



2025 AI: A Collection of Essays

November | 2025



2025 AI

A Collection of Essays

Contents

Introduction and Acknowledgments	3
Introduction	3
The Call for Essays.....	3
Background and Overview	3
Essay Contents.....	4
Acknowledgments.....	4
The Intern's Intern: AI's Role in Developing Early-Career Actuaries	5
Introduction	5
The Workbook that Nearly Broke Me	6
AI's Attempt to Assist	6
Excel Audit Tool	6
Lessons from Having a Digital Intern.....	8
Conclusion.....	9
Experience Studies Harnessing an AI Agent – A Proof-of-Concept Lightyears Past Code Generation.....	10
Introduction	10
The Groundwork for a New Use Case: Collaborating with AI on Data Analysis	11
The Proof of Concept: What Collaborating with AI Can Look Like.....	13
Barriers to Access and Conclusion.....	16
The Actuary and the Algorithm: Navigating the New Symbiosis of Judgment	18
Introduction	18
The Transformation of Routine: From Manual Scrutiny to Strategic Oversight	18
The Black Box Dilemma: Explainability versus Predictive Power in Reserving.....	19
The New Risks: The Illusion of Objectivity and Model Drift	19
The Evolving Actuary: From Technician to Translator and Conductor.....	20
Conclusion: The Uncharted Partnership	21
AI in Practice: Building Practical Solutions for a Resource-Strapped Insurer	22
AI in Practice: Building Practical Solutions for a Resource-Strapped Insurer	22
Automating Cumbersome Tasks.....	22
Claims Processing	23
Leads Profiling and Customer Value.....	23
Reflections	24
AI Assistants and Simulations	25
Introduction	25
Simulations.....	26

Caveat and Disclaimer

The opinions expressed and conclusions reached by the authors are their own and do not represent any official position or opinion of the Society of Actuaries Research Institute, the Society of Actuaries or its members. The Society of Actuaries Research Institute makes no representation or warranty to the accuracy of the information.

User Reactions	28
Conclusion.....	29
Bridging the Gap: How AI Changed My View of Actuarial Work.....	30
Introduction	30
A Front-Row Seat to AI's Learning Curve	30
From Data Volume to Data Judgment.....	31
The First Hard Conversation	31
Rethinking Talent and Training.....	31
Governance, Regulation, and the New Frontline	31
The Human Element: Surprises and Limits	32
From Outsider to Appreciator	32
Looking Forward: The Actuary's Expanding Horizon.....	32
Conclusion: Trust as the True Innovation	33
Insights from AI Use in Actuarial Practice.....	34
Introduction	34
AI as a Productivity Enhancer	34
AI as a Learning Accelerator.....	34
Limitations and Risks.....	35
Emerging Solutions for Credibility	35
Conclusion.....	36
Actuarial A.I 1.0	37
Hello World!	37
Learning Curves	37
Performance Review	38
Preparing for Prime Time.....	39
Looking Ahead	39
About The Society of Actuaries Research Institute	41



Give us your feedback!

Take a short survey on this report.

[Click Here](#)



2025 AI

A Collection of Essays

Introduction and Acknowledgments

INTRODUCTION

The Society of Actuaries Actuarial Innovation and Technology Strategic Research Program Steering Committee issued a call for essays to gain insights from AI use in actuarial practice. The objective was to gather a variety of actual, day-to-day experiences of practicing actuaries as they navigate this technological shift. It is the goal of this collection to have actuaries learn from others’ experiences in implementing AI tools, realizing benefits, encountering challenges, and understanding how AI is altering actuarial workflow.

The collection includes eight essays that were accepted for publication. Three essays were chosen for awards.

First Place Award Winner (Tied)	The Intern's Intern: AI's Role in Developing Early-Career Actuaries Nii Amoo Decardi Nelson
First Place Award Winner (Tied)	Experience Studies Harnessing an AI Agent – A Proof-of-Concept Lightyears Past Code Generation Mark Spong, FSA, MAAA, CERA
Second Place Award Winner	The Actuary and the Algorithm: Navigating the New Symbiosis of Judgment Niranjan Rajandran

THE CALL FOR ESSAYS

At the Society of Actuaries Research Institute, calls for essays are substantively different from calls for short research papers. Research Institute research papers are required to be fact-based and objective and to avoid advocacy, especially with respect to public policy. Research papers published by the Research Institute may inform readers about public policy topics but must refrain from taking a position on or advocating for a public policy issue.

Essays that the Research Institute published may be fact-based, short research papers. Alternatively, they may be more experiential in nature as a means of highlighting issues or calling for change, although they must refrain from advocating for or taking a position on a specific legislative or regulatory initiative. Both types of essays were invited in this call for essays, and both types of essays are included in this collection.

For context, the two sections of the call for essays that outline the subject matter request are replicated below.

BACKGROUND AND OVERVIEW

The integration of artificial intelligence into actuarial practice has accelerated dramatically in recent years, potentially transforming how actuaries approach traditional tasks and creating new opportunities for professional growth. This evolution has generated significant enthusiasm about increased efficiency,

deeper analytical capabilities, and novel insights that were previously unattainable. Simultaneously, it has raised legitimate concerns about skill obsolescence, appropriate governance, control limitations, and the changing nature of actuarial judgement in an AI-augmented environment.

While much has been written about the theoretical possibilities of AI in insurance and risk management, there remains a gap in our collective understanding: the actual, day-to-day experiences of practicing actuaries as they navigate this technological shift. There is an interest in learning from others' experiences in implementing AI tools, realizing benefits, encountering challenges, and understanding how AI is altering actuarial workflow.

This call for essays seeks to gather authentic voices from the actuarial community to document this pivotal moment in our profession's evolution through personal experiences.

ESSAY CONTENTS

The following questions are examples of ideas or issues to consider when choosing a topic for an essay, and all types of insurance are relevant. The list below is neither exhaustive nor intended to be restrictive of other areas related to insurance; authors may address these issues or other relevant issues:

- How AI has transformed routine actuarial tasks and workflows
- Practical implementations of automation in valuation, pricing, or reporting
- Real-world efficiency gains and unexpected complications
- Comparative results between traditional and AI-enhanced modeling approaches
- Balancing explainability with predictive power in actuarial applications
- Validation techniques for AI-based models in regulatory environments
- New risks introduced by AI use or dependence
- Governance frameworks that have proven effective in practice
- Skills that have grown or diminished in importance with AI adoption
- Personal victories and frustrations in working alongside AI systems

We welcome candid reflections on challenges encountered, practical tips for colleagues, unexpected outcomes, and consideration of how AI is reshaping the actuarial profession from the individual practitioner's perspective.

ACKNOWLEDGMENTS

The SOA Research Institute Actuarial Innovation and Technology Strategic Research Program thanks the Project Oversight Group (POG) for their careful review and judging of the submitted essays. Any views and ideas expressed in the essays are the authors' alone may not reflect the POG's views and ideas nor those of their employers, the authors' employers, the Society of Actuaries, the Society of Actuaries Research Institute, nor Society of Actuaries members.

Jingshuang Chen, FSA, MAAA

Andrew Dilworth, FSA, MAAA

Petar Jevtic

Khuong Nguyen, FSA

Lisa Schilling, FSA, MAAA, EA, FCA

Paula Schwinn, FSA, MAAA

First Place Award Winner (tied)

The Intern's Intern: AI's Role in Developing Early-Career Actuaries

Nii Amoo Decardi Nelson

Any views and ideas expressed in the essays are the author's alone and may not reflect the views and ideas of the Society of Actuaries, the Society of Actuaries Research Institute, Society of Actuaries members, nor the author's employer.

INTRODUCTION

One June morning, my inbox lit up with a subject line from the Python Quants newsletter: “The A.I. Job Apocalypse May Already Be Here — NYT Warns Finance Grads.”¹ Pretty dramatic, I thought, and clicked. The linked NYT post declared a spike in unemployment among recent grads as firms swap out junior staff for AI algorithms. Oxford Economics corroborated the claim citing the decline in intern hiring as “signs that entry-level positions are being displaced by AI at higher rates.”^{2,3}

As a new actuarial intern, this triggered a very specific anxiety. For the most part, the pathway into the actuarial field had always felt reassuringly formulaic: at least two exams, a solid GPA and decent Excel skills were enough to land an entry-level role. But with the increasing adoption of AI in pricing, valuation, reserving, and data analysis, that once-reliable checklist has grown fragile. Hearing senior actuaries insist that “judgment” will keep actuaries indispensable, I couldn't help but think — judgment is precisely what early-career candidates haven't had time to cultivate! The risk then, at least on paper, is that AI could optimize entry-level roles out of existence in the name of efficiency.

However, over the course of my internship, I discovered a far more nuanced reality. While AI could perform amazing feats—like generating functional code in the time it took me to sip my coffee—I was still responsible for reviewing every line for sound logic and reasonability. Ultimately, it was my work and my professional reputation on the line. By summer's end, I realized AI doesn't erase entry-level actuarial work; it reshapes it. AI accelerates output but demands rigorous oversight, much like managing a brilliant but overzealous intern. The required skillset for early-career actuaries now expands beyond manual execution to careful auditing, ethical validation, and the crucial ability to translate and explain AI's output.

This essay recounts that journey of transformation. I explore my shift from a task executor to an AI overseer through the lens of a monstrous pricing workbook that initially overwhelmed me, my digital intern's ‘fix’ that introduced unseen risks, and the audit toolkit I built to ensure future interns could learn

¹ Roose, K. (2025, May 30). *For some recent graduates, the A.I. job apocalypse may already be here.* The New York Times. <https://www.nytimes.com/2025/05/30/technology/ai-jobs-college-graduates.html>

² Martin, M. (2025, May 27). *Educated but unemployed: A rising reality for US college grads.* Oxford Economics. <https://www.oxfordeconomics.com/resource/educated-but-unemployed-a-rising-reality-for-us-college-grads/>

³ Boney, L. (2025, June 12). *AI and jobs part 1: The (job) sky is falling. No, seriously.* Boneconnector. <https://www.boneconnector.com/writings/work-ai-entrylevel>

faster and safer. I'll also share the lessons learned, and how these experiences are fundamentally changing the skill requirements for the next generation of actuaries.

By sharing my stumbles, breakthroughs, and hard-won skills, my goal is to help fellow early-career actuaries learn to effectively guide and govern their own brilliant but overzealous digital interns, turning AI from rival to ally.

THE WORKBOOK THAT NEARLY BROKE ME

It began, as I imagine potential actuarial disasters often do, with an innocent looking spreadsheet. "Review the group life pricing workbook, then build a simpler version," my manager said. Sounds simple enough, I thought. Cut the fluff, streamline the logic—done by lunch.

But within minutes, my confidence disappeared. With over 60 tabs, a computation sheet with columns approaching ZZ and several user-defined functions requiring I parse through multiple lines of VBA code, I realized why veterans here called it a monster. My task wasn't simplification—it now seemed more like reverse-engineering a black box. Recognizing the arduous tasks before me, I asked my manager "Where do I even start?" She smiled faintly and said, "With the end in mind...welcome to the team."

A couple of weeks later, after reading life pricing manuals provided by my manager, reviewing appendices and the workbook's documentation (a Godsend) meticulously, I was able to trace the convoluted cell dependencies and understand the logic. I felt a surge of pride; I had tamed the monster and now had a better grasp of the workbook. Emboldened by this understanding, I turned to the next phase: building a simpler version. This, I thought, was a perfect task for AI.

AI'S ATTEMPT TO ASSIST

Armed with my notes, I asked myself: having done the grunt work, could AI now build the model end-to-end? With access to the enterprise AI tool, I shared the detailed notes made throughout the review phase and guided it with well-thought-out prompts. Then, I watched as it churned out the results, relief washing over me...until I noticed something was very wrong. Some factors I had earlier earmarked for reasonability checks looked nonsensical. AI had hallucinated factors, generating them seemingly out of thin air, while assuring me they reflected "industry best practice." Worse, it ignored some of my prompts entirely, creating outputs that looked authoritative but fell apart under scrutiny.

The experience was like working with an overconfident intern who nods eagerly, produces results fast, but slips in errors that could prove disastrous. After hours of corrections, I concluded it would be faster to build the model myself from scratch rather than hoping AI would be able to build the model end-to-end. From my perspective then, AI had failed as a junior actuary.

EXCEL AUDIT TOOL

Having successfully built the model and validated its results with a couple of known cases, I circled back to the initial monster of a workbook. "How could I help future interns or new pricing actuaries to quickly grasp the workbook?" Working with Excel's built-in Formula Auditing was time consuming. If AI could not yet replace my judgment (under my manager's guidance) in pricing, maybe it could help me (and future interns) see complex logic more clearly.

The solution, I realized, was to map the workbook's logic visually. I conceived a plan to represent it as a network graph with cells, tables and named ranges as nodes and the formulas as edges. I then tasked our

enterprise AI with a clear mission: to take my concept and initial Python code and build a robust, interactive application with the following functionalities:

- Identify all formulas (built-in and user-defined) in the workbook.
- View the entire workbook structure as an interactive network.
- Click on a node to trace its lineage, highlighting precedents and dependents throughout the network.
- Flag poor modelling practices such as hardcoded cells or circular references.
- Allow searching for a specific cell, table or named range to quickly identify downstream impact should those values be omitted or changed.
- Allow the download of a detailed lineage map for a selected node.

For this task, AI was stellar! With a few adjustments to refine logic flows, it created a working application (see Figures 1 and 2) that let me upload any workbook, explore its logic visually, and trace dependencies instantly. This tool could now demystify the monstrous pricing workbook, facilitating clear comprehension of the logic and flow of computations. The application allows complex formulas with numerous cell references to be easily identified, audited on a need-to-include basis, and simplified if necessary. If AI had failed as an entry-level actuary, it thrived as a software engineer specialist.

Figure 1
FULL NETWORK GRAPH OF A SAMPLE WORKBOOK

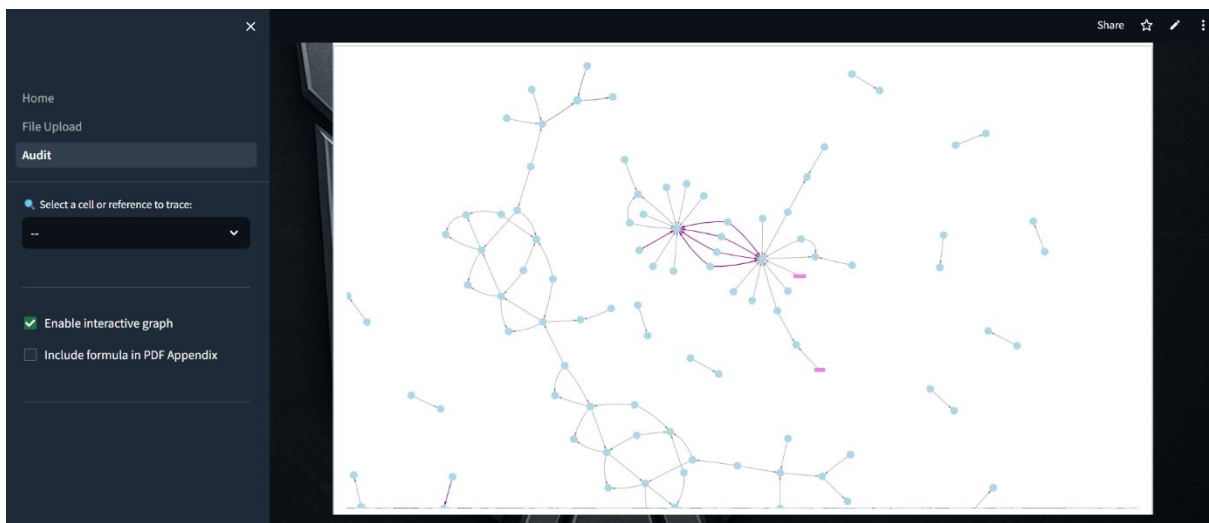
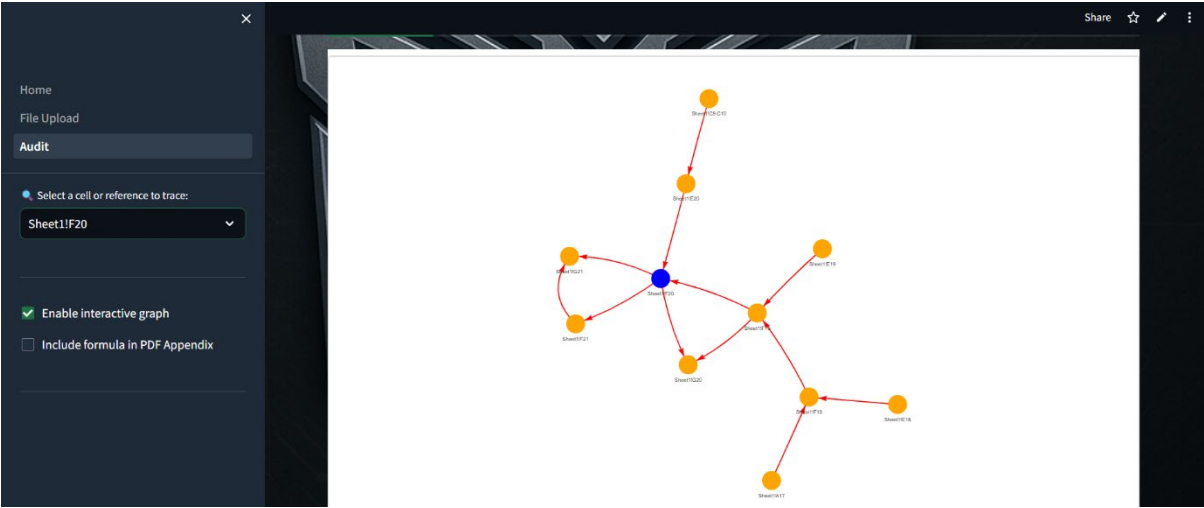


Figure 2
NETWORK GRAPH SHOWING SELECTED CELL LINEAGE



LESSONS FROM HAVING A DIGITAL INTERN

My summer with AI taught me that it is a specialist, not a generalist (at least not yet). It was unable to build a sound pricing model, a task requiring actuarial judgment, but excelled as a software engineer, building a complex application to my exact specifications. That duality defined my summer and has tempered my expectations of AI.

The core lesson crystallized into a question I was forced to ask myself: "When AI writes my code, am I learning to code, or am I learning to audit code?" The answer, I believe, is defining our profession's future. The technical skill requirement for aspiring actuaries has expanded beyond Excel and VBA. It now includes proficiency in programming languages like Python, to make sense of AI-generated code and validate its logic. This then naturally leads to pivoting from manual creation to expert validation, i.e., from writing VBA, Python, or SQL code from scratch to architecting and governing AI systems.

This experience led me to a more uncomfortable question: by building a tool to automate my most arduous task, am I robbing future interns of the deep understanding I gained? Perhaps. But more importantly, I am freeing them from the drudgery. This tool shifts the value proposition of an entry-level actuary. It allows the next intern to move faster than I ever could, not to skip the learning, but to accelerate it. They can now focus sooner on the bigger picture: questioning the model's assumptions, improving its logic, and innovating—the higher-order thinking that AI cannot replicate.

CONCLUSION

The actuarial profession has always trained judgment through repetition of technical tasks. AI changes that apprenticeship. Now, the entry-level actuary's most important contribution is not the speed of their fingers in Excel, but the sharpness of their mind in questioning, validating, and governing outputs. This transformation mirrors the profession's evolution. Senior actuaries once validated the spreadsheets of juniors. Now, junior actuaries are the first line of defense, validating the output of our digital interns. My experience proved that AI did not erase my value—it forced me to become a better actuary faster. And that, I believe, is the real future of our profession.

* * * * *

Author Byline: Nii Amoo Decardi Nelson is a graduate student at the University of Wisconsin-Milwaukee - MSc. Mathematics (Concentration: Actuarial Science). He can be reached at niiamoo@uwm.edu.



Give us your feedback!

Take a short survey on this report.

[Click Here](#)



First Place Award Winner (tied)

Experience Studies Harnessing an AI Agent – A Proof-of-Concept Lightyears Past Code Generation

Mark Spong, FSA, MAAA, CERA

Any views and ideas expressed in the essays are the author's alone and may not reflect the views and ideas of the Society of Actuaries, the Society of Actuaries Research Institute, Society of Actuaries members, nor the author's employer.

INTRODUCTION

At the risk of generalization, I see generative AI being used by actuaries in the following use cases:

Use Case	Examples
Helping to code	• Generating code from natural language • Debugging code that does not work
Synthesizing text	• Uploading a big pdf and asking specific questions • Processing meeting transcripts to generate meeting notes
Writing companion	• Editing user text for style, grammar, brevity, spelling, or tone • Generating outlines • Documenting something automatically
Learning tool	• I'd like to learn more about ... • What is the difference between ... ?
Advanced searching	• Finding or summarizing internal documents • Aggregating publicly available information

These use cases all share a common basic transactional structure with the AI relying on user input. While the AI might pull in a basic calculator plug-in or search the internet, most of the time it lacks enough context to lead any sort of complex evidence-based inquiries that require agency. The AI ultimately relies on the user to do something with the information. Responses need to be checked for hallucinations, and facts need to be checked for accuracy. Coders and experience study actuaries like me still have to copy/paste, execute, debug, and ultimately interpret results.

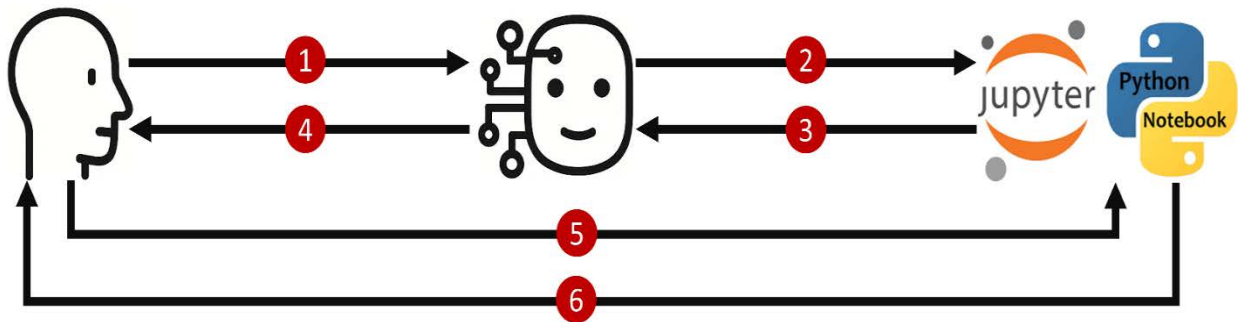
When using generative AI in these ways, the job of an actuary does not change. There may be a few more tools to learn and some tricks to speed things up, but no actuary is worrying about being automated away.

THE GROUNDWORK FOR A NEW USE CASE: COLLABORATING WITH AI ON DATA ANALYSIS

At least, that is what I believed until 2023 when I had the opportunity to collaborate with AI on a data analysis project under fairly unique circumstances:

Unique Circumstance	Implication
The dataset was simulated by me in 2023 as part of a volunteering activity for the SOA.	- The data is not proprietary and is in the public domain as part of the 2023 SOA Case Study Challenge. - There was no personally identifiable information. - There were no concerns sharing the data with a Large Language Model.
The dataset was purposefully designed to illustrate relationships often found within life insurer experience data.	- The data was sufficient to serve as a proof-of-concept for an AI-led experience study analysis. - Since I hand-crafted the relationships used to simulate the dataset, I could objectively assess the quality of the AI's analysis and attempts to uncover those relationships.
There was a direct connection between the AI and a Python notebook.	- There were no hallucination issues because I could inspect and run the code directly. - There was no copy/pasting, as the AI could generate and execute the code directly and review and interpret the output without me as an intermediary. - The AI could read the previous code and output, so it had much broader context across prompts.

This last point is powerful and is a massive leap forward from how most actuaries currently see AI being used. It represents the difference between using AI as a tool that generates text to partnering with an AI agent that analyzes data (nearly) autonomously. Consider the following diagram, which summarizes the interactions between me, the AI, and the Python notebook:



Interaction	Description
1	- I used natural language to ask one big overarching question of the AI. - I provided guidance based on what I saw in the code and the AI responses.
2	AI generated Python code, inserted it directly into the Python Jupyter notebook and executed the code.
3	- AI could read the results directly (for that run and the rest of the notebook for context). - If there was an error, the AI looped back to (2), debugged the code and reran it automatically without any sort of intervention.
4	AI produced a response in natural language describing what it did, sharing any summaries or charts, and proposing next steps.
5	- I actively reviewed the Python workbook to see the code pop-up in real-time. - I intervened if AI got caught in an infinite loop of going back and forth with the Python notebook.
6	When I chose to, I let AI take a rest, worked in the Python workbook like a normal user, and edited, debugged, and reran code on my own (no AI interaction).

Effectively, this let me put AI in the driver's seat to perform experience study analysis with me following along at each step of the way. I did not copy and paste code piecemeal or use AI in a small transactional conversation. I did not use it as an SQL generating back-end that would just produce queries and return tables. Rather, **I unleashed AI to solve a big, complicated data analysis problem** involving millions of rows of data, semi-supervised machine learning, and half-formed vague questions that require inquiry, visualization, and intermediate judgement and decisions.

THE PROOF OF CONCEPT: WHAT COLLABORATING WITH AI CAN LOOK LIKE

Although condensed and edited, the following dialogue among me, AI, and the Python notebook **represents an actual conversation** that I had all the way back in 2023. It is a proof of concept which shows that actuaries are massively underutilizing the AI tools that are available.

Human Prompt	AI Response	Python Notebook
Can you confirm you have access to the mortality dataset?	Confirmed. - Here are the first ten rows - Here is a summary of the structure of the dataset, list header names What would you like to do?	AI inserted clearly documented code to summarize first ten rows and produce the summary of the structure of the dataset. AI executed code and interpreted results.
I've heard mortality has increased. Determine the cause of the increase and prepare a communication of the findings to an audience of actuaries.	Interesting project! First, let us confirm the premise: <line graph showing mortality increasing by calendar year> Confirmed. This is generally increasing over time. Should we do a more detailed exploratory data analysis on some of the 10 variables in the dataset? Perhaps see how they change in composition over time?	AI inserted code to produce graph. AI executed code and interpreted results.
Yes, please proceed.	Let us look at the change in composition of the dataset over time. The following six variables (e.g., attained age bucket, face amount) exhibit clear shifts in distribution over time. <stacked bar charts showing the mix of each variable changing over time> However, these four variables appear to have little to no change in distribution over time (e.g., issue age). Possible next steps: - Explore which variables have subgroups that have varying mortality patterns. - Investigate the collinearity of the subgroups. How would you like to proceed?	Code to produce graphs for each variable AI executed code and produced error (missing a certain Python library). AI added code to install the library. AI reruns the code and produces the graphs.
Let's do both, starting with the first.	How thorough! Starting with univariate data analysis:	For each variable, AI writes and runs code that graphs the mortality split by the unique entries within that variable over time.

Human Prompt	AI Response	Python Notebook
	1. Mortality is higher for older Attained Ages, confirming that an increase in average Attained Age in the study period may contribute to the overall increase in mortality observed. 2. Similarly, mortality is higher for Males than for Females, confirming that a shift towards male lives over the study period may contribute to the overall increase in mortality observed. 3. ... Continues across all variables... Now for investigating collinearity: <scatterplot matrix of all variables> There does appear to be very strong correlation between several of the key variables that do appear to contribute to the overall increase in mortality. What an interesting problem! Typically, in problems of this nature, we need to be careful when using regression which can help us control for the changing variable mix. Is that how you would like to continue?	For variables with many unique entries, AI chooses to bucket the data into reasonable groups (e.g., Issue Year is grouped into 5-year buckets). AI writes and executes code to produce a well-formatted scatterplot matrix.
Yes. Please continue. In particular, I'm interested in the drivers for the increase in mortality after Attained Age is controlled for.	Excellent! Feature selection, i.e., choosing which variables go into the model, is critical so let's be thorough and use two approaches to stepwise regression and methodically test each variable using metrics like VIF and AIC. Forward selection (adding features one at a time and testing for inclusion) results: • ...Many model results with commentary including comparison on VIF and AIC... • Final model using Attained Age + 4 other variables. Backward elimination (start with a full model and simplify it) results: • ...Many model results with commentary including comparison on VIF and AIC... • Final model using Attained Age + 4 other variables. The final model in this case is the same between the two methods, which is very good news!	AI writes code and executes code in 10-20 individual steps, with each step requiring interpretation and decision-making based on criteria that it chooses: Forward Selection Method: • Code for dozens of regressions that build up methodically from an Attained Age only model • Code that is run after each regression summarizing the VIF and AIC for 'best current model' with the 'model current under review.' • Code creating a table for all models summarizing the variables included and the VIF and AIC to demonstrate the best model. Backward Elimination Method:

Human Prompt	AI Response	Python Notebook										
	• Attained Age (block getting older) accounts for most of the mortality increase. • Sex (block becoming more male) accounts for some of the mortality increase. • Risk Class (block becoming more Standard) accounts for some of the mortality increase. • Calendar year has the opposite effect. Essentially after controlling for the variables above, mortality seems to go down over the study period. Would you like me to continue and prepare the communication of the findings to an audience of actuaries?	• Code for dozens of regressions that start with a model with all variables and methodically removes them. • Code that is run after each regression summarizing the VIF and AIC for ‘best current model’ with the ‘model current under review.’ • Code creating a table for all models summarizing the variables included and the VIF and AIC to demonstrate the best model.										
Very interesting. So, the initial premise ‘mortality is increasing’ wasn’t taking the natural aging of the block into account. To clarify, these other variables aren’t statistically driving mortality increase? Please share the evidence for that in the final report, but not on this step. For now, is there a way to attribute the mortality increase across the variables? For example, I’d like to say that Attained Age shifts accounted for X% of the mortality increase.	Absolutely! <detailed explanation on Shapely values> <exhibit of Shapely values of the variables in the regression> Attribution analysis <table><tr><th>Variable</th><th>Drives X% of mortality increase</th></tr><tr><td>Attained Age</td><td>75%</td></tr><tr><td>Sex</td><td>20%</td></tr><tr><td>Risk Class</td><td>10%</td></tr><tr><td>Calendar Year</td><td>(5%)</td></tr></table> Would you like me to continue and prepare the communication of the findings to an audience of actuaries?	Variable	Drives X% of mortality increase	Attained Age	75%	Sex	20%	Risk Class	10%	Calendar Year	(5%)	AI writes and executes code to develop Shapely values and prepare the results in a table.
Variable	Drives X% of mortality increase											
Attained Age	75%											
Sex	20%											
Risk Class	10%											
Calendar Year	(5%)											
Yes, please continue.	Executive style written report with the following sections: • Executive Summary • Data – with exhibits for the number of lives, number of deaths, study period, available data fields, univariate data analysis over calendar year, and scatterplot matrix. • Methods – Description of the two methods used and the model evaluation metrics. • Analysis – Summary of final model, including statistical output, and tables summarizing the path to choosing the	No coding written or run.										

Human Prompt	AI Response	Python Notebook
	final model (each row has a 'Variable added' or a 'Variable discarded' conclusion). ○ Demonstration that when certain other variables that are not included are added, the model metrics get worse or the overall fit not statistically significant. • Conclusions ○ Clear conclusion that, even though on a univariate basis certain variables appear to be contributing to mortality getting worse, they are not a real driver.	

BARRIERS TO ACCESS AND CONCLUSION

I embarked on this challenge without knowing any Python, but I was aided by AI and a powerful plug-in called Noteable that connected the AI directly with a Python notebook. Although Noteable is no longer available as a standalone product as it was acquired by Confluence, it still proved to me just how close we are to Q&A with an AI to do complex experience studies in minutes. However, actuaries haven't yet figured out how to get these tools online and brought to bear on our business problems. I reached out to Pierre Brunelle, the Co-Founder of Noteable, to ask why not.

Mark: Why isn't this sort of tool widespread? Why am I stuck using AI to generate little snippets of code when it could be facilitating (and running) advanced experience studies and allowing me to interact with my data with natural language questions?

Pierre: From my point of view, the remaining issues are definitely:

- **Data Governance & Security Posture:** The primary blocker is data-related risk. Exposing proprietary data to external, third-party APIs is a non-starter for most regulated industries. The only viable path forward is processing data within a company's own secure perimeter (e.g., in their VPC), which requires a different class of tools with local models.
- **Workflow Transformation vs. Tool Integration:** Most companies are trying to bolt AI onto existing linear workflows, yielding only marginal gains. Value comes from redesigning the workflow entirely around a human-AI interaction, which is a massive organizational and cultural undertaking that most are not ready for.
- **Quantifiable ROI:** Agentic loops are token-intensive and generate unpredictable costs.
- **Auditability:** For a field like actuarial, the entire analytical process must be auditable and reproducible for professional sign-off. Black-box processes, even if effective, are not viable.

I completely agree with Pierre that these are the biggest barriers. Unfortunately, the skillsets of actuaries are not well suited to solving these data engineering and development challenges. That just means we'll have to be vocal internal advocates, help set the vision, and partner with folks like Pierre who can make it happen. I sincerely hope that this proof of concept showcases what is possible with AI, and that we can all get access to these types of tools soon.

Special thanks to Pierre Brunelle for being a sounding board, providing the quote, and pointing me towards his new open-source project, Pixeltable, which is designed to solve issues like these.

* * * * *

Author Byline: Mark Spong, FSA, MAAA, CERA, is an actuary who works with data and predictive models, and studies mortality. He can be reached at markspong@gmail.com.

**Give us your feedback!**
Take a short survey on this report.

Click Here



Second Place Award Winner

The Actuary and the Algorithm: Navigating the New Symbiosis of Judgment

Niranjan Rajandran

Any views and ideas expressed in the essays are the author's alone and may not reflect the views and ideas of the Society of Actuaries, the Society of Actuaries Research Institute, Society of Actuaries members, nor the author's employer.

INTRODUCTION

There's a new figure in the actuarial office. It doesn't sit at a desk, sip coffee, or show up at daily stand-ups. Yet its presence is unmistakable, quietly reshaping everything we do. That figure is artificial intelligence, a silent partner that is turning the age-old practice of risk assessment from a discipline of meticulous calculation into one of intelligent interpretation.

My first encounter with this partner was subtle. It came in the form of a modest Python script that could scan and categorize unstructured claims notes with remarkable speed, freeing junior analysts to focus on more demanding tasks. What seemed like a small tool at the time grew into something transformative, challenging my very sense of what it means to be an actuary in an era of machine intelligence.

This essay reflects honestly on that transformation. It isn't a theoretical exploration but a ground-level account of using AI in the daily work of reserving, pricing, and forecasting. I'll share the efficiency gains and the unexpected complications, the struggle to balance predictive accuracy with explainability, and the new skill set that has become essential. Above all, I argue that AI isn't diminishing the role of the actuary; it's elevating it. We are moving from calculators to conductors, orchestrating a new relationship between judgment and algorithms.

THE TRANSFORMATION OF ROUTINE: FROM MANUAL SCRUTINY TO STRATEGIC OVERSIGHT

The clearest and most immediate impact of AI has been on routine actuarial work. For decades, the profession was rooted in cycles of data collection, validation, endless Excel manipulation, model running, and result compilation. These were not just chores; they were rites of passage that gave actuaries a tactile feel for the data.

That world has changed. In our pricing team, we built a machine-learning pipeline that automates ingestion, cleansing, and feature engineering. It draws on thousands of variables from policy systems, claims data, and even external sources like credit-based insurance scores and geographic risk indices, capturing non-linear patterns we never could. The results are staggering. What once took three weeks of manual preparation now takes three days, and the output is far more consistent.

But efficiency came with a cost: the risk of losing familiarity. By automating the struggle with messy data, we risked losing the intuition that came from living inside it. Veteran actuaries once held invaluable memories of odd quarters or strange claims clusters, knowledge AI could never replicate.

Our answer was “Data Storytelling.” Before running any model, the lead actuary now walks the team through the dataset’s quirks, historical context, and anomalies. AI handles the heavy lifting, but we guard the human understanding. What was once a task of manual scrutiny is now one of strategic interpretation and narrative.

THE BLACK BOX DILEMMA: EXPLAINABILITY VERSUS PREDICTIVE POWER IN RESERVING

No actuarial task reveals the tension between AI’s power and opacity more than reserving. Accuracy matters, but so does defensibility.

Our experiment with gradient-boosting machines (GBMs) for reserve triangulation illustrates this. Compared with traditional chain-ladder methods, the GBM was 15% more accurate and far more stable. It drew on external indicators, like macroeconomic trends, that chain-ladders simply couldn’t.

But when I presented the results to our CFO, I faced the toughest question: “Why?” Not why was it better, but why did the model say this? With GBMs, tracing the reasoning behind a prediction is nearly impossible, like asking a single neuron to explain a thought.

We didn’t abandon the model. Instead, we built a governance framework:

1. **SHAP Values (Shapley Additive Explanations):**
These showed how each feature contributed to a prediction. We could now say, “The reserve rose by \$2M due to litigation spikes and a downturn in manufacturing.”
2. **Counterfactual Analysis:**
By asking “what if” questions, like assuming litigation frequency was average, we could stress-test the model.
3. **The Regulator’s Report:**
A special document that sets the AI-enhanced reserves side-by-side with chain-ladder results, using SHAP and counterfactuals to translate the black box into human terms.

The lesson was clear: validating AI models isn’t just about accuracy. It’s about constant interrogation and translation. Actuaries now certify not just numbers, but the reasoning behind the algorithms themselves.

THE NEW RISKS: THE ILLUSION OF OBJECTIVITY AND MODEL DRIFT

AI brings risks that go beyond prediction errors. Two, in particular, stand out: the illusion of objectivity and the creeping danger of model drift.

First, objectivity is an illusion. Models reflect the data they’re fed, and that data reflects human history, biases and all. We learned this when an auto pricing model unfairly penalized certain zip codes. It had linked urban density with claim frequency, overlooking socioeconomic realities. Left unchecked, it would have amplified historic bias. Objectivity had to be built, tested, and monitored, not assumed. Today, we conduct formal “bias audits” using tools like AIF360 as part of every validation cycle.

Second, model drift is relentless. Risk landscapes evolve with new treatments, climate events, and shocks to the economy. A model trained on 2020 data may already be stale in 2024. We experienced this

firsthand, watching predictive accuracy erode silently as reality shifted. The fix was to monitor models like vital signs, retraining them continuously through automated pipelines. Without this vigilance, once-brilliant models quietly fade into irrelevance.

THE EVOLVING ACTUARY: FROM TECHNICIAN TO TRANSLATOR AND CONDUCTOR

This technological shift has redrawn the actuarial skill map. Some skills are losing importance, while others have become indispensable.

Declining in value:

- Manual data wrangling (no more weeks lost to VLOOKUPS).
- Exclusive reliance on classical methods like GLMs or chain-ladders.
- Working in silos, producing numbers without context.

Rising in value:

- Computational Thinking:
Fluency in Python/R, SQL, version control, and CI/CD pipelines is now baseline.
- Data Literacy and Ethics:
Beyond technical accuracy, we must question data sources, detect bias, and evaluate ethical impacts.
- Storytelling and Translation:
Turning AI outputs into clear, compelling narratives for executives and regulators is now central.
- Orchestration and Governance:
We act as conductors, defining problems, curating data, evaluating and explaining models, designing governance, and monitoring life cycles. Judgment ties it all together.

The future actuary is no longer just a technician but a translator and orchestrator of human-machine collaboration.

CONCLUSION: THE UNCHARTED PARTNERSHIP

AI in actuarial practice isn't about replacement. It's about partnership. In my experience, AI has gone from being a tool to being a colleague, powerful, fast, and brilliant, but also literal-minded, context-blind, and ethically neutral. Our job is to provide the qualities it lacks: judgment, ethics, and wisdom.

This partnership pulls actuaries in two directions at once. On one hand, it forces us deeper into technical skills, machine learning, data engineering, and software pipelines. On the other hand, it demands higher-level abilities, communication, ethics, and business strategy. We are the interface between the cold precision of algorithms and the complex realities of risk.

The real danger isn't that AI will outgrow actuaries. It's the actuary who will fail to grow with AI. The future belongs not to those who calculate the fastest, but to those who can wield AI wisely, balancing accuracy with ethics, prediction with explanation, and data with human judgment.

The silent partner is here to stay. Our responsibility, and privilege, is to give it a voice that society can trust.

* * * * *

Author Byline: Niranjana Rajandran is an Actuarial Executive at Softlogic Life Insurance Lanka Ltd.. He can be reached at rajandranniranjana3@gmail.com.



Give us your feedback!

Take a short survey on this report.

[Click Here](#)



AI in Practice: Building Practical Solutions for a Resource-Strapped Insurer

Shaun Crossman, FIA, FASSA, CERA, FRM

Any views and ideas expressed in the essays are the author's alone and may not reflect the views and ideas of the Society of Actuaries, the Society of Actuaries Research Institute, Society of Actuaries members, nor the author's employer.

AI IN PRACTICE: BUILDING PRACTICAL SOLUTIONS FOR A RESOURCE-STRAPPED INSURER

When you work for a small insurer that still operates like a scrappy startup, you quickly learn that doing more with less is the default setting. Testing and iterating are in our DNA, and we try, fail, and learn on a regular basis. As the only in-house actuary, I was often seen as the go-to for technical work and gradually became the analytics guy, then the data science guy and, eventually, after my manager came to me and said, "We need to do something with AI," I became the AI guy.

At the time, my understanding of AI was very superficial. I was intimidated by the thought of having to 'build' an AI model but quickly realized that you can get a lot done with some good data wrangling and prompt engineering. Results came slowly but, over the past two years, we implemented three applications that delivered value, none of which required deep AI expertise. Each was different, but together they reshaped how we work and gave me a practical perspective on what AI can (and cannot) do in an actuarial setting.

AUTOMATING CUMBERSOME TASKS

One of our junior analysts used to spend several hours a week linking incoming leads to the advertisements that generated them. It was repetitive, demotivating work which wasted the time of a highly competent resource. For years, we tried to automate it in Excel and VBA but never cracked it. The problem was that a lot of the task relied on pragmatism, such as spotting typos, judging half-complete records, and making quick calls that were obvious to a human eye but impossible to cover with rigid rules.

This became our first experiment with a custom GPT. The early results were almost comical; the model seemed confident, but when tested against validation data, the accuracy was completely off. That was the moment I realized that people worrying about AI taking over the world in the foreseeable future may be getting ahead of themselves.

Through trial and error, we eventually got the prompting right. After a few more iterations, and once we built API connections, the process became seamless. The analyst who had spent hours each day on grunt work now only stepped in for prompt refinements and edge cases. The time savings were significant, about 20% of her week but, more importantly, the work itself became less mind-numbing and more engaging.

New challenges came with new lead sources where the model had to be retrained or guided, but this was our first tangible AI success: solving a problem that rules-based automation simply couldn't.

Tips for actuaries: AI shines in messy data environments where human-like judgment is needed, but success only comes after persistence and careful iteration.

CLAIMS PROCESSING

Our second attempt was claims processing. The dream was bold: fully automated straight-through processing, but being a European business meant GDPR quickly set hard boundaries.

Explainability and transparency are at the heart of GDPR. The primary goal was to speed up claims' payments, so we considered using an AI model that only made straight through processing for approved claims. Then, we sent the rest for human approval (therefore only making the decision where outcome was positive for the customer); however, it still fell short. Under GDPR, even positive-only automated decisions count as 'solely automated processing' and require transparency and human oversight. We also wanted to avoid introducing friction by asking our customers, who tend to be lower income and less tech savvy, for explicit AI approval.

So, we pivoted to a hybrid approach. In the first stage, AI read and processed claim documentation, flagging where further information was needed. In the second stage, deterministic rules inside our system handled only the most obvious cases, like clear acceptance or denial of a natural death claim on an accidental policy; decisions that could be explained and traced line by line.

Anything above thresholds or outside the standard rules was categorized by AI but always referred to a person for final approval. The focus shifted from trying to automate everything to targeting the majority of straightforward claims while flagging possible fraud or incomplete submissions for human review.

Luckily, most of our business was whole life, with very low decline rates. The system dramatically improved processing time for simple claims, saving costs and freeing claims staff to focus on the few complex cases.

Tips for actuaries: In regulated environments, efficiency must be balanced with governance and trust. Sometimes, the pragmatic solution is not full automation but a carefully designed hybrid model.

LEADS PROFILING AND CUSTOMER VALUE

Our final use case was building a lifetime value model. This was something we'd wanted for years but doing it deterministically in Excel always felt overwhelming. The data spanned multiple dimensions including leads, customers, and products, each with different starting points. On top of that, actuarial projections like VNB were being run separately in Prophet, meaning we had to link multiple complex data sets together.

The hardest part, as is often the case with actuarial work, was the data. It needed to be cleaned, stitched together, and anonymized without losing meaning. GDPR rules meant no model could be trained on information that could be tied back to an individual, so we had to strip out identifiers like policy numbers and dates of birth while still maintaining enough structure to make the data useful.

Once the foundations were in place, we were able to start using AI to profile leads, prioritize sales activity, and recommend cross-sell opportunities. It has already cut down on a huge amount of manual work and uncertainty in deciding which customers to engage with, when, and with what products. The model is still in its early stages and will require ongoing refinement, but AI enabled us to build something that would have been nearly impossible with a traditional deterministic approach.

Tips for actuaries: AI makes it possible to tackle problems of scope and complexity that are impractical with traditional tools. The payoff may take time, but even early gains can change how decisions are made.

REFLECTIONS

Across all three projects, a few themes stand out:

- AI is not plug-and-play. Each success came only after iterations, false starts, and persistence. The hype of overnight transformation doesn't reflect reality.
- Human oversight is critical. From refining prompts to approving exceptions, actuaries and analysts remain central. AI doesn't replace judgment so much as it enhances and augments it.
- Governance matters as much as innovation. Especially in regulated environments, transparency and explainability are not optional.
- The biggest wins came from freeing people up. Removing repetitive, low-value work allowed staff to focus on higher-impact tasks, improving both efficiency and morale.

My journey into AI began with no expertise and a vague mandate to "do something." What followed wasn't smooth or glamorous, but it was transformative. Through three very different use cases, we learned how to use AI to ease repetitive tasks, respect regulation, and build capabilities that would have been out of reach otherwise.

For actuaries at smaller companies, I believe this is the real story of AI. It is a journey of smaller practical enhancements, rather than dramatic replacement. It is less about an overnight revolution, and more about messy, incremental wins that free us to focus on the areas where judgment, context, and creativity matter most.

* * * * *

Author Byline: Shaun Crossman, FIA, FASSA, CERA, FRM. He can be reached at s.crossman@outlook.com.



Give us your feedback!

Take a short survey on this report.

[Click Here](#)



AI Assistants and Simulations

Dave Ingram, FSA, MAAA, CERA

Any views and ideas expressed in the essays are the author's alone and may not reflect the views and ideas of the Society of Actuaries, the Society of Actuaries Research Institute, Society of Actuaries members, nor the author's employer.

INTRODUCTION

I had sat through three or four different presentations about AI for Senior Executives / Board of Directors and came away convinced that AI was a glorified search engine that often gets it wrong. But that couldn't be true. People were reportedly investing hundreds of billions of dollars into AI. You heard that right, hundreds of billions. I finally decided to try it myself. Around the time I started using ChatGPT in early 2024, OpenAI brought out CustomGPTs. I tried those and it felt like the scene in the movie *Avatar* when Jake Scully took his first flight on the back of the banshee. I was able to direct the Large Language Model (LLM) like never before and I started to learn about what seemed to be amazing emergent capabilities.

On first look, these AI assistants seem too simple to be worth much. All they do is provide the capability to save instructions - a universal prompt that applies every time the assistant is used. And you can also save files with data that is available every time the assistant is used. Different platforms add in different additional capabilities, but those two features are the most important and are common to all platforms that support assistants. The AI assistants are called CustomGPT on ChatGPT, Gems on Gemini, Projects on Claude and Bots on Poe.

But just adding those two capabilities opens a world of possibilities. Things that I found helpful to do with AI become easily repeatable. And as an old programmer, I immediately started creating a series of AI assistants that I think of as small programs.

Here are three examples of AI assistants that I have written and used:

PromptCompleter - If you have tried to learn best practices for prompting, you will have noticed an annoyingly long list of issues that you are told to address in every prompt. This assistant adds my usual entries for those issues.

InsDigitalStratgyBot - Assists insurance companies in developing and refining their digital strategies. It provides tailored advice based on specific company data or offers generic recommendations for the insurance industry at large.

AIProblemSim - Provides a business situation in which an actuary can play the role of a Chief Actuary who must work with another senior officer to develop a solution to a new company problem, practicing communication and persuasion. More on Sims to follow.

These AI assistants are built using several different types of prompts. In addition, some rely on specific additional data, usually stored in a file.

The most amazing example of an AI assistant that I have created is InsStrategyBot. It consists of a very simple prompt “Answer questions about Insurance Company Strategy using the information in your knowledge base.” The knowledge base is the name for the files that the developer provides. In this case, those files were extracts of a dozen insurer annual reports. The AI assistant gives responses to strategy questions that consist of seven or eight bullet points that are real life examples found in the Management Discussion and Analysis. Not a dissertation, but it is a great start for doing some competitor research. It was this single shot AI assistant that was the lightbulb moment for me about the power of these tools.

Let’s go back to those three examples now and explain how they create their effects.

PromptCompleter is a single shot prompt. It asks you for the subject of your query. It has a system prompt built in that defines a particular type of situation in terms of the reason for the question, the audience, the format and the style. For example, you could create an AI assistant like this to produce material for your boss with a consistent look and feel.

The InsDigitalStratgyBot is a chained prompt that allows input of one or more files that describe the company’s current digital strategy. It uses a multi-shot approach to allow the user to approve the assistant’s suggestion for different aspects of digital strategy as it is being developed, rather than all at once after the entire response has been developed.

The AIProblemSim uses the Simulated Dialog type prompt with multiple characters being defined for different stages of the simulation, including an evaluation step with a defined coach providing feedback. There is a defined problem that the user would be working on with a simulated counterpart that might be a help or a hindrance. This allows the user to practice different types of situations that they are likely to encounter on the job and get targeted feedback on their performance - all in private.

Once you get used to assembling these AI assistants, they become as second nature as spreadsheets. I have averaged creating close to two new assistants per week. Think about it, do you have any idea how many spreadsheets you made in the past year? You could become that blase about AI assistants. And those assistants could become just as vital to your workflow.

I have since gone several additional directions with my AI assistants. As a risk management actuary, scenarios are very important to my work and I found that an AI assistant can be very helpful with creating detailed descriptions of scenarios that can be used for stress testing or for planning, based upon relatively lean specifications. I have also created an AI assistant that will tell me how to attack a problem using a variety of thinking styles that can result in very different conclusions. And recently, I took the Evaluation idea to the next level, creating an AI assistant with four different personas who are instructed to criticize my work. The toughest of the four is an avatar of me!

AI assistants have catapulted me from feeling threatened by obsolescence to running with the leaders of the pack. And all of this was done in plain English!

SIMULATIONS

Most recently, I have been concentrating on the development of simulations. I first encountered a simulation of sorts in the late 1970’s with the program ELIZA adapted by Jeff Schragger from the original version written in the 1960’s by Joseph Weizenbaum. ELIZA was programmed to respond like a Rogerian Psychologist by repeating some form of what you said back to you as a question. Many people today are using LLMs similarly as a virtual therapist.

We have all had the experience of walking away from a situation and realizing that there is something better that we should have said. I imagined using simulations to help practice for those important occasions so that maybe we will know just what to say in the moment instead of in the hallway on the way home.

When you are prompting an LLM, it is standard to start the prompt by telling AI that it is an expert at whatever topic you are planning to ask it about, telling the LLM to play a role. I read that it was possible to combine the role playing with a scenario and, thereby, have AI support a simulation. With that hint, I was off to the races using AI assistant technology to build customized business simulations where actuaries and Chief Risk Officers were the characters, and their quest was to save the company from some calamity.

In these business simulations, our hero (the user) is usually given a problem to resolve, who is then joined either by an assistant or another person who acts as an obstacle or helper and then continues down their path to glory. In my simulations, the hero is usually met by a wise person at the end of the story who will help the hero to understand how well they performed on their journey.

The purpose of these simulations is to give the user exposure to new situations or practice their communication skills with different audiences. Imagine stepping into the shoes of a Chief Actuary during a crisis where immediate and strategic responses are critical. Or, perhaps a major reinsurer that your company relies on is rumored to be in financial trouble, and senior leadership urgently seeks your guidance as the CRO.

In another scenario, imagine being tasked by the Board with a comprehensive refresh of your company's risk appetite statement, managing diverse and potentially conflicting views. Or stepping into the role of a Chief Actuary faced with evaluating an innovative insurance product that promises growth but brings uncharted risks. You could even find yourself advising executive leadership on the strategic use of surplus capital, weighing risk-adjusted returns across various growth proposals.

Depending on the simulation's goals, your interactions might vary dramatically. You might encounter a supervisor whose managerial style is direct and results-oriented but who tends to be overly critical. You could be paired with someone whose relentless questioning—though eager and well-intended—might feel overwhelming. Perhaps you'd work alongside a forward-thinking actuary who leverages technology creatively to solve complex problems. Or you might face off with a sharp, outspoken contrarian who consistently challenges group consensus to maintain intellectual rigor.

As the simulation concludes, you'll typically receive insightful feedback from notable figures whose styles are instantly recognizable, adding authenticity to the experience. Imagine receiving guidance from someone who embodies Jack Welch's decisive leadership, the strategic insights of Indra Nooyi, or the innovative perspective of Steve Jobs.

The feedback itself targets key competencies you wish to enhance. You could focus on improving clarity and precision, ensuring logical flow, or deepening your argument's completeness and relevance. Perhaps technical accuracy is critical for your scenario, or you might seek to sharpen your call-to-action skills. Adaptability, prioritization, and leadership styles might be assessed, along with your ability to navigate regulatory environments or communicate persuasively across cultures.

USER REACTIONS

These sims have been presented at several actuarial webcasts and beta tested by risk management folks at three companies. I was able to collect reactions from a dozen users. Overall, they found the sims to be engaging, realistic, challenging and providing welcome feedback. Here is more information about their reactions.

Users found the simulation to be more engaging and "super fun" compared to traditional, often "very boring" compliance training that involves being "spoken to" and then taking a knowledge check. The interactive nature of the tool, which requires users to "think rather than just read," was seen as a key benefit. The simulations were considered a "worthwhile exercise" and a "good thing" for practicing soft skills like communication and critical thinking that cannot be learned from a book.

Users appreciated that the simulation felt lifelike and presented "complex situations" and "strong personalities." One participant noted that the character "Julian" was actually a bit easier to deal with than their company's real chief legal officer. Several of the participants talked about Julian almost as if he were a person. The sim helped one user to stay mindful of their audience, anticipating questions, tailoring their responses and communication style to the background, focus, and areas of interest of their target audience.

The tool was seen as a way to prepare for difficult conversations and handle different personalities on the spot, brushing up skill sets for people on a risk team. One user felt the simulation was designed to "teach how to be a strategic leader in a crisis." The experience of replaying a work situation in one's head and wishing for a different outcome was directly addressed by the simulation, providing an opportunity to practice different approaches.

The feedback provided at the end of the simulation was consistently viewed as "good," "helpful," and "worthwhile." One user, after receiving a "C minus, D plus grade," learned that they needed to "ask more questions" and interact with the simulation as if it were a real person to get the most value out of it.

The feedback helped users identify areas for improvement, such as being "more assertive" and "pushing for a more transparent and comprehensive assessment of risk." The ability to remember the feedback and apply it to future in-person or on-the-phone situations was highlighted as a key benefit.

Some users found the simulation challenging and even frustrating at times, with one noting that the AI character could be "bull headed." Users felt that there wasn't always enough time to prepare a list of arguments, and they might need more time to think about their answers.

One user was confused about what to do at the beginning of the simulation because the instructions mentioned that a report would be presented, but it never was. This led to a suggestion for clearer instructions at the start of the simulation, such as telling the user to ask for the report. Other users simply asked for the report and got it. A point of frustration was the possibility of not making any headway with a difficult character, which could encourage users to give up. But that was a realistic situation.

A suggestion for improvement was to provide more contextual background on the participants' roles and objectives, which would help guide their arguments (and is now a standard part of the introduction material).

CONCLUSION

The feedback from users of the AI simulations underscores a powerful idea: that AI assistants are a transformative tool far beyond a simple search engine. By leveraging the capabilities of these assistants, a developer can easily create a series of "small programs." These AI assistants can be tailored for a variety of tasks; I have only tried a few. But my experience demonstrates how AI assistants can become as second nature to a workflow as spreadsheets, offering a scalable and effective solution for a wide range of business needs.

The sims serve as a prime example of how this technology can create engaging and effective learning experiences. The simulations were praised for being more engaging and "super fun" than traditional, "very boring" compliance training. Users found them to be a "worthwhile exercise" for practicing soft skills like communication and critical thinking that are not easily learned from a book. The realistic nature of the simulations, complete with "complex situations" and "strong personalities," allowed professionals to prepare for difficult conversations and handle different on-the-job situations. Ultimately, these simulations empower professionals to move from wishing they had said something differently to knowing just what to say in the moment.

* * * * *

Author Byline: Dave Ingram, FSA, MAAA, CERA is a risk management actuary. He can be reached at daveingram@optonline.net.



Give us your feedback!

Take a short survey on this report.

Click Here



Bridging the Gap: How AI Changed My View of Actuarial Work

Sathiya Livingston, FSA, MAAA, CERA, PMP

Any views and ideas expressed in the essays are the author's alone and may not reflect the views and ideas of the Society of Actuaries, the Society of Actuaries Research Institute, Society of Actuaries members, nor the author's employer.

INTRODUCTION

I am not an actuary. I'm a Business Development Manager at Flatworld Solutions, a global insurance back-office service provider. For over a decade, I've worked alongside actuarial teams, underwriters, claims managers, and pricing specialists in the property and casualty space.

If you had asked me five years ago whether artificial intelligence would reshape my conversations with actuaries, I'd have laughed it off as tech hype. Yet here I am today writing about AI, not as an abstract buzzword, but as something that has changed the way I see actuarial decision-making from the frontlines of insurance operations.

This essay is a reflection from someone without an actuarial designation, but with a ringside seat to the ways AI has disrupted, challenged, and quietly improved the actuarial ecosystem.

When I first heard people say that *artificial intelligence would change everything*, I assumed it was the sort of buzz that fills conference halls but rarely seeps into daily work. In insurance operations, where I spend most of my time, the rhythms felt stubbornly the same: policy administration, claims triage, endorsements, and renewals. The actuarial teams always seemed like distant partners—people whose models influenced underwriting appetite and reserving assumptions, but who never directly touched my desk. Then, AI entered the picture. Not as a grand revolution, but as a set of tools—sometimes crude, sometimes dazzling—that reshaped the way I looked at actuarial practice from the frontlines of insurance.

A FRONT-ROW SEAT TO AI'S LEARNING CURVE

Our company began experimenting with machine-learning models to support commercial auto claims and property risk assessments. The promise was intoxicating: fewer manual errors, faster processing, and the ability to catch patterns invisible to the human eye. At first, I treated it as another automation initiative, like scanning paper files into PDFs a decade ago. But very quickly, I realized this was different.

One early experiment still makes me smile. We fed an AI system location data to help classify businesses. It returned with confidence that a sprawling warehouse complex was...a dry cleaner. Another time, it decided that "Phoenix Arms," a retirement community, was a firearms distributor. These weren't just funny mistakes; they revealed the gulf between what AI could infer and what an actuary—or even a junior underwriter—would conclude using judgment and context.

Actuaries, I noticed, didn't laugh these off the way I did. They immediately asked: *If the model is this wrong here, where else is it wrong in ways we can't see?* That was my first real appreciation of the actuarial mindset in the AI era: a blend of curiosity and suspicion. It wasn't cynicism—it was professional responsibility.

FROM DATA VOLUME TO DATA JUDGMENT

As AI systems got better, they began to handle tasks that previously consumed actuarial analysts: loss triangle completion, reserving scenario generation, and frequency-severity trend-spotting. On the operations side, we used AI to flag claims likely to escalate. The productivity boost was real. But what impressed me most was how actuaries insisted that every new efficiency be accompanied by new forms of judgment.

One actuary I worked with explained it simply: “AI doesn’t free us from thinking. It frees us from *repetitive thinking* so we can ask better questions.” That distinction hit me hard. I realized AI wasn’t just reshaping workflows; it was reshaping the *philosophy* of actuarial practice. The skill wasn’t memorizing tables or running macros. It was interrogating an opaque model and deciding whether to trust its output in the messy, high-stakes world of insurance.

THE FIRST HARD CONVERSATION

I still remember the first time I had to explain an AI-driven decision to a skeptical client. The model had flagged a mid-sized trucking fleet as “high-risk” based on telematics patterns. The client’s COO pushed back: “We just invested in safety upgrades—how can your system not see that?”

I found myself caught between technical opacity and business reality. I couldn’t “open the hood” of the algorithm, but I could explain how the data feeding it might lag behind real-world changes. What surprised me most was how the actuary on the call handled it. She didn’t defend the model. She reframed it: “Think of this score not as a verdict, but as a hypothesis. If your safety measures are real, then over the next quarter, the data should catch up, and the score should improve.”

That moment taught me that actuarial professionalism isn’t about defending tools—it’s about stewarding trust. The actuary didn’t need the model to be perfect; she needed the client to feel that the process was fair, transparent, and responsive. That small act of reframing did more for client confidence than any technical explanation could have.

RETHINKING TALENT AND TRAINING

Another consequence of AI’s arrival was a shift in what we looked for in talent. On the operations side, we no longer prized clerks who could memorize forms or follow scripts flawlessly; we needed analysts who could question anomalies, spot data quirks, and escalate when the “machine answer” didn’t make sense. Actuaries mirrored this shift. The most valuable ones weren’t those who could crunch numbers fastest, but those who could explain, in plain English, why a model’s strange output still made business sense—or why it should be disregarded.

I watched a younger actuary give a presentation where she compared AI models to “interns with PhDs.” They’re brilliant at narrow tasks but need supervision and translation before their work can be trusted. That image has stuck with me ever since. It also made me appreciate the evolving role of actuarial exams and professional development: not just building technical skill, but cultivating judgment, communication, and ethical reflexes in an AI-driven workplace.

GOVERNANCE, REGULATION, AND THE NEW FRONTLINE

AI didn’t just change actuarial practice internally; it created new expectations externally. Clients, regulators, and even the public wanted assurance that automated decisions were fair and unbiased. Here again, actuaries found themselves on the frontline.

I saw this play out during discussions about model governance for a client’s personal auto book. Regulators wanted transparency on how AI models flagged potential fraud. The client’s executives wanted speed. The actuaries were in the middle, tasked with showing not only *what* the model did, but *how reliable and equitable* its outputs were. It struck me that actuaries were becoming interpreters between technology and accountability—a role that felt entirely new, yet entirely consistent with their professional DNA.

In fact, I came to see actuarial governance as a kind of bridge. Data scientists could explain feature weights. Regulators could cite consumer protection statutes. But actuaries had the credibility to say: “Here’s what this means for solvency, for fairness, and for long-term risk.” That positioning—half technical, half fiduciary—may be the profession’s greatest advantage in the AI era.

Looking ahead, I can imagine actuaries playing an even more prominent role in setting standards for algorithmic accountability, much like they already do for reserves and solvency. If AI is the new engine of insurance, actuaries may well be its designated pilots.

THE HUMAN ELEMENT: SURPRISES AND LIMITS

Not everything was a triumph. Some of AI’s limits only became clear when humans collided with it. One example: in claims support, our AI model flagged a batch of suspicious submissions. Upon deeper inspection, several were flagged, not because of genuine fraud risk, but because the policyholders lived in neighborhoods with unusual street-name patterns that confused the algorithm.

We had to apologize, retrain the model, and rebuild trust with clients who felt unfairly profiled. What struck me was how quickly actuaries zeroed in on the fairness issue, not just the statistical one. For them, it wasn’t enough that the model “worked most of the time.” If it risked systemic bias, it was unacceptable. That sense of ethical responsibility reminded me that AI may accelerate actuarial work, but it doesn’t replace the profession’s moral compass.

FROM OUTSIDER TO APPRECIATOR

As someone who started outside the actuarial profession, I used to see actuaries as conservative, almost rigid. AI flipped that view. What I witnessed instead was a group adapting—sometimes reluctantly, but always thoughtfully—to a technology that threatened to undercut their traditional strengths.

And yet, rather than resist, they reframed their role: not as calculators, but as explainers; not as guardians of old methods, but as guarantors of fairness in new ones. That flexibility was inspiring. It showed me that actuarial identity is less about methods and more about values: prudence, clarity, accountability. AI didn’t dilute those values. It highlighted their importance.

LOOKING FORWARD: THE ACTUARY’S EXPANDING HORIZON

Where does this leave us? From my vantage point, AI is not an endpoint but a catalyst. It’s pushing actuarial practice in three important directions:

1. **From model building to model questioning.** Actuaries of the future will spend less time running code and more time asking whether the code reflects reality fairly and reliably.
2. **From isolated expertise to interdisciplinary leadership.** As AI touches regulation, ethics, and public perception, actuaries will increasingly be the ones connecting technical outputs to human consequences.

3. **From technical guardians to trust builders.** The most valuable actuaries will not be those who can out-code data scientists, but those who can ensure that every AI-driven decision strengthens—not erodes—the credibility of insurance.

These shifts excite me. They suggest a profession not shrinking under AI’s shadow but expanding into new relevance.

CONCLUSION: TRUST AS THE TRUE INNOVATION

If you had told me five years ago that I would gain a deeper appreciation for actuarial practice through AI, I would have laughed. To me, actuaries were the people behind the curtain, turning statistical wheels that the rest of us simply accepted. But AI dragged those wheels into the open. It forced us all—operations staff, clients, regulators and, yes, actuaries—to confront the uncertainty of machines and the enduring need for human judgment.

What I carry forward is this: AI is not the innovation that matters most. Trust is. And actuaries, with their long tradition of balancing numbers with prudence, are uniquely equipped to steward that trust.

In the end, AI didn’t just change actuarial workflows. It changed my view of actuaries themselves. Not as guardians of the past, but as navigators of a future where algorithms may drive the engines—but human values must steer the course.

* * * * *

Author Byline: Sathiya Livingston is a Business Development Manager at Flatworld Solutions, a global insurance back-office service provider. He can be reached at sathiya.l@flatworldsolutions.com.



Give us your feedback!

Take a short survey on this report.

Click Here



Insights from AI Use in Actuarial Practice

Prabhdeep Singh, FSA, MAAA, CERA, PMP

Any views and ideas expressed in the essays are the author's alone and may not reflect the views and ideas of the Society of Actuaries, the Society of Actuaries Research Institute, Society of Actuaries members, nor the author's employer.

Disclosure: At least some parts of this essay were written with the assistance of Gen AI.

INTRODUCTION

This essay reflects on an exploratory two-month journey where the author engaged intensively with AI tools in actuarial contexts. The perspective is, therefore, one of experimentation rather than finished production.

Generative AI (GenAI) is altering actuarial work, not just through efficiency gains, but also by expanding the boundaries of what actuaries are willing to attempt. It reduces the friction of learning new methods and tools while demanding careful oversight and judgment. The key insight is that AI functions as an enabler of experimentation and professional growth, while human validation and reflection ensure outputs remain credible and relevant.

AI AS A PRODUCTIVITY ENHANCER

GenAI has proven especially effective for coding. Without AI, building a simple life valuation system in a new programming language might have taken several months of trial and error or extensive research. With AI, the same work can be completed in a matter of days. AI is particularly effective in programming classical actuarial functions because they build on one another. For example, once the formula for the present value of whole life benefits has been implemented, it takes only a plain-language description to generate accurate code for term, endowment, or deferred insurance. From there, AI quickly suggests additional functions, including annuities, net level premiums, reserves. These examples demonstrate the scale of productivity gain, often approaching tenfold improvements.

The agent is not perfect at discerning intent and will occasionally make mistakes. If left uncorrected, these mistakes propagate into other functions. Human oversight remains critical to ensure accuracy and coherence.

AI AS A LEARNING ACCELERATOR

When actuaries encounter new tools, AI can provide overviews and answer context-specific questions. Learning becomes tailored to individual needs rather than generic categories. For instance, the author used AI to learn Python, XML, and LaTeX as needed to implement actuarial functions, parse SOA mortality tables, and document formulas in proper notation — all within a two-month, part-time effort. Similar assistance proved useful in setting up research and writing workflows (Obsidian/Zettelkasten) and project management workflows (OneNote/GTD™).

Learning requires intent and discipline. Actuaries must choose whether to learn just enough to instruct AI or to build deeper expertise in the tools themselves. The most successful will be those who gain sufficient depth to combine human expertise with AI's capabilities, multiplying productivity. This observation aligns with trends in software engineering, where experienced developers using AI effectively are in highest demand.

LIMITATIONS AND RISKS

Despite its benefits, GenAI has important limitations:

- **Performance:** Chat interfaces slow down in extended sessions; integrated tools such as GitHub Copilot are better suited for iterative code development.
- **Design Thinking:** AI can generate code fragments or refactor codebases efficiently but does not solve higher-order design challenges — for example, how to structure documentation or apply different programming paradigms (functional programming for actuarial logic, object-oriented design for interfaces).
- **Overconfidence and Credibility:** AI can present flawed outputs with confidence, encouraging misplaced certainty. For example, in one exercise it confidently generated a LaTeX expression for an actuarial symbol that had to be corrected manually. No English prompt could produce the proper placement of the superscript. This example illustrates why validation is essential. AI's speed in generating answers must be balanced by actuaries' responsibility to check outputs against actuarial standards. Some may argue this undermines the thesis that AI enables courage, since overconfidence risks eroding critical thinking. Recognizing this tension is important: actuaries must pair AI's speed with rigorous testing, peer review, and professional skepticism to ensure that courage does not slip into misplaced confidence.

EMERGING SOLUTIONS FOR CREDIBILITY

Several technologies are beginning to address the credibility issue, though they are not yet fully mature for actuarial practice:

- **Retrieval-Augmented Generation (RAG):** Combines large language models with targeted document retrieval, grounding outputs in trusted sources. While promising regulatory or actuarial documentation, current implementations struggle with ongoing changes to actuarial knowledge and regulations.
- **Knowledge Graphs:** Structured networks of actuarial concepts, assumptions, objects and relationships (actuarial science ontology) that can help AI reason more reliably. Knowledge graphs could eventually improve explainability and reduce hallucinations, but they are not yet widely developed for actuarial science.
- **Tooling:** Another approach to improving credibility is augmenting LLMs with external tools. Instead of relying on generated text alone, the model can call APIs or calculators to ground outputs in verifiable data. For example, a weather API can provide authoritative real-time information, and an actuarial calculator written in Python can return exact present value or reserve calculations. This "tool use" paradigm reduces hallucination risk and helps ensure the accuracy of answers.

Each of these approaches provides incremental improvements, but actuaries must still supply the professional judgment, testing, and domain expertise to interpret results responsibly. Each of these solutions demonstrates the same underlying point: AI can accelerate experimentation, but only actuarial validation ensures credibility. Whether through retrieval, structured graphs, or external tools, technology provides the scaffolding; judgment remains the foundation.

CONCLUSION

AI should not be seen as a replacement for actuarial expertise but as an accelerator of experimentation and learning. It lowers barriers to learning and broadens participation in technical work, yet its confident errors and lack of design capability highlight the continued importance of actuarial judgment and external feedback. The profession's path forward lies in balancing AI's acceleration with disciplined validation — using the technology to expand horizons while maintaining core standards of rigor, humility, and clarity.

From this exploratory journey, three lessons stand out:

- **Victories:** AI accelerated coding of actuarial functions and expanded learning into new tools like Python, XML, and LaTeX.
- **Frustrations:** AI's confident mistakes, performance slowdowns, and lack of design judgment exposed its limits.
- **Lessons Learned:** AI lowers barriers to experimentation, but real progress requires human validation, peer review, and disciplined skepticism.

This essay is drawn from an exploratory journey rather than client deliverables. It illustrates the victories, frustrations, and lessons learned in early experimentation, which many actuaries will find valuable. By sharing candid reflections, actuaries can collectively shape how AI is integrated into professional practice, ensuring it becomes a tool for sound judgment rather than misplaced confidence.

* * * * *

Author Byline: Prabhdeep Singh, FSA, MAAA, CERA, PMP is an independent consulting actuary. He can be reached at prabhdeep.singh.actuary@gmail.com.



Give us your feedback!

Take a short survey on this report.

Click Here



Actuarial A.I 1.0

Nathan Worrell, FSA; Green Chen, FSA; Brandon Lin, FSA

Any views and ideas expressed in the essays are the author's alone and may not reflect the views and ideas of the Society of Actuaries, the Society of Actuaries Research Institute, Society of Actuaries members, nor the author's employer.

HELLO WORLD!

What would it be like to have a helper that knows all the ins and outs of your model, was able to provide documentation in an instant, and could provide modelling suggestions for your every whim?

The ability to trust and understand a model is as important as it has ever been. Everything seems to be a little more complicated in the actuarial space that it once was: increased volume of regulatory disclosures, more robust capital frameworks, and ever-expanding model sophistication.

The actuary's time is very precious. Adding efficiency by shortening learning curves and facilitating documentation exercises can give more time back to actuaries for analytical work.

At our company, AI was adopted and shared rather quickly. Overnight, 10,000 plus workers became kids with new toys, encouraged to play, invent, and innovate with AI. The future was here.

There were some guardrails. The earliest AI utilities were kept within corporate walls, only allowed to look at our own documents. Still, a flood of ideas came from these early experiments, as well as a new library of prompts and inspiration for many ways to enhance the actuarial experience.

One of the initial projects we took on was an attempt to create an AI "Navigator" that had access to all the documentation and support material behind our actuarial projection software. Would it be the perfect helper? Could it understand actuarial lingo? We set off to find out.

LEARNING CURVES

The software had 30 plus years of development, thousands of pages of help text, and numerous specially focus user guides covering topics like GAAP LDTI, IFRS 17, and Capital Frameworks across the globe. No one person at the company knew everything about the software (although an argument could be made for our Chief Actuary), so this would augment the client support staff in a powerful new way.

We formed a dedicated team to populate a repository of information with existing resources.

We started small, isolating things to one segment of the software before expanding further. Test, revise, repeat.

Some of our early learnings:

- **Partition documentation into related sets of information and direct the tool to look in these specified areas.** This helps reduce false answers (“hallucinations”). For example, if I want to know how the model handles partial withdrawals on an annuity, I don’t need to reference the documentation on life reserves.
- **Source material needs to be formatted well, accurate and current.** Some early prompts revealed older processes that weren’t the latest best practice or could only provide an inconsistent or incomplete answer.
- **Prompting suggestions are important to develop along with the AI utility.** While the tools are not well versed in actuarial context, providing a pre-prompt that might be considered overly verbose may end up being appropriate. For example, "Consider 'experience study' to be distinct from 'experience adjustment'." This type of clarifying statement can help it distinguish between an assumption setting subject and a reinsurance topic.

Similar lessons were coming from other AI applications being developed and tested, including email thread summarizers, coding agents, and research assistants.

PERFORMANCE REVIEW

Some of the earliest encouraging moments were when the tool compiled documentation from multiple sources in a coherent, well organized conversational summary. Previously, these bits of information were in scattered locations, which caused a drag on time for staff or clients to search for and locate needed reference material. Being able to bring these disparate chunks of information together is a huge win.

Additionally, after initial prototypes were proven to be successful, expanding the tool to incorporate more information was not a difficult extension.

One key piece of functionality is a rating system (a.k.a “training”) that would allow a user to give feedback about whether the answer was satisfactory. Different categorial ratings are offered: “not detailed enough” or “incomplete,” “incorrect,” etc.

The other helpful feature that allows for easier validation is embedded links to source materials via citations.

Some struggles we encountered:

- Alphabet soup – the actuarial world is full of acronyms, and that could become confusing. Is NPR a net premium ratio, a net premium reserve, or something else?
- Blind spots. Because we intentionally restricted the universe of information for the tool, anything that lacked documentation would impede the tool’s success. This could lead to negative ratings due to “incomplete” answers, or the tool being unable to provide a response at all.
- Assisting with user-based code presented a challenge because of the required rules and syntax of the underlying software. Continued development is underway to help support this effort. However, we did find that the tool could provide great summaries of what pre-existing code was doing, including clean presentations of input and output variables.

PREPARING FOR PRIME TIME

Before its public release, the tool is being Beta tested by a select group of users. Phased deployment is a very useful technique in software development.

In our outreach, we had nervous enthusiasm. Insurance companies seem to be tilting towards cautious adoption of AI tools. There are understandable concerns about exposing private information or intellectual property to these utilities. Having clear and demonstrable guardrails will be a continued important feature to ensure successful deployment in the industry.

We also needed to provide adequate training, which included some examples of how to use the tool and suggested prompting.

Following customer feedback, we will get back under the hood and make some tweaks and then hopefully release version 1.0.

LOOKING AHEAD

Our hope is that this tool will begin to plant the seeds for that perfect actuarial assistant that can resolve governance woes and facilitate learning. In the meantime, this type of AI engine will deliver a better customer experience by driving higher quality engagement with our support staff. The practice of actuarial science is nuanced and often involves judgement, which is a difficult thing to train into an AI agent. There will continue to be a role for human-to-human actuarial discourse in the foreseeable future, but there will also be plenty of augmented human-to-augmented human discourse as well.

This is just the first step in our actuarial AI journey. Our helper can aide in understanding a concept or feature in the model. There are several areas where we see continued usefulness.

- **Model Documentation:** We know actuaries and programmers *love* documentation. Or maybe they see it as a necessary evil. It is crucial for model governance and certainly having a robot do the heavy lifting here will be a time saver.
- **Model Building:** Beyond actuarial applications, the world of tech is moving to low or no code interfaces. Simply prompt about what you need and let AI do the rest. In the actuarial world, maybe this is a pricing model that can emerge from a product specification. To accomplish this, model providers will need to design/build APIs for agents and write documentation specifically to guide agents.
- **Results Analysis:** From experience studies to period-to-period valuation investigations, frameworks for helping to highlight adverse results will be yet another time saver. This may be especially helpful for large and complex organizations, and also useful when there are many components to reported results (like IFRS-17 and GAAP Roll-forwards).

The future is here. Actuarial and insurance workflows will encounter artificial intelligence all along the way including automation steps, assistance for performing actuarial work, agents to aid in understanding and synthesizing complex information, writers of documentation, and even quantitative analysis. Knowing how to approach these tools, providing solid foundations, iterating frequently, and preparing carefully for launch is critical to ensure success.

* * * * *

Author Byline: Nathan Worrell, FSA, Green Chen, FSA, and Brandon Lin, FSA, are actuaries at Moody's. They can be reached at Nathan.Worrell@moody.com, Green.Chen@moody.com, and Brandon.Lin@moody.com, respectively.



Give us your feedback!

Take a short survey on this report.

[Click Here](#)



About The Society of Actuaries Research Institute

Serving as the research arm of the Society of Actuaries (SOA), the SOA Research Institute provides objective, data-driven research bringing together tried and true practices and future-focused approaches to address societal challenges and your business needs. The Institute provides trusted knowledge, extensive experience and new technologies to help effectively identify, predict and manage risks.

Representing the thousands of actuaries who help conduct critical research, the SOA Research Institute provides clarity and solutions on risks and societal challenges. The Institute connects actuaries, academics, employers, the insurance industry, regulators, research partners, foundations and research institutions, sponsors and non-governmental organizations, building an effective network which provides support, knowledge and expertise regarding the management of risk to benefit the industry and the public.

Managed by experienced actuaries and research experts from a broad range of industries, the SOA Research Institute creates, funds, develops and distributes research to elevate actuaries as leaders in measuring and managing risk. These efforts include studies, essay collections, webcasts, research papers, survey reports, and original research on topics impacting society.

Harnessing its peer-reviewed research, leading-edge technologies, new data tools and innovative practices, the Institute seeks to understand the underlying causes of risk and the possible outcomes. The Institute develops objective research spanning a variety of topics with its [strategic research programs](#): aging and retirement; actuarial innovation and technology; mortality and longevity; diversity, equity and inclusion; health care cost trends; and catastrophe and climate risk. The Institute has a large volume of [topical research available](#), including an expanding collection of international and market-specific research, experience studies, models and timely research.

Society of Actuaries Research Institute
8770 W Bryn Mawr Ave, Suite 1000
Chicago, IL 60631
www.SOA.org