

Editorial: Generative AI in maintenance decision-making hype, hazard, or genuine capability?

1. The new conversation in the maintenance office

Generative artificial intelligence (AI) has moved from research curiosity to boardroom agenda in under three years. Large language models can now draft technical reports, summarize regulatory documents and hold plausible diagnostic conversations about equipment failures. In asset-intensive industries, such as oil and gas, power generation, mining, manufacturing and maintenance teams are beginning to ask whether these capabilities can solve persistent operational pain points including slow access to equipment knowledge, inconsistent work-order quality and overburdened planners who spend more time documenting than deciding. Early deployments already report measurable gains; Siemens, for example, has reported that its Industrial Copilot for maintenance reduces reactive maintenance time by an average of 25% in pilot use cases (Siemens, 2025).

The enthusiasm is understandable. But maintenance is not marketing copy. A hallucinated paragraph in a blog post is an embarrassment; a hallucinated torque value in a turbine reassembly procedure is a safety incident. Hallucination, which is the generation of plausible but factually incorrect content, remains a critical barrier to reliable deployment of large language models in any domain (Rawte *et al.*, 2023). The consequences are amplified in safety-critical environments where AI outputs can trigger physical interventions. This editorial examines where generative AI is already delivering value in maintenance, where its hazards are most acute and, more importantly, what governance posture organizations must adopt to use it responsibly. I argue that the decisive factor is not the technology itself but the organization's capacity for what is known as skeptical intelligence: a structured, teachable discipline of interrogating AI outputs before they reach the field.

2. Where generative AI is already being applied

Fault diagnosis and troubleshooting. Several early adopters now use large language models as conversational interfaces to maintenance knowledge bases. A technician facing an unfamiliar alarm code can describe the symptoms in natural language and receive a ranked list of probable causes drawn from historical work orders, OEM manuals and reliability databases (Thompson, 2025). This is not predictive analytics in the traditional sense; it is knowledge retrieval accelerated by natural language understanding. When the underlying corpus is well curated, the results can significantly reduce diagnostic time, particularly for less experienced technicians or those working on aging, poorly documented assets.

Work-order drafting and report generation. Maintenance planners routinely spend hours translating condition-monitoring findings into structured work orders. Generative AI can draft these documents in seconds, pre-populating fields such as failure mode, recommended action, required parts and estimated duration from templates trained on the organization's own maintenance history. A recent survey on generative AI for predictive maintenance in manufacturing confirms that automated work-order generation and SOP drafting are among the most documented and immediately productive use cases (Li *et al.*, 2026).

Knowledge retrieval from unstructured sources. Decades of maintenance intelligence sit locked in free-text fields, scanned inspection reports and engineering change notices that no keyword search can reliably surface. Retrieval-augmented generation (RAG) architectures are proving effective at extracting actionable information from these unstructured repositories by answering questions like "How was this failure mode addressed last time it occurred on a



similar asset?" with contextualized, source-referenced responses. [Chiang et al. \(2025\)](#) demonstrated a RAG-based interactive industrial knowledge management system achieving a mean reciprocal rank of 88% and recall of 85% in technical-service queries, underscoring the practical viability of this approach.

3. The hazards: what can go wrong

Hallucination in safety-critical contexts. Large language models generate text that is statistically plausible, not factually guaranteed. In maintenance, this distinction is dangerous. A model may confidently recommend a lubricant grade that does not exist, cite a standard clause that was revised three years ago or suggest an inspection interval that contradicts the asset's risk profile. Unlike a human expert who signals uncertainty through hedging language or a pause, a generative model delivers incorrect information with the same fluency as correct information. [Diab et al. \(2025\)](#) have mapped the gap between existing LLM benchmarks and the demands of hazard analysis in safety-critical systems, concluding that current evaluation frameworks are insufficient to guarantee reliable performance in these contexts. In safety-critical environments, such as pressure vessels, rotating machinery and electrical switchgear, an undetected hallucination can propagate through a work order into a physical intervention with serious consequences.

Erosion of institutional knowledge. There is a subtler risk. Over-reliance on AI-generated answers gradually erodes the diagnostic reasoning skills of the workforce. If technicians habitually accept the model's first suggestion without independent verification, the organization loses the very expertise that makes human oversight meaningful. This creates a dependency loop in which the people responsible for validating AI outputs become progressively less capable of doing so ([Thompson, 2025](#)).

Data privacy and intellectual-property exposure. Cloud-hosted generative models require that prompts and often the documents attached to them leave the organization's network. Maintenance data frequently contain proprietary process parameters, equipment configurations and failure histories that represent significant competitive and safety-sensitive intellectual property. Without robust data-governance controls, the convenience of a cloud-based AI assistant may come at an unacceptable cost.

4. Governance and trust: the case for skeptical intelligence

The hazards outlined above are real, but they are not arguments for rejection. They are arguments for a specific governance posture that is called skeptical intelligence. Skeptical intelligence is neither technophobia nor uncritical adoption. It is a structured organizational capability and a discipline of systematically interrogating, validating and contextualizing every AI-generated output against domain expertise, equipment history and operational evidence before that output is permitted to influence a maintenance decision.

Embedding skeptical intelligence in practice. Skeptical intelligence must be operationalized, not merely declared. This means establishing concrete review gates: no AI-drafted work instruction reaches the field without sign-off by a qualified reviewer who has cross-checked the recommendation against the asset's maintenance history and the applicable technical standard. It means defining confidence thresholds below which the model's output is flagged for mandatory human re-analysis rather than presented as a recommendation. And it means requiring that every AI-generated output carry a provenance trace: which data sources were consulted, which model version produced the response and what retrieval context was used. Without this traceability, review becomes theater rather than governance. The International Association of Privacy Professionals' analysis of hallucination governance reinforces this point, arguing that the most effective mitigation pattern treats the model not as an oracle but as a generator operating inside a verification loop ([IAPP, 2025](#)).

Accountability when AI-drafted instructions lead to error. Skeptical intelligence is not only a quality filter; it is also a governance layer that creates an auditable chain of responsibility. When

an AI-generated recommendation is reviewed, approved and executed, the review record documents who validated the output and on what basis. If the recommendation subsequently proves incorrect, the accountability trail is clear, not because blame is the objective, but because learning requires knowing where the process failed. Did the model hallucinate? Did the reviewer miss a contradiction? Was the underlying data stale? Skeptical intelligence transforms post-incident analysis from finger-pointing into structured improvement process.

Regulatory alignment. Emerging regulatory frameworks, most notably the EU Artificial Intelligence Act (Regulation 2024/1689), require meaningful human oversight of AI systems deployed in high-risk contexts (EU AI Act, 2024). A maintenance organization that has embedded skeptical intelligence including documented review protocols, defined accountability and traceable decision logs, is far better positioned to demonstrate compliance than one that relies on informal, ad-hoc checking. Skeptical intelligence, in this sense, is not merely good practice; it is a regulatory readiness strategy.

5. From pilot to practice: what maintenance leaders should do now

Begin with low-risk, high-frequency tasks. Report drafting, knowledge search and administrative summarization are ideal starting points: the consequences of error are manageable, the volume is high enough to demonstrate efficiency gains, and the outputs are already subject to human review in most organizations. These early use cases build familiarity and reveal the model's failure patterns in a controlled setting (Li *et al.*, 2026).

Cultivate skeptical intelligence as a workforce competency. This is not a one-time training module. It requires ongoing investment in critical-thinking skills applied specifically to AI outputs. Technicians and planners should be trained to ask: Does this recommendation align with what I know about this asset? Can I verify the cited source? What would I do differently if the model were unavailable? Organizations that treat skeptical intelligence as a measurable competency that is assessed, coached and rewarded, will build the human infrastructure that responsible GenAI deployment demands.

Scale to diagnostic and prescriptive roles only after skeptical-intelligence protocols have been tested and embedded. Moving generative AI into safety-critical decision support before the organization has demonstrated that it can reliably catch the model's errors is a governance failure waiting to materialize. The sequence matters: capability first and then confidence and expanded scope.

6. Closing reflection: genuine capability, conditionally

Generative AI is neither pure hype nor existential hazard. It is a genuinely powerful capability whose value in maintenance decision-making depends almost entirely on the governance discipline that surrounds it. The organizations that will extract lasting benefit are those that institutionalize skeptical intelligence, not as a slogan, but as a set of enforceable protocols, traceable review processes and a workforce that is trained to question fluent machines with the same rigor it applies to any other source of technical advice.

The current literature is heavily weighted toward algorithmic performance such as retrieval accuracy, response latency and benchmark scores. Far less is known about how skeptical intelligence functions in practice on the maintenance floor, whether it can be reliably measured and matured over time and what organizational structures best sustain it. Until these questions receive the empirical attention they deserve, the field will remain long on technological promise and short on operational proof.

Mohamed Ben-Daya

Acknowledgments

LLMs were used to identify appropriate references structure the editorial and extensive editing.

References

- Chiang, W.H., *et al.* (2025), "Application of retrieval-augmented generation for interactive industrial knowledge management via a large language model", *Computers in Industry*, Vol. 167, 104262.
- Diab, S., *et al.* (2025), "From hallucinations to hazards: benchmarking LLMs for hazard analysis in safety-critical systems", *Safety Science*, Vol. 183, 106781.
- EU AI Act (2024), "Regulation (EU) 2024/1689 of the European Parliament and of the Council laying down harmonised rules on artificial intelligence", available at: <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng> (accessed 6 April 2026).
- IAPP (2025), "Hallucinations in LLMs: technical challenges, systemic risks and AI governance implications", *International Association of Privacy Professionals*, available at: <https://iapp.org/news/a/hallucinations-in-llms-technical-challenges-systemic-risks-and-ai-governance-implications> (accessed 6 April 2026).
- Li, X., *et al.* (2026), "Generative AI for predictive maintenance for manufacturing: a survey", *Procedia CIRP*.
- Rawte, V., Sheth, A. and Das, A. (2023), *A Survey of Hallucination in Large Foundation Models*, arXiv preprint arXiv: 2309.05922.
- Siemens (2025), *Siemens Expands Industrial Copilot with New Generative AI-Powered Maintenance Offering*, Press release, March 2025, available at: <https://press.siemens.com/global/en/pressrelease/siemens-expands-industrial-copilot-new-generative-ai-powered-maintenance-offering> (accessed 6 April 2026).
- Thompson, E. (2025), "Shaping the future of industrial maintenance: generative AI, condition monitoring, and the rise of the citizen maintenance worker", *Transformer Technology*, March 2025, available at: <https://transformer-technology.com/article-hub/shaping-the-future-of-industrial-maintenance-generative-ai-condition-monitoring-and-the-rise-of-the-citizen-maintenance-worker/> (accessed 6 April 2026).