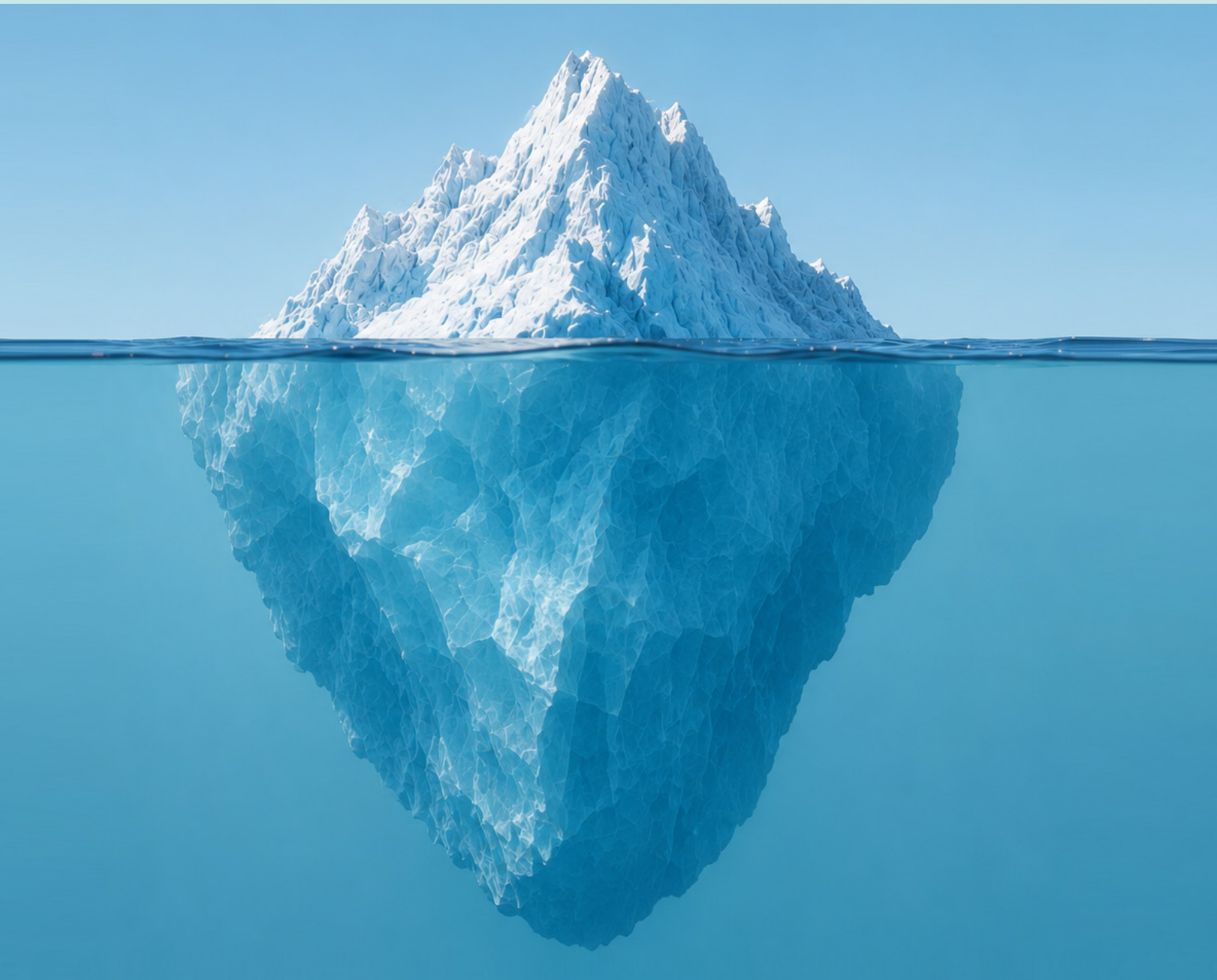


Simple because it matters.

ERGO

From Tribal Knowledge to Self-evolving Agents

Capturing, Governing, and Scaling
Institutional Knowledge in Financial Enterprises



In collaboration with

Google Cloud

Deloitte.



ERGO is shaping the future of insurance through innovation, emerging technologies and digital transformation. The ERGO Innovation Lab, based at the Merantix AI Campus in Berlin, explores new technologies and develops forward-looking services jointly with respective business owners and technology delivery units that challenge the status quo and help transform insurance, risk and finance.



Google Cloud is the premier hyperscaler for AI, accelerating innovation by bridging cutting-edge research with secure, scalable enterprise applications. Google Cloud provides a comprehensive, full-stack AI ecosystem, ranging from custom-built silicon and industry-leading Gemini models to seamless integrations across its SaaS and PaaS portfolios. By partnering closely with top-tier financial services organizations, Google Cloud co-creates tailored solutions that drive meaningful business transformation and set new industry standards.



Deloitte provides leading professional services to nearly 90% of the Fortune Global 500® and thousands of private companies. Legal advisory services in Germany are provided by Deloitte Legal. Our people deliver measurable and lasting results that help reinforce public trust in capital markets and enable clients to transform and thrive. Building on its 180+-year history, Deloitte spans more than 150 countries and territories. Learn how Deloitte's over 470,000 people worldwide work together every day to make an impact that matters at www.deloitte.com/de.

„We can know more than we can tell.“

Polanyi (1996, p.4)

In many enterprises, some of the most valuable knowledge is never formally documented anywhere. Instead, it can be found in conversations, habits and the accumulated experience of individuals.

Much of this knowledge is shared informally through interactions. This form of practical know-how is commonly referred to as “tribal knowledge”: knowledge that develops through experience and is passed on through practice rather than written records (Polanyi, 1958; Watts, 2026).

However, tribal knowledge has two defining characteristics: it is largely undocumented, or documented only sporadically at the individual level and rarely shared across the organization, particularly in a form accessible to enterprise systems; and it continues to evolve over time.

Enterprise Knowledge at Risk

In the coming years, a significant number of senior employees in Germany are expected to retire (GDV, 2026; Statistisches Bundesamt, 2025). This demographic shift is not unique to Germany; the retirement of the baby boomer generation is a global phenomenon with significant implications for organizations worldwide (McKinsey, 2026).

Without proper documentation, a significant share of tribal knowledge leaves the enterprise over time. Rebuilding this expertise from scratch leads to avoidable productivity losses and inconsistent quality. In addition, knowledge loss in enterprises persists due to staff turnover, limited knowledge transfer between senior and junior colleagues and initiatives where new processes replace old ones without preserving the insights and knowledge accumulated through their use.

The Rise of Agentic Expertise

Today’s advancements in artificial intelligence offer a way to address this knowledge loss or gap (Cai et al., 2026; Gundurao et al., 2026).

Particularly in the insurance industry, Agentic AI has already been implemented at different levels of maturity, tool use and autonomy (Corradi et al., 2026).

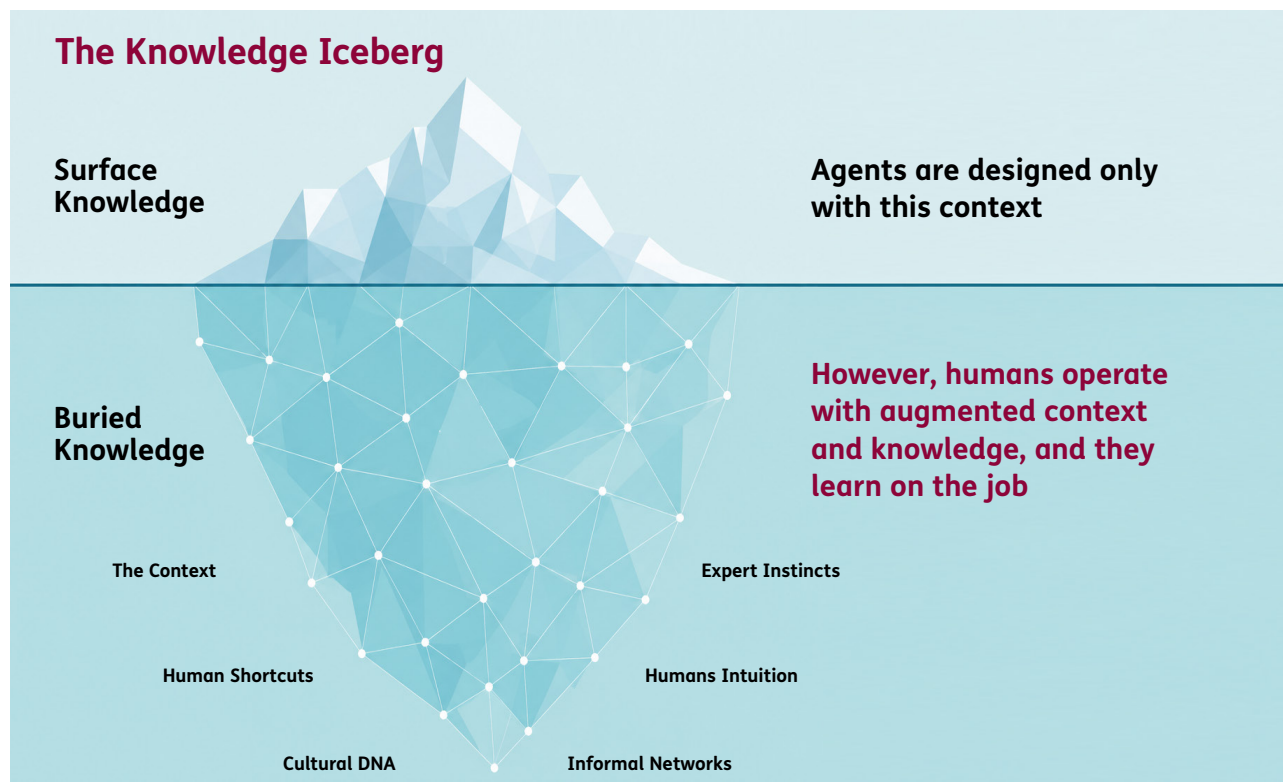
The development of agentic AI opens new ways to capture, maintain, and activate otherwise invisible knowledge. Ultimately, we believe self-evolving agents can significantly reshape how institutional knowledge is scaled and executed.

Table of Contents

The Challenge of Capturing Hidden Expertise	5
The Concept of Self-evolving Agents	7
Self-evolving Agents in Financial Enterprises	10
Self-evolving Agent Use Cases in Insurance	10
Enterprise Readiness	13
Where Enterprises Stand: A Framework for Agent Maturity	16
Case Study: AI Content Engine	18
The Path Forward	22
Bibliography	25

The Challenge of Capturing Hidden Expertise

The vision of autonomous, Self-evolving systems transforming hidden expertise into scalable value is promising. However, this concept entails addressing a few technical but also enterprise structural challenges to close a knowledge gap.



Most knowledge used at work is hidden below the surface. Humans tap into this rich, contextual knowledge while learning on the job, while traditional agents rely only on documented surface knowledge.

Five Key Challenges for Knowledge Capturing

1. Catastrophic Forgetting

When trained on new data, AI models tend to overwrite previously learned information. Without a dedicated memory architecture, continuous knowledge accumulation is difficult (Cai et al., 2026; Ven et al., 2025).

2. Lack of Context/Misconception

Agents often lack clear context on the data they are given. Thus, agents can develop recurring logical errors and a misconception problem (Agarwal et al., 2026).

3. Drift

Business conditions, data distributions and the regulatory landscape change regularly. Previous knowledge becomes outdated and agents may act on this (Gulli, 2025).

4. Expert-to-Engineer Translation Gap

Often, domain knowledge sits with business experts while implementation responsibility sits with technical teams (Zimmermann, 2026c). Both use different terminology, assumptions and success criteria that can get misunderstood in communications. Selecting the appropriate experts for agent training is often a complex task in itself (Zuin et al., 2025).

5. Legacy Constraints

Most enterprises run on legacy systems that are not capable of supporting continuous AI learnings (Lorica, 2026). According to an MIT study by Challapally et al. (2025), 95% of generative AI implementations fail due to a structural learning gap between AI models and the organization they operate in.

Rethinking Human-in-the-Loop

In conclusion, the implementation of agentic AI into enterprises raises the question of technical maturity, organizational structure and trust in the AI (Calipo et al., 2026). Without a doubt, collaboration in the form of “Human-in-the-loop” (HITL) is fundamental to capture and preserve tribal knowledge. However, the traditional HITL setup is not sufficient. Feedback from a single expert is not enough to capture knowledge that is representative of the organization as a whole. The agent needs to be in regular exchange with several experts to learn from and request the needed information or monitor exceptional cases or workflows of its environment (Estrada, 2025; Zimmermann, 2026c).

The Concept of Self-evolving Agents

We define self-evolving agents as autonomous systems that continuously learn on their own through interaction with environments, aiming to perform better while preserving safety. In particular for enterprise settings, the self-evolving agent should start by drawing from the richest source of expertise available. This includes formal policies and procedures as well as tribal knowledge such as shortcuts, preferences, exceptions, decision patterns, and informal ways of working. However, even if tribal knowledge often contains the most valuable context, much of it is rarely documented (Fang et al., 2025; Zimmermann, 2026b).

Principles of Self-evolving Agents by Zimmermann (2026b)

Value from Day One

A self-evolving agent's return on investment does not depend on full autonomy at launch. In the beginning, they can support experts by retrieving relevant documents, assisting with complex reasoning, and making their work transparent.

Later on, as the agent learns, it can gradually take on more autonomous decisions, while continuing to deliver productivity improvements at each stage.

This means self-evolving agents must create measurable productivity gains from day one, even before they reach high levels of autonomy.

Autonomous Evolution

The target architecture should follow a Subject-Matter-Expert (SME)-in-the-loop and developers-out-of-the-loop pattern. Self-evolving agents learn primarily through interaction with domain experts, e.g. using a chat interface, and enterprise context, and not through constant software engineering intervention. If every improvement requires developer involvement, the economic value of self-evolving agents is quickly reduced.

Efficient Expert Feedback

Human attention is scarce, and knowledge workers cannot be expected to become full-time agent trainers. Self-evolving agents must therefore extract high-quality learning signals from minimal expert input. The goal is to move beyond capturing feedback literally and instead extract the underlying expert judgment reason.

Progressive Trust Through Failure

Early-stage agents will not be perfectly accurate, but each failure can become a learning opportunity if it is captured, analyzed, and translated into improved behavior. In this sense, failures become valuable learning signals. We might move to the era where failures are the new gold. Over time, validated corrections and repeated improvements create the basis for stronger trust and greater autonomy in the agent.

Human-Guided and Governed Autonomy

Self-evolving agents should progress toward autonomy under human supervision and transparent governance. Domain experts remain responsible for determining when an agent is ready to move from recommendation to partial execution and, eventually, higher levels of delegation. Similar to onboarding a new employee, agents can be granted greater autonomy as their knowledge, reliability, and task quality improve under supervision. We envision that this trust model (human-agent) will be similar to human-human.

Enterprise Knowledge as an Editable Learning Layer

In-context learning is the most practical and flexible approach for agents to learn and unlearn. Knowledge should not be permanently built into the AI model. Instead, it should be in a separate, editable knowledge layer that the organization can review, correct, and update. This allows agents to continuously build the enterprise's operational "DNA".

Continuous Model Uplift

Because the local knowledge remains in an editable layer, rather than being fixed in the model itself, each model upgrade can improve the agent's reasoning, planning, and interaction quality while continuing to use the organization's validated knowledge base. This separation of model and knowledge layer allows enterprises to adopt the latest GenAI capabilities while preserving their accumulated domain learning.

Self-evolving Characteristics

Based on the stated principles, we derive five essential self-evolving agent characteristics (Fang et al., 2025; Ven et al., 2025):

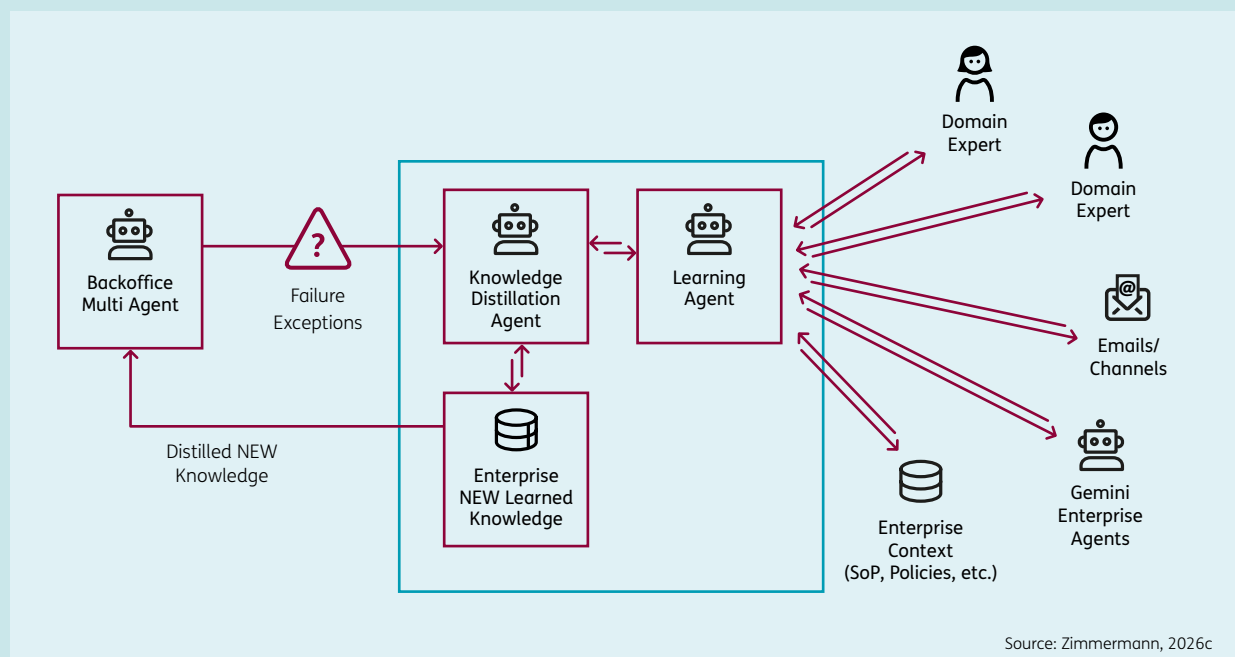
1. Adaptation

The agent must be able to detect and adapt to new environments and situations in near real time, without requiring extensive offline retraining. As conditions, data or user change, the agent should be able to update its behavior autonomously.

2. Exploiting Task Similarity

The agent should be able to recognize when tasks or situations are similar and reuse what it has already learned. This means that learning from one task can help the agent perform better on related tasks or learn similar tasks faster in the future.

How a Self-evolving Agent workflow can look like based on Google's Agent Gym:



Google's Agent Gym is an experimental learning environment in which agents learn from operational experience and are gradually enabled to reach higher levels of autonomy. Here, friction is turned into learnings by a knowledge distillation agent. The unexpected outcomes, exceptions, or incorrect decisions are systematically analyzed and checked against existing rules, policies, and

domain expert input. Afterwards, knowledge is turned and translated into new enterprise knowledge and reintegrated into the operational agent. In this way, a governed learning loop emerges that enables continuous improvement without losing transparency, governance, or human control.

3. Task Agnosticism

The agent should be able to identify the task context without relying on explicit labels from humans or external systems. This is particularly important in dynamic workflows, where tasks may change without clear signals or where no predefined task catalogue exists.

4. Noise Tolerance

Typically, traditional AI is trained on carefully cleaned and prepared datasets. However, data in real enterprise environments can arrive out of context, is incomplete, outdated, duplicated, inconsistent, or stored in different formats. A self-evolving agent must be able to work with imperfect data in daily operations, while still producing safe, stable, and reliable outputs.

5. Resource Efficiency and Sustainability

Self-evolving agents must strike the right balance between performance and resource efficiency. This includes thoughtful approaches for memory management, e.g. selective retraining and prioritization of what knowledge to keep, compress or discard.

What This Chapter Means in Practice

Self-evolving agents require both the right technical characteristics and the right operating principles. They must be able to adapt to changing environments, recognize related tasks, infer context, and handle unstructured data. However, these capabilities only create value when they are embedded in a governed enterprise setup.

This means self-evolving agents should not start as fully autonomous systems. They first support domain experts and over time, extract judgment by learning from feedback, operational decisions, exceptions, and corrections, while relevant learning moments are captured in an editable enterprise knowledge layer.

Because enterprise knowledge is not always stored as clean, queryable data, Self-evolving agents in particular need an infrastructure that feeds them with critical knowledge.

As reliability and task quality improve, supervision can gradually be reduced and more autonomy can be granted. The goal is to extract expert judgment, and make their knowledge scalable for the enterprise. In this way, operational experience can be transformed into governed institutional knowledge.

Self-evolving Agents in Financial Enterprises

Agentic AI has already been implemented by a few insurers, particularly in areas such as claims, underwriting, marketing and sales, customer service, research and development (Munich Re, 2026). Many insurance cases deployed agents that follow predefined rules or execute multi-step tasks independently within a controlled environment, such as a chatbot that guides a customer through a service journey and acts on their behalf (Corradi et al., 2026).

In some enterprise environments, such agents perform well, when rules, processes and outputs are stable and clearly defined. Yet, even this stability has its limits: Google notes that even deterministic, hard-coded flows often require intervention as novel cases and situations arise.

Self-evolving agents become relevant in environments where automation reaches its limits and human judgment is needed to handle complexity and dynamic changes. We define four core conditions where self-evolving agents make sense:

4 Self-evolving Agents Conditions

1. New knowledge is continuously generated, or existing knowledge is drifting, which causes information to be outdated after time.
2. Human experts repeatedly override or augment standard processes.
3. Operational behavior itself is a source of learning.
4. A clear optimization signal or target metric is available to indicate whether new knowledge actually improves performance.

Hence, self-evolving agents are useful, when they can constantly absorb new knowledge, learn from recurring exceptions, and must show impact by measuring whether learning leads to performance improvements.

Since no system can anticipate every current and emerging business scenario, self-evolving agents should not be

designed for absolute autonomy. Their value lies in a governed learning loop, where operational experience is captured, validated, and translated into improved agent behavior over time.

Without meeting these prerequisites, self-evolving agents risk creating unnecessary complexity and governance challenges without delivering tangible business value.

Self-evolving Agent Use Cases in Insurance

This chapter identifies exemplary agentic AI use cases along the insurance value chain defined by the Munich Re (2026) Tech Trend Radar. The potential use cases have been discussed with selected domain experts and show the impact of self-evolving agents that capture and scale tribal knowledge within insurance.

Rather than starting from a technology perspective, we will take a novel approach in identifying use cases: Each case is derived from any of the three questions Q1-Q3 (p. 11) to have the proposed self-evolving agent conditions being met. This ensures that use cases are selected because they address a concrete business problem, create measurable value, fit the operational reality of insurance organizations, or create the highest ROI impact.

Impact of Self-evolving Agents on the Insurance Value Chain



Three Questions to Identify Self-Evolving Use Cases

Q1

Fragmented knowledge

Where in the value chain is tribal knowledge fragmented, manual, or difficult to access?

Q2

Recurring exceptions

Where do recurring exceptions, contextual decisions, or undocumented operational practices shape outcomes?

Q3

Measurable improvement

Where can measurable improvement be achieved through automation, assistance, and continuous learning?

Example for Self-Evolving Use Cases in Insurance

Product Design & Pricing

Product lifecycle knowledge and discontinuation logic

Insurance products are regularly modified, repriced, or discontinued, but the reasoning behind these decisions, why a product was withdrawn from a market, why a coverage extension was removed, or why a segment was repriced out of appetite, is rarely preserved. Self-evolving agents can capture product lifecycle decisions as they are made, building institutional memory around what has been tried, what worked, and what did not, preventing the organization from repeating costly product experiments.

Q2

Underwriting

Senior underwriter knowledge transfer

Senior underwriters carry knowledge of risk intuition that is rarely documented and difficult to transfer. Self-evolving agents can engage directly with these experts to absorb and preserve their underlying reasoning, interpretation and judgment. In a recent study, such systems demonstrated knowledge recall of up to 94.9% (Zuin et al., 2025).

Q1

Q2

Q3

Referral and exception pattern learning

As underwriters assess, refer, or decline risks, they build practical knowledge that helps guide future decisions. Agents that log these patterns build a live, self-updating guide to non-standard risks that can replace ad hoc email threads and undocumented desk notes.

Q1

Q2

Q3

Sales & Distribution

Sales Agent & Broker Coaching at Scale

Successful broker interactions are rarely documented. Self-evolving agents can learn from high-conversion sales interaction, extract successful objection handling, and continuously coach the wider salesforce.

Q1

Q2

Q3

Optimized content engines

The rise of GEO brings a change into the marketing content landscape. Self-evolving agents can internalize branding guidelines and build optimized content for AI search. They can continuously learn which content structures, formats, and messages are surfaced by generative search systems and adapt recommendations accordingly. This enables marketing teams to create content that remains brand-consistent while becoming more discoverable in AI-driven search environments.

Q2

Q3

Renewal risk re-scoring

Renewal underwriting is often manual and experience-dependent, especially when changing claims development, loss patterns, and portfolio shifts require judgment beyond standard pricing rules. Self-evolving agents can learn continuously from these signals to support more consistent renewal decisions, scalable re-pricing, and portfolio steering over time.

Q2

Q3

Tailored Risk Coverage Recommendations

Matching complex commercial risks to the right coverage structures is experience-intensive and depends on judg-

Q1

Q2

Q3

ment accumulated across many placements. Agents that learn continuously from past underwriting decisions can suggest optimal coverage structures for new and non-standard risks, improving consistency and speed of placement.

Policy Administration

Complex query self-resolution

Q1

Q2

Q3

Policy administration teams handle back-office decisions outside standard workflows: endorsement conflicts, mid-term structural changes, and multi-policy interdependencies requiring system-level resolution. Agents learn from resolved cases within policy and billing systems, reducing the need to escalate every non-standard configuration to senior staff.

Risk Management

Regulatory & compliance knowledge capture

Q2

Q3

Regulatory interpretation varies by market, line of business, and individual expert. They change regularly and the agent continuously updates its interpretation framework as new rulings, guidance, or market-specific decisions emerge. While compliance is governed centrally, the knowledge it captures spans the entire value chain.

Customer Engagement

Proactive churn detection & retention knowledge

Q1

Q2

Q3

Experienced retention agents develop informal intuition for which signals indicate a customer may need additional support or a more suitable product. Self-evolving agents continuously learn from resolved retention cases and proactively trigger the right intervention, improving both customer outcomes and retention rates.

Conversational exception handling knowledge base

Q1

Q2

Q3

Customer-facing agents must navigate coverage disputes, mid-term change requests, and multi-policy questions live, in conversation, while managing customer expectations in real time. Self-evolving conversational agents accumulate resolution patterns specifically from these live interaction contexts, continuously improving their ability to handle novel exceptions autonomously without transferring the call.

Real-time agent coaching with live knowledge surfacing

Q1

Q3

Senior customer service agents develop highly effective resolution and communication patterns that junior staff cannot access at the moment. Self-evolving agents learn continuously from top-performer interactions, identifying what distinguishes high-quality resolutions from standard

ones and updating their guidance models accordingly, so that the knowledge embedded in those interactions compounds over time rather than remaining with the individual.

Claims

Claims handler expertise codification

Q1

Q2

Q3

Experienced adjusters apply consistent judgment across standard and mid-complexity cases when interpreting policy coverage and estimating reserves. Self-evolving agents learn from historical resolutions to surface structured decision support for these cases, with human adjusters retaining judgment and sign-off throughout.

Complex liability triage

Q1

Q2

Q3

Large loss and liability claims involve non-routine decisions shaped by jurisdiction, policy history, and negotiation precedent. Agents that learn from comparable resolved cases improve both speed and consistency of complex reserving.

Fraud pattern evolution

Q2

Q3

Unlike static rule-based systems, self-evolving agents adapt dynamically as fraud techniques change, continuously updating detection logic from newly observed patterns within governed validation boundaries.

Data & Analytics

Continuous model validation & pricing signal evolution

Q1

Q2

Q3

Pricing and reserving models are validated quarterly at best, while the risk landscape shifts continuously. Self-evolving agents that are built on a robust data hub can monitor model drift in real time, recalibrate assumptions as loss patterns change, and flag emerging signals without waiting for a scheduled review cycle transform analytics from a periodic activity into a living capability.

Portfolio intelligence

Q1

Q3

Insurance portfolios contain many useful patterns across products, risks, and customer groups. Self-evolving agents can continuously analyze how the portfolio develops over time, helping leadership identify and take strategic decisions.

3 Clusters for High-potential Use Cases

The first is **knowledge preservation**, especially present in underwriting, product development, and risk management, where intuition, judgment, and interpretation built over decades leave with the individual unless an agent captures it first.

To do this, a foundational knowledge base or system must exist as a prerequisite, providing the structured environment into which the self-evolving agent can collaborate with experts and capture, organize, and retrieve accumulated expertise.

Knowledge preservation is focused on making irreplaceable expertise durable.

The second cluster is **operational intelligence**, which works well across sales, policy administration, customer engagement, and claims, where thousands of daily interactions and exception resolutions contain valuable learning signals.

Agents that learn continuously from this operational behaviour compound their capabilities at a rate no structured knowledge programme could replicate.

The third cluster is **continuous recalibration**, across the entire value chain, but particularly valuable in data and analytics as well as risk management. Pricing models, fraud detection logic, and risk thresholds all degrade as the environment changes over time. Self-evolving agents can continuously adjust tools to reflect changing conditions, helping maintain their accuracy and effectiveness over time.

The three clusters show that the most valuable goal of a self-evolving agent in insurance is to collect knowledge from different sources and ensure that this knowledge survives and scales within the enterprise over time.

Enterprise Readiness

To successfully deploy AI in enterprises, companies must understand the technical risks that come with AI implementation, the regulatory landscape and the organizational requirements.

Understand New Threats and Risks

The most critical AI-related risks for financial institutions include persuasive misinformation, hallucinations, discrimination and bias amplification, data security threats, and cybercrime (Hirz et al., 2024).

Self-evolving agents add a further layer of complexity because of their dynamic characteristics and changing behavior over time: they can drift or stagnate as environments, policies, and artifacts change, causing outputs to diverge from business, regulatory, or ethical expectations (Gullí, 2025).

Hence, enterprise readiness requires more than just model performance: it depends on appropriate technical and organizational measures to manage risk, ensure compliance, and build trust (Calipo et al., 2026).

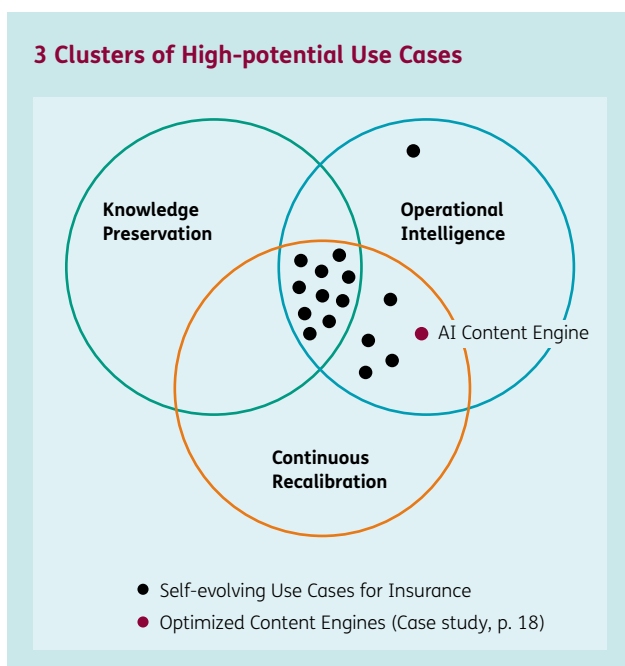
This chapter outlines the key requirements that financial enterprises need to consider before deploying agentic and self-evolving AI systems at scale.

Legal & Regulatory Framework

Over and above established IT Service Management (ITSM) and DevSecOps practices like Site Reliability Engineering (SRE) or the Software Development Lifecycle (SDLC), a large number of regulatory requirements either specify specific implementation details of the above frameworks or apply to the use of AI systems within insurances in a general fashion.

Establishing a coherent approach for enterprise ready AI solutions across domains and jurisdictions is complicated by the fact that some regulations aim at specific industries, like insurances in particular or financial services in general, while others target the use of artificial intelligence or data privacy across industries. In the European Union, the Digital Operational Resilience Act (DORA) or the Insurance Distribution Directive (IDD) are examples for sector-specific legislation.

Similar standards in the US are set by the Federal Trade Commission (FTC). Additional country specific legislation



within the EU or state regulations in the US like those of the New York Department of Financial Services (DFS) further complicate the picture.

Similarly, the EU AI Act, the EU Data Act, the EU General Data Protection Regulation (GDPR), the Colorado AI Act (CAIA) or the California Consumer Privacy Act (CCPA) are examples of domain-specific regulations that apply to all sectors and have significant implications for the use of AI in insurance.

In this highly-regulated, cross-jurisdictional, multi-domain environment it is important to understand not only the technical challenges and potential failure modes, but also the local regulatory differences in order to put appropriate technical and organizational measures (TOMs) in place to manage the risks involved.

Governance & Accountability

Because self-evolving agents continuously learn and change dynamically, oversight needs to be built in from the start, not added on later. This raises a need for transparency: who is accountable when an agent's behaviour shifts or a new risk emerges, and who is responsible if that decision turns out to be wrong.

Oversight works best when it is built into existing risk and control structures rather than set up as a separate process, with a clear focus on identifying high-risk use cases that require independent review. This is especially important given that regulation around agentic AI is still evolving, which means organizations need to actively manage this uncertainty rather than wait for clearer rules.

Identity and Access

Identity and access management deserves particular attention when it comes to self-evolving agents. As these agents grow more capable over time and accumulate more knowledge, a natural question follows about whether they should also be granted broader access to tools and APIs. This raises the further issue of who decides when that access is warranted, and what must be proven first before it is granted. In particular situations, agents require an identity of their own, similar to employees, with clear identification when they access or handle sensitive data.

Observability

Besides the classic structured logs that capture the entire “chain-of-thought” of agents, e.g. which tool is called, which data was received, reasoning for the next steps, confidence scores for decisions (Gullí, 2025) etc., implementing self-evolving agents requires transparency on three criteria:

Exception Observation: The agent recognizes a failure or a case outside its current knowledge or standard process, which is the trigger for the entire learning episode.

Expert Observation: Agents must capture the human judgment extraction process, whether that is a single expert, multiple experts, or it can also be enterprise sources and artifacts.

Learning Observation: As self-evolving agents continuously learn, it is required to capture the actual act of modifying, or integrating that new information into its persistent knowledge base.

Data and Privacy

Once agent systems process sensitive customer, employee, and business information across multiple tools and systems, implementing robust data privacy controls is critical. Clear safeguards prevent unauthorized access, unintended data exposure, or use of data beyond its intended purpose (Albada, 2025).

Agent systems should therefore process only the information required for a specific task, and ensure that personal or sensitive data is protected through encryption, access restrictions, monitoring, and regular deletion of unnecessary logs, cached data, and outputs (Albada, 2025).

KPIs & Measurements

KPIs are necessary measures to tell whether continuous learning is actually making the agent better or quietly steering it in the wrong direction.

To address this, agentic systems need mechanisms for drift and degradation detection (Gullí, 2025), that is, the ability to continuously monitor performance, real-world signals (sometimes hidden from the agent) and trigger corrective action before the system stagnates or deteriorates (Zimmermann, 2026c).

Additionally, the impact of the self-learning agent on applicable, enterprise outcomes should be taken into consideration as well and comparison between old vs new learnings can be measured.

Change Management

BCG's (Bedard & Beauchene, 2026) research outcomes state that 10% of implementation success comes from algorithms themselves, 20% from the technology required to implement them, and 70% from empowering people to capitalize on the technology.

Organizations should therefore implement AI right into valuable enterprise-wide priorities, rather than treating

them as isolated technology projects, i.e. leaders should set strategic priorities and ensure that AI investments directly support measurable business outcomes (Bedard & Beauchene, 2026).

At the same time, successful adoption requires more than just deployment: employees need a clear narrative, trusted guardrails, and targeted upskilling to understand

how AI changes daily work and creates value (Bedard & Beauchene, 2026; Calipo et al., 2026).

Ultimately, companies need an AI-enabled operating model that clearly defines how humans and AI systems work together, including roles, governance, workflows, and responsibilities (Bedard & Beauchene, 2026; Calipo et al., 2026; Gundurao et al., 2026).

What Enterprises Must Have in Place

- Legal & Regulatory Framework
- Governance & Accountability
- Identity Access Management
- Observability Criteria
- Data & Privacy Safeguards
- Clear KPIs and Business Outcomes
- Change Management Plan

What This Chapter Means in Practice

Trustworthy AI deployment can be achieved through disciplined governance, policies, controls, capabilities, and ongoing monitoring. This can include mandatory human review for outcomes, traceable transparency for decisions, and explicit validation steps before decisions are acted upon.

On the technical side, organizations must establish core technical safeguards such as data residency, output filtering, strict access controls, stronger identity verification, sandboxed tool execution, and tightly limited permissions.

After deployment, governance must continue through drift detection, benchmarking, clear performance targets, alignment audits, and safety guardrails. This ensures that agent autonomy can develop within controlled and compliant boundaries while preserving performance and trust enterprise. Like this, operational experience can turn into governed institutional knowledge.

Where Enterprises Stand: A Framework for Agent Maturity

Since “agent” and “agentic” are overloaded terms, we describe the characteristics of agentic AI solutions in a multi-dimensional framework in order to be able to clearly separate different aspects that may or may not be realized in concrete systems.

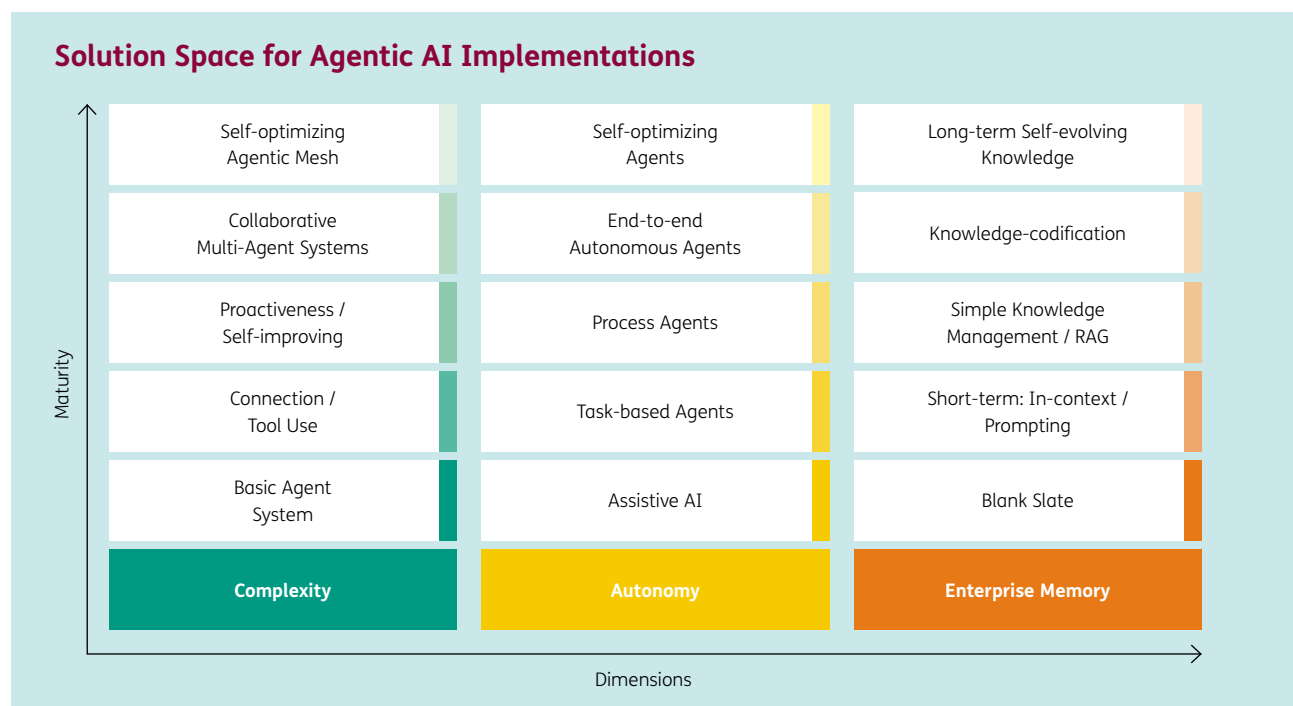
Solution Space for Agentic AI

We use three technical dimensions to describe how “agentic” a system is:

- **Complexity:** How many steps and agents are involved, and how fixed is the workflow?
- **Autonomy:** How heavily is HITL involved in agents’ decision-making and actions?
- **Memory:** How does the agent store and update enterprise knowledge over time?

In addition to the three technological dimensions complexity, autonomy, and enterprise memory (or „statefulness“), there are many additional organizational dimensions like access to sensitive or classified data, the legal or commercial significance of the actions or decisions of the agent or the number of users with access to the services provided by the agent. These additional dimensions are also of crucial importance to the governance of agentic systems, but remain out-of-scope for our current discussion.

Each of the three technical dimensions has its own maturity level and the extent to which each of these are realized in an implementation signifies how “agentic” any concrete AI system is.



A structured view of agentic AI along key dimensions, illustrating the path from basic agents to self-evolving enterprise AI.

Complexity

The complexity dimension describes the degree to which an agentic workflow is pre-determined by the design of the implementation. At the most basic, a single agent could replace a traditional software component by leveraging generative AI to solve a task that would be challenging to implement. More sophisticated agentic solutions leverage several such agents with or without using external tools to achieve a certain outcome. In all of these cases the workflow is pre-determined, likely by a traditional workflow automation tool. At the far end of

the spectrum, the workflow management is also agentic. This leads to an emergent workflow, i.e. the agentic workflow manager decides autonomously which path to take instead of choosing from a limited set of pre-arranged options. Several of these systems may, of course, be combined into a collaborative multi-agent mesh, and later evolve to a self-evolving one. Systems like this provide the high degree of flexibility that allows for the optimization of results that is enabled by self-evolving agentic systems that are described below.

Autonomy

While “autonomy” in relation to the freedom to complete (self-) assigned tasks is associated with agentic systems, and even more so since the arrival of agentic-harnesses, it is neither necessary nor desirable in all cases, or even legal e.g. in the case of the management of critical infrastructure or certain HR and medical tasks prohibited under the EU AI Act.

The least autonomous systems will hence employ purely assistive agents that support human operators much like chat interfaces that have become ubiquitous in the GenAI revolution. In a next step, agents might be used to complete simple tasks, but always with a HITL to verify and approve the agentic actions. More sophisticated systems have an access management system that is risk-based and decide automatically which decisions need HITL and when a fully-autonomous system is acceptable from both commercial and legal/governance point of view.

Enterprise Memory

Memory storage and retrieval mechanisms are essential for agentic AI and in particular self-evolving agents to behave coherently over time. For agents in general, we can distinguish between short-term and long-term memories (Cai et al., 2026; Gullí, 2025; Maragheh, 2025). In between, we are adding RAG and knowledge codification for self-evolving agents.

Knowledge codification describes the stage at which enterprise knowledge is no longer only retrieved from documents or experts, but translated into structured, reusable, and operationally relevant assets. It actively converts that knowledge into a structured, reusable asset, something like a formal Standard Operating Procedure or a written policy document. The value lies in making tribal knowledge explicit, scalable, and easier to govern.

Self-evolving enterprise memory represents a more advanced maturity stage. Here, enterprise knowledge continuously adapts, is refined over time through feedback, interactions, changing policies, and new organiza-

tional experience. Instead of remaining static, the memory layer adapts to new realities while preserving traceability and governance.

Financial enterprise memory needs to evaluate and distinguish different types of memory sources and usages. On the one hand, the question arises on which enterprise-specific data and type of information should be memorized and prioritized, while on the other hand, how a self-evolving agent can unlearn data that is outdated or even classified as unimportant. Looking at the characteristics of a self-evolving agent (Fang et al., 2025; Ven et al., 2025), the agent needs to stay resource efficient and sustainable at all times when it comes to memory.

What This Chapter Means in Practice

For enterprises, self-evolving agents only create value when they are tied to real work and clear outcomes. This means:

Choosing the Right Problems

Use self-evolving agents where knowledge changes frequently, experts often override rules, operations generate useful learning signals, or learning success can be measured.

Focus on Three Value Clusters

1. Knowledge Preservation
2. Operational Intelligence
3. Continuous Recalibration

Governance, Compliance and HITL

Treat expert overrides and exceptional cases as structured learning episodes. New insights must be validated and reintroduced under clear human oversight and imposed regulations.

Done well, self-evolving agents do not replace human experts. They ensure that their knowledge is captured, reused, and continuously improved across the enterprise.

Case Study: AI Content Engine

The insurance industry is currently undergoing a structural transformation when it comes to marketing, brand awareness and customer acquisition due to the rise of AI Search (Schmolke & Hosseini, 2025). Generative Engine Optimization (GEO) impacts the Sales & Distribution value chain of insurance directly (Munich Re, 2025) and is located in the intersection of the Operational Intelligence and Continuous Recalibration use case cluster.

The AI Content Engine

To address this major shift, ERGO, in collaboration with Google Cloud and Deloitte, developed the AI Content Engine.

This multi-agent system automates the end-to-end creation of multimodal marketing assets, comprising text, images, video, audio podcasts, and complete HTML landing pages, from a single primary keyword instruction.

The operational architecture leverages Google Cloud and Gemini models to reduce asset creation cycles from weeks to minutes, aiming to scale content output. The target baseline for this strategic scale is driven by empirical performance values: extensively integrated generative AI architectures reduce overall content production costs by an average of 41%, while allowing enterprises to exceed predefined revenue goals by 22% (Deloitte, 2025).

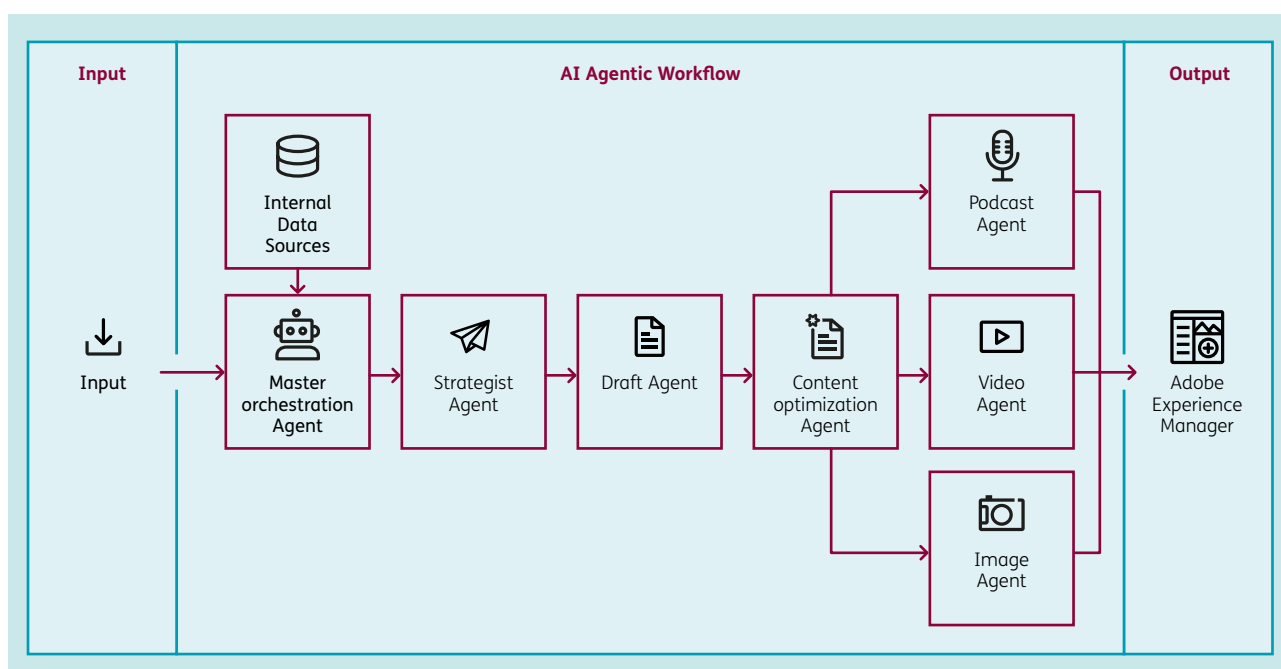
The Architecture

The system architecture is implemented as an agentic workflow deployed entirely on the Google Cloud Platform, while protecting company knowledge and intellectual property.

The process is steered by a central Master Orchestrator Agent. This orchestrator combines the primary user inputs and enriches them via integrated internal data sources.

These internal data sources combine product information, related topics and market knowledge, while hiding the technical processing steps of the processing components.

The combined information is then passed to the main group of agents, which is structurally partitioned into a Strategist Agent, a Drafter Agent, and a Content Optimizer Agent.





Einfach, weil's wichtig ist.

[Versicherungen & Finanzen](#)
[Service](#)
[Kontakt](#)

 Suche
  Berater
  Log-in

 **0800 / 3746 095**
7-24 Uhr (gebührenfrei)

[Home](#) > [Ratgeber](#) > [Zahnimplantate](#)

Zahnimplantate

Ihr Weg zu einem strahlenden Lächeln

Der Verlust eines Zahnes kann die Lebensqualität stark beeinträchtigen. Zahnimplantate bieten eine moderne und dauerhafte Lösung, um fehlende Zähne ästhetisch und funktional zu ersetzen. Erfahren Sie hier alles Wichtige rund um das Thema Zahnimplantate, von den Kosten über den Behandlungsablauf bis hin zur passenden Absicherung durch die ERGO.



Umfassender ERGO Ratgeber zu Zahnimplantaten: Erfahren Sie alles über Kosten, Behandlungsablauf, Vorteile und wie eine Zahnzusatzversicherung Sie finanziell absichert.

[Jetzt informieren](#)
[Beitrag berechnen](#)

- Zahnimplantate sind künstliche Zahnwurzeln, die fest im Kieferknochen verankert werden und eine dauerhafte, ästhetische sowie funktionale Lösung für fehlende Zähne bieten.
- Der Behandlungsablauf umfasst mehrere Phasen, von der sorgfältigen Planung über den chirurgischen Eingriff bis zur Einheilung und dem Einsetzen des endgültigen Zahnersatzes.
- Die Kosten für Zahnimplantate variieren stark und hängen von Material, Aufwand und individuellen Gegebenheiten ab. Eine Zahnzusatzversicherung der ERGO kann die finanzielle Belastung erheblich reduzieren.
- Zahnimplantate verbessern die Lebensqualität durch ein natürliches Kaugefühl und ein selbstbewusstes Lächeln. Sie erfordern jedoch eine sorgfältige Mundhygiene und regelmäßige zahnärztliche Kontrollen.

[Was sind Zahnimplantate und wann sind sie sinnvoll?](#)
[Der Ablauf einer Zahnimplantat-Behandlung: Schritt für Schritt](#)
[Kosten für Zahnimplantate: Was Sie wissen müssen](#)
[Zahnimplantate und Ihre Krankenversicherung: Wer zahlt was?](#)
[Vorteile und Risiken von Zahnimplantaten im Überblick](#)
[Die richtige Pflege für langanhaltende Freude an Ihren Implantaten](#)

Was sind Zahnimplantate und wann sind sie sinnvoll?

Zahnimplantate sind künstliche Zahnwurzeln, die in den **Kieferknochen** eingesetzt werden, um fehlende Zähne zu ersetzen. Sie bestehen meist aus Titan oder Keramik und verwachsen fest mit dem Knochen. Auf diesen Implantaten können dann Kronen, Brücken oder Prothesen befestigt werden, die sich wie natürliche Zähne anfühlen und aussehen.

Sie stellen eine hervorragende Lösung dar, wenn Sie einen oder mehrere Zähne verloren haben. Sie bieten eine dauerhafte, stabile Alternative zu herkömmlichen Brücken oder herausnehmbaren Prothesen. Implantate verhindern zudem den **Knochenabbau** im Kiefer, der nach Zahnverlust oft einsetzt, und erhalten so die natürliche Gesichtsstruktur.

Aufbau eines Zahnimplantats

The AI Content Engine delivers a fully branded ERGO landing page.

These core components direct the downstream production pipelines for image, video, and audio assets, ultimately converging within an HTML Generator module before the complete asset passes through automated security guardrails and is transferred directly to the Adobe Experience Manager production frontend (ergo.de). In the end, a human can confirm or edit the results.

Infrastructure & Governance

The setup is built on a robust, fully managed foundation using Google Cloud Run for hosting, Agent Engine for multi-agent orchestration, and Cloud Storage for asset management. This provides a highly solid starting point for the enterprise scale.

SEO Fundamentals for GEO

Aligning this automated infrastructure with modern search mechanics requires adherence to official search platform frameworks rather than attempting to exploit speculative generative engine optimization trends.

According to Google's official AI optimization guidance (Google, 2026) on maximizing visibility for generative AI features on Google Search, successful indexing and visibility within AI-driven search models rely fundamentally on applying foundational search engine optimization (SEO) practices. This strategy prioritizes the production of valuable, non-commodity content that reflects direct expertise

and deep topical value instead of recycling existing web content.

Ensuring Long-Term Relevancy through Self-evolving Agents

The content engine must ensure high semantic relevance, close alignment with search intent, and structured data delivery, which remain the core criteria for discovery within generative search architectures. Rather than executing search engine-first manipulation, the architecture is tuned to produce people-first information that satisfies user intent while supporting text assets with high-quality, context-aware images and videos.

To ensure long-term relevance, the engine incorporates a self-evolving agent infrastructure that continuously adapts to branding and content output quality automatically from expert feedback.

Testing Phase

The self-evolving agent concept is validated through a methodical testing and validation process designed to demonstrate performance optimization through continuous evolution (Zimmermann, 2026c). The empirical testing design incorporates a sample size of 60 dedicated test cases for the evaluation phase, alongside a prior training phase consisting of 25 training cases designed to baseline the learning mechanics.

Methodology: Measuring State vs. Evolving

To measure changes in system performance, the evaluation methodology freezes and compares two distinct operational states: the Stable Version represents the technological status quo, executing asset generation fully automatically based on the static system prompt.

Conversely, the Evolving Version utilizes the active state of the long-term memory layer, though it remains frozen during the active test execution to prevent live configuration drift or the concurrent opening of secondary feedback channels. This methodology establishes the empirical baseline to capture the Exception Observation, where the agent recognizes a failure or a case outside its current knowledge or standard process, serving as the definitive trigger for the entire learning episode.

SME-in-the-Loop

The operational steering of this evolution follows the principle of the human expert in the loop (SME-in-the-loop). Reviewers inspect the generated landing pages at the end of the pipeline and manually log an Expert Rating using a standardized checklist based on a five-star scale.

This review interface functions as the mechanism for the Expert Observation, where agents must capture the human judgment extraction process. In the beginning this can be a single expert, multiple experts, or enterprise sources and artifacts. The unstructured feedback and manual text edits provided by the editor are analyzed for generalizability, distilled into a machine-readable lesson card, and stored within a Firestore database layer.

At runtime, these lesson cards attach dynamically to the system instructions of the Strategist Agent. This dynamic ingestion executes the Learning Observation, fulfilling the requirement that as self-evolving agents continuously learn, the architecture must capture the actual act of modifying or integrating that new information into its persistent knowledge base without software engineering intervention.

3 Critical KPIs for Success

The quantitative and qualitative performance evolution of the AI Content Engine is measured across three central key performance indicators before and after SME-in-the-loop intervention to monitor the variances between the stable baseline and the evolving configurations. Empirical data gathered over the test cases ($n = 60$) validates the practical efficacy of the self-evolving framework, demonstrating clear operational shifts across the core evaluation metrics:

Topical Depth Index: Measures search visibility and content authority through a data-driven approach that evaluates conceptual relationships and entities rather than traditional keyword density. The target state is defined by the primary keyword and subtopics constructed via search engine result page analysis to identify up to 5 essential subtopics. The individual subtopics coverage scores are assigned a value of 1.0 for complete definition, 0.5 for peripheral mention, or 0.0 for missing gaps.

In benchmarking the two architectures, Agent A (Stable Baseline) achieved an average index of 80.67% with a median of 80.00%. Agent B (Self-evolving Agent) exhibited an advanced performance lift, yielding an average topical depth index of 85.33% and a median of 90.00%. This performance allows Agent B to sit comfortably above the 80% threshold required for strategic GEO readiness.

From a statistical standpoint, a paired t-test reveals that this ~5% performance lift falls just short of strict mathematical significance ($p > 0.05$) due to high volatility (a standard deviation of 18.31% for Agent A and 17.99% for Agent B) driven by extreme individual outliers. Operationally, however, the 90.00% median confirms that the self-evolving architecture consistently dominates the majority of production runs.

This structure ensures that the generated assets comply with official search engine guidelines, which mandate high-value, people-first information demonstrating direct domain expertise for discovery within AI-driven search features, rather than relying on algorithmic manipulation (Google, 2026).

Expert Rating: Evaluates general text quality and stylistic compliance based on direct internal reviewer ratings from one to five stars recorded in the task inbox, anchoring the core framework criteria of human responsibility and explicit accountability (Deloitte AI Institute, 2026).

Reflecting a clear statistical correlation ($r \approx 0.68$), the analysis establishes that enhanced topical coverage directly corresponds to superior human evaluation. Consequently, internal reviewers awarded Agent B a higher average score of 3.25 stars compared to the 3.12 star baseline, established by Agent A, citing its capacity to comprehensively address sophisticated subtopics over spinning generic marketing copy.

Readability Score: Tracks shifts in textual clarity using localized Flesch Reading Ease scripts to ensure that automated stylistic variations do not degrade comprehension, satisfying the requirement for explainability and technical algorithmic reliability (Deloitte AI Institute, 2026).

Both systems demonstrated exceptional stability, maintaining consistent readability marks of 29.18 for Agent A and 29.35 for Agent B.

This uniform performance indicates that the architecture successfully integrates dense, complex insurance terminology without generating unreadable walls of text. While the standard Flesch script is not entirely optimized for native German syntactic structures, it remains mathematically valid for this study as a means of assessing relative variance between the two testing states rather than absolute linguistic readability.

The empirical evaluation of the AI Content Engine demonstrates that transitioning from a static multi-agent setup to an adaptive, self-evolving architecture yields clear operational value.

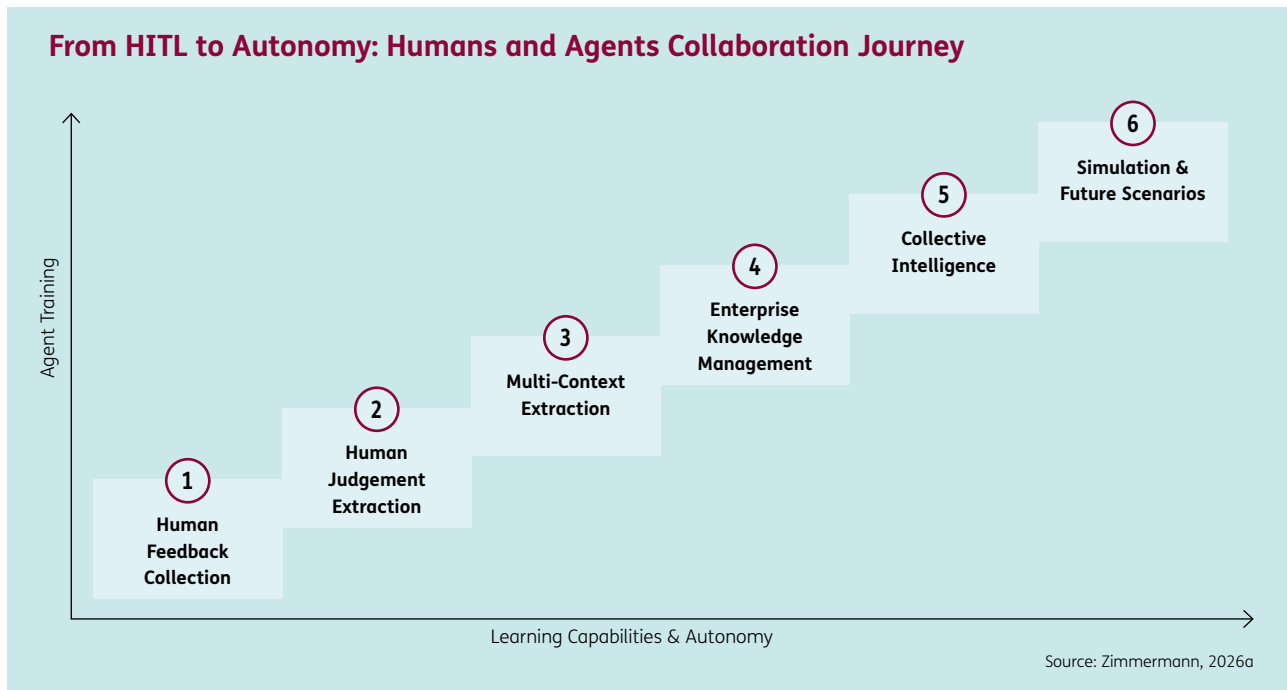
By establishing a structured SME-in-the-loop validation process, the engine achieved an average Topical Depth Index of 85.33% and a 90.00% median, crossing the critical threshold required for generative engine optimization readiness. A strong statistical correlation ($r \approx 0.68$) proves that automating the internalization of expert edits into structured lesson cards directly improves qualitative results. However, the susceptibility to contextual drift, evidenced by individual performance drops to 30%, highlights the permanent necessity of human oversight and strict enterprise guardrails.

Ultimately, the case study validates that continuous, in-context learning loops successfully transform transient tribal knowledge into scalable and auditable corporate assets without requiring ongoing software engineering intervention.

KPI / Metric	Agent A (Stable Baseline)	Agent B (Self-Evolving Agent)	Key Operational Insights
Avg. Topical Depth Index	80.67%	85.33%	Agent B demonstrates advanced content authority, sitting comfortably above the 80% benchmark required for generative engine optimization (GEO) readiness.
Median Topical Depth	80.00%	90.00%	The 90.00% median proves that the self-evolving agent dominates the vast majority of production runs despite strict statistical significance falling slightly short ($p > 0.05$).
Avg. Expert Rating (1–5 Stars)	3.12	3.25	Backed by a strong statistical correlation ($r \approx 0.68$), internal reviewers gave higher ratings to Agent B because it covers complex subtopics instead of generic copy.
Standard Deviation (Volatility)	18.31%	17.99%	High volatility is driven by extreme individual outliers, such as Agent A dropping to 10% or Agent B dropping to 30% due to context drift.
Avg. Readability Score	29.18	29.35	Both agents show steady readability, proving the models successfully integrate deep domain knowledge without creating dense walls of unreadable text.

The Path Forward

Looking ahead, agentic autonomous systems will evolve not in a single leap, but through a series of stages in which machine learning and human expertise become increasingly integrated.



The Staircase Journey to Collective Intelligence

Early deployments will centre on collecting human feedback that flows back to R&D for prompt optimizations. Later on, human judgment extraction is where agents engage directly with domain experts to surface the reasoning behind decisions rather than simply recording outcomes.

As these systems mature, they will draw on broader organizational context, consulting multiple experts, enterprise policies, and external knowledge sources simultaneously. Here, expertise is systematically observed and captured.

Over time, agents will evolve a useful knowledge management capability, that is able to append, update, and retire its own knowledge entries in response to new information, regulatory change, or shifting business priorities.

This progression represents one of the most significant shifts in how financial enterprises will organize and preserve institutional expertise over the coming years, and understanding where a given organization sits on this

staircase is an essential first step in planning any agentic AI programme.

From there, knowledge becomes a collective intelligence system: agents become interconnected across workflows, allowing learning agents to share expert pools and draw on each other's learning rather than operating in isolation.

Beyond Collective Intelligence

The self-evolving agents described throughout this staircase learn primarily from expertise held by experts within the enterprise. This is a powerful and currently underexploited source of learning, but it is not an inexhaustible one.

As the self-evolving agent continues to observe expert overrides, exception resolutions, and recalibration decisions, it will approach the limit of what its human experts can teach it.

At this stage, the agent may begin to use the captured knowledge as the foundation for a more advanced capability: in order to continue improving, it must begin to

generate additional learning opportunities itself, for example by simulating future scenarios, creating synthetic cases, generating novel concepts, and testing alternative strategies in controlled environments.

In doing so, the agent moves from absorbing existing expertise to extending it through structured exploration. Early research from Google's Agent Gym already points towards a future in which agents become the "student outsmarting the master" in certain domains.

Learnings

A practical takeaway from deploying self-evolving agents in enterprise settings is that the strongest initial ROI may come from an agentic overlay rather than from changing core systems, SaaS platforms, or systems of record.

In this model, agents sit on top of existing processes and reduce repetitive knowledge work performed by junior and mid-level employees. Examples of tasks often involve structured reasoning, document review, comparison, preparation, and recommendation, but do not require the strategic judgment of senior executives yet.

As a result, self-evolving agents can more realistically encode and improve entry- and mid-level knowledge work before expanding toward higher-complexity decision domains.

Final Reflections

Tribal knowledge has long been one of the least visible assets in organizations. Self-evolving Agents change the way this knowledge can be recorded and used, but it also requires stability, responsibility and control in an enterprise environment.

In that sense, the real innovation lies not only in more capable models, but in how we design the long-term interaction between human expertise, organizational knowledge, and adaptive AI systems.

Authors

Thuy Emilia Dang

Innovation Developer
ERGO Group AG

Michael Zimmermann

Office of the CTO
Google Cloud

Tim Oberbauer

Manager
Deloitte

Dr. Matthias Lein

Head of AI Foundations
ERGO Group AG

Acknowledgement

This white paper was made possible by the support, constructive ideas and professional feedback of numerous other contributors. We would therefore like to express our sincere thanks to all those who have made a significant contribution to the creation and improvement of this white paper through their valuable technical input, careful proofreading and constructive feedback.

Special thanks to our colleagues who led and built the AI Content Engine. Their work formed the outcomes that enabled the insights presented in this paper:

Uwe Scharrer

Search Marketing Expert
ERGO Group AG

Dr. Raphael Kronberg

Tech Lead AI
ERGO Group AG

Florian Barth

Manager
Deloitte

Manpreet Kaur

Senior Consultant
Deloitte

We would also like to thank the following individuals for their support:

Hanbing Ma

Head of Innovation
ERGO Group AG

Maximilian König

Director
Deloitte

Georgina Neitzel

Head of Innovation Lab
ERGO Group AG

Fynn Gebers

Insurance Industry Executive
Google Cloud

Image credit: Some images in this brochure were generated with the assistance of AI.

Bibliography

- Agarwal, S., Biswal, A., Zeighami, S., Cheung, A., Gonzalez, J., & Parameswaran, A. G. (2026). Arming Data Agents with Tribal Knowledge (Version 2). arXiv. <https://doi.org/10.48550/ARXIV.2602.13521>
- Albada, M. (2025). Building Applications with AI Agents: Designing and Implementing Multiagent Systems. O'Reilly Media, Inc.
- Bedard, J., & Beauchene, V. (2026, February 4). AI Transformation Is a Workforce Transformation [BCG]. <https://www.bcg.com/publications/2026/ai-transformation-is-a-workforce-transformation>
- Cai, Y., Hao, Y., Zhou, J., Yan, H., Lei, Z., Zhen, R., Han, Z., Yang, Y., Li, J., Pan, Q., Huai, T., Chen, Q., Li, X., Chen, K., Zhang, B., Qiu, X., & He, L. (2026). Building Self-Evolving Agents via Experience-Driven Lifelong Learning: A Framework and Benchmark (arXiv:2508.19005). arXiv. <https://doi.org/10.48550/arXiv.2508.19005>
- Calipo, A., Winklhofer, R., Sattelhak, D., & Barbarella, F. (2026). A moment to lead: The foundations Asia Pacific life insurers need to scale agentic AI with confidence [Whitepaper]. Deloitte. <https://www.deloitte.com/content/dam/assets-zone1/apac/en/docs/collections/2026/agentic-ai-scale-for-life-insurers-in-apac.pdf>
- Challapally, A., Pease, C., & Chari, P. (2025). The GenAI Divide: State of AI in Business 2025. MIT Nanda. https://mlq.ai/media/quarterly_decks/v0.1_State_of_AI_in_Business_2025_Report.pdf
- Corradi, D., Durmus, S., Franco, G., Khoury, J., Mishra, S., & Koh, E. (2026). Competitive Advantage in the Agentic AI Era. <https://www.bcg.com/assets/2026/white-paper-competitive-advantage-in-the-agentic-ai-era-ap-insurers.pdf>
- Deloitte. (2025, December 16). GenAI is ready now. Are your marketing operations? Getting the most out of human-AI collaboration. <https://www.deloittedigital.com/us/en/insights/research/genai-human-marketing-operations.html>
- Deloitte AI Institute. (2026). The AI Dossier: A selection of high-impact use cases across six major industries. <https://www.deloitte.com/nz/en/services/consulting/research/artificial-intelligence-dossier.html>
- Estrada, S. (2025, August 18). MIT report: 95% of generative AI pilots at companies are failing. FORTUNE. <https://fortune.com/2025/08/18/mit-report-95-percent-generative-ai-pilots-at-companies-failing-cfo/>
- Fang, J., Peng, Y., Zhang, X., Wang, Y., Yi, X., Zhang, G., Xu, Y., Wu, B., Liu, S., Li, Z., Ren, Z., Aletras, N., Wang, X., Zhou, H., & Meng, Z. (2025). A Comprehensive Survey of Self-Evolving AI Agents: A New Paradigm Bridging Foundation Models and Lifelong Agentic Systems (Version 2). arXiv. <https://doi.org/10.48550/ARXIV.2508.07407>
- GDV. (2026, March 25). Demografischer Wandel: Versicherer setzen auf Digitalisierung. <https://www.gdv.de/gdv/themen/digitalisierung/demografischer-wandel-versicherer-setzen-auf-digitalisierung-197170>
- Google. (2026). AI Optimization Guide. <https://developers.google.com/search/docs/fundamentals/ai-optimization-guide?hl=de>
- Gulli, A. (2025). Agentic Design Patterns: A Hands-On Guide to Building Intelligent Systems. Springer Nature Switzerland. <https://doi.org/10.1007/978-3-032-01402-3>
- Gundurao, A., Krishnakanthan, K., Walsh, R., Kaniyar, S., & Catlin, T. (2026, April 29). Can agentic AI (finally) modernize core technologies in insurance? [McKinsey & Company]. <https://www.mckinsey.com/industries/financial-services/our-insights/can-agentic-ai-finally-modernize-core-technologies-in-insurance>
- Hirz, J., Kivisaari, E., Miehe, P., Senatore, C., Tautan, B., & Toraldo, F. (2024). What should an actuary know about artificial intelligence? [Paper]. Actuarial Association of Europe. <https://actuary.eu/wp-content/uploads/2024/01/What-should-an-actuary-know-about-Artificial-Intelligence.pdf>
- Lorica, B. (2026, May 19). Stop upgrading your LLM. Start fixing your data. Gradient Flow. https://gradientflow.substack.com/p/stop-upgrading-your-llm-start-fixing?img=https%3A%2F%2Fsubstack-post-media.s3.amazonaws.com%2Fpublic%2Fimages%2F-9da51d32-ab7e-4f94-bcd1-0c46f4cfbe4a_1345x1020.jpeg&open=false
- Maragheh, R. Y. (2025). The Future is Agentic: Definitions, Perspectives, and Open Challenges of Multi-Agent Recommender Systems. University of Illinois Urbana Champaign. <https://arxiv.org/html/2507.02097v1>

McKinsey. (2026, April 15). Boomers and the business baton. <https://www.mckinsey.com/featured-insights/week-in-charts/boomers-and-the-business-baton>

Munich Re. (2025). 2025 Tech Trend Radar. Munich Re. <https://www.munichre.com/en/company/innovation/tech-trend-radar-2025.html#1278543003>

Munich Re. (2026). 2026 Tech Trend Radar (From Trends to Transformation at Scale - What Matters Now? Discover Our Expert Perspective for Insurance). Munich Re. <https://www.munichre.com/en/company/innovation/tech-trend-radar-2026.html#1278543003>

Polanyi, M. (1958). Personal Knowledge: Towards a Post-Critical Philosophy [Digital]. Routledge & Kegan Paul Ltd. <https://martin.kcvs.ca/phil334/papers/polanyi-m-personal-knowledge-towards-a-post-critical-philosophy.pdf>

Polanyi, M. (1996). The Tacit Dimension. Routledge & Kegan Paul.

Schmolke, L.-M., & Hosseini, H. (2025). From SEO to LLM Discovery: Insights from the Insurance Sector to Navigate the Next Chapter of Search. <https://www.ergo.com/de/newsroom-ergo/medieninformationen/2025/20250627-ergo-whitepaper-llm-search>

Statistisches Bundesamt. (2025, September 3). 12.9 million economically active people will reach statutory retirement age in the next 15 years. Federal Statistical Office. https://www.destatis.de/DE/Presse/Pressemitteilungen/2025/08/PD25_N048_13.html

Ven, G. M. van de, Soures, N., & Kudithipudi, D. (2025). Continual Learning and Catastrophic Forgetting (pp. 153–168). <https://doi.org/10.1016/B978-0-443-15754-7.00073-0>

Watts, M. (2026, June 3). Documenting Tribal Knowledge: Ensuring Continuity in Operations. ISIXSIGMA. <https://www.isixsigma.com/dictionary/tribal-knowledge/>

Zimmermann, M. (2026a). HiTL to Autonomous [Google Presentation].

Zimmermann, M. (2026b, April 6). Agent Gym First Principles [Google Doc].

Zimmermann, M. (2026c, April 23). Self-evolving agents: Scale domain expertise with in-context learning. <https://www.youtube.com/watch?v=bRy7bW41Pow>

Zuin, G., Mastelini, S., Loures, T., & Veloso, A. (2025). Leveraging Large Language Models for Tacit Knowledge Discovery in Organizational Contexts (arXiv:2507.03811). arXiv. <https://doi.org/10.48550/arXiv.2507.03811>