

Multi-Agent Advisory Networks: Redefining Insurance Consulting with Collaborative Agentic AI Systems

¹Balaji Adusupalli,

DevOps architect - Small commercial Insurance , Balaje.adusupalli.devsecops@gmail.com, ORCID: 0009-0000-9127-9040

Abstract

AI was designed to make human lives easier. However, as more data became available and neural networks became more complex, the technology became prohibitive rather than helpful. Today, to utilize untamed pre-trained machine learning models, an AI research group would require at least a few hundred GPUs to handle hundreds of billions of parameters. To work well, large machine learning models require massive amounts of computational, financial, and energy resources.

The Multi-Agent Advisory Network (MAAN) is an AI model with supervised learning that represents an answer to the expensive requirements of pre-trained models. MAAN can maintain, optimize, share, and combine its knowledge, resulting in high-quality collaborative performance across diverse tasks. MAAN uses a splitting neural network that splits its processing and memory demands among multiple parallel agents. This means that each agent is smaller and lighter and that their larger capacities are unlocked via multi-agent collaborative learning. In contrast to most multi-agent research that leverages evolution or reinforcement learning techniques for training agents, MAAN uses easy-to-scale widely used supervised learning. While MAAN uses zero reinforcement learning techniques for multi-agent cooperative learning, it does not use zero for joint learning, reducing the overall computational demands. This allows MAAN to navigate the concept of the scalable quality generalization of smaller and easier-to-train image and language processing models. In summary, MAAN is a lightweight, flexible, agile, and efficient deep learning model based on a loose and parallel-based agent architecture.

Keywords: AI, Neural Networks, Machine Learning, Multi-Agent Systems, Supervised Learning, Computational Efficiency, Pre-Trained Models, GPU Requirements, Collaborative Learning, Splitting Neural Network, Parallel Agents, Scalable AI, Reinforcement Learning, Zero Reinforcement Learning, Model Optimization, Memory Efficiency, Image Processing, Language Processing, Scalable Quality Generalization, Deep Learning.

1. Introduction

The increasing availability and use of electronic records, together with recent advances in artificial intelligence and machine learning, have created unprecedented opportunities in cognitive systems for solving complex decision problems. Driven by these opportunities, research in AI and machine learning has made significant progress toward creating artificial agents that can understand, reason about, and interact with complex, real-world environments to solve problems important for individuals, organizations, and society. These agents often embody uniquely deep knowledge and skilled problem-solving capabilities for specific interests. Hence, it is increasingly common that AI problem-solving agents are consulted by or collaborate with a group of human problem-solving

agents. For example, leading-edge AI systems win games, answer routine actuarial questions, and partner with physicians to provide diagnostic recommendations, all while interacting with human counterparts.

One promising way of aligning the information used to train AI systems with human needs is by mining electronic records from these large organizations and by creating AI platforms that learn and update their knowledge of the organization's identity. The use of electronic records in client-specific contexts allows AI agents to update policies and make predictions that are uniquely tailored to individual clients. This paper constructs a framework and proposes implementations to develop AI advisory systems for the insurance consulting practice. The goal is not to replace insurance consultants, but to add to the value they deliver to clients: while human consultants specialize in building insights across a wide range of industries and expose clients to a broad set of considerations, the AI's ability to provide innovative recommendations unique to the client context using large volumes of internal insurance records is high. The framework can be implemented on large loss models and for severity modulations of simulations in other areas of insurance.

1.1. Overview of the Study Framework

This architecture provides a comprehensive guide to real-world systems with their incorporation of individual artificial agents, who would coordinate together to surface complex patterns and trends in collective behavior. The derived technology, Multi-Agent Ensembles, is capable of supporting the underwriting and brokerage side of the insurance continuum in performing multi-functional tasks like advisory, concurrency control, and knowledge acquisition in dynamically complex, uncertain, and distributed environments. Our approach is illustrated with detailed real-life experimentation using property and casualty insurance data.

We propose the concept of a Collaborative Multi-Agent Artificial Network that would address data and decision complexity, lack of decentralized control, and knowledge-centric cooperation in insurance organizations. The Generalized Multiple Agent Framework, based on the interplay of AI, mathematical data mining, computational risk theory, and meta-heuristic search techniques, operates as a manager of several autonomous, adaptive, and distributed AI working models. Distinct individual AI models could integrate strategies including combining diverse expert knowledge and judgment, modular reasoning, representation of inconsistent empirical measurement probabilities, and rule-based uncertainty. Such a system softens the real-world complexities of the insurance domain, in which some system instances must work independently and yet need to be employed, aggregating at the task of interest, namely decision-making under uncertainty.

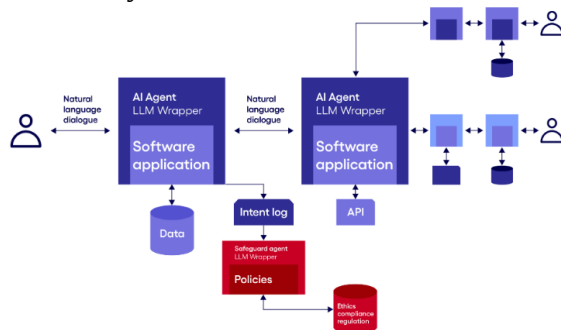


Fig 1 : Multi-agent AI is set to revolutionize enterprise operations

1.2. Research Objectives and Questions

Given the potential of recent advances in AI technologies for digitally augmenting the availability, range, and quality of advisory consulting services, we set out to explore and investigate the hypothesis that AI technology can be used to redefine the delivery of insurance consulting services. However, we believe that the research perspective should be broader to explore both the potential contributory properties of these digitally augmented AI tools and their dynamic interaction within the broader context of organizational and situationally defined man-machine processes within real professional consulting experiences. Subsequently, the research question involves an exploration of AI structures designed to incorporate the utilization of systems of many AI tools, including examining and interpreting analysis structures optimized to understand the areas in which such systems may perform well or poorly in real client consultation settings.

2. Background and Motivation

Today's insurance consulting is far from perfect. Insurance consultants and operations personnel readily acknowledge that there exists a massive knowledge gap in the industry. Insurers, as well as their clients, frequently ask for advice on how to best proceed, but where can they turn for reliable, up-to-date, on-demand assistance? Insurance carriers exclusively generate industry-summarizing documents. Insurance brokers answer questions with their best educated guesses. The rating experts rarely deal with client-specific situations. There currently exists no single source of granular or macro-level, real-time, comprehensive, consultative information or services for the actual insurance industry. Insurance professionals currently obtain shared or external access to pertinent information through time-consuming searches and contact with industry-wide peers or industry-specific data repositories. Unfortunately, the high level of noise in all topics of interest and the increasing complexity of the documents further exacerbate these tasks. With current technology, organizations need a team of thousands to coordinate and provide industry-summarizing consultative information.

In days of yore, government agencies had the same information-providing challenge—provide a single source of comprehensive, consultative information while maintaining current information, quality, relevance, security, and privacy. They solved the problem by constructing relevant, collaborative, information-providing agents. Those agents then created the first operational artificial intelligence systems and the machine learning required to make these systems competitive and of human quality; finally, they functioned as integrated parts of the government organizations' information delivery systems, supporting operations and not just providing summative or descriptive data. The proposal is the introduction of similar truly collaborative, familiar flavored, functionally integrated agents to address the consulting that insurance companies need so badly.

Equation 1 : Collective Decision Optimization in Multi-Agent Systems

$$O_t = \sum_{i=1}^n w_i U_i$$

where
 O_t = Overall system optimization,
 w_i = Weight assigned to agent i ,
 U_i = Utility function of agent i ,
 n = Total number of agents.

2.1. Key Drivers of Change in Insurance Consulting

Drawn from the above-discussed context and looking into currently addressed problems by insurance consultancies, four main factors have been identified to significantly reshape the insurance consultancies from present to future: expertise, especially in rapidly expanding unconventional coverage; custom coverages, especially for unconventional risks that promise a larger loss; dynamic pricing responsive to high-volume, low-value risks; and integration and coherence assurance in large future M&A transactions, even for massively complex firms' organic growth. Before further explaining these factors in the next section and substantiating their value, it is crucial to connect the dots from these factors that are redefining insurance consulting. A triad of the following questions: Why are expert consultancies valued? Why insurance? Why AI?

The convergence of the expanding expertise to remain best differentiable in an increasingly feature-barrier market, the rise of unconventional risks and associated custom solutions to make "larger loss" events less frequent or less expensive, the need for a broader and deeper transfer of insurance capital for high-volume, low-value risks, and the expansion of particularly complex and predominantly large firms are key underlying forces driving current tensions and disruptions in the business model of the conventional insurance industry and the consulting industry it supports.

2.2. Factors Influencing Transition in Insurance Consulting

What are empirical drivers facilitating collaboration and enacting value exchange? Provide a brief picture of a multi-agent decision-making model for a network, involving intelligent mobile agents linking multi-user elements with information, knowledge, and expertise exchange functions. Describe the advisory web proposed as the proto-network.

The intensity of relations with each relationship type and its indicators are found.

There is evidence of acceleration in the process of integration in many kinds of business networks, including services such as providing consulting services on insurance. Thus, insurance consulting firms are facing significant changes and challenges. Insurers too are facing a rapidly changing and deregulated market. They must respond quickly to competitive forces and should adapt to changing economic conditions. They are concerned with expanding markets through product differentiation and value-added services. To provide these value-added services, insurers hire consulting firms.

The value of insurance consulting services is based mainly on the expertise of the insurance consultant. The purpose of the study is to understand how integration can be accelerated between key stakeholders in the relationship to create collaborative teams more suited to the consultants' operating model. Furthermore, such an understanding can suggest procedures for consultants determining strategy and facilitate better results for all intervention agents.

3. Conceptual Framework

1. Definition of Advisory Network A multi-agent-based information system dedicated to providing consulting services in specific fields is called an advisory network. An advisory network refers to a collection of advisors, underwriting experts, or consultants, who concentrate their expertise, talent, and advice on different problems. We denote as an advisor the abstraction of an agent specialized in a task that makes it capable of providing particular advice or guidance over relevant matters to the user of a service. An advisor can be a professional who provides expert advice in a particular area, which is commonly associated with subject matter experts, including management consultants, academic advisors, educational, technical, scientific consultants, financial planners, insurance underwriting experts, foreign advisors, marketing consultants, political consultants,

lifestyle and health advisors, etc. An advisory network can address both general advising in any context and scenario and offer advice according to specific advising issues. The concept of advising or advisory is tightly associated with the process of providing counsel in specific fields and activities, the activity of lending advice, and the discipline that studies the principles of advising.

2. Advisory Network Definition and AI Confluence An advisory network system is an information technology service that offers professional consulting to constituents, covering a wide range of policy areas and subjects. The objective of the advisory network is to provide credible, high-level, important guidance to facilitate and ensure informed decision-making. The purpose of creating advice-seeking systems is to provide technological assistance to different actors in the search for professional guidance. It has been developed by taking advantage of significant developments in knowledge acquisition as well as in expert system reasoning strategies and by making use of modern communication and advanced display capabilities. Information system consultants are a type of expert system that addresses issues associated with developing information systems. The emergence of artificial intelligence as a field of research and development was plagued during the 20th century by infighting between the diverse species in the AI field.

3.1. Defining Multi-Agent Systems

A multi-agent system (MAS) is a set of intelligent agents that interact to solve problems that are beyond the capabilities of any single agent. A distributed problem has to be split into sub-problems, each of which is potentially better solved by different agents. In typical agent problem domains, it is the nature of real-time data intake that makes it difficult for a centralized system to scale up to address a worldwide audience that exceeds billions of claims and adheres to a variety of human language activation modalities, which adds yet more complexity to the system. While tailoring to collectivity, the importance of a user-specific approach is not diminished.

The importance of interaction between agents and environments is emphasized. The agents act in environments of interest; therefore, a generalized agent design replaces the environment with a copy of itself so that the design can be applied to any agent-environment system. In turn, an environment contains other agents that interact with the design, allowing flexibility in the use of agent models. As the environment has autonomous features that the design or concept of each agent does not have, the design will have to observe at least some other agents and at least some other similar designs that have knowledge that the agent can learn from. To put it uniquely, an environment is everything in the universe of the agent in existence that is not in the agent. An example of this previously would be language-word learning.

3.2. Understanding Advisory Networks

An advisory network is a network of agents designed to take over the functionality of a specific consultancy task. For instance, it can be a network of doctors for diagnosis, insurance agents, or financial planners for risk and financial advice. Here, we are particularly interested in applying multi-agent systems to convey the expertise needed to aid and consult on risk and well-being assessment tasks that take place during insurance consulting. Below, we summarize the three kinds of extensions that are characteristic of the insurance consulting task, making it different from other multi-agent tasks, and also what the key consequential effects are from each characteristic. Upon doing so, we aim to define a novel consulting problem to be tackled by an advisory network, combining various other bridges and technologies developed for different purposes, while

preserving its distinctive properties and enabling it to interact effectively with human consultants naturally and productively.

Firstly, consultants and clients often wish to establish and maintain some confidence in the hidden network and evaluate the impacts of expected future changes. These properties are essential for an agentic advisory network to be both beneficial and frequently used. Secondly, the adverse moral hazard and adverse selection related to various contractual relationships and consultancy tasks should be studied in a different light from a variety of design spaces. Thirdly, in consultancy tasks, the consulting process and result evaluation tend to be of paramount importance. Whether aiming to create a knowledge base, work together with existing human consultant-broker teams, or establish an independent consulting ecosystem, it is key that we address all those or related concerns.

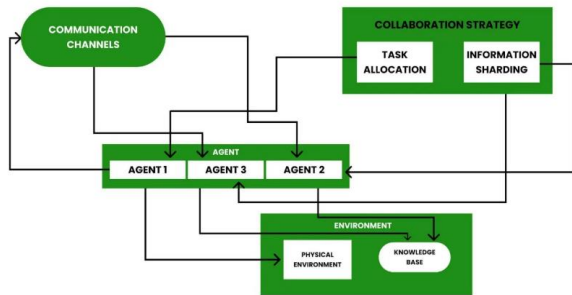


Fig 2 : Understanding Multi-Agent AI Frameworks

4. Literature Review

A picture is emerging in the literature that describes the possible far-reaching impact of multi-agent AI on organizations and businesses. There are claims that AI systems will be able to perform many tasks that are currently performed by humans, and therefore, a significant number of human roles will be replaced. However, there is also much research highlighting that AI systems do not exhibit the same capabilities as humans and are not designed to operate in the ways that humans do. Whether AI systems are reaching human-level intelligence either already or on the near-term horizon recently appears debatable, or at least less clear-cut than previous publications suggested. Understanding how multi-agent and other AI systems will develop over time is therefore of wide interest and importance. Consideration of the potential for multi-agent AI is a part of that landscape. Understanding the limitations and opportunities of AI from a business and strategic viewpoint is therefore a very important aspect of AI strategy.

Much less attention has been paid by academia to the impact of AI on advisory roles and the implications of either turning advice into data and then transacting with that data or transforming advisory roles through AI-enabling processes and technologies. The purpose of this paper, therefore, is to explore the emerging and potential impact of AI on the service sector, using insurance as a case study, where the service is the provision of risk management and insurance advice. Data from a novel design science approach was used in implementing a prototype and real-world learning about natural language chatbot-driven advisory systems, seeking to capture flows of exogenous impacts on design and understand client behavior. Additionally, data from a more traditional literature review and desk research is presented.

4.1. Existing Models of Insurance Consulting

To understand the existing and potential models of Insurance Consulting (IC), we first need to understand the traditional definition of IC, since that is what we are endeavoring to redefine with MAAIS. IC has been a core function in the insurance industry since its inception. Insurance consulting is fundamentally problem-driven, and there are three major problem scenarios where insurance consulting expertise is required. Scenario 1: Specialist Risk Expertise is required to advise insurance underwriters or brokers in the screening, selection, action, or placement components of the compliance and fortification cycle. In most insurance consulting situations, it is claimed that the IC capital has the knowledge that the broker needs, but the broker ultimately makes the decision to sell the insurance product. Though the broker may be part of a sizable broking consultancy with some selling skills, they are not a specialist in identifying and valuing risks. This specialist, who practices the art of persuading underwriters to accept uncovered risks, cannot practically sell protection without extensive training, which could be up to 20 years. This knowledge is held by a small number of specialist brokers. Scenario 2: Special Industry Risk NGO or Government Entity help is required in existing insurance or non-insurance disputes, claims, accumulation, or market-making negotiations, valuing either individual or group property insurance covers, typically covering more impoverished risks, such as parametric or microinsurance. Once more, this is not the broker's game, as very different consulting skills and techniques are required. Scenario 3: A basic technical or more holistic, near mini-relational, security or risk evaluation and advisory report will be required. The first purchase situation occurs with the low-cost promoter buttons affixed to spaces in standardized insurance schedules. The registered conclusion "value listed" declares the sum and pulse of insurable interest provided when promoting a business or other property, but brokers rarely acquire it until the due diligence stage of the transaction.

4.2. AI in Advisory Roles

Artificial intelligence (AI) has achieved superhuman levels of performance in fields as diverse as playing games, computer vision, and conversational recommendation systems. Unlike human advisors, there is no evidence that these AI systems can generalize the experience gained on one task or one domain to a different task or domain. Whether they can accomplish it or not, human beings also learn continuously from communication and relatedness with others, working in organizations, teams, and socio-economic systems. Consequently, AI systems designed to assist, advise, or manage tasks or knowledge often lack the necessary flexibility, understanding of the perspectives of different involved parties, human-related skills, as well as the ability to efficiently negotiate and achieve compromise solutions while solving conflicting objectives. AI in an advisory role can assist with training and productivity enhancements, lowering the occurrence of bias, and minimizing discrimination. The performance can be monitored near real-time, and the boundaries of engaged AI agents are formally defined. This presents a novel, end-to-end multi-agent, and blackboard AI system with an explicit, quantitative measure of the individual performance and respective global level of collaboration among the different assisting AI agents and within the entire deliberation team.

AI in an advisory role refers to intelligent systems designed to assist, advise, suggest, educate, or provide consulting services to human users, individuals, or organizations. The intelligent agent, whether created through knowledge or machine learning, can identify patterns or regularities among existing data and make forecasts or assist with suggestions. The agent could act autonomously, where the advisory agent performs decisions and actual task execution on behalf

of the human user, or the agent can act in an authority role, where the advisory agent only makes suggestions or offers recommendations to the human user, so the user will have information to be used in conjunction with their knowledge and responsibility while making the ultimate decision. Many practical decision-making problems could benefit when an advisory AI agent assists the human decision-maker, based on the extensive knowledge of history and the capacity to learn current and emerging patterns. Indeed, a range of professional judgment tasks that include prediction, problem-solving, and diagnosis, performed through the application of decision heuristics and rules or machine learning algorithms could benefit from the intelligent assistance of highly competent AI agents.

4.3. Collaborative AI Systems

This section envisions multi-agent settings and associates a select number of stakeholders (agents) to an assigned role or task, like advisers. The stakeholders are programmed such that they can talk, discuss, and establish agreements among themselves. A recent manifestation of Personal AIs would be an excellent amalgam of multiple roles under multi-agent settings; for example, an aggregator, an adviser, and a companion for a user. The purpose of this section is to explain how multi-agent systems can be used to design collaborative environments where a stakeholder, whether an AI or a human, can work and consult with fellow stakeholders.

This section convenes various ideas and projects, both past and present, from the multi-agent area and the emerging consultation genre. Recommender systems have employed the idea of similar agents to improve sellers' services. Another related idea is mentioned for mobile agents while determining network traffic. A project promotes the pooling of resources in an underserved category by advertising competitors' products. Metadata provides an organization that can offer a means of selecting human assistants without lexicographically searching. An aspect that differentiates a project from the field of recommendation systems is that it concerns a professional multi-stakeholder data platform, not just data that can be used to generate revenues for sellers.

5. Methodology

The system should be able to process both opinions from different advisors, company knowledge, and goals, which are partly affected by the most important factors in insurance underwriting. Therefore, we propose a multi-agent multi-task end-to-end deep learning model-based advisory network with explicitly enhanced argumentation and mitigated risks. First, the model decomposes the preference from the composite large, complex, and uncertain prerequisite Company-Advisor-Fact tri-vector in a multi-task fashion, so that it can handle both opinions coming from potentially different advisors and conflicting company goals effectively. Only by combining opinions from different sources can we promote the impartiality, precision, and diversity of advisory systems consistently. Then, a novel Fact-Enhanced Pooling Method is proposed to assist the model in resolving the inherent risk based on the company's fact position, which reflects the relationship between companies and advisors, and the preference of advisors' opinions. Last, we present Class-aware Metric Learning for Collaborative Argumentation, which benefits greatly from the deep neural network, to further mitigate the resulting risk by considering the class information associated with both advisors and their opinions. In addition, we also contribute extensive empirical evidence to demonstrate the efficiency of our model, with concrete implementations based on real-world insurance. The superior results across a variety of real-world business settings consistently support the multi-faceted benefits of the proposed model.

The design and process of the model are as shown: First, the model consists of parties, including a large company and a diverse group of advisors who assess the composite prerequisite Company-Advisor-Fact with potentially different preferences or goals. In this study, our experiments include five different advisors with four possible preferences. Moreover, the company may change its stance on different fact positions, holding various preferences in different scenarios. Additionally, we choose insurance underwriting as the underlying implementation scenario in conjunction with real-world data for proof of concept. Second, a multi-agent and multi-task structure is proposed to simultaneously help the model master each particular advisor's preference and the company's preference, so that risk can be mitigated by skillfully predicting the preference assignment of each advisor. Third, the developed Fact-Enhanced Pooling Method is utilized to attune the risk via fact and multiple preferences. Fourth, to further reduce the risk, the innovative Class-aware Metric Learning for Collaborative Argumentation is introduced, which is facilitated greatly by the deep neural network by considering more details. Then the beliefs of the advisory result are proposed, which are needed to mitigate the risk caused by the argumentation collaboration.

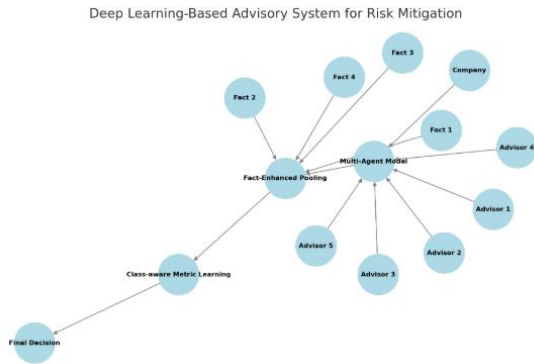


Fig 3 : Deep Learning-Based Advisory System for Risk Mitigation

5.1. Research Design

We use the terms insurance experts and agents interchangeably in this paper as they both represent knowledgeable professionals in the insurance domain. We leverage qualitative empirical research methods where practitioners, as domain experts, provide us with rich contextual data, as well as offer insights and nuances into their practice. We used semi-structured interviews in this exploratory and interpretive investigation of insurance consulting. Interviewees varied from insurance companies of differing sizes, as larger companies could provide us with the senior executives' perspectives of the general insurance industry. The use of differing views from the internal industry landscape encourages the representation of a rich and deep analysis of the topic of interest, which is an essential aspect of case research. Interviews were conducted selectively in one-on-one as well as group settings and were supplemented further with industry-specific conference participation. Group settings were initiated where subject matter experts of consultants and employees from both large and small-scale insurance companies participated in the discussions. Each interviewee was asked a preparative open questionnaire and some allowed prompting when necessary. All communications or discussions pursued were recorded, transcribed, and then subjected to the iterative process of qualitative analysis. Theoretical sampling drives our data collection, ensuring an in-depth analysis that focuses on interesting or novel problems and the plausibility of the questions.

Equation 2 : Risk Prediction Model Using Agentic AI $R_p = \alpha X + \beta Y + \gamma Z$
 where

R_p = Predicted risk level,

X = Customer claim history,

Y = Market volatility factor,

Z = External economic variables,

α, β, γ = AI-determined weight parameters.

5.2. Data Collection Techniques

The data required to validate the insurance consulting model consisted of unprocessed data from a variety of sources such as claim logs, contact center call-in data, assessment reports, surveillance videos, and user interactions. Since the dataset utilized was private and confidential, data associated with each state cannot be shared. Furthermore, the number of validations per user input, as well as the input text provided, varies depending on the nature of the input text. The average validation count corresponded to tasks with similar levels of complexity. Three general guidelines were put forward by the experts: anytime the AI agent correctly predicted the task, the first guideline corresponded to the expert recommending the output rendered as the best advice for the user; the second guideline was associated with letting the AI deliver the advice as long as the quality was acceptable; and the third guideline related to when the user wanted to contact a human assistant.

The data, consisting of prescription narratives, ICD-9-CM, and HCPCS codes, as well as discharge dates and outcomes, was collected in random order from all of the notes and prescriptions that were created over seven months in a large dental practice. The prescription notes and details related to each treatment prescribed were created and recorded during the patient's treatment. A compliant distributed framework was utilized to process the case data. The data was received by applying the anonymization client to the data. The anonymization process processed all of the data that was returned from the database. Upon finishing the data anonymization process, the anonymization client provided four new data files: three temporary files in CSV format and the final anonymized file in ZIP format, to allow them to be moved to a secure location.

5.3. Analytical Methods

Our model is termed Collaborative Agentic AI Model (CAAIM) as it is a Multi-Agent System that is composed of a diverse ensemble of AI specialists. The objective is to create a decision support system where the strengths of different AI sub-specialists can be leveraged to effectively propel data to solutions to a much higher level. Our implemented case concerns the complex problem of optimizing the product mix of an insurance consultancy. The consultancy offers specialized fields of expertise that allow the entity to execute projects in high-resource B2B lines of business, and they would like to maximize the collective net return from projects in their multi-year pipeline. The high-resource nature of projects means that the consultancy might execute most K projects, each project driven by the outcomes produced by AI systems that specialize in each B2B line of business. The selection of the B2B line of business that the consultancy accepts also involves strategic long-term planning. Subsequently, we also created a demonstration case to show the advantage of using CAAIM on a dataset where the true underlying data structure is not known a priori to the agent.

Specifically, the unique collaborative problem-solving of three agents that investigate in a region of the same multi-feature data points provides significant efficiency gains. A full summary of all the agents that comprise the CAAIM model is presented in this section. The emphasis on insurance consulting, as a rapidly growing complex real-world domain with numerous interfaces, is that there is a personal bias or interest in strategies that are better handled if the appropriate bio-inspired strategies are utilized in solving the model. Our bio-inspired model simulates and learns through collaboration, adaptation, distributed decision-making, self-organization, and other techniques that incorporate a variety of attributes such as investigative ability and the capacity to form ad-hoc, specialized investigation sub-groups.

6. The Role of Multi-Agent Systems in Insurance Consulting

The paper addresses problems involved in decision support systems designed for insurance companies and independent insurance agencies. We propose a hybrid architecture of the system that utilizes a multi-agent model showing potential for reasoning under situations of fundamental uncertainty. We analyze the structure of the insurance consulting process to suggest a proper set of functionalities in computer decision support systems. Our model is based on the idea of a "fractal" structure of both the modeled insurance market and the multi-agent system. Which of the insurance companies provides the best set of insurance products presents itself as a multi-agent competition model. We tackle the multi-agent behavior issue by representing real insurance salesmen as competing agents. The payment fee would be an agreed percentage of profit from implemented proposals. Proposed multi-agent simulation: therefore, the insurance consulting process is modeled by simulating a selection of decisions within multiple fractured insurance companies. Real experts represent the modeling agents. The simulation process itself is carried out through the multi-agent market of insurance agencies, which comes to life in the software multi-agent environment. The results of the decisions made by the leading experts can be compared to the knowledge of professionals from different districts. Our model demonstrates the positive elements of unpredictability typical of regional markets. At the same time, performed studies show a high degree of decision security significance in regional insurance markets.

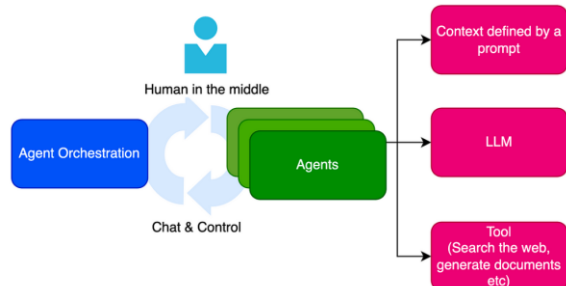


Fig 4 : Multi-Agent System

6.1. Agentic Behavior in Consulting

While customers and agency executives often come into the relationship armed with a mass of data, few come to the process with a defined problem. Indeed, in most fields of consulting, identifying the issue or problem to solve can often consume far more time and effort than the problem-solving that follows. Moreover, the nature of the documented problem has a significant impact on whether, how, and where the consultants are used. Current research into management consulting in the insurance sector finds many segments where an automated analytics service

would be a valuable adjunct to the consultancy or replace the consultancy entirely. Examples include routine benchmarking of KPIs or underwriting operations, detecting patterns or trends in portfolio performance, or even checking for common errors.

In the commercial world, one of the most frequently articulated visions for AI products is a vision of agents tirelessly working through data for humans. A digital assistant that can scour a stream for vital information—a personal assistant that retrieves information and helps users with various tasks on request. These domain-driven applications of AI create enormous value from data by being responsive to the data. They have two obvious modes of delivery: a consumer can call upon the AI agent, or, more commonly and valuable, the AI agent calls out when it detects something. Imagine the soft mental state that develops when people are freed from the need to stay focused: how much deep thinking may replace the ordinary? More work on different aspects of agentic AI is being done to understand how collaboration affects the level of (and ability to handle) agentic behavior.

6.2. Collaboration Among Agents

RDMEN manages its knowledge base based on concepts, which represent a learning structure for agents to reason about the world, act within it, and communicate among themselves and to users. Such concepts define the roles of agents and also the way they should interact. The analysis of the communicative roles of agents and the design of communicative structures of the knowledge base was previously developed for MAA technology.

In it, the concept that they supported was patented, and agents had specific roles and goals. Both SKS and its MAA technology development were preserving the collaboration of human expertise. Later on, others showed that the expansion of a human who processes knowledge to a community of knowledgeable agents offers the potential for building more flexible and trusted knowledge-based systems. Their theory-driven approach not only explains expert experience at any level of a social organization but can lead to a new generation of consulting and advisory systems that draw upon community knowledge. The latter paper also presented a potential dialogue machine technology that utilizes virtual agents from different commercial autonomous agent systems as subjects to integrate multi-agent behavior. With such virtual agents, one may work to understand, support, and leverage evolved group-level cognitive strategies in the field of communication science through a partnership for known agents and multi-agent systems.

Thus, the proposed solution is collaborative and competes with alternative technologies through its common knowledge base shared by many humans, and it also learns.

In this work, approaches for multi-agent advisory AI were exposed, where agents are designed to be subservient to human experts, be virtually invisible during their failure recovery support, and employ human expectations to enhance individual autonomous agents. A multi-agent decision engine design space is outlined, facilitating the creation of a distributed AI platform that contains multiple collaborating agents capable of learning. It is then shown how, after system deployment, structural complexity in the collaborative properties of multiple agent systems, that serve as agents of stakeholders, can grow and adapt to the changing requirements of long-lived socio-technical problems. Such adaptation has been proven to provide significant benefits against a variety of performance metrics, which are crucial for the long-term operation of mission-critical systems in both the public and commercial sectors. Finally, I based the agent implementation around the architecture, which has a focus on robustness, efficiency, verifiability, and scalability.

7. Case Studies

We discuss four case studies, each motivated by an aspect of our client motivation scenario. Our scenario can be extended by associating customer service calls and other customer interactions with each aspect of enforcement. Each case study formulates a compliance rating problem with the same enforcement aspect but with different sets of potential partners.

Health compliance:

Problem statement: An insurance company wants to help its corporate clients improve their employees' health. As part of its HR services, the corporate client regularly conducts wellness-related activities such as health fairs, yoga classes, nutrition classes, and smoke-out initiatives. Increasing compliance maximizes certain insurance agreement benefits for both the clients and our clients. The corporate client reports these events to us weekly. A variety of health service providers visit the employees at the client company regularly, but not necessarily weekly. Our goal is to use the health-related outcomes information and the wellness activity reports to motivate adherence to wellness activities and measure wellness support effectiveness by assigning a compliance rating to non-compliant or less effective clients.

7.1. Real-World Applications of Multi-Agent Systems

In the past, intelligent systems have been composed of one agent or monolithic modules. Recently, large intelligent systems have been proposed that are composed of multiple agents. Agents can share the workload of performing a huge amount of computation, as well as concurrently discover more elucidative solutions. The framework of multi-agent systems (MAS) has become a prevalent paradigm in computational intelligence, thanks to the decentralized feature, self-driven capability, and ability to balance centralized systems. The multiple-component network structure also allows the system to maintain a certain level of operability if one of the units experiences failure. Starting with this concept, this section reviews the real-world applications of multi-agent systems in various domains, such as finance, manufacturing, medicine, and so on. The examples are classified into the following categories.

Insurance Consulting: Multi-agent systems integrate different and diverse expertise from artificial intelligence. A variety of business processes from different sectors, like advisory and consulting, can be supported by intelligence. In the insurance industry, consultancy aims to detect fraudulent or exceptional claims. The consultancy agents are attorneys, claim adjusters, and law enforcement agents. Law enforcement agents perform data mining and text mining using legal knowledge and evidential exceptions. They are designated as machine learning or data mining agents that discuss, collaborate, or default by reasoning about the process of handling claims with exceptional hints. Claim adjusters, from their nondetection of any exceptions, are defaulters of the detectors of fraudulent behavior. If the defaulters do not meet their duties properly, the productivity of the organization can be harmed. Some exceptions are detected by claim adjusters and are passed on to other agents, who deliver them as cases to an attorney. The attorney provides guidelines to medical investigators. The roles of the agents are not fixed and can vary due to the interaction caused by the collaborative and reasoning dialogues carried out by the agents about the claimed exceptional behavior. The claim adjuster is responsible for the administrative enforcement of the nondetectable predicted exceptional activities.

7.2. Comparative Analysis of Traditional vs. AI-Driven Consulting

Traditional consulting involves proposing recommendations based on the requirements and expectations of the client. Rehabilitation centers are at the forefront of the treatment program because they outsource consulting services for patient treatment to other enterprises. That inevitably affects a decline in the quality of information during the exchange, and it is also bound to interaction structures through knowledge pipelines designed using the “one-to-many” principle, caused by the desire to search for the most reliable predictions. AI-driven consulting advantages within the scope of deep trust-value-long-term partnerships implementing “many-many” knowledge pipeline interaction are for clients, the companies to whom the implementation of long and expensive patient treatment is outsourced, and their clients – the largest insurance companies.

8. Challenges and Limitations

In this work, we demonstrated the potential of utilizing an array of collaborative intelligent agents in decentralized control systems to decrease the technical vulnerabilities of a service. Although several studies have demonstrated the successful implementation of collaborative control in other complex systems, computer experiments, and simulation models cannot fully cover the complexity and the issues mentioned above. Thus, our future work will mainly consist of the development and real-world experimentation of a physical network of agents. However, before doing that, we have to overcome several challenges in transparency, interpretability, reasoning, and scalability of our approach. In our existing system, agent tasks, reasons for the suggested solutions, as well as the global problem-solving process are hidden from the user. To trust and use the system, human interaction needs to understand what is happening inside. To the best of our knowledge, no existing frameworks offer the combination of human interpretability tools that our system provides in one place.

AI logic that enables end users and regulators to comprehend, query, and audit the system functionality is very complex and is not connected to current research in reinforcement learning, agent-based systems, or related fields. From our perspective, initializing this research trend is crucial not only for our field but also for the development of trustworthy and responsible AI systems. Currently, the recommended solutions are agent-driven and precede industry competitor research and negotiations, which may result in information leaks. In our future work, we plan to investigate ways to provide the end user with more control of the system's reasoning process while maintaining and protecting the confidentiality of the curated data. Our system relies on handcrafting the sequences for solving the client's problems or their primitive versions. After reaching several layers during the iterative process, our system promptly realizes that it is stuck and informs the user. The user's task is to define successive needs in more detail or to perform work related to a particular issue.

Because we have a preliminary version of the solutions to almost all of the defined goals, our system does not deliver the actual values of the corresponding intermediate layers during a single interaction. Realization of the above-mentioned advantages on the scalability front requires careful experiments and engineering of a system that integrates reinforcement learning and natural language processing techniques to handle situations where the developed system is unable to deliver a solution for a particular aspect of the composite problem, and decision making, instead of following the suggestion rapidly, may lock in the user to get data until the system is trained or competitive. A looming policy for disclosure ensures that the available competitive software for related problem-solving is developed, implemented, and on the market.

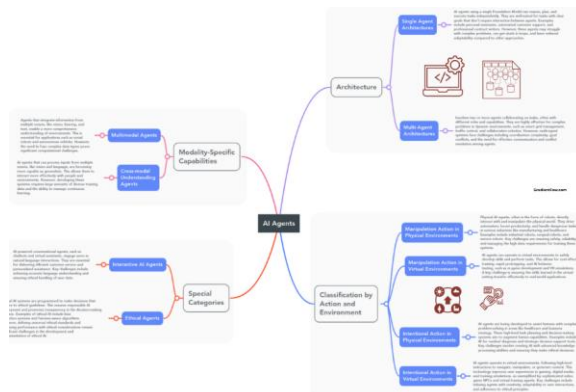


Fig 5 : Agentic AI: Challenges and Opportunities

8.1. Technical Challenges

In MANI, expert agents cannot dictate their positions or be designed to act solely on their interests. The whole must be greater than the sum of its parts for that wide-ranging potential to come to fruition. Accomplishing that ideal will demand collaboration among our network's agents. Building deep representations of the policyholder or organization as the composite set of network features will simplify this task, but complications will still arise. The 2D-ConvNet receives and interprets images and produces a new transformed visual representation of the set of diseases or conditions; this transformed representation would reside in the new dimension created by the introduction of the new type of modeling agent and would generate an adversarial gradient if audited.

Agents need to impact not the insured policyholder as independent players but rather the joint interest or defined function of the insured policyholder and the insurance company. This example hints at the many collaborations and partnerships that might need design. Some of these functions, like the "fight household fire risk" one mentioned earlier, represent joint interest or multi-actor functionality. Acting in concert, they generate jointly backward propagated gradients of interest. Other functions, though, are geared to serve interactional partners. For example, if one model foresees a "low risk of a subcutaneous abscess", and mimics observed radiological findings well, then the models overlap enough to be adversaries. That function does represent a domain-specific knowledge needed to direct the operations of an insurance company, but it produces adversarial gradients that dissolve any possibility of gaining useful knowledge about the materials or methods used to conduct or evaluate the model training.

8.2. Ethical Considerations

Agents powered by advanced AI components are capable of reasoning, making decisions, and carrying out complex operations, as well as having a significant amount of influence on themselves and their outcomes. When building such systems, we are mandated to always prioritize safety, reliability, and robustness, and to design measures to ensure that the influence AI systems have on themselves and third parties is by human values. This is important for users to trust these systems and, of course, to ensure that we are responsible stewards of the technology. Collaborative agentic systems will also be difficult solutions given their multifaceted nature, and it certainly will take some time for our AI research and development to be mature enough to fully address it.

For agents to align with human values, it is important not only that they are constructed using values-aligned techniques and reasoning, but also include and rely on important feedback from stakeholders both to assure people that their overall interests are properly optimized and to ensure

strategic alignment. This feedback for AI models can be rather abstract and mainly involves a human user determining relevant weights for a series of feedback events. This is, of course, difficult to both detect and agree on, even though some methods for doing this have been demonstrated, especially concerning reinforcement learning. These methods depend on relevant feedback concepts being kept in the utility and action tuples that underlie the models. Because of these difficulties in modeling human feedback, multi-agent models must not work in isolation but are instead coupled to better-grounded social or user theories to facilitate applying partial reinforcement learning.

8.3. Market Acceptance

By far the most important critique or requirement is the lack of bias of the insurer-consulting agent – a key characteristic of the consulting profession. Traditional broker intermediaries often had business relationships of a personal nature with the top layers of insureds that helped correct natural institutional insurer bias – i.e., wanting to take as much money as possible from the insureds while transferring as little risk as possible in return. Insureds were already concerned about the possible bias of reinsurers, but at least they had, in general, a professional counterpart in the single broker. Now insureds are faced with the possibility of a double bias, with human bias being augmented by AI agent bias.

It seems likely that AI-based advisers and insurance consulting will play an increasing role in the overall risk and insurance milieu. The actual intensity and timing of that increase is, of course, affected by the specific economic and demographic environment of each line of insurance. Similarly, the acceptance of AI in consulting will vary between its various possible applications, with pricing quotation advice representing the least AI risk, perhaps being the first accepted, especially by insurance buyers. Finally, the different key insurance industry players – insurers, empowered by their access to historical insurance business data and power over their brokers, or allying reinsurers with their unique insurance business perspective – are likely to be forced into using various combinations of AI techniques in their relationships with their clients.

9. Future Directions

In this paper, we've made a strong proposal for collaborative agentic AI, specifically realized within the use case of insurance consulting. Given our focus, our approaches are consultative and advisory, extending the potential for recommender systems. However, the potential extensions and mapping to new fields are exceedingly wide. For most of these other applications, there is a high value on interpretability, debate, and predictability - all touchstones of our consultative approach. Moreover, insurance is a popular area for AI research. Over our experience, we've found important and noteworthy connections to many different research areas. Here, we propose just a few additional areas that might be inspired by this work. Prediction isn't everything. These systems are not solely about predictions. We use our agent network to quickly and efficiently find where a prediction may be wrong, perform live analyses, carry out robustness checks, and begin to understand when a prediction is called for and when a claim or prediction needs to be further checked. This is an application area where AI benefits people, specifically if the AI system isn't just predicting things but entering the discourse. Efforts to improve the interpretability of recommendation systems, through visualization, explanation, and controversy testing, must continue. The agent network is built to use interpretability throughout the system. Future research can do better in using more of our agents in our recommendation process and better understanding

and exposing the arguments and considerations of the agents in the recommendable system to benefit the underwriter.

9.1. Innovations in Agentic AI

The technological landscape is shifting with AI that is designed to work like people and with people. It is the dawn of multi-agent systems, whose definition and principles are more than robotics or game AI. While AI's teleological conception is on a path guided by its human originator, AI can change our society and its foundations. AI's real impact is on a higher layer, a moral one that questions our humanity. Multi-agent systems, both virtual and physical, provide the most essential form of AI through their interactivity with the agents, just like humans do. One of the design principles of intelligent and autonomous entities is that they adapt their behaviors and roles constantly under dynamic environmental conditions. Assuredly, adaptable behavior and nonlinear benefits to the behaviors of insurance entities would help both insurers and policyholders to minimize adversities.

Advisory AI in insurance was trained with knowledge, rules, regulatory compliance, financial knowledge, and logical reasoning. Furthermore, insurance advisory AI using reinforcement learning combined with logic and rule-based approaches with complex properties can initiate interactions with policyholders through frequent learning and provide in-depth knowledge. Participatory interaction from the policyholder while applying reinforcement learning would develop, design, and assist in improving the policyholder's ability. This AI autonomously enhances knowledge, learns new information, captures the learning process, and adapts to real-world dynamic environmental conditions. Moreover, the roles of the insurers play a crucial part in how such entities behave and prioritize their behaviors under dynamic real-world insurance conditions. Frequent policyholder interactiveness helps to understand better ways of helping and responding to the real-world agent. The AI traces results, and negligence reflects information and background about agent activities.

9.2. Potential for Industry Transformation

The potential for domain-specific applications is huge because many knowledge-intensive and highly specialized sectors will require these solutions. Insurance is one of these sectors. Insurance is an information society product par excellence that has seen relatively less transformation by AI compared, for example, to Wall Street. Over time, insurance services could become free to the end user because of the potential for low prices that AI will exhibit. These potential massive changes mean that it is appropriate to think about redefining insurance consulting in both product and research directions. This is especially the case since insurance is a major area of the economy in its own right, and it also has significant input and funding attached to it from other areas of economics such as healthcare economics, disaster economics, and the pensions arena.

Retraining is required for the next generation of consultative skills that will need to accompany these transformations in problem-solving and intellectual property. This paper makes an exploratory foray into this thinking exercise—what could be called the ultra-long-term future of private insurance and private protection research that is driven and connected by themes in the literature that are multidisciplinary and that are practiced not just by one consultant but by virtual teams of human-aided AI advisers centered around the Insurance Knowledge Portal that we introduce in this study and are suggesting that the sector adopts. This future will be inspired by these themes. Both the protection provided by and the choice of insurance will be transformed by the kinds of AI-managed economies that could already be seen as the origin of future

developments. The paths to that future by private insurance could be very liberating. While both the bailout and compliance-based models of future financial regulation involve significant state involvement, and therefore have less futuristic innovation content, if all economic activities were subject to these types of state surveillance, then this could lead to less imaginative searching for new problems like those inherent in catastrophes and conducted at a European insurance firm today, deemed to be solved historically by the contractual innovation facilitated by accident insurance as a form of horizontal protection for a vertical loss.

10. Implications for Practitioners

This paper has very clear implications for AI practitioners in the insurance industry, as well as for other professionals — for example, management consultants, medical practitioners, lawyers, and accountants — whose practice involves consulting activities. For consulting professionals, we argue that they have a great deal to learn from AI researchers about developing tailored AI systems for supporting consultative processes without aspiring to create fully autonomous, agentic AI systems. For AI professionals, the effectiveness and efficiency in the development and operation of AI systems can be significantly improved by a better understanding of consulting expertise and its development in the corresponding professional domains.

It can be expected that many professional domains can benefit from a better and deeper understanding of the role of professional expertise in consulting processes. This deeper understanding will allow professional bodies and organizations to rethink the future of their professions in the face of AI systems. The constantly increasing complexity and fragmentation of human knowledge as we progress in the twenty-first century causes extensive professional specialization in the form of ever-finer sub-domains within the traditional professions. On the one hand, this strategy naturally sustains the high level of expert service in a vast range of areas; but on the other hand, this also promotes an artificial distinction between professional and client capabilities and knowledge, further eroding the right of professionals to undertake requested tasks without the clients' interference.

10.1. Best Practices for Implementation

Across the organizational hierarchy and different centers of interest, we reflected on a set of directives to inform agents about the ethical and efficacious execution of their goals. It is highly recommended to implement feedback loops from the consumer, audit logs of agent interaction, and address inadequate risks through user safeguards, protective upper bounds, and user-definable comfort zones. It is also very important to clearly and carefully delineate the role of the agent and properly communicate what the agent is doing, to suggest a course of action rather than perform that action.

A few positive child-parent directives could specify template structures, autonomy sensitivity, minimized and bounded divergent education, precautionary and explicit context removal, and the implementation of client-centric guarding acts; it is also important to ensure proper handling of discordant directives. Additionally, directives could offer added ethical value and lucrative commercial procedures via typical, encouraged, and value-earning attitudes, alternative regression behaviors, and transparent decision-making. Given this model and initial insights into significant factors for actions, the next chapter begins the task of promoting the conversation between clients and a large set of insurance applications of our model framework, such as analyses of corporate behavior, product design, customer segmentation, and reputation and trust management.

10.2. Training and Development Needs

Consultants who have put in years of business experience or teams with tenured expertise in any industry vertical possess an accumulated wealth of problem-solving experience. If some of the most valuable agents are retiring, how can their replacements be trained? I'd argue that the question of "who is replaceable" is constrained by a single-agent model and that in a truly powerful advisory context, a retiree can be replaced with the retained knowledge, experience, and protocols of the entire corporation and all its contributing agents; a network uniquely qualified to contribute its deep, diverse, and richly intellectual value at any moment in the life experience of the eventual user. No single agent will be squandered or lost, but the collective will only grow in value as the diversity of negotiated risk experiences is expanded. The development of individual powerful learning agents comes under the category of on-the-job training and experience, with the most powerful real-time feedback and knowledge extraction processes already in place. When teleoperation or telesupervision is required, the architect is available to reshape the network structure dynamically to provide the desired life experience as a mentor or guide. This capability drastically reduces the training period, since the professional learning environment is always present.

Equation 3 : AI-Driven Policy Recommendation Efficiency

$$E_p = \frac{S_c}{T_r}$$

where
 E_p = Efficiency of policy recommendations,
 S_c = Successfully converted recommendations,
 T_r = Total recommendations issued.

11. Policy Recommendations

Following the careful analysis and design considerations needed to unleash MANs' disruptive capabilities, we reemphasize our belief in the importance of addressing regulation and governance as key open research challenges. By doing so, the insurance sector can catalyze positive disruption, stimulating the development of innovative and truly value-creating products, and delivering significant economic efficiency improvements not just in the insurance sector, but also cascading across various other sectors that use commercial insurance as a key input. We strongly believe that solving the important applied economic and societal challenges that commercial insurance is naturally designed to address, utilizing appropriate AI systems, trained in a manner respectful of the specific RegTech, solvency, and systemic consequences, could become a powerful solution to societal problems perceived to have public good characteristics that might otherwise arrogate public markets to solve through command-and-control regulatory policy intervention.

In developing and stimulating wider use of advisory AI insurance agency networks, positioning also wider insurance regulatory sandboxes and possibly even net-benefit-of-the-regulation considerations, we think the proposed MAN framework could provide a unifying, generative, and powerfully innovative approach to addressing applied economic and societal questions, whilst ensuring the much-needed early deployment of emergent AI technology guided by a framework dominated by the society-wide public interest consideration. Advances should enable a collaborative global resurgence in developing, governing, and deploying AI systems for maximizing public good. With careful collaborative design, execution, and oversight, we believe that AI systems have the potential to collaborate in addressing the world's significant public good challenges more effectively than our current inefficient economic and societal systems can.

11.1. Regulatory Considerations

In designing multi-agent advisory networks for insurance consulting, we have been mindful of regulatory considerations. Any AI-powered advisory system must be transparent and explainable. Data sources must be verifiable, and any use of personal data must ensure privacy norms. The advisory system should not rely on biases or perverse incentives. It must comply with laws and regulations. Ethical considerations must be kept in mind. Specifically, in the context of insurance consulting, navigating a multi-agent network must not recommend inappropriate or fraudulent actions, make claims not supported by facts or reasonable forecasts, or involve selling and buying activities where the recommendations do not support the proponent's objectives. It must encourage actions that are, on balance, best for the proponent considering the proponent's circumstances and long-term objectives. It must aim to preserve financial value and prevent destruction or diminution of value. It must not act for personal financial gain at the expense of the client, make recommendations subject to any conflict of interest, or fail to identify and address conflicts of interest.

The European Parliament called on the Commission to assess the feasibility and suitability of a separate legal person status for AI systems to enable AI agents to have the same obligations, rights, and legal status as humans can. Such a status would allow individual AI systems to conduct business in their own right and their capacity as "employees", "entrepreneurs", or "employers", be held liable where things go wrong, and own the intellectual property they create. India's National Strategy for AI called for "dignity and legal status for AI systems appropriate for their role in human society". Such evolving regulatory and policy outlooks can embrace multi-agent networks, recognizing them not as independent but as interdependent advisory avatars of the agent that the multi-agent network represents, where communication with human agents and the ultimate decision rights rest with the human agents. This way, legally, it will be the natural person who will be responsible for justifying the advice, and not the multi-agent network itself.

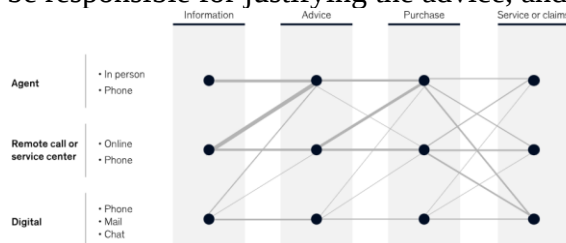


Fig 6 : The revolution in multi-access insurance

11.2. Standards for AI in Consulting

Artificial intelligence (AI) technologies are not only potentially disruptive but also, in many cases, create liability if they cause negative external effects and there is a clear cause. This is especially true for services such as insurance consulting, where there is a strong focus on minimizing future risks for clients. However, independently from liability aspects, businesses are interested in achieving a certain level of quality and transparency when applying AI to these services. To that end, businesses call for a form of certification system or quality label, similar to established standards in the digital product realm. Certification systems provide companies with useful information about the quality and reliability of AI technologies. In addition, sensors support buyers in identifying products that meet certain standards. This ensures a minimum level of transparency and safety. Existing initiatives focused on an ethical perspective and the goal to develop "human-centric" AI, but fell short of addressing practical aspects such as compliance with licensing laws.

Even if standards are proposed, it is not easy to apply them in practice on commercial AI systems because checks cannot be planned and need to be repeated frequently. Changes to the data basis as well as model adaptation regularly need to be reviewed, and this involves subjective considerations such as the appropriateness of relevant past events trained on, the quality of observed underwriting and risk management decisions, etc. Subjectivity and fuzziness in model checks indicate a need for a hybrid approach combining heuristic elements with formal testing and monitoring. This raises underlying questions: What are common elements for different types of AI in consulting systems that can be tested in a standardized way, such as transparent decision rationales, the usage of audited data or open model structures, and what issues are too dependent on business models, regulations, or the particular use case and cannot be properly addressed by a general standard for service-providing AI?

12. Conclusion

With the advent of big data and deep learning methods, AI systems have achieved great success in solving various application issues. A new research trend is to integrate distributed learning AI systems into multi-agent collaborative models such that a coherent and learnable social fabric among AI systems can be created. In the modeling scenarios, the human-communication-oriented natural language is adopted for agent-to-agent communication. Specifically, we propose the concepts of multi-agent advisory networks, which contain collaborative dialog agents and policy processors, serving various applications. We provide a theoretical framework of multi-agent advisory networks and implement a practical application for insurance consulting. We demonstrate the practical use of AI systems that collaborate smoothly through experiments on real-world datasets.

In particular, as a demonstration of how to apply the model in an industrial scenario, the scheme consolidates an insurance consulting process that contains knowledge transfer, underwriting new business, and pricing models among client, broker, risk engineer, reinsurance department, and actuary. The model achieved impressive performances in the tasks, and the interactive dialogues show that it can provide advisors with recommended risk-level pricing. As a future work, the practical process of the proposed model application will go further to integrate it with an execution system to achieve the ‘zero-touch’ consulting experience. In addition, a task of further study will be to construct insurance consulting models for each subservice and professional center. At the same time, incorporating inputs such as monitoring signals and diary records, enhancing collaborative learning for advisor parity, understanding, and the implementation of the practical business processing assist serviceability.

12.1. Summary of Findings and Implications

This paper reports on auditable, highly interpretable, and highly collaborative AI systems—MAANs—that are currently absent. To remedy this, we propose Multi-Agent Advisory Networks (MAANs): systems embodying agents in network form, with each agent a network, where agents represent the player institutions and the network edges facilitate communication, competition, coordination, and cooperative development in both learning and utilizing AI software-led networks. Investigating the feasibility and implications of building, deploying, regulating, and managing MAANs is vital to understanding what kind of AI software we want to be designed and tested so that the AI software modifications, performance assessments, parameter value

adjustments, assurances, and safety measures when AI software operates avoid disregarding qualitative information from player institutions in partly autonomously advising on extremely important tasks prone to high-severity errors. Moreover, the construction and critiques of MAANs require rethinking the aforementioned levels separately and in conjunction. Hence, we examine the desiderata and interconnectedness of players, institutions, and software in agent and network form from these three separate yet interrelated viewpoints. The full realization of the research institute will have direct implications for several existing AI systems. AI system research and development will likely diverge and converge in different places, scholars will switch careers that allow them to keep their ethical opinions, and the most valuable products will not always have the highest sales volume. We argue that since AI systems designed in light of MAANs may provide more sophisticated predictions, safety, and interpretations, more public transparency, and more human collaboration, developing ethical MAANs is not only the right thing to do—it is the consequentialist thing to do.

12. References

- [1] Kalisetty, S., & Ganti, V. K. A. T. (2019). Transforming the Retail Landscape: Srinivas's Vision for Integrating Advanced Technologies in Supply Chain Efficiency and Customer Experience. *Online Journal of Materials Science*, 1, 1254.
- [2] Sikha, V. K. (2020). Ease of Building Omni-Channel Customer Care Services with Cloud-Based Telephony Services & AI. Zenodo. <https://doi.org/10.5281/ZENODO.14662553>
- [3] Siramgari, D., & Korada, L. (2019). Privacy and Anonymity. Zenodo. <https://doi.org/10.5281/ZENODO.14567952>
- [4] Maguluri, K. K., & Ganti, V. K. A. T. (2019). Predictive Analytics in Biologics: Improving Production Outcomes Using Big Data.
- [5] Ganesan, P. (2020). PUBLIC CLOUD IN MULTI-CLOUD STRATEGIES INTEGRATION AND MANAGEMENT.
- [6] Siramgari, D., & Korada, L. (2019). Privacy and Anonymity. Zenodo. <https://doi.org/10.5281/ZENODO.14567952>
- [7] Polineni, T. N. S., & Ganti, V. K. A. T. (2019). Revolutionizing Patient Care and Digital Infrastructure: Integrating Cloud Computing and Advanced Data Engineering for Industry Innovation. *World*, 1, 1252.
- [8] Somepalli, S. (2019). Navigating the Cloudscape: Tailoring SaaS, IaaS, and PaaS Solutions to Optimize Water, Electricity, and Gas Utility Operations. Zenodo. <https://doi.org/10.5281/ZENODO.14933534>

- [9] Ganesan, P. (2020). Balancing Ethics in AI: Overcoming Bias, Enhancing Transparency, and Ensuring Accountability. *North American Journal of Engineering Research*, 1(1).
- [10] Somepalli, S., & Siramgari, D. (2020). Unveiling the Power of Granular Data: Enhancing Holistic Analysis in Utility Management. Zenodo.
<https://doi.org/10.5281/ZENODO.14436211>
- [12] Ganesan, P. (2020). DevOps Automation for Cloud Native Distributed Applications. *Journal of Scientific and Engineering Research*, 7(2), 342-347.
- [13] Vankayalapati, R. K. (2020). AI-Driven Decision Support Systems: The Role Of High-Speed Storage And Cloud Integration In Business Insights. Available at SSRN 5103815.
- [14] Ganti, V. K. A. T. (2019). Data Engineering Frameworks for Optimizing Community Health Surveillance Systems. *Global Journal of Medical Case Reports*, 1, 1255.
- [15] Sondinti, K., & Reddy, L. (2019). Data-Driven Innovation in Finance: Crafting Intelligent Solutions for Customer-Centric Service Delivery and Competitive Advantage. Available at SSRN 5111781.
- [16] Pandugula, C., & Yasmeen, Z. (2019). A Comprehensive Study of Proactive Cybersecurity Models in Cloud-Driven Retail Technology Architectures. *Universal Journal of Computer Sciences and Communications*, 1(1), 1253. Retrieved from <https://www.scipublications.com/journal/index.php/ujcsc/article/view/1253>