

As Physical AI Market Grows, New Research Identifies a Fundamental Safety Gap

An emerging discipline proposes that autonomous machines need safety infrastructure the tech industry has not yet built and may not yet recognize it needs.

ABU DHABI, UNITED ARAB EMIRATES, April 13, 2026 /EINPresswire.com/ -- The artificial

“

The AI industry has spent a decade making machines intelligent. The next decade will be defined by whether we make them trustworthy. That requires a new discipline, not just better engineering.”

*George Bancs, Founder,
Synthan Sciences*

intelligence industry is in the middle of its most consequential pivot. After a decade defined by software - language models, image generators, recommendation engines - the center of gravity is shifting toward machines that exist in the physical world. Autonomous vehicles, humanoid robots, surgical systems, industrial drones, and warehouse automation are no longer prototypes on conference stages. They are entering production. They are entering public spaces. And they are entering an environment with almost no standardized safety infrastructure designed for them.

This is the gap that a growing body of research and industry analysis now identifies as one of the most significant blind spots in modern technology development. The broader AI market - which encompasses autonomous systems, robotics, and intelligent machines operating in real-world environments - is projected to reach \$1.81 trillion by 2030, according to Grand View Research, with physical AI representing one of its fastest-growing segments. Yet the safety frameworks governing these systems remain fragmented, inconsistent, and in most cases, borrowed from adjacent fields that were never designed for the challenge.

The gap is not theoretical. It is already producing consequences. Autonomous vehicle recalls, industrial robot incidents, and drone-related regulatory disputes have made headlines across multiple continents. Each incident exposes the same underlying problem: the industry has invested heavily in making machines capable, but comparatively little in making them trustworthy.

THE SOFTWARE SAFETY PARADIGM AND ITS LIMITS

The current discourse around AI safety is dominated by concerns about software: bias in

algorithms, hallucination in language models, misinformation generated by image synthesis tools. These are legitimate issues. But they represent only one dimension of a much larger problem.

When artificial intelligence operates exclusively as software, its failure modes are informational. A chatbot gives a wrong answer. A recommendation engine surfaces inappropriate content. A credit scoring model produces biased outcomes. These failures cause harm, but they occur within digital boundaries. They can be patched, rolled back, and contained.

Physical AI operates under fundamentally different constraints. When an intelligent machine moves through the world - driving a vehicle, assembling a product, interacting with a person - its failure modes are kinetic. They involve mass, velocity, and force. They cannot be patched in real time. A wrong decision by an autonomous system operating at highway speed or on a factory floor has consequences measured in injuries, structural damage, and, in the worst cases, fatalities.

Despite this fundamental difference, the safety paradigm applied to physical AI systems is largely an extension of software safety thinking. Risk assessment methodologies, testing protocols, and certification frameworks have been adapted from software development and traditional manufacturing safety - two fields that predate the existence of machines capable of autonomous decision-making in unstructured environments.

This is not a failure of intent. It is a failure of framework. The tools being used were designed for a different category of problem.

THE EMERGENCE OF [SYNTHANITY](#)

Into this gap has emerged a new conceptual and research framework that some observers believe could reshape how the industry thinks about [physical AI safety](#). The framework is called synthanity - a term coined by researcher and author George Bancs to describe the comprehensive study, governance, and integration of synthetic intelligent beings into human society.

Synthanity is not a product or a single technology. It is a proposed discipline - an intellectual



George Bancs, Author and Founder of Synthan Sciences

architecture that attempts to address the full spectrum of challenges created by intelligent machines operating alongside humans. It spans science, law, and culture, arguing that the safe integration of physical AI cannot be solved by engineering alone. It requires new legal categories, new social contracts, and new cultural norms.

The framework is detailed across a three-volume series, The [Syncyclopedia of Synthanity](#), which lays out the theoretical foundations, legal governance structures, and cultural implications of a world in which autonomous machines are not tools but participants in shared physical space.

What distinguishes synthanity from existing AI safety discourse is its scope and its starting point. Where most industry safety efforts begin with a specific technology - autonomous vehicles, surgical robots, industrial automation - and work outward to identify risks, synthanity begins with first principles about what it means for an intelligent entity to operate in the human world, and works downward toward specific applications.

This inversion of perspective is deliberate. The framework argues that technology-specific safety approaches produce technology-specific solutions that fail to generalize. The autonomous vehicle industry develops its own safety standards. The surgical robotics industry develops different ones. The warehouse automation industry develops others. There is no shared vocabulary, no common architecture, and no interoperability between these safety regimes.

Synthanity proposes that what is needed is not more technology-specific safety work but a foundational layer of safety infrastructure that applies across all physical AI systems - what the framework describes as a universal trust architecture for autonomous machines.

CHALLENGING THE INDUSTRY CONSENSUS

The implications of this framework extend beyond academic theory. If synthanity's core thesis is correct - that the AI industry's approach to physical safety is fundamentally fragmented and structurally insufficient - then the current trajectory of autonomous systems deployment carries risks that are not being adequately priced by markets, regulators, or the public.

This is a challenging position for an industry that has invested billions in existing safety approaches. It suggests that the problem is not a lack of effort or resources but a misalignment of paradigm. The industry is solving the wrong problem well, rather than solving the right problem at all.

Several industry dynamics support this concern. First, the speed of physical AI deployment is accelerating faster than the development of safety standards. Humanoid robots are entering commercial pilot programs in manufacturing, logistics, and healthcare. Autonomous delivery vehicles are operating on public roads in multiple countries. Intelligent drones are being integrated into emergency response, agriculture, and infrastructure inspection. In each case, the technology is being deployed into environments where it interacts with humans, under safety

frameworks that were designed before the technology existed.

Second, the regulatory landscape remains fragmented. Different jurisdictions apply different standards. There is no international equivalent of what the internet has in ICANN or what aviation has in ICAO - no global body setting interoperable safety standards for autonomous physical systems. This fragmentation creates regulatory arbitrage, where companies can deploy in jurisdictions with the least stringent requirements, and systemic risk, where an incident in one market has no mechanism for informing safety practices in another.

Third, the talent and research pipeline for physical AI safety is thin. The overwhelming majority of AI safety research funding and academic attention is directed toward software AI safety - alignment, interpretability and bias mitigation. Physical AI safety - the study of how to make autonomous machines safe in shared physical environments receives a fraction of the attention and resources, despite arguably posing more immediate and more severe risks.

RESEARCH AND PRACTICE

One of the more unusual aspects of the synthanity framework is the relationship between its published research and its practical application. George Bancs, who developed the framework and authored the Syncyclopedia series, is also the founder of Synthan Sciences, an Abu Dhabi-based startup that is building commercial safety infrastructure for physical AI systems.

This dual positioning - researcher and entrepreneur - is itself a commentary on the state of the field. In established disciplines, the path from theoretical framework to commercial application typically runs through academic institutions, government research programs, and standards bodies before reaching the private sector. In physical AI safety, those intermediary institutions are only beginning to organize around the problem. The gap between identifying the need for a new discipline and having the institutional infrastructure to develop it has created space for private-sector actors to move faster than the traditional knowledge pipeline.

Synthan Sciences has developed a proprietary multi-layer safety architecture that spans hardware-level safety components, communication and behavioral protocols, and identity and certification frameworks for autonomous systems. The company operates under the Abu Dhabi Global Market (ADGM) regulatory framework and is currently preparing for a seed funding round.

The company represents one approach to translating the theoretical work of synthanity into deployable technology. Whether it becomes the dominant approach or one of many will depend on how quickly the broader industry recognizes the structural gap that the framework identifies and how rapidly institutional responses - from regulators, standards bodies, and academic institutions materialize.

THE QUESTION OF TIMING

The central tension in the synthanity thesis is one of timing. If the framework is right that physical AI safety requires a fundamentally new paradigm - not incremental improvements to existing approaches - then the window for building that paradigm is narrowing as deployment accelerates.

History offers instructive parallels. Cybersecurity emerged as a discipline only after the consequences of its absence became impossible to ignore. The early internet was built with almost no security infrastructure, and the cost of retrofitting security into systems designed without it has been measured in trillions of dollars and counting. Aviation safety, by contrast, developed its international regulatory infrastructure relatively early in the technology's lifecycle, and the result has been one of the safest transportation systems ever created.

Physical AI sits at a similar inflection point. The infrastructure choices made in the next several years - whether the industry builds a unified safety paradigm or continues with fragmented, technology-specific approaches - will determine the safety trajectory of autonomous systems for decades.

The emergence of synthanity as a framework, regardless of whether it becomes the prevailing paradigm, signals that the intellectual groundwork for this transition is being laid. The question is whether the industry, regulators, and the public are ready to engage with it before the deployment curve makes the problem orders of magnitude harder to solve.

The Syncyclopedia of Synthanity, the three-volume series that details the framework, is available on Amazon and through major booksellers worldwide. For more information about the research and its practical applications, visit synthansciences.com.

Media Contact

Agentic Press Holdings

[email us here](#)

This press release can be viewed online at: <https://www.einpresswire.com/article/905240402>

EIN Presswire's priority is source transparency. We do not allow opaque clients, and our editors try to be careful about weeding out false and misleading content. As a user, if you see something we have missed, please do bring it to our attention. Your help is welcome. EIN Presswire, Everyone's Internet News Presswire™, tries to define some of the boundaries that are reasonable in today's world. Please see our Editorial Guidelines for more information.

© 1995-2026 Newsmatics Inc. All Right Reserved.