

How to Identify and Operationalize GenAI Use Cases Which Really Help You

Christoph Bergen, HMS Analytical Software GmbH, Heidelberg, Germany

Luis Wirth, HMS Analytical Software GmbH, Heidelberg, Germany

ABSTRACT

Generative Artificial Intelligence (GenAI) has emerged as a transformative technology. Distinguishing its value within regulated and complex domains such as the pharmaceutical industry remains challenging. This paper presents an introduction to guide decision-makers in evaluating where and how to apply GenAI effectively. We begin by contrasting classical, discriminative AI approaches with generative models, outlining the types of problems best suited for each paradigm. Through real-world pharmaceutical examples, we illustrate how GenAI models can be leveraged and highlight infrastructure and processes required to translate a model into a productive and adopted solution. Furthermore, we discuss strategies for operationalizing GenAI within enterprise environments, ranging from self-service AI platforms to distributed systems of GenAI tools.

INTRODUCTION

Since the introduction of ChatGPT at the end of 2022, the category of generative AI use cases has attracted a great deal of attention. Many promises are made on the basis of examples that work well in a demo environment, yet the requirements in the pharmaceutical industry are complex. An informed and nuanced understanding of the capabilities and appropriate use cases of AI, with a focus on generative AI in this paper, is important. The following article aims to provide an introduction and guidance.

In the pharmaceutical and life sciences but as well in other industries, this development coincides with increasing demands for data-driven efficiency, regulatory compliance, and innovation under constrained resources. Traditional AI methods have already demonstrated their value in areas such as process optimization, predictive maintenance, and anomaly detection. However, their applicability is often limited by data availability, domain-specific model training requirements, and narrowly defined problem statements with uncertainty about model performance. Generative approaches mitigate several of these constraints by leveraging pre-trained foundation models that can be adapted to diverse contexts.

Despite these opportunities, the integration of GenAI into (regulated) enterprise environments remains a complex challenge. Beyond technical implementation, the identification of use cases that provide verifiable business or scientific value demands a structured evaluation of both feasibility and benefit within domain-specific constraints.

The objective of this paper is therefore twofold. First, it aims to differentiate the omnipresent term “AI”. By delineating the conceptual boundaries between classical and generative AI and to characterize the resulting implications for model development and deployment. Second, it provides a structured framework for identifying and operationalizing GenAI use cases within the pharmaceutical context. This includes the consideration of organizational, technical, and regulatory factors that influence scalability and sustainability. By synthesizing these perspectives, the paper seeks to contribute to a more systematic understanding of how generative AI can be responsibly and effectively applied in complex domains. The insights presented are largely grounded in the authors’ practical experience and observations from real-world projects and collaborations.

DIFFERENTIATION BETWEEN DISCRIMINATIVE AND GENERATIVE AI

Classical, or discriminative, AI approaches are primarily focused on analyzing and predicting outcomes based on existing data. These models are designed to classify, detect, or forecast by learning the relationships between inputs and outputs. Common use cases include predictive maintenance, fraud detection, tabular data classification, and time series forecasting. For instance, a model may predict equipment failures in a pharmaceutical production line or identify fraudulent transactions in billing systems (discriminate between fraudulent/ not fraudulent or failure/ non

failure). Typically, these solutions require use case-specific training data and result in models optimized only for a well-defined task that typically cannot generalize to other tasks.

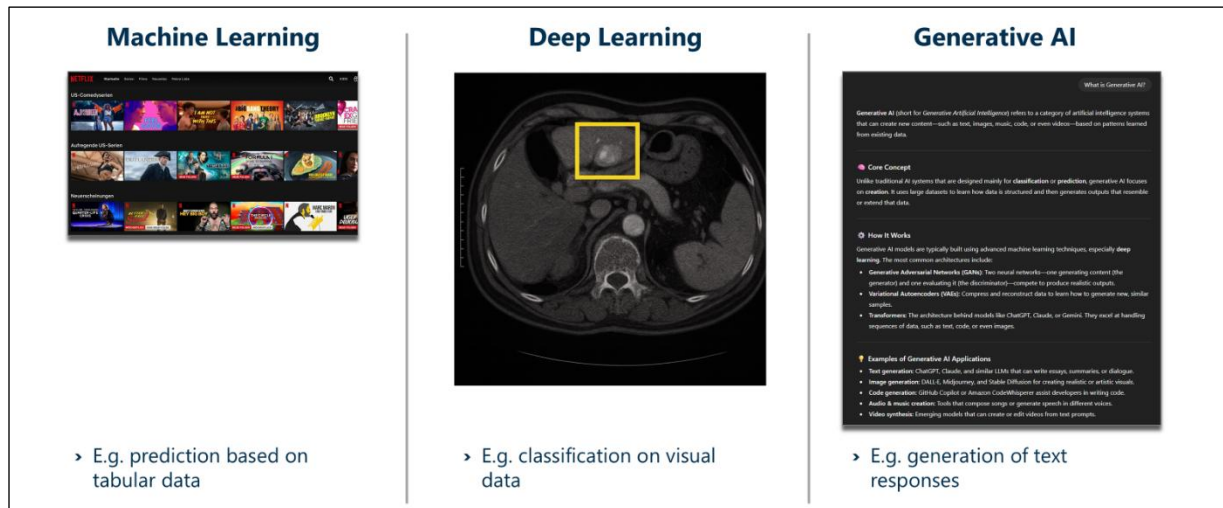


Figure 1. illustrates three archetypal use cases to clarify the vocabulary commonly used in connection with Artificial Intelligence. Machine Learning is a subfield of Artificial Intelligence, while Deep Learning is a specialized form of Machine Learning that uses deep neural networks as learning algorithms. Generative AI, which typically relies on deep learning approaches for model training, refers more to the nature of the model's output - it generates new, not predefined or bounded outputs (e.g., yes/no; blue/red; bounding box coordinates). In theory, every Generative AI model is also a Machine Learning model, as well as a Deep Learning model. However, in current discussions, the terms are often used imprecisely and tend to describe applications from a historical development perspective. Machine Learning is, although inaccurately, often associated with use cases involving primarily tabular data, while Deep Learning is linked to data modalities such as image, video, audio, and static text. Generative AI, in turn, is typically associated with generating text and images.

In contrast, generative AI models are creating new data or content given an input, they are not restricted to predefined output value ranges. Instead of mapping inputs to predefined outputs, generative models learn underlying patterns and structures from large amounts of data and use this knowledge to generate novel outputs such as text, images, or audio. This enables applications like chatbots, natural language translation, and automated literature review systems. Importantly and in contrast to the case-specific training of classical AI models, generative AI approaches often leverage large, pre-trained foundation models that can be adapted to a wide range of tasks without the need for domain-specific training data.

While there is some overlap - for example, in recommender systems and literature search applications - the fundamental distinction lies in the role each paradigm plays within a workflow. Classical AI excels at decision support and classification in structured environments, whereas generative AI introduces possibilities for content creation, automation, and interaction. Classical AI can support rigorous, rule-based processes such as quality control, while generative AI can enhance innovation and efficiency, for instance, by drafting regulatory documents or generating patient communication materials. This complementarity makes understanding both paradigms essential for selecting the right approach to a given challenge.

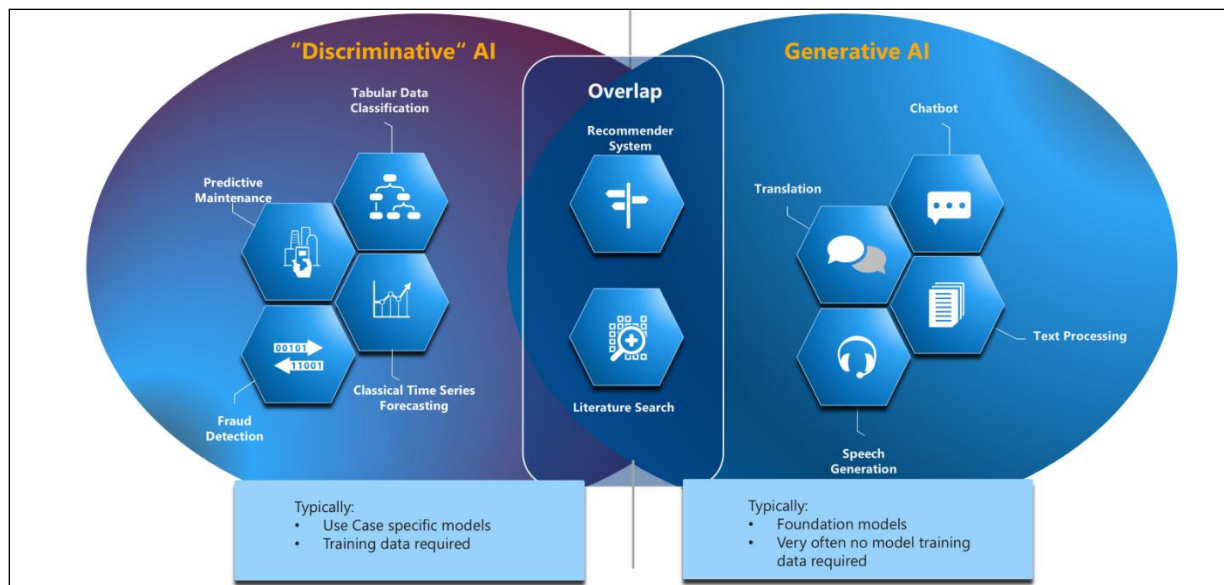


Figure 2 illustrates the distinction and overlap between discriminative and generative AI paradigms. Discriminative AI models are typically tailored to specific use cases and require customized training data suited to the application. In contrast, generative AI generally built on foundation models that require little to no additional training data. Between these two paradigms lies an overlap, with use cases such as recommender systems and literature search, where both discriminative and generative modeling techniques can be applied.

DIFFERENCE IN MODEL TRAINING AND SELECTION

Another key differentiating factor between classical and generative AI in practical use lies in the way an AI model is selected and deployed for a specific application.

In classical AI, model selection is typically application- and data-specific. A model is developed or trained to be precisely tailored to the given dataset and problem. The most effort is usually required for data collection, curation, and preparation. This is followed by the model selection process, which is based on experimental evaluation (e.g. using cross-validation and performance metrics such as accuracy or F1-score) and requires technical expertise in machine learning. The goal is to optimize a clearly defined objective criterion based on a limited dataset.

In contrast, generative AI is usually based on large, pre-trained foundation models (e.g. large language models (LLMs), multimodal foundation models) that already possess broad world knowledge and versatile capabilities. Here, "selecting" a model means less about training or comparing a set of models and more about identifying a suitable foundation model and adapting it to the specific application. Adapting in this context means through prompt design, fine-tuning, retrieval augmentation, or by building a workflow of multiple LLM calls (orchestration). As a result, the focus shifts from model training and architecture to system integration and contextualization. For further reading, [1] provides a more detailed discussion of these aspects.

IMPACT ON USE CASES

This change in how generative AI systems are built and deployed compared to classical AI systems has an impact both on economics and the practicality of implementing AI solutions. Instead of investing in model development and data preparation, the focus has shifted to leveraging existing foundation models. Use cases can now be realized more efficiently, as feasibility assessments and prototyping are much faster due to the availability of pre-trained models and standardized, API-based interfaces. These developments lower entry barriers and accelerate the transition from experimentation to production. As a result, integration into existing IT systems has also become simpler.

This paradigm shift has led to what can be described as a *"shift to the right"* in the feasibility-benefit landscape of AI use cases. Many applications that were previously considered infeasible - due to high training costs, limited data availability, or complex integration requirements - are now accessible.

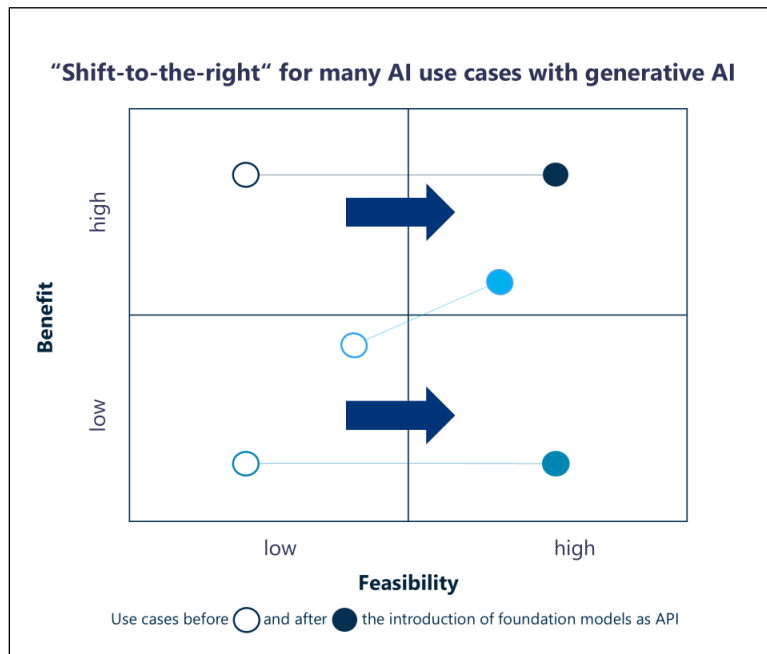


Figure 3 shows how use cases that had not been implemented before became implementation targets with positive RoI after gaining access to API-based generative AI model capabilities. API stands for Application Programming Interface and in this context refers to the ability to access a generative AI model through an interface without investing effort in training, deployment and hosting of the model - typically through cloud-based offerings.

These use cases have become strong candidates for early implementation, often described as "low-hanging fruit". This trend is ongoing: as model capabilities continue to grow and frameworks, tools, and solution understanding improve, the range of feasible and beneficial applications continues to expand.

THE MANY VIEWS ON GENERATIVE AI USE CASES

Generative AI can be integrated into organizational ecosystems at different levels, each representing a unique perspective on its role, purpose, and interaction with stakeholders. Table 1 illustrates these perspectives, ranging from back-end automation to direct customer-facing applications.

At the **system level**, GenAI operates as a largely invisible component, automating background processes such as cleaning unstructured text data or extracting relevant information into standardized formats. In this scenario, direct model interaction with end users is minimal or non-existent, with systems typically connected through REST API calls to model endpoints. One example of such a system could be in pharmacovigilance where unstructured text data from safety reports, patient feedback or trial notes could be transformed into structured data (e.g. transformation into drug name, event type, severity, ...).

The **process view** situates GenAI as a functional element embedded within business workflows. Here, it automates repetitive manual tasks, adds quality control steps, or generates intermediate data products. While primarily backend-driven, these use cases may expose limited outputs to users through dashboards or workflow interfaces. For example, such a system could generate an initial version of textual descriptions of statistical outputs, potentially based on templates and domain-specific terminology, to serve as a starting point for a human medical writer.

Moving closer to direct **user engagement**, the user view highlights GenAI as an interactive assistant supporting daily tasks. Examples include semantic search within knowledge bases, coding assistants, and retrieval-augmented generation (RAG) systems. These tools are typically delivered via chatbots, custom interfaces, or integrations within existing software platforms, providing immediate value to users. For example, exploratory data analysis (also integrating standard plots via existing libraries) could be provided within a chatbot interface for non-programming experts.

Finally, in the **customer view**, GenAI directly generates value for end customers, such as personalized summaries, Q&A systems, or FAQ responses. These use cases offer a entry point for organizations to integrate GenAI into products and services. An example could be a medical information chatbot which answers questions about drug dosing, contraindications, or storage conditions, importantly based explicitly on verified content.

By distinguishing these four perspectives (system, process, user, and customer), the layered approach helps identify the appropriate level of technical integration, interaction model, and exposure to stakeholders, providing a starting point for GenAI use cases to be effective and embedded within organizational processes.

VIEW TYPE	APPLICATION WITHIN VIEW CONTEXT	EXAMPLE USE CASES	INTERACTION WITH THE MODEL
System View	Automated processing Preprocessing with a human in the loop	Cleaning unstructured text data, e.g., extracting relevant information into a database schema in pharmacovigilance	REST calls to model API endpoint, no exposure to end users
Process View	Automation of tedious manual tasks or additional quality gate Gen AI data product of the process	Generate an initial version of textual descriptions of statistical outputs	REST calls to model API endpoint, some exposure via UI components, deep dive chatbot
User View	Assistant tool in daily work Semantic search in knowledge base	Copilot-like-support, chatbots, RAG systems, coding-assistants, exploratory data analysis via Chatbot	Chatbots, custom systems, software built-ins
Customer View	Generation of content for customer (e.g., summaries) Interface to (parts of the) applications (e.g., FAQ)	Low barrier entry point into product (e.g., data analysis). QnA / FAQ Answering	Chatbots, software built-ins

Table 1 provides an overview of different application views for integrating generative AI. It classifies use cases into four perspectives: System, Process, User, and Customer View.

OPERATIONALIZATION

For pharmaceutical companies, the difference between pilot and scale determines whether GenAI is “just” a productivity tool or a strategic differentiator. Moving from experimentation with GenAI to systematic, scalable deployment in enterprise environments requires more than a successful pilot. Traditionally, “operationalization” refers to the process of bringing an AI model or system into operations (into “production”). However, succeeding in this process involves more than just technical implementation. **Operationalization** encompasses the organizational, technical, and strategic decisions that determine whether GenAI efforts evolve into sustainable capabilities or remain isolated prototypes.

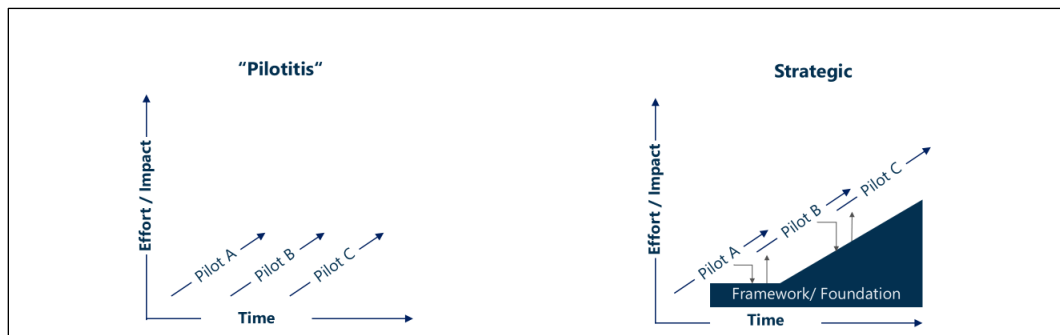


Figure 4: The above illustration is based on the book “Elements of Data Strategy” by Boyan Angelov [2].

These choices span leadership and governance, distribution of responsibilities, sourcing and technology models, regulatory alignment, resource readiness, and value delivery. The following section addresses each aspect in as much detail as possible within the constraints of the paper's length. It provides an overview of the various elements and, in particular, the different types of options available for companies to pursue. This helps the reader navigate their own organization and understand available options when planning to implement a GenAI initiative.

A central question is where generative AI initiatives originate. Many pharmaceutical companies start bottom up, with domain experts or teams creating lightweight solutions to address workflow bottlenecks. These efforts tend to fit team needs and gain quick adoption, but without shared standards they often cannot be scaled. In this context, *scaled*

means expanding a solution beyond its initial pilot team or use case. For example, deploying it across multiple therapeutic areas, geographies, or business units. Top-down initiatives, by contrast, usually begin with C-level priorities and dedicated budgets. They bring resources and strategic clarity but risk being disconnected from daily operations and slowed by bureaucracy. Increasingly, firms pursue a hybrid model: executives set direction and secure resources, while domain teams propose and test use cases within a clear governance and platform framework. This balances relevance with scalability. At one point, organizations usually opt for building or buying a platform to facilitate GenAI use cases. Early adopters often begin with multiple, scattered pilot projects, each with bespoke implementations. While these initial efforts create valuable learning experiences, they also tend to increase maintenance costs and lead to silo solutions. Despite these challenges, early experimentation holds significant value. It allows the organization to gain first-hand exposure and practical experience with GenAI technologies. Moreover, generative AI systems depend heavily on integration with the organization's existing IT ecosystem. Relevant data sources such as knowledge management systems, document repositories, databases, and proprietary software libraries that already model business processes must be connected and made accessible to the AI systems. Use cases are then layered onto this shared foundation, enabling reuse, faster scaling, and more efficient delivery. Once compliant data sources and interfaces are in place, they can be applied consistently across products and geographies. This transition shifts the organization's approach from project-by-project delivery toward a capability-driven operating model.

One very important aspect that is often overlooked is talent and change management. Beyond technical infrastructure, employees need training and support to integrate GenAI into their work. Successful firms create upskilling programs, introduce roles such as AI engineers or AI solution architects, and form cross-functional product teams pairing domain experts with technical specialists. They also frame GenAI as an assistant rather than a replacement and highlight early wins to build confidence and adoption. In particular, it also covers conveying the weaknesses or other behaviors of these technologies compared to traditional IT systems, since the use of generative AI is by nature non-deterministic.

Portfolio management ensures resources are focused on the right use cases. Organizations define clear evaluation criteria, considering feasibility, regulatory risk, business value, and change readiness. Projects are categorized into quick wins, strategic bets, or deprioritized experiments. Stage-gated investment models control scale-up, while success is measured through both business outcomes (time saved, errors reduced, outcomes improved) and technical robustness (explainability, reproducibility, monitoring). Sharing these results strengthens future selection and builds trust across stakeholders.

Beyond vertical decision-making hierarchies, firms must also consider the horizontal organizational structure, which determines where in the company ownership and oversight should reside. Early GenAI projects often sit in central innovation or AI/Data Science units. As adoption grows, firms either centralize capabilities in Centers of Excellence (CoEs) or distribute them into business functions such as regulatory affairs, medical affairs, or manufacturing. Centralized CoEs support standardization, reuse, and risk management, but may struggle to meet varied business demands quickly. Decentralized models respond faster but often duplicate work and fragment governance. A federated structure has become common in large pharma: a central team sets standards, maintains shared infrastructure, and validates models, while business units adapt and own their implementations. This approach preserves both domain ownership and enterprise-wide consistency.

Technology and sourcing choices are equally important. Companies can build custom solutions, buy external tools, or combine the two. Building in-house integrates closely with proprietary data and workflows, ensures privacy, and creates opportunities for IP, but requires major investments in infrastructure and talent. Buying external tools accelerates deployment and reduces development overhead but can limit flexibility and raise compliance concerns. Many firms adopt a hybrid path, building differentiated internal capabilities while relying on external APIs or services for foundational components. Multi-cloud strategies are often used to retain flexibility and reduce vendor risk.

Risk and compliance must be embedded from the start. In pharma, trial-and-error approaches are not viable. Companies implement credibility frameworks to assess outputs, validate models under GxP conditions, and ensure decision traceability. Governance typically includes three layers: domain teams manage first-line compliance, cross-functional committees review high-risk implementations, and C-suite or board members oversee strategy. Automated safeguards such as hallucination detection, source validation, and audit trails are built into pipelines, easing review processes with regulatory, legal, and medical teams.

These dimensions describe how pharmaceutical companies operationalize GenAI. Hybrid leadership, federated implementation, platform-first development, hybrid sourcing, embedded governance, talent investment, and disciplined portfolio management together help companies turn GenAI from isolated pilots into enterprise-wide capabilities.

OUTLOOK

This paper primarily focuses on the identification, adoption and implementation of Generative AI. The next evolutionary step, which is already beginning to emerge, is Agentic AI. While Generative AI use cases are typically characterized by the application of an LLM with a clearly defined solution path (if complex logic is required at all), Agentic AI represents a step towards more autonomous systems designed for more complex tasks. These systems are given a task objective and, for example, the tools to solve it, rather than being provided with the specific solution path. This allows for the modelling of more heterogeneous, non-linear challenges. The building block of an agentic system is the so-called agent. This agent can perceive, decide and act. It operates goal-oriented and autonomously within the framework of an overarching strategy. LLMs form the core of agents, acting as the decision-making and orchestration engine, and controlling the agent's behavior.

Hand in hand with the development of agentic AI comes a drive for standardisation. For example, access to a knowledge database should not be redeveloped for every AI use case, or in a further step, there should be a possibility to make these resources generally available to all newly emerging agents. In this context, standards such as MCP (Model Context Protocol [3]) are being created, and timely implementation should be considered.

CONCLUSION

This paper examined the conceptual and practical differences between discriminative and generative artificial intelligence, with a specific focus on the identification and operationalization of GenAI use cases in the pharmaceutical domain. By contrasting these paradigms, we have shown that generative models extend the range of feasible applications beyond traditional, data-specific approaches. Through pre-trained foundation models and adaptable interfaces, GenAI enables the automation and augmentation of tasks that previously required extensive domain-specific data or manual effort or were not feasible at all.

The discussion of operationalization emphasized that successful deployment of GenAI systems requires a multidimensional framework that integrates technical, organizational, and regulatory considerations. The transition from isolated experiments to sustainable enterprise capabilities depends on the establishment of coherent governance structures, the availability of compliant infrastructure, and the systematic alignment of GenAI initiatives with business processes.

In highly regulated environments such as the pharmaceutical industry, adherence to quality, traceability, and validation standards is a prerequisite for the responsible use of generative systems. Embedding risk management, auditability, and human oversight within the AI system ensures that GenAI deployments meet industry requirements. Furthermore, the role of human expertise remains central - both in curating input data and in critically evaluating and processing AI outputs.

From a broader perspective, generative AI represents a significant shift in how computational systems interact with knowledge and language. The emerging evolution towards agentic architectures may further extend these capabilities by enabling autonomous goal-oriented behavior.

In conclusion, the systematic identification and operationalization of GenAI use cases require not only technical proficiency but also organizational maturity and domain-specific understanding. A rigorous and structured approach allows organizations to integrate generative methods responsibly and effectively, transforming them from isolated technological experiments into stable and scientifically grounded components of enterprise decision-making.

Based on the author's experience, the adoption of generative AI is not about identifying the right use case or whether to implement it, but about determining the best way to do so. Companies still at the starting line risk being left behind. The differentiator is increasingly not just the adoption of the technology, but concrete experience with it. Generative AI is moving from experimental use to becoming a baseline requirement for rapid decision-making and efficiency through automation.

REFERENCES

- [1] Huyen, Chip. AI Engineering: Building Applications with Foundation Models. O'reilly, 2025
- [2] Boyan, Angelov. Elements of Data Strategy: A Framework for Data and AI-Driven Transformation. 2023
- [3] Anthropic. Introducing the Model Context Protocol. Accessed October 27, 2025.
<https://www.anthropic.com/news/model-context-protocol>

CONTACT INFORMATION

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

1st. Author Name: Christoph Bergen

Company: HMS Analytical Software GmbH

Address: Grüne Meile 29, 69115 Heidelberg, Germany

Email: Christoph.Bergen@analytical-software.de

Website: <https://www.analytical-software.de/en>

2nd. Author Name: Luis Wirth

Company: HMS Analytical Software GmbH

Address: Grüne Meile 29, 69115 Heidelberg, Germany

Email: Luis.Wirth@analytical-software.de

Website: <https://www.analytical-software.de/en>

Brand and product names are trademarks of their respective companies.