

Proceedings of the 35th European Safety and Reliability & the 33rd Society for Risk Analysis Europe Conference
 Edited by Eirik Bjorheim Abrahamsen, Terje Aven, Frederic Boudier, Roger Flage, Marja Ylönen
 ©2025 ESREL SRA-E 2025 Organizers. Published by Research Publishing, Singapore.
 doi: 10.3850/978-981-94-3281-3_ESREL-SRA-E2025-P2957-cd

Risk Management and Uncertainty of Artificial Intelligence in a High Hazard Industry

Elisabeth Lootz

Norwegian Ocean Safety Authority (HAVTIL). Norway. E-mail: Elisabeth.Lootz@havtil.no

Linn Iren Vestly Bergh

Norwegian Ocean Safety Authority (HAVTIL). Norway. E-mail: Linn.Iren.Bergh@havtil.no

Bjørnar Heide

Norwegian Ocean Safety Authority (HAVTIL). Norway. E-mail: Bjornar.Heide@havtil.no

The Norwegian Ocean Safety Authority (HAVTIL) has observed rapid changes involving Artificial Intelligence (AI). Currently we have limited knowledge of possible consequences, and more uncertainties related to the development, deployment and maintenance of AI compared to other technologies. Deployed in high-hazard contexts such as the petroleum industry, AI can introduce new safety risks that must be mitigated to ensure prudent operations. Meticulous risk management throughout the AI's lifecycle, understanding new uncertainties, will be essential for both providers and deployers. The possibilities AI represents are vast, and optimism among the players is tremendous. However, the speed of development of AI solutions, prediction risk, unexpected or malicious use, new discipline experts unfamiliar with the potential of major hazards offshore can make it challenging to develop sufficient risk understanding and necessity of reliable AI predictions in a high hazard industry. We assert that the concept of uncertainty in our regulations will be more essential in risk management of safety to provide knowledge-based decisions to ensure safe operations when new technologies such as Artificial Intelligence (AI) is introduced.

Keywords: Artificial Intelligence, Risk Management, Uncertainty.

1. Introduction

Artificial Intelligence (AI) represents a resource and can help to reduce safety risks. Technology developments are fundamental prerequisites for the continuous improvement in the petroleum industry, emphasizing the potential for new technology to contribute to increased efficiency and safety (Bergh and Teigen, 2024b; Markussen et al 2024). This assumes however that the companies understand and actively follow up the potential hazards when new technology is developed and adopted.

HAVTIL's main risk issue for 2025 is "AI is also a risk factor", which underscores the importance of reliable and safe AI in a high-hazard industry.

In 2015 the Norwegian petroleum regulation clarified that risk should be understood as the consequence of the activities, with associated uncertainty. HAVTIL also published a

Memorandum on the risk concept in 2016 and a Memorandum on risk management in 2018.

This clarification is in line with international standards, such as ISO 31000: 2009, and is not specific to Norway. However, in practice we sometimes observe that a narrow risk concept is used, with a strong focus on probabilities and consequences. This was the background for the need for clarification in the regulation in 2015.

In this paper we explore the following issues:

How can the risk concept from the 2015 petroleum regulation provide a good foundation for managing new digital technologies such as AI risk?

What role does the increased emphasis on uncertainties and knowledge strength of risk assessments play?

2. Method

This paper is based on analysis of data and information on risk management, risk identification, risk mitigation and formal qualification of new digital technologies such as AI-solutions, from follow up activities since 2018. It includes findings from several studies initiated by HAVTIL to gain insight and new knowledge on digitalization and AI-risks such as Ernsten et al Safetec (2021) and Markussen et al. DNV (2024). We have also gathered information by participating in conferences, collaboration forums such as The Safety Forum Norway and International Regulators Forum (IRF), where HAVTIL has headed the work on digitalization since 2020. We participated in a tripartite collaboration on a report on working environment and safety concerns related to the introduction of AI and other new technologies (Safety Forum 2022). We have closely followed the research discourse on AI, automation, Human-AI collaboration and safety risks. Research and studies have been utilized to prepare for risk-based audits and dialogue meetings within the industry. In addition to studies, observations in audits from 2018 and onwards, and information summaries from series of meetings with deployers and providers in the same time span, form the basis of findings presented in this paper.

3. The concept of risk

New technologies can help to improve safety and efficiency. However, we sometimes observe that a narrow risk concept is used, with a strong focus on probabilities and consequences. Such a risk concept rests on assumptions such as relatively stable and well-known processes, where incident reports can serve as a good basis for risk management.

An example is a simplified view sometimes taken of the annual HAVTIL reports Trends in risk level in the petroleum activity (RNNP 1999-) which at its core measures the number and severity of annual precursor events, to inform about the risk level of the Norwegian petroleum industry. The RNNP process has been used since 1999 to measure how the level of risk is developing in Norway's oil and gas industry.

There is a long tradition of such simplifications in the Norwegian petroleum industry. Simplifications can be useful, but we clarified the risk concept to avoid oversimplifications. Even in the Risk Level Project, the clarified concept is used because the precursor events do not give enough information about the uncertainties and possible surprises of future rare catastrophic events in the industry (Vinnem et al 2006).

Furthermore, the clarified risk concept seems especially valuable in risk management when new technologies such as AI are introduced in industries such as the petroleum activities, where the potential consequences can be large. HAVTIL's Memorandum on risk management states that before decisions are taken, issues relating to health, safety and the environment (HSE) must be adequately identified and risk assessed. The framework regulations § 15 requires a sound health, safety and environment (HMS) culture where knowledge of, involvement in, commitment to and engagement with safety must be a core value and shape decisions in every part of risk management. Management responsibilities and decisions will be key elements at all levels of the business to develop a sound HSE-culture. A core principle in risk management entails that a high degree of uncertainty or great potential consequences calls for a cautious approach.

3.1. Examinations of risk management

HAVTIL enforces a technology-neutral, performance- and risk-based regulatory regime. The two perhaps most central conclusions in the Engen report (2013) evaluating the Norwegian Petroleum regime are that the Norwegian regime has proven to be robust in the face of significant technological and structural changes. And given this, they recommended continuing regulations that are mainly functional and linked to industry standards.

These conclusions are based on the fact that the regime has built-in mechanisms for handling new technology – such as digital technology – especially when the industry has very limited experience and involves major changes in organization and cooperation models. This is also highly relevant for AI.

These mechanisms include that the operator must be able to document that the same level is achieved using new solutions such as AI; qualification and use of new technology and new methods, and risk management and the precautionary principle, where the recognized good professional practice has not been determined.

The facilities regulations § 9 on qualification explicitly states; “Where the petroleum activities entail use of new technology or new methods, criteria shall be drawn up for development, testing and use so that the requirements for health, safety and the environment are fulfilled. The criteria shall be representative of the relevant conditions of use, and the technology or methods shall be adapted to already accepted solutions. The qualification or testing shall demonstrate that applicable requirements can be fulfilled using the relevant new technology or methods”. This regulation refers to DNV-RP-A203 that categorizes new technology according to levels of uncertainty.

3.2. Regulations on AI

Initial reviews from the authorities indicate that the HAVTIL regulations are applicable concerning the follow-up of AI solutions as they stand today. They contain relevant basic requirements for responsible operations, risk management, and qualification of new technology before deployment, which are important for a holistic and responsible development and use of AI solutions.^a The newly adopted AI Regulation (EU AI Act) is also founded on a risk-based approach, categorizing systems into risk categories (unacceptable risk, high risk, limited risk, and minimal/no risk). The prerequisites for compliance of a system are dependent upon the degree of risk it carries. The majority of stipulations within the AI regulation are directed towards systems deemed high-risk. Classification of an AI system as high-risk under the EU's Regulation hinges on several factors, such as its potential impact on health, safety, and

fundamental rights. In the offshore industry and energy sector, we expect that utilization of AI will be considered a «high-risk application of AI» as energy production, transportation, and distribution fall under the AI Regulation definition of critical infrastructure.

Three key points sum up the HAVTIL's expectations, based on regulations and research of the industry's increased use of AI in their operations; 1) The companies must assess vulnerability and risk from an integrated perspective which includes human, technological and organizational (HTO) factors; 2) Each company must take ownership of and manage the risk related to the implementation of new systems and technological solutions; and 3) Involving and training employees is crucial for promoting expertise and risk understanding. (Bergh and Teigen 2024).

As these types of technologies develop, there is an increased need for standards and best practices in response to the challenges and opportunities associated with their growing use of AI in industries and society at large (Markussen et al 2024).

4. Risk management: Learning from incidents

The overall historic development in the Risk Level Project on the Norwegian continental shelf, as mentioned above, related to several predefined major hazard scenarios such as well control incidents, hydrocarbon leaks and structural incidents, has shown a solid positive trend from 1991 to 2023. The underlying positive trend in the major hazard indicators indicates that the industry has progressively improved its management of technical, operational and organizational factors that affect risk (HAVTIL RNNP 2023). Risk-informed decision-making, aimed at reducing risk through increased knowledge, has been central to risk management practices. Reported near-misses with major-accident potential have been reported and investigated both by the players^b themselves and HAVTIL to gain insight into the causes of these incidents. Players are actively learning from

^a However, the regulations may lack references to norms and standards that can provide sufficient guidance for the use of AI. Therefore, it will be important for us to follow up the industry's work on standardization.

^b Players are responsible operating companies and contractors with duty to comply with regulations in the Norwegian petroleum activities.

incidents to prevent similar occurrences in the future. This is achieved among other things by increasing maintenance, improving work processes, updating operational procedures, using scenarios from investigated incidents for training purposes to raise awareness of risks related to personal injuries and major hazards, as well as training scenarios in emergency preparedness exercises. Studies based on investigation reports, supplemented by survey data from Risk level in the Norwegian petroleum industry (RNNP 2010; 2011; 2014 and 2022) and interviews with experts have provided increased in-depth knowledge and insight into patterns of causes and possible risk-reducing measures. The findings from these studies have been communicated to different stakeholders, both operating companies, contractors and technology providers that have the responsibility, but also skills and resources to mitigate identified risks. Scientific papers have been published addressing and elaborating on specific risk topics found in these studies such as on HSE-culture and structural integrity (Jensen et al 2015) and on hydrocarbon leaks caused by inadequate human-machine interfaces in the process plants (Lootz et al 2012). The knowledge has provided insights into patterns of causes, and therefore also information on how to mitigate risks. HAVTIL has systematically utilized these studies in our follow-up activities of the players, requesting risk-reducing measures. There is reason to argue that this systematic learning from incidents within the companies, and on an industry level, has significantly contributed to a more knowledge-based risk management and a positive effect on the overall risk level.

5. AI in a high hazard industry

Compared to other technologies, we have less knowledge of the consequences of introducing AI into our high-hazard industry. We have less knowledge of, and experience with the technologies themselves as incidents are not reported, investigated and studied.

There are various AI models, some already in use, others under development. We observe digital twins, robots, drones, both aerial and subsurface, image recognition for corrosion detection subsurface or in plants, automated crane operations, domain-specific AI solutions in

production optimization, energy usage reduction, condition-based maintenance, well operation planning, well operations, and Managed Pressure Drilling. These are typical AI-solutions within the industry (Markussen et al 2024). Most of these AI solutions provide decision support for operators or engineers, while some systems used for control and monitoring start to take more direct control over certain processes, albeit under human supervision, such as automated drilling operations. As of now, there are limited examples of fully automated AI-systems operating without human surveillance. One example, however, is autonomous robots for subsurface inspections (Markussen et al 2024).

The safety risks posed by AI are different from traditional industrial risk factors and conventional IT solutions. Unreliable predictions or incorrect generated output are a main safety risk in a high hazard context. Non-transparent or non-explainable outputs make it difficult, and potentially impossible, for users to interpret predictions and intervene if necessary (e.g. Endsley 2023, Markussen et al 2024).

Different factors in multiple development stages can cause AI to provide incorrect predictions. AI models can be trained on data that is insufficient for training purposes, or not adequately representative of the world one wants insight into. Data can be biased or change over time, impacting the functionality and reliability of the system in ways that can be challenging to understand and detect.

An important driver of risks in AI models is their complex life cycles involving maintenance and follow-up needs. This includes aspects related to data collection, model retraining, testing, implementation, and integration verifications (Suresh and Gutttag 2019). Traditional risk monitoring involves regular inspections, incident reporting, and safety audits. Once hazards are identified and mitigated, they remain relatively stable if the environment does not change, as described above. AI, however, requires continuous monitoring due to potential model drift and data changes.

AI models and the contexts in which they are used are often complex, involving different systems, sometimes distributed organisations, making it demanding to detect and respond to errors as they

unfold. A lack of training data on rare incidents, edge cases, can lead to inability to detect such cases if they occur (Markussen et al 2024). In addition to safety risks related to AI predictions and reliability, there are also uncertainties related to unexpected, unintended or malicious use of AI. Incidents reported from healthcare and transport underline the safety risks both inherent in the AI systems, but also from unexpected use, such as over-trust in Tesla cars ability as self-driving.

New research points to the fact that there are few formal prospective risk assessments published in the literature on AI and that it is also unclear what controls are required to eliminate or mitigate potential safety risks (Salmon et al 2024).

AI thus represents more uncertainties than many other technologies due to a weaker knowledge basis of what safety risks in different AI-systems might represent. All development, deployment and maintenance phases can potentially introduce new uncertainties that need to be understood and mitigated (Suresh and Gutttag 2021). AI-models developed for one context might not be reliable in another. HAVTIL has highlighted the need to address both inherent and contextual AI risk factors related to the potential for non-reliable outputs from AI-systems (Bergh and Teigen 2024).

As of now, we do not know what new scenarios might arise with the introduction of AI into our industry. We do not know how failures in these new technologies might affect well-known high hazards scenarios as underlying causes. As numerous sources during an AI solution's life cycle can introduce or increase risk, uncertainties in all AI life phases must be handled

(Suresh and Gutttag 2021; Towhidul Islam Tonmoy et al 2024). We have little statistical experience data, and there seem to be limited research efforts on prediction risks and deterioration of models over time related to different AI-models and -solutions, serving as decision support for which models will provide reliable output.

According to regulations, establishing a risk picture, describing the context for use, work processes and users when starting an AI-project provides a critical starting point for identifying risks, assessing risk and mitigating risks across all phases. For instance, choosing a transparent AI-

model in an early development phase is essential if the AI will collaborate with humans to support their decision-making in a safety-critical setting. While deep neural networks improve predictions, they paradoxically make them less transparent and explainable for those who are supposed to monitor and intervene if AI fails, making these tasks harder to perform correctly (Endsley 2023).

6. Follow up observations so far

AI is being developed rapidly in a competitive environment. Companies are expressing concern about the fear of missing out, either becoming outdated by new technologies constantly being developed or not making a profit due to high development costs without or with delayed deployment. A representative example of what we hear in dialogue with providers and deployers comes from a CEO in a major operating company on the Norwegian continental shelf. He stated in an interview that he believes the biggest risk with AI is not using it, as actual AI deployment will determine winners and losers within the industry (Havtil 2024). He also advocated for using AI surveillance of safety barriers. He explained: "Programmers traditionally use a fixed logic for each safety barrier, but an AI agent can handle thousands of such logics automatically". In our dialogue with the players, we observe notable variations in risk perceptions among leaders and the developer communities related to potential consequences of AI.

Recently, several new disciplines have entered our high-risk industry. During meetings, audits, and conferences, we've met managers and experts who are working on AI solutions for an industry they're not familiar with. New technology and methods shall be qualified before deployment according to regulations. Audits on automated operations show a lack of awareness of safety risks related to AI. We meet developers and managers that are unfamiliar with requirements of qualification of new technology before deployment.

Many experts and leaders claim that their systems will be safe because there will always be humans in the loop to detect and intervene if problems arise. However, available knowledge and research on human, organizational and operational factors that can influence system performance such as

Bergh (2024a); Endsley (2023); Ernstsén. et al (2022); Kirwan et al (2024) seems to be unknown among various developers and leaders we talk to. Some argue that user-representation during development will provide necessary knowledge on relevant risks factors. In projects we do to some extent meet user representatives providing information on the context and operations where these solutions are to be implemented. However, we observe that they are sometimes involved late in the project, or they focus primarily on normal operations, rather than how the system will perform as a barrier in a failure-, hazard- and accident situation. Safety delegates are frequently not included in digital technology projects, even when these projects entail significant technical and organizational changes within the company. Human factors experts with knowledge on human capabilities and limitations in Human-AI collaboration are to a limited degree involved in such projects. We observe insufficient assessments of system performance, on transparency and situational awareness to support decision-making, as well as possible under-trust or over-trust and workload, particularly when humans are expected to monitor and intervene. Engaging in meaningful conversations on safety risks when introducing complex technologies in the petroleum activities with managers and developers, has sometime been challenging. Our observations suggest that topics related to safety risks, and major hazards, have not been adequately addressed in all digital technology projects, even though user representatives have been involved. Through our follow-up, we observe great variation in developers' knowledge of high-hazard operations and major accident risks. The lack of knowledge effect risk assessments being performed. So far, we have not found competence requirement regarding risk assessments, or plans to familiarize with offshore activities, in these development communities. However, this is not a uniform scenario. We have also encountered AI developers with strong knowledge of risk management, major hazards, and activities in the petroleum industry, who are highly motivated to manage risks and qualifying new solutions before deployment. As an example, Reliable Augment Generate (RAG) models based on CoPilot or ChatGPT subtracting unstructured

data from company data sources are developed and highly regarded by some developers and players. Other developers strongly oppose the plans of implementing RAG models in high-risk contexts due to hallucinations or minor changes in output that can distort safety-critical information. This perspective on risk is supported in research literature such as by Alkaissi and McFarlane (2023). Incidents caused by hallucinations in LLMs, providing nonsensical text within high-risk contexts such as health care and aviation are also highlighted by OECD.

Many AI solutions have already been deployed through purchased solutions. In our follow-up of digital technologies, we see examples of providers not always informing buyers about the digital solution being used, when solutions are changed, or updates are made. Documentation on data usage and versions of models in use are not always provided. Sometimes it can be unclear whether it is the provider or the deployer who will maintain the systems over time. We also hear examples where providers of solutions are using data from deployer without ensuring the data's quality, representativeness, or labelling. Roles and responsibilities related to data input are not always clearly defined. Based on information from audits and other follow up activities, there appears to be limited contact between the developer communities and health, safety and environment professionals or risk managements experts within companies.

Most of the AI-systems that are developed will be decision support tools. Developers, managers and sales personnel are expressing high expectations with regards to human ability to monitor AI-performance over time, and to detect and intervene if AI fails. We are informed that humans in the loop will keep the systems safe. In conferences and meetings, we have during the last few years addressed the issue of possible over-trust in human capabilities, asking for assessments of human-AI collaboration using research-based knowledge on human capabilities and limitations and measurements of system performance. Research on human limitations, such as inability to monitor over time without losing situational awareness, or decreased ability to intervene over time (Endsley 2023; Naikar et al 2023), seems to be somewhat unknown in parts of

the developer communities, not at least among decision-makers and managers. So far, we have barely seen increased complexity included in the risk picture or part of risk assessments. In our follow-up, we observe that risk management is limited to a single digital system, with developers focusing mainly on a single human-machine collaboration. This approach reduces the problem to simple, dyadic components and relationships, thereby limiting the scope of risk assessments. Central control rooms, collaboration rooms, driller cabins, and functions and tasks placed outside these work environments, consist in many cases of various individuals or teams, who interact with each other and a large array of technologies. This complexity is, as of now, hardly addressed in the projects we have followed up. Research such as Naikar et al (2023) and Endsley (2023) call for a more holistic approach where all systems must be addressed, and research-based knowledge on capabilities and limitations of both AI and humans must be utilized in designing for Human AI collaboration. (Bergh et al 2024a)

7. Discussion

As a regulator, we have sought to gain insight into the possible risks associated with various AI models by following the research discourse, initiating new research for increased knowledge, and instigating a dialogue with both providers and deployers of AI. The concept of uncertainty has particularly proven to be vital in fostering insightful discussions, raising awareness of safety risks, and sharing learnings with the players. It has enhanced comprehension of the need to provide knowledge-based decisions and a cautionary approach to ensure safe operations. Our main concern as a regulator is that limited use of risk management systems, understanding and knowledge of safety risks, and not qualifying new technology can lead to unreliable solutions being deployed in a high-risk environment. We do acknowledge that increased accuracy in AI predictions and generated texts might be provided in the future but, there are still challenges with accuracy in predictions and generated output. Currently, the accuracy in many models may show 70%, 80% or 90% during testing. There are risks related to human-AI collaborations that are

still not assessed and mitigated. We register limited attention, so far, on maintenance needs to avoid model drift or requirements of and plans related to competency, training and exercise activities for users of AI.

Another concern is the lack of historic knowledge into AI-risks and possible causes of incidents. As of now, there is a lack of criteria for reporting of AI-related incidents in our industry. No incidents are being reported to HAVTIL, and we have not seen any systematic reporting within companies. So far, no investigation of AI incidents has been conducted in our industry nationally or internationally. Consequently, no systematic studies on the causes of such incidents have been provided to be used as a basis for risk mitigation. The lack of knowledge represents uncertainties related to AI and AI-incidents. Other agencies such as OECD has also identified incident reporting as one of their main priorities, pointing to the need for data on past and present AI risks as important evidence for making informed decisions about AI governance. They have implemented OECD AI Incident Monitor (AIM) to measure developing trends of recurring incidents but also to identify early signals to prevent future incidents for learning purposes.

Historically, risk management of both organizational, operational and technical changes has proven to be a fruitful way to identify, monitor and mitigate risks. Using learnings from previous incidents to mitigate risks have proven imperative. Before deployment of new technology, meticulously testing based on test criteria and piloting before full deployment are vital.

When new uncertainties due to AI are introduced, it will be even more important to identify and understand potential risks at an early stage, Bergh et al (2024b); Heide and Ersdal (2018).

Based on follow-up activities, we see that descriptions of contexts and risk pictures to understand what can go wrong here if introducing AI, is not sufficiently covered in companies' risk management. We see examples where the consequences related to personal injuries and major hazard scenarios due to incorrect predictions are not adequately addressed when planning for deployment offshore.

In most of AI-development projects, working environment and safety risk are not adequately identified, assessed and mitigated. As a regulator, that is a concern. Acquiring more knowledge, meticulous risk management, the cautionary and precautionary principles will be imperative when uncertainty is high.

8. Conclusion

In this paper, we argue that the new concept of risk in the 2015 petroleum regulation is especially valuable in risk management when new technologies such as AI are introduced in a high-hazard industry. This risk concept places more emphasis on uncertainties and the knowledge strength of the risk assessments. As emphasised in this paper low knowledge strength regarding safety impact of AI calls for a thorough risk management and a precautionary approach.

In our various follow-up activities, addressing uncertainty, has been particularly valuable in opening meaningful dialogue on AI-related safety risks, also with representatives from new disciplines in our industry. Currently, uncertainties related to AI are high, and compliance with risk management requirements in some of the projects we have followed up is inadequate. As a regulator, it is our responsibility to share these observations and remind providers and deployers of regulatory requirements, stating that risk management in a high-risk context is a precondition for a license to operate.

References

- Alkaissi H, McFarlane SI (2023) Artificial hallucinations in ChatGPT: implications in scientific writing. *Cureus*. <https://doi.org/10.7759/cureus.35179>
- Bergh, L.IV., Teigen, S.K., and Dørum, F. (2024a). Human performance and automated operations: a regulatory perspective. *Ergonomics*, vol 67).
- Bergh, L.I.V.; Teigen, K.S. (2024b) AI safety: A regulatory perspective ICAPAI
- Endsley, M.R. (2023) Ironies of Artificial Intelligence. *Ergonomics*, 66:11, 2023. <https://doi.org/10.1080/00140139.2023.2243404>
- Engen, O., A.; Hagen, J.; Lindøe, P., H.; Kringen, J.; Selnes, P., O.; Kaasen, K.; Vinnem, J. E.; Hansen, K., Solberg, A.; Teig, I., L. Tilsynsstrategi (2013)
- Ernstsen, J., Aulie, E. G., Sætren, G. B., Stenhammer, H. C., & Phillips, R. (2021). Et menneskesentrert perspektiv på kognitiv teknologi i petroleumindustrien. 2021.
- Heide, B., Ersdal, G. (2018) Revitalization of risk management in the Norwegian petroleum sector in Safety and Reliability – Safe Societies in a Changing World, - Haugen et al. (Eds) Taylor & Francis Group, London
- Jensen, R.; Antonsen, S. Haugen, S.; Vinnem, J., E.; Heggland, J.E.; Lootz, E. (2015) Soft Causes with Hard Consequences: The role of organisational factors in the creation of structural integrity in Safety Science Monitor <https://airisk.mit.edu>
- Kirwan, B., Venditti, R., Giampaolo, N., Sánchez, M. (2024). A Human Centric Design Approach for Future Human-AI Teams in Aviation <http://doi.org/10.54941/ahfe1005464>
- Lootz, E., Hauge, S., Kongsvik, T., & Mostue, B. A. (2012). Human factors and hydrocarbon leaks on the Norwegian Continental Shelf in Contemporary Ergonomics.
- Markussen, C.; Agrell, C; Van der Meulen, M., DNV (2024) Responsible use of artificial intelligence in the petroleum sector responsible-use-of-artificial-intelligence-in-the-petroleum-sector-in-norwegian-only.pdf
- Norwegian Ocean Safety Authority (HAVTIL 1999-2023) Trends in risk level in the petroleum activity (RNNP) Main reports
- Naikar, N., Brady, A., Moy, G. M., Kwok, H. W. Designing human-AI systems for complex settings: ideas from distributed, joint and self-organising perspectives of sociotechnical systems and cognitive work analysis. *Ergonomics* 2023
- Salmon, P.M., King, B.J., Elstak, I., McLean, S., Read, G.J.M (2024) Tomorrow demons: A scoping review of the risks associated with emerging technologies. *Ergonomics*. <https://www.tandfonline.com/doi/full/10.1080/00140139.2024.2416554>
- Suresh, H. and Gutttag, J. (2021) A Framework for Understanding Sources of Harm throughout the Machine Learning Life Cycle EAAMO21, October 5-9, 2021-, NY, USA <https://doi.org/10.1145/3465416.3483305>
- Towhidul Islam Tonmoy; S.M. Mehedi Zaman; S.M. Vinija Jain; Anku Rani; Vipula Rawte; Aman Chadha; Amitava Das (2024) A Comprehensive Survey of Hallucination Mitigation Techniques in Large Language Models
- Vinnem, J.E., Aven, T., Husebø, T., Seljelid, J., Tveit, O.J (2006) Major hazard risk indicators for monitoring of trends in the Norwegian Offshore petroleum sector in Reliability Engineering System Safety, volume 91 ELSEVIER