



CONSIGLIO NAZIONALE
DEGLI ATTUARI



CONSIGLIO ORDINE
NAZIONALE DEGLI ATTUARI

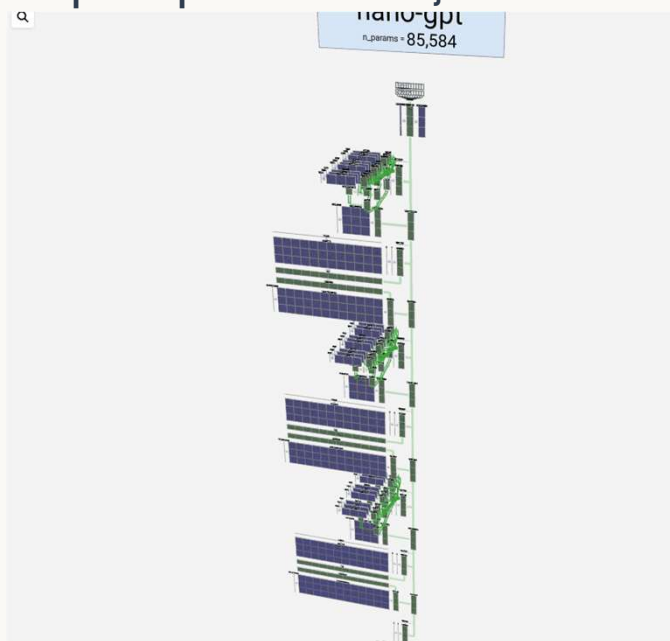
AI: Opportunità, sfida e nuovi stimoli per la professione Attuariale

Claudio Senatore Reso – Ordine degli Attuari – Task Force AI

Università degli Studi di Verona -Pomeriggio Attuariale – 30 Gennaio 2026

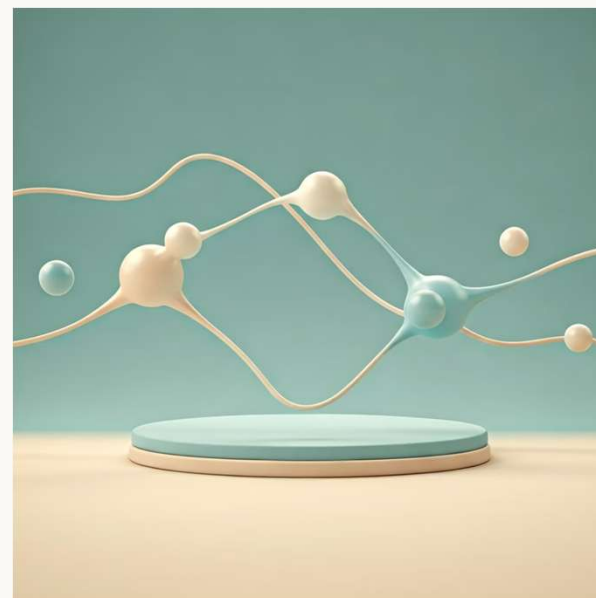
Lo stupore come motore della conoscenza: AI e Ricerca

Un principio condiviso, due domini



Nei modelli di Intelligenza Artificiale

- Lo stupore è formalizzato come $\text{surprisal}(x) = -\log p(x)$
- Misura lo scarto tra aspettativa del modello e dato osservato
- L'apprendimento consiste nel ridurre sistematicamente questo scarto
- Lo stupore è il segnale che guida l'aggiornamento del modello



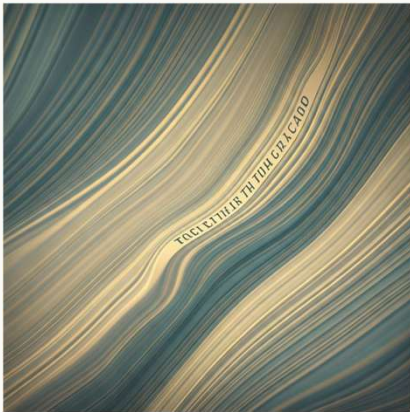
Nella ricerca

- Lo stupore emerge quando l'osservazione contraddice un'ipotesi implicita
- Da Platone ad Aristotele: lo stupore segnala un limite della comprensione
- La ricerca nasce non dalla conferma, ma dallo scarto tra teoria e realtà
- Lo stupore orienta la formulazione di nuove domande e modelli

Il costo della persuasione scalabile: quando l'apparenza di competenza sostituisce la competenza reale

Gli LLM possono dare 'apparenza di sapere' senza comprensione

📄 **Domanda guida:** quando un discorso è conoscenza e quando è solo persuasione?



Il rischio non è produrre risposte: è scambiare fluidità per verità



Nel mito di Theuth (scrittura): memoria delegata → illusione di sapienza



Con GenAI: la persuasione diventa scalabile → cresce il valore di verifica e competenza

“

Intorno a ogni cosa, uno solo è il principio per quelli che devono prendere buone decisioni: bisogna conoscere la cosa su cui si devono prendere decisioni, altrimenti per necessità si sbaglia tutto. Ma i più non si accorgono di non sapere che cosa sia l'essenza di ciascuna cosa; e, come se lo sapessero, non si mettono d'accordo all'inizio della ricerca e, proseguendo, ne subiscono le naturali conseguenze, perché non si accordano né con sé stessi né fra di loro (Fedro 237b7 – 237c8).

” aio 2026

IA e nuovo modo di fare ricerca/lavoro

Localizzare, interpretare e sintetizzare ciò che altri fanno (in modo verificabile)

Localizzare: Dove sta l'informazione (paper, norme, report, codice, dataset)

Interpretare: Contesto, assunzioni, qualità della fonte, limiti

Sintetizzare: Ricombinare ciò che serve per decisioni e ricerca (senza confondere forma e verità)

Anche in ambito attuariale cresce la co-autorialità: più competenze integrate, più collaborazioni

Esempio (ASTIN Bulletin):

1985 ~ 1-2 autori/paper → 2024 spesso 3-4

Mega-team:

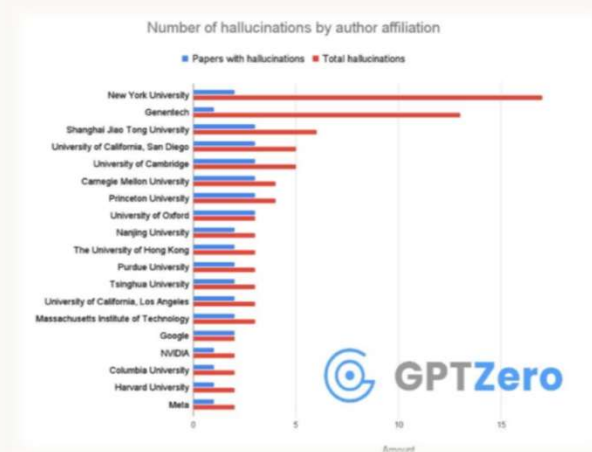
Higgs/ATLAS (2012) ~2.932 autori | LHC (2015) 5.154 | COVIDSurg (2021) 15.025

Denario: Assistente scientifico end-to-end

Denario è un sistema AI multi-agente per la ricerca scientifica end-to-end. Genera idee, verifica la letteratura esistente, sviluppa e esegue piani metodologici, analizza dati, crea grafici e redige articoli scientifici. Questo dimostra come l'AI possa orchestrare automaticamente molte fasi della ricerca interdisciplinare, combinando conoscenze da diverse discipline.

Theorizer: Sintetizzatore automatico di teorie

Theorizer è uno strumento AI che sintetizza automaticamente migliaia di teorie analizzando l'intero corpus di ricerca di un campo scientifico. A differenza di strumenti che riassumono singoli articoli, Theorizer identifica regolarità—pattern ricorrenti in molteplici studi—ed esprime teorie testabili strutturate in tre componenti: LAW (legge qualitativa o quantitativa), SCOPE (condizioni di validità ed eccezioni), ed EVIDENCE (prove empiriche tracciate da articoli specifici). Ogni teoria include nome e descrizione contestuale, contenendo tipicamente una o due leggi.



La miglior partenza è quella radicata nella realtà

L'Insurance Judgment Experiment da «Noise, a Flaw in Human Judgment» – Kahneman, Sibony & Sunstein (2022)

Impostazione sperimentale

- Dirigenti di compagnie assicurative hanno preparato descrizioni dettagliate di casi per sottoscrittori e liquidatori sinistri
- Gli underwriter hanno valutato 4 casi di assicurazione contro gli infortuni sul lavoro
- I liquidatori hanno valutato 6 casi di sinistri danni auto
- Lo studio ha raccolto 86 valutazioni da 48 underwriter e 113 valutazioni da 68 liquidatori sinistri
- I partecipanti hanno valutato i casi in modo indipendente e non erano informati che lo studio analizzasse la variabilità del giudizio

Risultati chiave

- I dirigenti prevedevano solo un 10% di variabilità tra i giudizi dei professionisti (il 15% era la seconda stima più frequente)
- Realtà per gli underwriter: differenza mediana del **55%** tra giudizi su casi identici
- Realtà per i liquidatori sinistri: differenza mediana del **43%** tra giudizi
- In metà dei casi, le differenze hanno superato questi valori mediani

Impatto reale

- Esempio: a fronte di un premio di 9.500 \$ assegnato da un underwriter, un altro ha tipicamente fissato 16.700 \$ (non 10.500 \$)
- La coerenza dei dirigenti nella valutazione professionale è stata fortemente sovrastimata
- Questo rumore non riconosciuto genera **impatti finanziari rilevanti** ed **esperienze cliente incoerenti**

IA in assicurazione: evidenze 2024–2025 e opportunità per l'attuariato

Adozione in accelerazione e nuovi ruoli per la professione

1 Adozione in accelerazione

Nel 2025 ~9 assicuratori su 10 sono già in fase di valutazione/uso della GenAI; ~55% è in early o full adoption e la spesa dichiarata in GenAI è circa raddoppiata in 12 mesi.

2 Dove impatta "operativamente"

I casi d'uso più citati nel settore riguardano workflow automation, underwriting/pricing decision support, claims (triage, documenti) e attività di knowledge-work "a valle" (ricerca, sintesi, drafting) — cioè aree dove l'attuario lavora su evidenza, ipotesi e controlli.

3 Evoluzione della funzione attuariale

In un report 2025, oltre il 70% dei rispondenti indica che gli attuari sono coinvolti anche in Data Science & Analytics, Cloud & Data Strategy, Generative AI e modernizzazione dei sistemi finance (non solo "modeling" tradizionale).

4 Benefici misurabili, ma la differenza la fa la scala

I "copilot/assistenti" possono aumentare la produttività dei knowledge workers di oltre il 30%, ma solo ~7% degli assicuratori ha scalato l'AI a livello enterprise; circa 2/3 restano in fase di pilot/POC (opportunità: chi industrializza prima crea vantaggio).

5 Outlook (trend enterprise)

Gartner stima che entro il 2028 l'agentic AI possa supportare ~15% delle decisioni operative quotidiane e comparire in ~33% del software enterprise; allo stesso tempo, avverte che >40% dei progetti agentic potrebbe essere cancellato entro il 2027 se non agganciato a valore e controlli.



IA Adozione – Utilizzo – Sviluppo

Flussi di sviluppo

1

1° Ondata

Gamification: Drag&Drop, Point&Click

2

2° Ondata

Gamification: Copilot, come assistente

3

3° Ondata

Gamification: Copilot, come agente

4

4° Ondata

Real World Models?

Funzionalità Macro



Funzionalità Micro

 **SAS Claims Insight Use Case:** Estrae pattern semantici dai sinistri per analisi attuariale e segmentazione del rischio

Scopo: la causa n.1 dei fallimenti (insieme ai Dati)

Ottimizzare perfettamente il bersaglio sbagliato



Obiettivi multipli non gerarchizzati

Conducono al caos e alla dispersione di risorse, rendendo difficile definire una direzione chiara per il sistema IA.



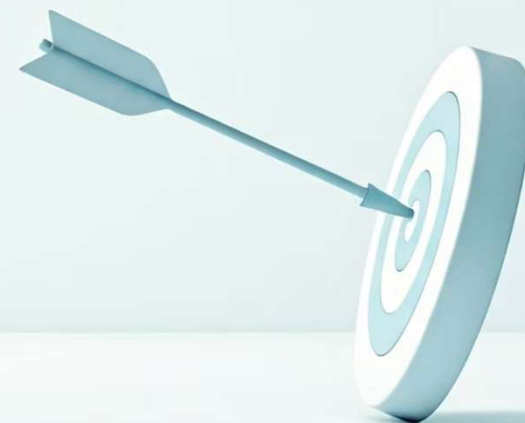
Trade-off non dichiarati

Portano a scelte incoerenti e a sistemi che non performano come previsto in situazioni critiche, poiché le priorità non sono esplicite.



Esempi di trade-off critici

- Accuratezza vs inclusività
- Velocità vs affidabilità
- Breve vs lungo periodo



Conoscenza: previsione $\neq$ comprensione $\neq$ decisione

Serve disciplina di validazione

- ☐ Correlazione $\neq$ causalità
- ☐ Confidenza dichiarata $\neq$ confidenza meritata
- ☐ Output pulito può essere fragile fuori campione
- ☐ Servono: robustezza, drift, calibrazione dell'incertezza, protocolli di verifica



Quando la mappa sostituisce il territorio

Misurare meglio non sostituisce pensare meglio

KPI e proxy diventano definizioni operative del concetto

Si ottimizza il numero, non l'esperienza reale

Serve: definizione ricca dell'obiettivo → poi metriche coerenti

"The map is not the territory."

A. Korzybski, Science and Sanity (1933)

Quattro capacità per sistemi affidabili

Da output a outcome



Scopo (teleologia)

Cosa vogliamo ottenere davvero? (priorità, trade-off)

AI Act: "scopo previsto / intended purpose" → lo scopo dichiarato guida obblighi, responsabilità e condizioni d'uso



Conoscenza (epistemologia)

Cosa conta come evidenza affidabile? (incertezza, validazione)

AI Act: "risk management + performance" → dimostrare affidabilità con dati adeguati, test/valutazioni, robustezza e monitoraggio



Realtà (ontologia)

Quali categorie usa il sistema per "vedere" il mondo?

AI Act: "data governance + trasparenza" → dati rappresentativi e controllati, limiti/contesto esplicitati, informazione chiara all'utente su capacità e vincoli



Etica

Come gestiamo casi nuovi e scelte delicate?

AI Act: "diritti fondamentali + controllo umano" → prevenire danni, garantire supervisione umana e tracciabilità/accountability nelle decisioni critiche

IA in assicurazione: etica, responsabilità umana e progettazione consapevole

L'etica dell'IA non è nelle macchine, ma nella responsabilità umana di progettazione e uso.

Posizioni nel dibattito: una mappa rapida

	IA come Strumento e Agente di Impatto L'IA produce effetti eticamente rilevanti, ma la normatività resta in chi definisce obiettivi, dati, vincoli, controlli e decide come usarne l'output. Qui l'espressione "AI etica" è spesso una scorciatoia linguistica.
	Machine Ethics Filone che prova a incorporare vincoli e valutazioni etiche nel design (rule-based, learning-based, value alignment), distinguendo tra "sistemi con considerazioni etiche implementate" e "agenti morali pieni".
	Agenti Artificiali e Responsabilità Distribuita Alcuni autori discutono se gli agenti artificiali possano essere "moral agents" in senso funzionale. La governance mira comunque a mantenere responsabilità e controllo chiaramente assegnati e tracciabili.

Implicazioni operative per assicurazioni e attuari

		
Chi risponde? Tracciabilità e Supervisione Serve una catena di responsabilità esplicita (provider/deployer, owner di processo, validatore, escalation) e una chiara supervisione umana sulle decisioni.	Etica della Progettazione (ex ante) Gestione dei rischi, qualità dei dati, robustezza, documentazione e logging sono cruciali per ridurre bias e possibili errori prima del "go-live".	Etica dell'Uso (in esercizio) Definire limiti d'uso, monitorare il "drift", effettuare audit e controlli su casi ad alto impatto (pricing, claims triage, frodi) ribadendo che "l'output non è una decisione finale".

Riferimenti: G. Ryle, The Concept of Mind (1949) – “category mistake”; Regolamento (UE) 2024/1689 (AI Act) – obblighi, human oversight, tracciabilità, responsabilità; J. Stenseke (2022) – panoramica su machine ethics; L. Floridi & J.W. Sanders (2004) – “morality of artificial agents”.



Perché gli attuari possono essere centrali

Governare incertezza, metriche e decisioni

Ragionare sotto rischio/incertezza/ambiguità

(non solo "predire")

Better be convex than smart → meglio progettare decisioni/prodotti che beneficino della volatilità e limitano le perdite, invece di puntare tutto su previsioni "brillanti".

Costruire metriche

che non tradiscono i concetti (prima definire bene l'obiettivo → poi KPI coerenti)

Stress test, robustezza

sensibilità, scenari, monitoraggio drift e degradazione delle performance

Governance

obiettivi, evidenza, categorie operative, controlli e responsabilità

Iniziative IA nella professione attuariale

Un panorama globale ed europeo in rapida evoluzione

ONA – Ordine Nazionale Attuari

- Task Force "Intelligenza Artificiale" (2025): area Innovazione
- Revisione del Codice Deontologico
- Corso FAC-SIA 2025: "AI generativa e prompt engineering" (Mar 2025)
- Corso FAC 2024: Machine Learning per riserve sinistri (Gen 2024)
- Webinar storici su data/AI (dal 2021)

ONA Task Force AI componenti:

- *Francesca Liguori*
- *Danilo Aulicino*
- *Manuel Caccone*
- *Claudio Senatore Reso*

AAE – Actuarial Association of Europe

- AI & Data Science Working Group (AI-DS WG)
- Discussion paper "What should an actuary know about AI?" (Gen 2024)
- Webinar "AI Regulation in Europe" (Giu 2025): focus EU AI Act
- Webinar "AI education for actuaries" (Giu 2025): da ML a GenAI
- Linee guida CPD Data Science (incluso AI)
- GitHub con Use Case

AI&Data Science Working Group membri italiani:

- *Manuel Caccone*
- *Claudio Senatore Reso – Vice Chair*

IAA – International Actuarial Association

- AI Task Force (AITF): guida iniziative globali su rischi/opportunità AI
- Statement of Intent (SOI): documento quadro per "responsible use"
- AI Summit Singapore (Apr 2024): kick-off attività AITF
- Interim Report 2024: workstream e risultati iniziali
- Paper "AI Governance Framework" (Nov 2025)
- Paper "Documentation of AI Models" (2025)
- Webinar su AI Agents e EU AI Act

AI Use Case Working Group membri italiani:

- *Manuel Caccone*
- *Claudio Senatore Reso*