

Intelligenza artificiale neuro-simbolica per l'ingegneria strutturale

Adriano Castagnone ¹, Giuseppe Nitti ^{1,2,*}

¹ S.T.A. DATA Srl, Via Saluzzo, 71 – 10126, Torino, Italia

² Politecnico di Torino, Corso Duca degli Abruzzi, 24 - 10129 Torino, Italia

* Recapiti: castagnone@stadata.com (A. Castagnone), giuseppe.nitti@polito.it (G. Nitti)

Abstract

Questo articolo affronta le opportunità e i rischi derivanti dall'integrazione dei large language models (LLM) nel campo dell'ingegneria strutturale. L'uso esclusivo degli LLM risulta inadeguato in questo ambito, poiché la loro natura probabilistica può generare allucinazioni e imprecisioni, inaccettabili in un settore a sicurezza critica che richiede calcoli rigorosi. Per risolvere questo dilemma, proponiamo l'adozione dell'intelligenza artificiale neuro-simbolica (NSAI), un approccio ibrido che bilancia l'intuizione neurale con il rigore simbolico. L'architettura NSAI impiega un sistema di interrogazione intelligente per arricchire le richieste dell'utente e delegare le operazioni critiche ad algoritmi deterministici esterni. Questo sistema garantisce affidabilità e conformità normativa, come esemplificato dal caso studio del chatbot 3Muri, un assistente intelligente basato su NSAI per software di analisi strutturale.

Parole chiave: *Intelligenza artificiale Neuro-simbolica, Large language models, Ingegneria strutturale, Architettura ibrida, AI spiegabile*

1. Introduzione

1.1 Contesto e motivazione

Il rapido avanzamento delle tecnologie di intelligenza artificiale ha trasformato numerosi settori industriali. Tra queste innovazioni, i large language models (LLM) sono emersi come strumenti particolarmente trasformativi, dimostrando capacità senza precedenti nella comprensione, generazione e ragionamento sul linguaggio naturale [1, 2]. Questi modelli, addestrati su vasti corpora testuali, hanno mostrato notevoli abilità nel condurre conversazioni simili a quelle umane, rispondere a domande complesse e assistere in compiti creativi e analitici.

Il settore dell'ingegneria strutturale, che per sua natura richiede precisione e risultati inequivocabili, ha naturalmente riconosciuto il potenziale di questi potenti sistemi di IA. Ingegneri e ricercatori hanno immaginato applicazioni che spaziano dalla verifica automatica della conformità normativa agli assistenti di progettazione intelligenti. Tuttavia, l'adozione degli LLM nell'ingegneria strutturale si scontra con un ostacolo insuperabile: il problema delle allucinazioni, ovvero la tendenza a generare output che appaiono plausibili ma sono di fatto errati o completamente inventati [3, 4].

Mentre tali errori potrebbero essere tollerabili in conversazioni casuali o scrittura creativa, rappresentano una minaccia esistenziale quando applicati a calcoli strutturali da cui dipende la vita delle persone. In un settore definito da parametri quali la responsabilità legale, le

regolamentazioni severe (Eurocodici [5, 6], NTC [7]) e la sicurezza critica [8], la tolleranza per gli errori computazionali è zero [9, 10] (figura 1).

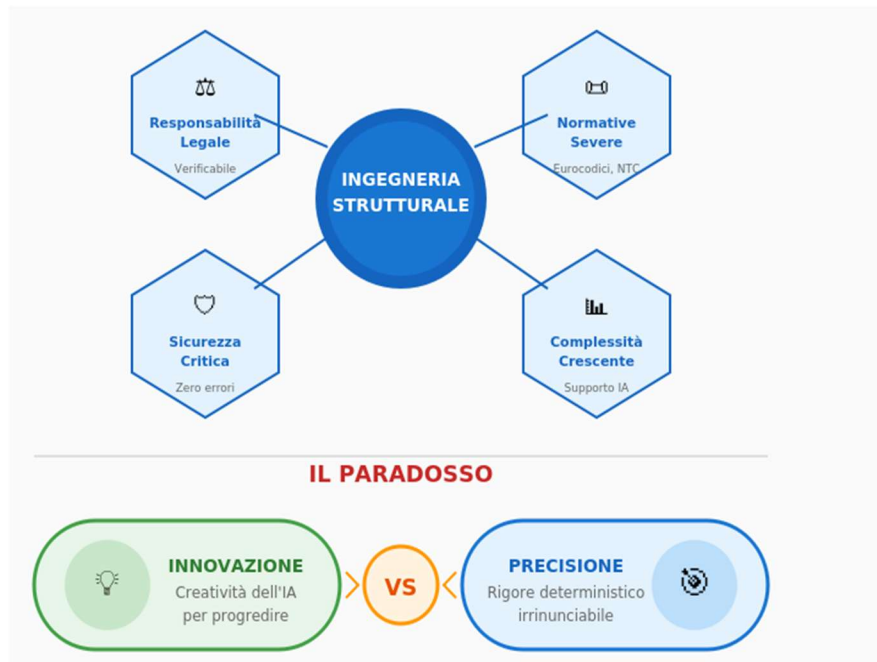


Figura 1: I vincoli dell'ingegneria strutturale e il paradosso innovazione/rigore

1.2 Il Paradosso fondamentale

Questi vincoli generano un paradosso fondamentale al cuore dell'adozione dell'IA nell'ingegneria strutturale. Da un lato, esiste una necessità genuina di innovazione alimentata dall'IA per aiutare gli ingegneri a gestire progetti sempre più complessi e navigare requisiti normativi estesi. Dall'altro, la natura probabilistica degli attuali sistemi di IA li rende inadatti all'applicazione diretta a calcoli dove è richiesta certezza [11, 12].

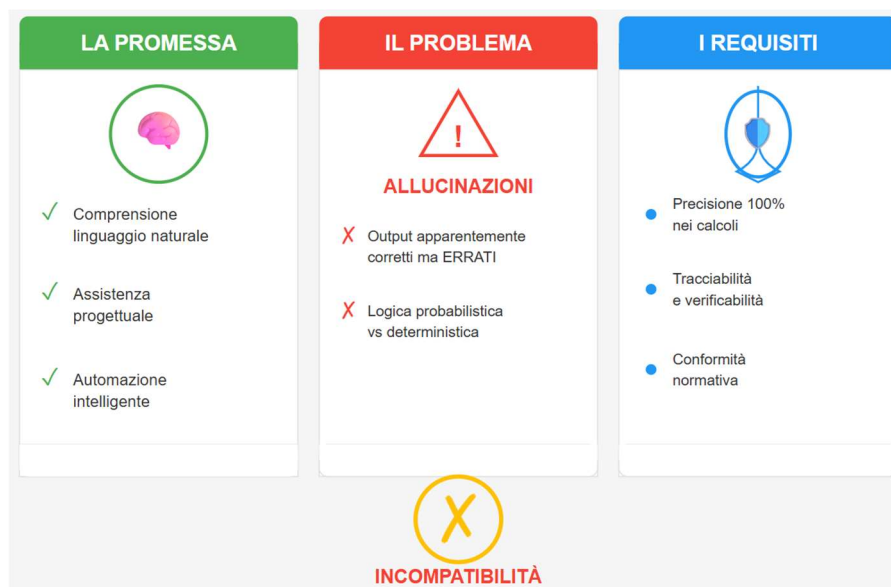


Figura 2: Il dilemma nell'uso degli LLM nell'ingegneria

L'equazione "LLM + Ingegneria = Rischio" (figura 2) sintetizza questo dilemma. La logica intrinseca degli LLM è probabilistica: generano output basati su pattern statistici appresi dai dati di addestramento, non attraverso ragionamento logico deterministico. Ciò contrasta nettamente con la natura deterministica richiesta dall'ingegneria strutturale, dove i calcoli devono produrre risultati identici indipendentemente da quante volte vengano eseguiti [1, 10].

Per risolvere questo dilemma, è necessario un nuovo paradigma, capace di sfruttare le capacità creative e interpretative dell'IA neurale preservando il rigore deterministico essenziale per l'ingegneria. Questo paradigma è l'intelligenza artificiale neuro-simbolica (NSAI) [13, 14].

2. Il ponte neuro-simbolico: fondamenti teorici

2.1 Panoramica del Paradigma

l'intelligenza artificiale neuro-simbolica rappresenta un paradigma trasformativo che mira a fungere da "ponte intelligente tra l'intuizione neurale e il rigore simbolico" (figura 3). Piuttosto che considerare gli approcci neurali e simbolici come alternative in competizione, l'NSAI li riconosce come capacità complementari che, opportunamente integrate, possono raggiungere risultati impossibili per ciascun approccio isolato [13, 14].

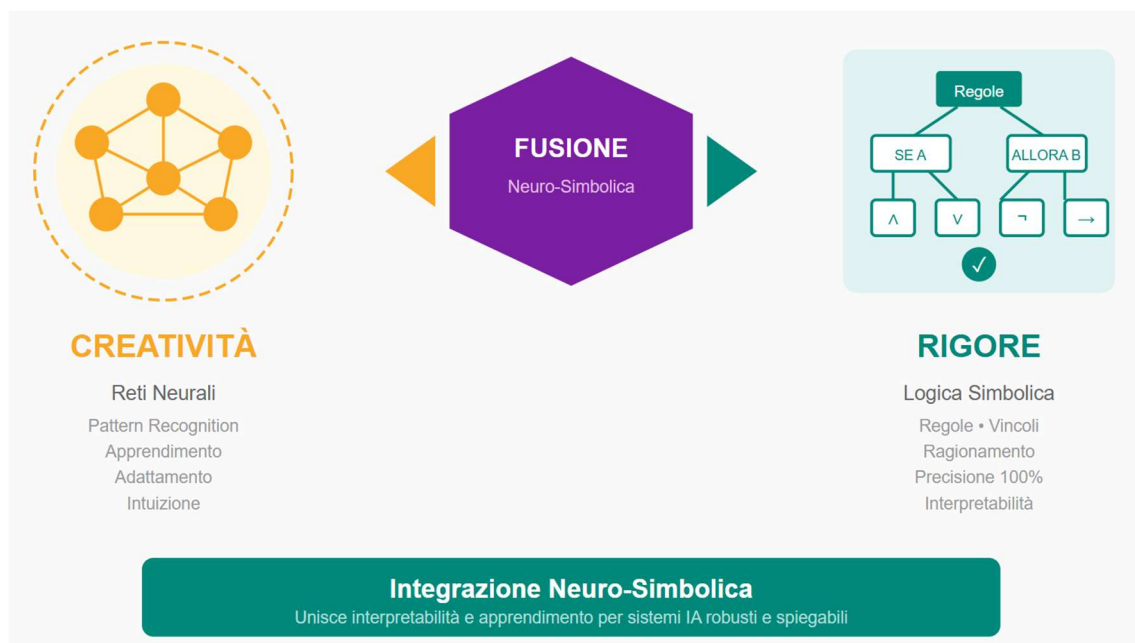


Figura 3: L'IA neuro-simbolica fonde creatività neurale e rigore logico

L'intuizione fondamentale alla base dell'NSAI è che l'intelligenza umana stessa opera attraverso l'integrazione di ragionamento intuitivo e analitico. Gli scienziati cognitivi hanno riconosciuto che il problem-solving umano coinvolge sia un rapido riconoscimento intuitivo di pattern (pensiero Sistema 1) sia un'analisi logica deliberata e lenta (pensiero Sistema 2). Le architetture NSAI cercano di replicare per via computazionale questa cognizione a doppio processo.

2.2 La componente neurale: creatività

La componente neurale di un sistema NSAI, tipicamente implementata mediante large language models o altre architetture di deep learning, fornisce capacità essenziali [15, 16]:

- **Comprensione del linguaggio naturale:** i modelli neurali eccellono nell'interpretare il linguaggio umano in tutta la sua variabilità e ambiguità. Gli ingegneri possono descrivere problemi con le proprie parole, usando terminologia di dominio o descrizioni informali.
- **Riconoscimento di pattern:** i modelli di deep learning possono identificare pattern complessi in dati ad alta dimensionalità, abilitando capacità come il riconoscimento di elementi strutturali in schizzi disegnati a mano.
- **Ragionamento flessibile:** i modelli neurali possono impegnarsi in ragionamento approssimativo e analogico che imita l'intuizione umana, suggerendo approcci progettuali basati sulla similarità con progetti passati.

Tuttavia, queste capacità comportano limitazioni fondamentali: output probabilistici (non deterministici), rischio di allucinazioni, opacità (modelli "black box") e precisione numerica limitata [17]. Queste limitazioni rendono le componenti neurali inadatte all'applicazione diretta a calcoli a sicurezza critica.

2.3 La componente simbolica: rigore

La componente simbolica di un sistema NSAI fornisce il rigore e l'interpretabilità che mancano alle componenti neurali [18, 19]. L'IA Simbolica opera su rappresentazioni esplicite della conoscenza usando logica formale, regole e relazioni strutturate:

- **Ontologia di dominio:** una rappresentazione formale dei concetti dell'ingegneria strutturale (travi, colonne, carichi, materiali) e delle loro relazioni.
- **Protocolli di verifica:** procedure formali per verificare che i progetti soddisfino i requisiti specificati, codificando i metodi prescritti dai codici normativi.
- **Regole di conformità:** codifiche esplicite dei requisiti normativi da fonti come Eurocodici e NTC, specificando calcoli, fattori di sicurezza e limiti.
- **Algoritmi deterministici:** algoritmi matematici per l'analisi strutturale che producono risultati identici per input identici.

La componente simbolica fornisce precisione deterministica (100%), trasparenza completa e conformità garantita [20]. Tuttavia, i sistemi simbolici sono rigidi, richiedono input in formati precisamente specificati e non possono gestire l'ambiguità del linguaggio naturale.

2.4 L'Integrazione: raggiungere la sinergia

L'intuizione chiave dell'integrazione neuro-simbolica è che i punti di forza di ciascuna componente affrontano precisamente le debolezze dell'altra [17, 21]. Le componenti neurali

gestiscono il front-end disordinato e ambiguo dell'interazione ingegneristica, mentre le componenti simboliche gestiscono il back-end preciso e verificabile.

Questa integrazione crea sistemi che ottengono risultati cruciali per l'ingegneria strutturale:

- **Affidabilità garantita:** ogni output subisce verifica formale dalla componente simbolica [22].
- **Spiegabilità completa:** ogni decisione può essere tracciata attraverso la catena di ragionamento simbolico [21].
- **Conformità normativa:** integrazione diretta con i requisiti normativi codificati [5, 6, 7].
- **Riduzione degli errori:** la doppia verifica neuro-simbolica intercetta errori che potrebbero sfuggire a ciascuna componente singolarmente.
- **Interazione naturale:** gli utenti possono interagire in linguaggio naturale ricevendo risultati rigorosamente verificati.