

Apri la flangia: se l'allucinazione dell'IA rischia di uccidere

□ IL PUNTO DI VISTA DI CLAUDIO DELAINI

“Claudio, ma l'IA ci fa risparmiare un sacco di tempo!”

Qualche settimana fa ero in videochiamata con un cliente. Mi fa vedere tutto orgoglioso il loro nuovo sistema: un'intelligenza artificiale che risponde ai manutentori in tempo reale.

“Guarda quanto è veloce! Gli chiedo come sostituire una guarnizione e in 10 secondi mi dà la risposta completa.”

Gli dico: “Bellissimo. Ma dimmi una cosa: quando l'IA ti dice ‘sostituisci la guarnizione’, ti dice anche di spurgare prima il sistema? Di verificare la pressione residua? Di aspettare il raffreddamento?”

Silenzio. Poi: “Beh, quello è implicito, no?”

Ed eccoci al problema. Per noi esperti è implicito. Per l'IA no. E per un manutentore junior che usa quel sistema per la prima volta, ancora meno.

Il mio incubo ricorrente è questo: che l'intelligenza artificiale diventi come quei manuali di istruzioni tradotti male dal cinese. Tecnicamente corretti, ma praticamente letali. Perché ti dicono “cosa fare” senza dirti “come farlo in sicurezza”.

E la cosa più insidiosa? L'IA non ti dà risposte palesemente sbagliate che ti fanno scattare l'allarme. Ti dà risposte verosimili, sensate, quasi giuste. E proprio per questo sono pericolose: perché ti fidi.

Il consiglio che do sempre: tratta ogni risposta dell'IA come tratteresti il

Apri la flangia: se l'allucinazione dell'IA rischia di uccidere

consiglio di uno stagista al primo giorno. Può essere utile, può essere corretto, ma tu devi validarlo prima di agire. Sempre. Anche se hai fretta. Anche se sembra ovvio. Anche se l'hai usata 100 volte prima.

Perché quella 101esima volta, quella risposta incompleta che non specifica le precauzioni di sicurezza, potrebbe essere l'ultima.

Oggi vorrei parlarti di un caso che mi è stato raccontato da un collega, e che mi fa venire i brividi ogni volta che ci penso. Non solo perché è particolarmente cruento, ma perché rappresenta esattamente il tipo di incidente che potrebbe succedere domani, in qualsiasi fabbrica che sta implementando sistemi di intelligenza artificiale per assistere i manutentori.

E la cosa più inquietante? In quel caso specifico l'IA non c'entrava nulla. Era stata una decisione sbagliata umana. Ma l'ipotesi da cui parto è, invece: e se fosse stata l'intelligenza artificiale a dare quel consiglio mortale?

Il caso della flangia

Vi racconto la storia come me l'hanno raccontata.

Siamo in un impianto industriale. Un manutentore ha un problema tecnico: il valore della pressione non appare più sul monitor di controllo. Deve intervenire su una flangia – un giunto di collegamento tra due tubi o elementi di un sistema in pressione.

Oggi, in teoria, un manutentore in questa situazione potrebbe interrogare un sistema di intelligenza artificiale per la ricerca di guasti. Magari un “cervello aziendale” – uno di quei mega software che assorbono le competenze dei tecnici esperti e le rendono disponibili a distanza.

Immaginiamo dunque che il manutentore immetta la richiesta: “Senti, il valore della pressione non c'è più a monitor. Cosa può essere?”

La risposta dell'IA potrebbe essere tecnicamente corretta: “Apri la flangia per verificare il sensore”.

Dal punto di vista logico, ha senso. Se il sensore di pressione non dà più lettura, può essere guasto, può essersi staccato, può essere ostruito. Per verificarlo, devi aprire.

Apri la flangia: se l'allucinazione dell'IA rischia di
uccidere

Ma c'è qualcosa che il software di assistenza "intelligente" non dice.

Non dice: "Stai attento che il sistema potrebbe essere in sovra-pressione. Aprila piano piano, con le dovute precauzioni, in modo tale che non ci sia il pericolo di espulsione violenta del contenuto."

La conseguenza reale

Nel caso specifico che mi è stato raccontato è successo questo: quando l'operatore ha aperto la flangia – senza le dovute precauzioni, senza verificare prima lo stato reale della pressione residua nel sistema – è partito un piccolo proiettile.

Quel proiettile gli ha bucato la giugulare.

Si tratta un infortunio mortale, reale, documentato. Non è un'ipotesi teorica, non è uno scenario da manuale. È successo davvero.

Ripeto, in quel caso non era stata l'intelligenza artificiale a dare il consiglio sbagliato. Era stata una decisione umana sbagliata, presa probabilmente per fretta, per superficialità, per mancanza di esperienza.

Apri la flangia: se l'allucinazione dell'IA rischia di uccidere

Ma ora fate questo esercizio mentale con me: e se invece fosse stata l'IA a suggerire "apri la flangia" senza specificare le precauzioni?

Il problema delle informazioni mancanti

Qui arriviamo al cuore del problema dell'intelligenza artificiale applicata alla manutenzione e alla sicurezza.

L'IA impara da dati storici. Impara da procedure scritte, da manuali, da interviste con tecnici esperti. Ma c'è una categoria di informazioni che spesso manca nei dati scritti: il contesto di sicurezza implicito.

Quando un manutentore esperto dice a un collega "apri la flangia", in quella frase ci sono un sacco di cose implicite:

- *Ovviamente* verifichi prima se c'è pressione residua
- *Ovviamente* aspetti che il sistema si sia raffreddato se necessario
- *Ovviamente* usi i DPI appropriati

Apri la flangia: se l'allucinazione dell'IA rischia di uccidere

- *Ovviamente* lo fai gradualmente, non di colpo
- *Ovviamente* ti posizioni lateralmente, non di fronte

Queste cose un esperto le sa senza doverle dire esplicitamente. Fanno parte del suo know-how tacito, di quella conoscenza che ha sviluppato con l'esperienza.

Ma quando quel know-how viene trasferito a un software attraverso interviste, procedure scritte, manuali, spesso questi aspetti di sicurezza impliciti non vengono catturati. Perché sembrano ovvi. Perché nessuno pensa di doverli specificare.

E quindi il software impara “apri la flangia” senza imparare il sottotesto “con queste 7 precauzioni”.

I software “seri” con link alle fonti

Qualcuno dirà: “Ma nei software seri, sotto la risposta ci sono i link alle fonti. Puoi cliccare e vedere da dove viene l'informazione.”

Apri la flangia: se l'allucinazione dell'IA rischia di uccidere

È vero. Nei sistemi ben progettati, ogni suggerimento dell'IA è accompagnato da riferimenti alla documentazione originale. Magari il link al manuale del macchinario, alla procedura interna, alla norma tecnica.

Ma ci sono due problemi:

1) Il manutentore può non cliccare sul link.

Perché? Perché ha fretta. Perché dopo N. volte che il sistema gli dà risposte corrette, si fida. Perché pensa "io ho 30 anni di esperienza, so fare il mio lavoro, questa risposta mi sembra sensata, eseguo".

E quella volta che la risposta è incompleta o ambigua, lui non se ne accorge.

2) La fonte può essere in una lingua diversa.

Se il manuale originale è in tedesco, in cinese, in inglese tecnico specialistico, quanto è probabile che il manutentore italiano clicchi, legga, comprenda nel dettaglio?

Teoricamente dovrebbe farlo. Praticamente, tendenzialmente si fida della sintesi che gli ha dato l'IA nella sua lingua.

Apri la flangia: se l'allucinazione dell'IA rischia di uccidere

Guarda i video su intelligenza artificiale e sicurezza in fabbrica sul mio canale YouTube



Di chi è la colpa? La zona grigia della responsabilità

Ed ecco che arriviamo alla domanda da un milione di euro: se l'intelligenza artificiale dà un consiglio sbagliato e qualcuno muore, di chi è la colpa?

Apri la flangia: se l'allucinazione dell'IA rischia di uccidere

Facciamo l'elenco dei possibili responsabili:

- Il manutentore

Si potrebbe dire: il manutentore deve avere la competenza per interpretare se la risposta dell'IA è valida o no. Deve validare l'informazione prima di agire.

Ma se il manutentore avesse già tutta quella competenza, probabilmente non avrebbe bisogno di interrogare l'IA. La usa proprio perché ha un dubbio, perché non è sicuro, perché vuole una conferma.

È colpa del manutentore se si fida di uno strumento che il datore di lavoro gli ha messo a disposizione come supporto ufficiale al lavoro?

- Il datore di lavoro

Il datore di lavoro ha messo a disposizione dei lavoratori uno strumento – il sistema di IA – senza validare le risposte che dà. Perché non può validarle: sono generate dinamicamente, sono sempre diverse, cambiano a seconda della query.

Apri la flangia: se l'allucinazione dell'IA rischia di uccidere

Quindi ha messo a disposizione un'intelligenza artificiale che, secondo la legge 132/2025, dovrebbe essere "sicura, affidabile, trasparente". Ma non ha modo di garantire che lo sia effettivamente.

Ha fatto una valutazione del rischio dell'uso di questo strumento? Come si valuta il rischio di uno strumento che può dare risposte imprevedibili?

- Il fornitore del software

Il fornitore potrebbe dire: "Noi abbiamo fornito il software con tutti i disclaimer. Abbiamo specificato che è uno strumento di supporto, non sostitutivo della competenza umana. Abbiamo messo i link alle fonti. Non possiamo essere responsabili di come viene usato."

E avrebbe ragione, fino a un certo punto.

- Il fornitore del motore di IA sottostante

Ma il fornitore del software – magari un'agenzia italiana che sviluppa soluzioni personalizzate – a sua volta si appoggia a un motore di intelligenza artificiale estero. Watson di IBM, OpenAI, Anthropic, Google, servizi cinesi.

Apri la flangia: se l'allucinazione dell'IA rischia di
uccidere

E quindi c'è una cascata di mancanza di informazioni: io uso un software che si appoggia a un motore che non controllo, basato su algoritmi che non conosco, addestrato su dati che non ho visto.

Come si attribuisce la responsabilità in questa catena?

La risposta: non lo sappiamo

La verità è che siamo nel Far West. Non c'è giurisprudenza consolidata. Non ci sono precedenti. Non ci sono linee guida chiare su come attribuire la responsabilità quando un sistema di IA dà un consiglio che porta a un infortunio mortale.

Il primo caso farà da apripista. E fino ad allora, tutti – costruttori, utilizzatori, fornitori, manutentori – operano in una zona grigia di incertezza legale.

Leggi anche: [AI e responsabilità giuridiche: sfide e normative in evoluzione](#)

Apri la flangia: se l'allucinazione dell'IA rischia di uccidere



Procedure generate dall'IA: quando la carta tradisce la realtà

Il problema non riguarda solo i sistemi di assistenza in tempo reale. Riguarda anche i documenti che l'IA ci aiuta a produrre. In particolare: le procedure operative.

Come si fanno le procedure (metodo tradizionale)

Vi racconto come faccio io.

Quando devo scrivere una procedura operativa per un cliente, non mi siedo alla scrivania a inventarla. Vado in fabbrica. Filmo come l'operaio fa il lavoro. Poi parlo con lui: cosa sbaglia? Cosa è giusto? Perché lo fa in quel modo? Perché non lo fa in quest'altro?

E insieme – io con la mia esperienza su tanti contesti diversi, lui con la sua esperienza su quella macchina specifica – creiamo la procedura.

Non pretendo di saper fare il lavoro al posto suo. Gli racconto quello che non mi convince, gli dico i miei motivi basati su infortuni che ho visto altrove, su errori comuni, su rischi che lui magari non considera perché ci è abituato.

Si fa una trattativa. Si ragiona insieme.

Questo processo richiede tempo. Richiede di andare in fabbrica, di osservare, di dialogare, di correggere bozze, di validare sul campo.

Apri la flangia: se l'allucinazione dell'IA rischia di uccidere

Il metodo IA: prompt engineering

Il “nuovo” metodo, che molti colleghi nella realtà già utilizzano, è molto più veloce.

Apri ChatGPT, Gemini, Claude. Scrivi: “Fammi la procedura di come si usa una pressa da 200 tonnellate per stampaggio lamiera”.

E il sistema ti genera una procedura. Ben formattata, con passaggi numerati, con avvertenze di sicurezza standard, con riferimenti normativi generici.

In 30 secondi hai una procedura di 3 pagine.

Il problema? È una procedura generica.

Non tiene conto di:

- Le specificità di quella pressa in particolare (età, modifiche fatte nel tempo, peculiarità)
- Il contesto specifico di quella fabbrica (spazi, altri macchinari)

Apri la flangia: se l'allucinazione dell'IA rischia di uccidere

vicini, layout)

- Le capacità e i limiti degli operatori che la useranno (esperienza, formazione, lingua)
- I rischi specifici di quel luogo di lavoro (presenza di sostanze, microclima, organizzazione)
- Le interferenze con altre lavorazioni

È una procedura che va bene “in generale” per “una pressa generica” usata da “operatori generici”.

Come si controlla una procedura generata dall'IA?

La risposta dovrebbe essere ovvia: devi andare davanti all'operaio con quella procedura e dirgli: “Guarda, l'ho scritta con l'aiuto dell'IA. Leggila. Questa cosa è fattibile qui? È realizzabile con questa macchina? C'è qualcosa che non va? Cosa manca?”

Apri la flangia: se l'allucinazione dell'IA rischia di uccidere

Ma se non lo fai – se ti limiti a generare la procedura, stamparla, metterla nel raccoglitore aziendale e considerarti “a posto” – stai creando un problema enorme.

C'è un principio fondamentale che in questo mestiere non si può dimenticare: la carta deve corrispondere alla realtà.

Se ho una procedura scritta che è troppo diversa da come si lavora veramente, quella procedura è non solo inutile, è pericolosa. Perché:

- Crea l'illusione della conformità (“abbiamo la procedura!”)
- In caso di infortunio, evidenzia la distanza tra prescritto e reale
- Nessuno la segue, quindi diventa carta straccia
- L'operaio perde fiducia anche nelle procedure serie

E l'intelligenza artificiale – proprio perché offre la scorciatoia facile e veloce – sta accelerando la distanza tra la carta e la realtà.

Sempre più aziende hanno procedure bellissime generate dall'IA. Sempre

Apri la flangia: se l'allucinazione dell'IA rischia di uccidere

meno aziende hanno procedure che corrispondono a come si lavora davvero.

Un caso reale

Vi racconto un altro caso reale che dimostra un limite strutturale dell'intelligenza artificiale.

Siamo in un'azienda che produce cosmetici. C'è un IBC – Intermediate Bulk Container, un grande contenitore di plastica – riempito al 30% con alcol.

Un operaio deve fare un'operazione che richiede di accedere all'interno del contenitore. Mette dentro il braccio.

Cosa succede? Il corpo umano genera carica elettrostatica. Mettendo il braccio dentro un'atmosfera con vapori di alcol in concentrazione sufficiente, ha creato le condizioni per un'esplosione.

Risultato: IBC esploso. E il braccio non sta benissimo.

Apri la flangia: se l'allucinazione dell'IA rischia di
uccidere

Dove era la procedura giusta?

Ora, la procedura corretta per fare quell'operazione esisteva. Era stata scritta. Era disponibile in azienda, validata, aggiornata.

Ma non era classificata come "procedura di sicurezza".

Era nelle procedure del GMP - Good Manufacturing Practice. Era scritta per garantire che il prodotto cosmetico non venisse contaminato dai microbi presenti sul corpo umano. Spiegava esattamente come un operatore deve comportarsi quando deve mettere le mani dentro un contenitore con prodotto: guanti specifici, pulizia, movimenti lenti per evitare cariche elettrostatiche.

Il modo giusto di lavorare che avrebbe evitato l'esplosione era scritto in una procedura di qualità del prodotto, non in una procedura di sicurezza delle persone.

Flessibilità umana vs compartimenti stagni

Qui sta il punto: un essere umano competente fa questa connessione mentale.

Apri la flangia: se l'allucinazione dell'IA rischia di uccidere

Uno che conosce bene l'azienda dice: "Guarda, in quella procedura lì sulla qualità c'è scritto come evitare la carica elettrostatica. Non è classificata come sicurezza, ma risolve anche il problema dell'esplosione. Usiamo quella."

Ma un'intelligenza artificiale fa questa connessione trasversale?

Dipende da come è strutturato il sistema. Se l'IA accede solo alle procedure di sicurezza, non vedrà mai quelle di qualità. Passa a compartimenti stagni.

Anche se tecnicamente ha accesso a tutti i documenti, potrebbe non cogliere la rilevanza di una procedura GMP per un problema di sicurezza ATEX (atmosfera esplosive), perché sono categorizzate in modi diversi.

La flessibilità mentale umana – che dice "questa informazione sta nel posto sbagliato, ma è quella che mi serve" – non è scontata in un sistema algoritmico che ragiona per categorie e tag.

Cosa fare concretamente

Dopo averti raccontato tutti questi casi inquietanti, è giusto provare a dare qualche indicazione pratica.

1. Validazione umana obbligatoria

Mai fidarsi ciecamente di una risposta dell'IA, specialmente quando riguarda operazioni con rischi per la sicurezza.

Anche se hai fretta. Anche se l'hai usata 100 volte e ha sempre dato risposte corrette. Anche se sembri paranoico agli occhi dei colleghi.

Quella 101esima volta potrebbe essere quella sbagliata.

2. Contestualizzare sempre

Una procedura generata dall'IA è un punto di partenza, mai un punto di arrivo.

Deve essere contestualizzata: andare sul campo, verificare con chi fa davvero quel lavoro, adattare alle specificità di quella situazione.

Il tempo che "risparmi" usando l'IA per generare la bozza deve essere reinvestito nella validazione accurata, non trattenuto come profitto netto.

Apri la flangia: se l'allucinazione dell'IA rischia di uccidere

3. Documentare i limiti

Nei documenti generati con l'aiuto dell'IA, è onesto scrivere:

- “Questa procedura è stata redatta con il supporto di strumenti di IA e validata in data X con l'operatore Y”
- “Eventuali modifiche operative devono essere segnalate per aggiornamento procedura”

Non (solo) per coprirsi legalmente, ma per creare consapevolezza che quel documento ha dei limiti e va mantenuto aggiornato.

4. Formare alla diffidenza costruttiva

Gli operatori che usano sistemi di IA per assistenza devono essere formati a:

- Verificare sempre i suggerimenti, specialmente su operazioni critiche

Apri la flangia: se l'allucinazione dell'IA rischia di uccidere

- Cliccare sui link alle fonti quando disponibili
- Segnalare risposte che sembrano incomplete o ambigue
- Non sentirsi stupidi a chiedere conferma a un collega più esperto

Conclusione: il verosimile può uccidere

L'intelligenza artificiale non dà risposte sbagliate nel senso di "completamente inventate" (anche se le allucinazioni esistono). Più spesso dà risposte verosimili ma incomplete.

"Apri la flangia" è tecnicamente corretto. Ma manca il contesto di sicurezza.

Una procedura generica per l'uso di una pressa è verosimile. Ma manca la specificità di quella pressa in quella fabbrica.

Una connessione tra procedura GMP e sicurezza ATEX è logica per un

Apri la flangia: se l'allucinazione dell'IA rischia di
uccidere

umano. Ma potrebbe non essere ovvia per un algoritmo.

E il verosimile può generare un infortunio grave.

Perché quando qualcosa sembra sensato, plausibile, ragionevole,
abbassiamo la guardia. Non verifichiamo con la stessa attenzione. Ci
fidiamo.

Ed è proprio in quei momenti che succede il disastro.

Leggi il libro "Intelligenza artificiale: i rischi nelle
fabbriche. Guida per chi vuole implementare l'IA negli
stabilimenti industriali"

Lo trovi in vendita su Amazon

[Vai](#)