$100 billion initiative (per hai 2025 report)

It's not a matter of if AI is regulated, but when, so it's critical that businesses start implementing trustworthy AI practices now.



EU AI ACT

One of the most comprehensive regulatory frameworks for AI is the EU AI Act. This legislation categorizes AI systems into **four risk levels**: unacceptable risk, high risk, limited risk, and minimal risk. High-risk AI systems, such as those used in healthcare, finance, and critical infrastructure, are subject to stringent requirements, including transparency, accountability, and human oversight. The EU AI Act aims to ensure that AI systems are safe, trustworthy, and respect fundamental rights.

US AI REGULATION

In the United States, AI regulation is more fragmented, with various federal and state agencies issuing guidelines and policies. The National Institute of Standards and Technology (NIST) has developed a framework for managing AI risks, focusing on principles such as fairness, transparency, and accountability. Additionally, the Federal Trade Commission (FTC) has issued guidelines on AI and machine learning, emphasizing the importance of preventing bias and ensuring consumer protection.

OTHER REGIONS

Other regions, such as Canada, Brazil, Australia, South Korea, and India, are also developing their own AI regulatory frameworks. These regulations often focus on similar principles, including transparency, accountability, and the prevention of bias. For example, Canada's Directive on Automated Decision-Making requires federal institutions to assess the impact of AI systems on individuals and implement measures to mitigate risks.

Human on the Loop Concept

Even with advancements in machine learning, deep learning, and generative AI, there is still a clear distinction between what humans and machines do well. Humans use common sense, intuition, creativity, empathy, and versatility. Machines take on tasks that humans could perform if we had all of the time in the world and didn't get fatigued: processing large data sets – not only consuming, but learning from massive amount of data, performing complex calculations, and automating tasks which can be performed without human intervention or assistance.

While the advancement in generative AI suggests that machines will be able to take over the world, let's remember how machines generate their output. In large language models, when given a prompt, machines grab the most probable next word, line of code, image, etc.. Some marvel at the LLMs creativity. But the creativity occurred when humans put information into the dataset. A machine doesn't understand what it means to have a broken heart but will finish the sentence "When he left me, he broke my..." with "heart" because millions of people have written that.

A Large Language model is operating without common sense and true intuition. Because this approach is effectively guessing, LLMs only assemble and present what humans have previously rendered.

The concept of "human on the loop" refers to a collaborative approach where humans and AI systems work together, with humans overseeing and intervening in the AI decision-making process when necessary. This approach aims to combine the strengths of both humans and machines to achieve better outcomes, particularly in critical sectors such as pharmaceuticals and clinical trial submissions.

HUMAN ON THE LOOP: ENSURING PEOPLE ADOPTION OF AI

The transition from "human out of the loop" to "human on the loop" involves several considerations to ensure people adoption of AI at scale. The key is to find the right balance between human oversight and machine autonomy. This balance is crucial for improving processes and ensuring that AI systems are trusted and accepted by users.

Many people have asked if AI is going to take away their jobs. Of course, this is complex to predict the future, but what is seen is that AI is being deployed to **complement the work of humans**.

- Making sense of patterns and trends surfaced from large data sets
- Interpreting results from complex calculations
- Assessing impacts of simulated scenarios
- Handling exceptional cases apart from those that can be automated

BENEFITS OF HUMAN ON THE LOOP

1. **Enhanced Decision-Making:** By combining human intuition and expertise with AI's data processing capabilities, decision-making can be significantly improved. Humans can provide context and judgment that AI systems may lack, leading to more accurate and reliable outcomes.
2. **Bias Mitigation:** Human oversight can help identify and mitigate biases in AI models. While AI systems can process large amounts of data, they may inadvertently learn and perpetuate biases present in the data. Human intervention can help ensure that decisions are fair and unbiased.
3. **Transparency and Trust:** Human involvement in the AI loop can enhance transparency and build trust in AI systems. Users are more likely to trust AI systems when they know that humans are overseeing and validating the decisions made by the AI.
4. **Ethical Considerations:** Human oversight is essential for addressing ethical considerations in AI decision-making. Humans can ensure that AI systems align with ethical standards and societal values, preventing potential harm and ensuring that AI is used responsibly.

CHALLENGES OF HUMAN ON THE LOOP

1. **Scalability:** One of the main challenges of the "human on the loop" approach is scalability. As AI systems become more complex and widespread, it may be difficult to ensure that humans can effectively oversee and intervene in all AI decisions.
2. **Training and Expertise:** Effective human oversight requires training and expertise in both the domain and the AI systems being used. Organizations must invest in training programs to ensure that their employees have the necessary skills to oversee AI systems.
3. **Balancing Autonomy and Oversight:** Finding the right balance between machine autonomy and human oversight can be challenging. Too much human intervention can slow down processes and reduce the efficiency gains offered by AI, while too little oversight can lead to errors and biases going unchecked.

SAS VISION ON AI REGULATION FRAMEWORK

SAS, a leader in analytics and AI, has developed its own vision for an AI regulation framework. This vision emphasizes the importance of responsible innovation, ethical AI practices, and robust governance to ensure that AI technologies are used for the benefit of society.

RESPONSIBLE INNOVATION

SAS believes that responsible innovation begins with responsible innovators. This means that organizations must prioritize ethical considerations and ensure that their AI systems are designed and deployed in a manner that respects human rights and promotes fairness. SAS emphasizes the need for transparency, accountability, and human oversight in AI systems to build trust and ensure that AI technologies are used responsibly.

AI GOVERNANCE

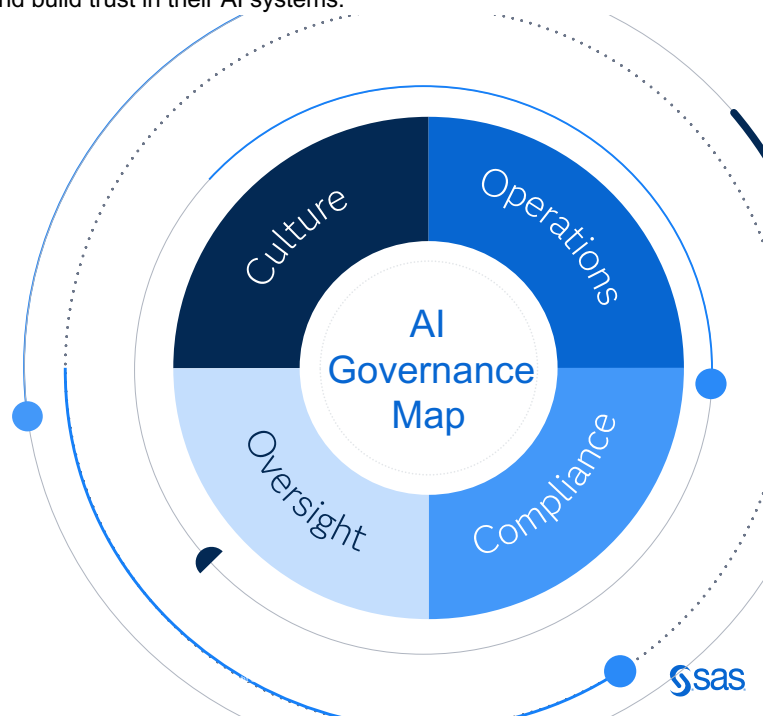
AI governance is a critical component of the SAS vision for AI regulation. This involves establishing clear frameworks and guidelines to ensure that AI systems are used ethically and responsibly. SAS advocates for the implementation of robust governance practices, including bias mitigation, transparency, and accountability. By doing so, organizations can build trust in their AI systems and ensure that they are used for the benefit of society.

STRATEGIC CONSIDERATIONS

SAS emphasizes the need for strategic considerations in the development and deployment of AI systems. This includes addressing issues such as data quality, regulatory compliance, and stakeholder engagement. By proactively planning for these challenges, organizations can ensure that their AI systems are used effectively and responsibly. SAS also highlights the importance of continuous monitoring and evaluation of AI systems to identify and address potential risks. The SAS vision for AI regulation framework underscores the importance of responsible innovation, ethical AI practices, and robust governance. By prioritizing these principles, organizations can harness the full potential of AI while ensuring that it is used for the benefit of society. The SAS vision provides a comprehensive framework for organizations to navigate the complex regulatory landscape and build trust in their AI systems.

AI Governance Map

Your guide to Trustworthy AI



Conclusion

The "human on the loop" concept emphasizes the importance of human oversight in AI decision-making processes. By combining the strengths of both humans and machines, this approach aims to achieve better outcomes, enhance transparency, and build trust in AI systems. However, it also presents challenges related to scalability, training, and finding the right balance between autonomy and oversight.

As AI continues to evolve, the collaboration between humans and machines will become increasingly important. Organizations must invest in training and governance frameworks to ensure that AI systems are used responsibly and ethically. By addressing the challenges and leveraging the benefits of the "human on the loop" approach, organizations can harness the full potential of AI while ensuring that it is used for the benefit of society.

The journey from "human out of the loop" to "human on the loop" is a critical step in the responsible adoption of AI. By prioritizing human oversight and ethical considerations, organizations can build trust in AI systems and ensure that they are used to enhance human capabilities and drive progress. The future of AI lies in the successful collaboration between humans and machines, where each complements the other's strengths to achieve better outcomes.

ACKNOWLEDGEMENT

RECOMMENDED READING

CONTACT INFORMATION

Your comments and questions are valued and encouraged. Contact the author at:

Author Name: Olivier Bouchard

Company: SAS

Address: Rue de la Branche, 77000 Evry Gregy Sur Yerres, FRANCE

Work Phone: +33663471356

Email: olivier.bouchard@sas.com

Website: <https://www.sas.com/>

SAS and all other SAS Institute Inc. product or service names are registered trademarks or trademarks of SAS

Institute Inc. in the USA and other countries. ® indicates USA registration.

Other brands and product names are trademarks of their respective companies.