



La teoria della pertinenza e gli LLM La metafora tra comunicazione naturale e artificiale: questioni aperte e premesse metodologiche

Relevance Theory and LLMs The Metaphor Between Natural and Artificial Communication: Analogies and Open Issues

Maurizio Maione

 0009-0000-3337-7881

Università Guglielmo Marconi

 00j0rk173

maurizio.maione63@gmail.com

 10.54103/2531-5994/31700

Abstract

[It.] Secondo la *Teoria della Pertinenza*, la comunicazione si giustifica in virtù della capacità dei parlanti di attivare e riconoscere reciprocamente le intenzioni e di legarle a determinati indizi o contesti. Si tratta di un modello teorico risolutivo rispetto a quei fenomeni che possano eventualmente mettere in crisi i processi di comunicazione/comprendimento: l'uso dell'implicito, il cosiddetto linguaggio figurato - nella fattispecie, il sistema delle metafore - e, infine, gli *atti linguistici indiretti*. In questa prospettiva, la *Teoria della Pertinenza* si configura come un modello teorico efficace nella misura in cui riconduce l'attività linguistico-comunicativa a processi cognitivi come le inferenze e il *mindreading*.

I *Large Language Model* sono i modelli linguistici dell'IA generativa che comprendono e generano testi nel linguaggio umano. L'attivazione degli LLM incoraggia il ricorso ad un modello teorico come quello della pertinenza come si desume dal dispositivo del *prompting*. In primis, l'attivazione degli LLM via *prompting* sembra offrire tutti gli elementi per applicare la teoria della pertinenza agli LLM e, quindi, per stabilire analogie sempre più forti tra la comunicazione umana (naturale) e quella tra l'IA (non strutturata cognitivamente!) e gli utenti umani.

L'obiettivo di questo contributo è quello di verificare se sia legittimo parlare di *mindreading* artificiale e se sia opportuno ripristinare il carattere naturale associandolo al solo utente umano. Dal punto di vista metodologico, si tratta di individuare un nucleo teorico in grado non solo di stabilire il confronto tra la comunicazione artificiale via LLMs e la comunicazione naturale, ma anche di predisporre e programmare uno studio di natura anche sperimentale. Il nucleo teorico è quello del *continuum* tra le espressioni metaforiche e quelle letterali, da inscrivere all'interno di quello più generale tra *implicito* ed *esplicito*. In questa sede, si stabilisce il peso e la forza metodologica di questo nucleo teorico, mentre si rinvia ad uno studio successivo la possibilità di configurarlo come un *caso di studio* da verificare all'interno di una ricerca incentrata soprattutto sull'analisi dei *corpora*. In tal senso, sarà proficuo esaminare una determinata situazione limite: l'azione degli LLM rispetto sia alla traduzione delle metafore morte da una lingua all'altra sia alla gestione parziale delle metafore nuove. Il minor peso esercitato dagli elementi lessicalizzati e la difficoltà a trovare una soluzione alternativa ad ampio spettro rendono

palesi alcune asimmetrie tra le due forme di comunicazione senza escludere affatto la possibilità di nuove indagini.

Parole chiave: Rilevanza, suggerimento, lettura del pensiero, metafora, implicatura.

[En.] According to Relevance Theory, communication is justified by the ability of speakers to mutually activate and recognize intentions and to link them to specific cues or contexts. This is a theoretical model that provides a solution to those phenomena that might potentially undermine communication and comprehension processes: the use of the implicit, so-called figurative language – in this case, the system of metaphors – and, finally, indirect speech acts. From this perspective, Relevance Theory emerges as an effective theoretical model insofar as it traces linguistic-communicative activity back to cognitive processes such as *inference* and *mindreading*.

Large Language Models (LLMs) are generative AI linguistic models that understand and generate texts in human language. The activation of LLMs encourages the use of a theoretical model such as Relevance Theory, as can be inferred from the device of prompting. First and foremost, the activation of LLMs via prompting appears to provide all the elements needed to apply Relevance Theory to LLMs and, consequently, to establish increasingly strong analogies between (natural) human communication and that between AI (which is not cognitively structured!) and human users.

The aim of this paper is to examine whether it is legitimate to refer to ‘artificial mindreading’ and whether it is appropriate to establish its natural character by associating it strictly with human users. From a methodological perspective, the aim is to identify a theoretical framework enabling not only a comparison between artificial communication via LLMs and natural communication, but also the design and planning of a study that may also be experimental in nature. This theoretical framework is based on the continuum between metaphorical and literal expressions, which forms part of the broader continuum between the implicit and the explicit. Here, we shall establish the significance and methodological strength of this theoretical framework, whilst deferring to a subsequent study the task of applying it to a case study to be tested within research focused primarily on corpus analysis. In this regard, it will be useful to examine a specific borderline situation: the role of LLMs in relation to both the translation of dead metaphors from one language to another and the partial handling of new metaphors. The lesser influence exerted by lexicalised elements and the difficulty in finding a broad-spectrum alternative solution highlight certain asymmetries between the two forms of communication, without in any way ruling out the possibility of further investigation.

Keywords: Relevance, prompting, mindreading, metaphor, implicature.

1. Introduzione

L'intelligenza artificiale generativa segna ormai molteplici settori della nostra vita; demonizzarne la diffusione è inutile e, oltretutto, genera una forma di oblio di quella che rimane la caratteristica principale della storia umana: la capacità dell'uomo di ricorrere a risorse o strumenti esterni come *estensioni* di varia natura della propria mente e del proprio corpo. La scrittura, la stampa, il computer sono esempi eloquenti di siffatte

estensioni. L'estensione non va intesa come simulazione bensì come un supporto/risorsa di cui l'uomo si avvale nel tentativo di ridurre la complessità della realtà. L'IA rientra tra le estensioni e l'espressione "intelligenza artificiale" è una metafora efficacissima che accentua, da un lato, il rapporto tra le due forme di intelligenza, artificiale e naturale, dall'altro, la natura protesica della prima che si riassume perfettamente nell'interazione tra l'utente umano e quello che è, al momento, l'esito più interessante delle ricerche ingegneristiche intorno all'intelligenza artificiale: la costruzione dei *Large Language Model*, dei modelli linguistici che configurano questa interazione come una comunicazione linguisticamente strutturata e ben articolata dal punto di vista lessicale, morfologico e sintattico.

Il presente saggio si propone di valorizzare l'interazione tra l'utente umano e gli LLM sussumendola alla Teoria della Pertinenza che Sperber e Wilson formulano in vista di una maggiore comprensione della comunicazione naturale orale, superando il modello teorico di Grice, incentrato esclusivamente sul rapporto tra l'intenzionalità dei parlanti e il codice/lingua di riferimento. Sperber e Wilson individuano come prioritari i nessi che intercorrono tra l'intenzionalità stessa e i molteplici indizi di natura contestuale che stabiliscono sia il livello ottimale della comprensione da parte dell'interlocutore sia quello della comunicazione naturale intesa nella sua interezza. L'*experimentum crucis* può essere individuato nella gestione delle *implicature* e nelle modalità di ricognizione delle stesse in vista di eventuali *esplicature*. Come si vedrà, la ricognizione delle implicature è *in primis* presente in entrambe le comunicazioni ma con modalità di esecuzione del tutto diversificate. Il terreno di verifica della diversa ricognizione delle implicature è offerto dalla metafora; situazioni simili sono però individuabili anche nell'attività di traduzione e nella ricognizione degli atti linguistici indiretti. Nel presente lavoro, si affronterà metodologicamente la questione della metafora, rinviando quindi ad uno studio successivo la disamina delle metafore anche in relazione ai diversi *corpora* di riferimento.

In conclusione, si metterà in evidenza la presenza di un'asimmetria cognitiva che risolve sostanzialmente il confronto tra le due forme di comunicazione pur all'interno di una disamina che si snoda a partire da alcune simmetrie suggerite dalla *teoria della pertinenza*, come quella che intercorre tra l'attività dell'utente degli LLM, segnata dal *prompting* – che è di fatto un dispositivo della pertinenza – e l'attività del parlante all'interno della comunicazione naturale.

2. La teoria della pertinenza e la comunicazione naturale

Secondo la *Teoria della Pertinenza (Relevance Theory)* (Sperber & Wilson, 1995, 2012), la comunicazione si giustifica in virtù della capacità dei parlanti di attivare e riconoscere reciprocamente le *intenzioni* e, soprattutto, di legarle a determinati indizi o contesti. Il riferimento alle intenzioni (intenzionalità) mostra l'attivazione dei processi cognitivi in vista di un equilibrio tra le istanze degli stessi e quelle della comunicazione in senso stretto. Non a caso, Sperber e Wilson individuano due principi contestuali

l'uno all'altro: il *principio cognitivo della pertinenza* e il *principio comunicativo della pertinenza*. In base al primo, la cognizione tende a essere orientata all'ottimizzazione; in base al secondo, invece, ogni atto di comunicazione esplicita comporta una presupposizione della propria pertinenza ottimale (Sperber & Wilson, 1995, pp. 260-272). Sono due principi congiunti: il parlante ricorre a comportamenti che rendano manifesta al suo interlocutore una determinata intenzione; tali comportamenti sono pertanto definiti ostensivi e la comunicazione che li attiva è chiamata *ostensive-inferential communication* (*comunicazione ostensivo-inferenziale*) (Sperber & Wilson, 1995, pp. 49-54). In *Meaning and Relevance*, gli autori integrano la precedente posizione dilatando tanto il comportamento ostensivo del parlante quanto piuttosto il processo cognitivo dell'interlocutore. La comunicazione inferenziale prevede infatti che il parlante adotti in modo ostensivo un determinato comportamento (ad esempio un gesto mimico o l'emissione di un segnale codificato) tale da attivare nel destinatario (attraverso il riconoscimento o la decodifica) una specifica struttura concettuale o idea (Sperber & Wilson, 2012, pp. 36-37). Il comportamento ostensivo implica in effetti diversi *stimoli ostensivi* come azioni gestuali, discorsi, suoni, enunciati, pensieri, ricordi che attirano l'attenzione dell'interlocutore e trasmettono un determinato messaggio secondo una procedura "anomala" in quanto gli stimoli ostensivi menzionati non sono affatto riconducibili ad un codice (Sperber & Wilson, 2012, pp. 98-101); per questa ragione, l'interlocutore non può far leva sulla sola possibilità di decodifica e deve quindi ricorrere a risorse di diversa natura secondo una "ricognizione" attivata e scandita da processi inferenziali non circoscritti alle istanze convenzionali, che regolano i significati degli enunciati, ma aperti ai multiformi significati ed intenzioni che i parlanti riferiscono in termini non convenzionali ai multiformi stimoli ostensivi (Sperber & Wilson, 2012, pp. 4-7). Perché questi significati e intenzioni vengano efficientemente elaborati, soprattutto in termini di comprensione, non è sufficiente che i parlanti si attengano ai soli "verbal inputs"; l'interlocutore va oltre la decodifica linguistica dell'enunciato del parlante-emittente: seguendo il principio del minimo sforzo, egli procede all'integrazione di quanto è esplicito con elementi impliciti di varia natura, pur di pervenire ad un'interpretazione dell'enunciato tale da realizzare le istanze e le aspettative della pertinenza (Sperber & Wilson, 2012, pp. 39-43). Si tratta di un processo di comprensione, in cui l'integrazione del livello esplicito dell'enunciato passa attraverso un vero e proprio arricchimento dello stesso mediante un parallelo livello implicito di inferenze ed interpretazioni (Sperber & Wilson, 2012, p. 40). La comprensione di un enunciato non dipende quindi né dalle sole componenti esplicite dello stesso né da una necessaria azione di decodifica. Efficace è la seguente sintesi: "(...) given the inferential nature of comprehension, the words in a language can be used to convey not only the concepts they encode, but also indefinitely many other related concepts to which they might point in a given context" (Sperber & Wilson, 2012, p. 43). La forza del livello implicito delle inferenze è strettamente connessa alla possibilità di *indicare* e di far leva su un determinato *contesto* che è il cardine della pertinenza. A questo punto, è effettivamente chiara la natura pragmatica della teoria della pertinenza: essendoci "more concepts in our minds than words in our language" (Sperber

& Wilson, 2012, p. 43), il ricorso al contesto stabilisce, da un lato, la pari dignità di tutti i concetti o pensieri presenti nelle menti dei parlanti in interazione, dall'altro, la necessità di diverse forme di coordinamento che la comunicazione attiva in vista dell'ottimizzazione semantica e della pertinenza in termini di processi intersoggettivi.

Questo non è il luogo per una disamina della Teoria della pertinenza completa, articolata ed esaustiva ma i nuclei concettuali presi in considerazione – funzionali agli obiettivi sopra menzionati – sono sufficienti a confermare la forza e le potenzialità delle ricerche di Sperber e Wilson. Si tratta di un modello teorico risolutivo – ai fini del presente contributo – rispetto a quei fenomeni che possano eventualmente mettere in crisi i processi di comunicazione/comprendimento: l'uso dell'implicito, il cosiddetto linguaggio figurato – nella fattispecie, il sistema delle metafore – e, infine, gli *atti linguistici indiretti*. In questa prospettiva, la *Teoria della Pertinenza* si configura come un modello teorico efficace nella misura in cui riconduce l'attività linguistico-comunicativa a processi cognitivi come il *mindreading* o *teoria della mente*: la capacità di riconoscere gli stati mentali altrui, a partire proprio dalla possibile ricognizione *situazionale* delle intenzioni (Kravchenko et al., 2024; Chandaria, 2024; Adornetti & Ferretti, 2025; Forceville, 2025).

3. I *Large Language Model* nella prospettiva della pertinenza

La teoria della pertinenza è in grado di giustificare, descrivere e spiegare la comunicazione naturale la cui complessità rende insufficienti modelli teorici come quello di Paul Grice incentrati sulla funzione prioritaria del codice e, quindi, sulla possibilità di risolvere l'implicito riconducendolo a determinate azioni che i parlanti performano a partire dal codice e in vista di una ricognizione dell'implicito che ha quindi la sua chiave di accesso esclusivamente nel codice stesso. Secondo Sperber e Wilson, la comunicazione non si risolve nel codice ma in processi cognitivi orientati verso l'elaborazione di diversi fattori cognitivamente rilevanti come la teoria della mente altrui, le intenzioni, il contesto e, infine, quei concetti che sono sempre più intesi come l'esito del coordinamento tra i concetti dei singoli parlanti e che, quindi, devono essere restituiti ad una dimensione collettivo-intersoggettiva. *Mutatis mutandis* – sorprendentemente – il modello della pertinenza può essere applicato anche all'interazione tra gli utenti naturali e quell'*intelligenza artificiale generativa* che si configura ormai come LLM, *Large Language Model* (modelli linguistici di grandi dimensioni).

I *Large Language Model* sono i modelli linguistici dell'IA generativa che, in base a determinate reti neurali artificiali, sono addestrati su enormi (*large*) set di dati testuali di diversa provenienza per apprendere continuamente dagli stessi e, soprattutto, per comprendere e creare testi nel linguaggio naturale umano. Ovviamente, l'addestramento in questione implica una complessa procedura interna che si risolve poi nella produzione linguistica: in base a parametri probabilistici, gli LLM stabiliscono relazioni tra parole per poi generare sequenze nuove come se fossero opera diretta dell'uomo (De Caro &

Giovanola, 2025, pp. 26-46). I testi generati sono anche funzionali a inferenze logiche e a vere e proprie catene di ragionamenti. L'addestramento è continuo in quanto si adatta a quello che è, in sostanza, un altrettanto continuo – se non proprio inesauribile – processo di integrazione e rinnovo dei dati testuali. In tal senso, la produzione testuale è variegatissima ma anche esposta ad alcuni rischi di errore – peraltro interessanti in vista del confronto con la comunicazione naturale – dovuti proprio alle *relazioni probabilistiche tra parole* sopra menzionate. I rischi si condensano nel fenomeno delle cosiddette *allucinazioni* che non sono altro che testi che, sebbene configurati come risposte apparentemente corrette a determinate richieste dell'utente umano, sono tuttavia connotati da informazioni false, imprecise e inaffidabili. Ovviamente, non si tratta di un dato allarmante visto che, nella stessa interazione umana, i parlanti sono costantemente esposti a informazioni imprecise e, talvolta, inaffidabili.

L'attivazione degli LLM incoraggia il ricorso ad un modello teorico come quello della pertinenza: i testi generati dagli LLM sono "incanalati" dal cosiddetto *prompting*, l'insieme di domande, istruzioni o esempi, che l'utente digita ed invia all'IA in vista di testi o informazioni con un livello di efficacia il più elevato possibile. Il *prompting* non è affatto un'azione marginale; anzi, costituisce un insieme di strategie che valorizza sia le "risposte" degli LLM sia l'identità cognitivo-culturale dell'utente (Roncaglia, 2023). In questa prospettiva, il *prompting* è strategico e non assimilabile a sporadiche richieste: il *prompting strategico* implica più richieste da ricondurre ad un unico tessuto connettivo ed è funzionale alla produzione di testi connotati dalla coerenza tra il contesto e le informazioni sollecitate.

In *primis*, l'attivazione degli LLM via *prompting* sembra offrire tutti gli elementi per applicare la teoria della pertinenza agli LLM e, quindi, per stabilire analogie sempre più forti tra la comunicazione umana e quella tra l'IA e gli utenti umani (Branda, 2026). Gli LLM rispondono bene ai *prompting* degli utenti, vale a dire, *intercettano* le intenzioni degli stessi mostrando così di "gestire" – anche se indirettamente – il *mindreading*, la lettura delle menti altrui, da ascrivere propriamente agli utenti che se ne avvalgono esplicitamente nell'interazione con gli LLM. L'efficacia cognitiva delle diverse "azioni" degli LLM si definisce soltanto in relazione alle richieste e/o aspettative dell'utente e non come un dato oggettivo oppure come una funzione strutturata che "emerge" dalle reti neurali artificiali. Gli LLM non sono di fatto strutturati cognitivamente. Le reti neurali artificiali non sono assimilabili alle reti neurali cerebrali che sono definite biologicamente. In un saggio incentrato sulla questione delle metafore, attribuire agli LLM la capacità di *cogliere i prompting* degli utenti e di *gestire il mindreading* significa ricorrere al linguaggio metaforico inteso come quell'unica risorsa cognitiva in grado di comprendere una circostanza nuova o inusuale; un primo segnale del ruolo pervasivo delle metafore non solo nel linguaggio ordinario ma anche in una disamina scientifica.

I testi generati non solo tengono conto della prospettiva degli utenti ma sono anche incentrati sul rapporto costante tra intenzioni e contesti; nel caso in cui qualcosa non funzioni o non sia coerente con quanto già prodotto, in base ad un nuovo *prompting*, gli LLM si autocorreggono dando forma a relazioni tra parole più corrette ed efficaci in quanto più esposte al controllo del contesto; controllo esercitato dall'IA ma verifica-

to e confermato dallo stesso utente umano. Il *prompting* esercitato e reiterato su una determinata informazione o su un determinato significato linguistico è un *prompting* orientato verso l'ottimizzazione della pertinenza e, nella misura in cui questo obiettivo avanza nella sua realizzazione e gli LLM rispondono coerentemente, l'interazione tra utente umano e IA si trasforma *tout court* in un oggetto di analisi pragmatica. In ogni caso, la pertinenza rimane appannaggio esclusivo dell'utente e il *prompting* ne diventa un dispositivo di controllo.

Il *prompting* che l'utente umano performa rispetto all'IA riproduce in parte la situazione del parlante nella comunicazione naturale. Il *prompting* assume anche la funzione che, nella comunicazione naturale, viene normalmente attribuita agli "stimoli ostensivi"; ma, mentre questi ultimi possono presentare strutture e modalità non linguistiche, il *prompting* può avvalersi di perifrasi che facciano riferimento indiretto a quegli stimoli ma sempre e necessariamente all'interno di un atto linguistico come un enunciato, una domanda oppure una richiesta (ordine). La reazione positiva agli stimoli ostensivi da parte dell'interlocutore della comunicazione naturale conferma il livello di pertinenza attivato dal parlante/emittente secondo modalità ugualmente riscontrabili nella risposta degli LLM all'utente umano. Nell'interazione tra l'utente umano e gli LLM, l'*incipit* è dunque fornito da un *prompting* che prevede uno scenario implicito che gli LLM definiscono sempre più esplicitamente in termini di maggiore pertinenza, anche in relazione agli eventuali errori *in itinere* e alle eventuali forme di *allucinazione*. In questa prospettiva, emerge un'istanza comune alle due forme di comunicazione: l'esistenza di un *continuum* tra implicito ed esplicito che, come si vedrà più avanti, Sperber e Wilson identificano come un elemento strutturale del linguaggio e, segnatamente, della comunicazione naturale orale.

Il *continuum* tra esplicito ed implicito concorre di certo a quell'indeterminatezza che, per diverse ragioni, si configura come una componente strutturale dell'intera attività linguistico-comunicativa ma che una più o meno radicata tradizione di studi e ricerche ascrive soprattutto a determinate attività come l'uso delle metafore e degli atti linguistici indiretti (Ervás & Gola, 2016; Littlemore, 2019; Kövecses, 2020; Ruytenbeek, 2021; Garello, 2024; Jucker, 2024).

Data la natura metodologica del presente saggio, si tralascia per il momento quella che potrebbe essere una *disamina corpus-based*, fondamentale sia per perfezionare e meglio esplorare la tassonomia delle metafore e degli atti linguistici indiretti sia per conferire maggiore precisazione scientifica al confronto tra la comunicazione naturale e l'interazione tra gli utenti umani e gli LLM.

4. L'indeterminatezza e gli LLM: la pertinenza e il prompting

Come si è già visto, la nozione di pertinenza è connessa a quella di *mindreading*: soltanto la possibilità di coordinare tra diversi concetti o significati elaborati individualmente ed associati a determinati contesti garantisce il raggiungimento della pertinenza.

za che è attivata dal parlante ma confermata dall'interlocutore. In vista del confronto tra comunicazione naturale e interazione tra utenti umani e LLM, un primo obiettivo è senz'altro quello di verificare fino a che punto sia efficace il ricorso al *mindreading artificiale*, soprattutto in virtù del fatto che si tratta di un'espressione metaforica (come lo è la stessa espressione "intelligenza artificiale"!) funzionale alla sola valorizzazione di eventuali analogie con la comunicazione naturale e non all'individuazione del presunto processo cognitivo artificiale corrispondente al *mindreading* naturale. In effetti, quello che si tende impropriamente a chiamare *mindreading artificiale* non è altro che un comportamento degli LLM da ricondurre ai molteplici repertori testuali su cui essi si addestrano. Il comportamento non implica un reale processo cognitivo e non può nemmeno esemplificarlo; il "come se" rientra soltanto nell'atteggiamento euristico dell'eventuale osservatore e studioso. D'altra parte, è necessario osservare che si tratta di comportamenti talvolta del tutto inaspettati per tipologia, velocità ed efficacia; il che va imputato ed ascritto non alla funzione reale di un *mindreading artificiale*, bensì alla capacità degli LLM di gestire, associare e organizzare un numero elevatissimo di informazioni distribuito tra molteplici repertori. Sebbene molti siano gli studi che hanno tentato o tentano di verificare l'attivazione di comportamenti che corrispondano al *mindreading*, non emergono tuttavia risultati in grado di confermare – seppur minimamente – la presenza di processi cognitivi simili a quelli del parlante naturale (Strachan et al., 2024; Chandaria, 2024). Alla luce di siffatte osservazioni, la pertinenza si configura come quel dispositivo cognitivo che determina esclusivamente la comunicazione naturale e che motiva l'attività cognitivo-comunicativa del parlante attivandone realmente il *mindreading* senza di cui sarebbe del tutto impraticabile la ricognizione delle intenzioni e del contesto. Rispetto alle azioni degli LLM, è possibile invece ascrivere il *mindreading* ai soli utenti umani che lo definiscono e modulano in base ai risultati di quelle azioni sollecitate via *prompting*. A questo punto, emerge palesemente una prima asimmetria: mentre nella comunicazione naturale la pertinenza è parte integrante di un processo cognitivo che si configura propriamente come *mindreading*, nell'interazione tra l'utente e gli LLM, la pertinenza è soltanto rilevata dall'utente in relazione a determinati *output* degli LLM. Ora è utile ritornare alla questione del *continuum* e della metafora che ne è l'istanza maggiormente presa in considerazione. La disamina del controllo del *continuum* tra implicito ed esplicito e del *continuum* tra espressioni metaforiche e quelle letterali potrà fornire qualche elemento più dirimente rispetto alle analogie tra la comunicazione utente-LLM e la comunicazione naturale (Eragamreddy, 2025; Allott & Textor, 2026).

Una prima questione riguarda l'origine testuale degli LLM: sebbene costante e diversamente articolato, l'addestramento degli LLM rimane ancorato ad informazioni testuali che fanno capo alla scrittura. La straordinaria capacità di interagire con l'utente umano, di *coglierne le intenzioni*, di rispondere ai suoi quesiti e di realizzare compiti di diversa natura richiesti dal medesimo si spiega in virtù di fonti testuali e, quindi, anche se algoritmicamente, della scrittura. La comunicazione artificiale, attivata dagli LLM, ha la sua matrice nella scrittura. La pertinenza che i *prompting* degli utenti umani attivano è, da un lato, interessante perché gli LLM elaborano gli indizi e li contestualizza-

no, dall'altro, l'attivazione/elaborazione degli indizi avviene in relazione ad un sistema neurale artificiale il cui perno non è affatto la comunicazione orale bensì la scrittura. Sebbene il richiamo alla teoria della pertinenza sia efficace nella descrizione e spiegazione dell'interazione tra utente e intelligenza artificiale generativa, non è tuttavia pienamente risolutivo in quanto le sue premesse teoriche riguardano *in primis* la comunicazione orale. Da questo punto di vista, il rapporto tra *prompting* e "performance" degli LLM presenta – va ribadito – evidenti asimmetrie rispetto al rapporto tra intenzionalità ed indizi nella comunicazione naturale. Si tratta di un primo elemento che ridimensiona – ma senza azzerarle – le analogie tra comunicazione artificiale e quella naturale ma che indica come obiettivo quello di prevedere LLM in grado di approssimarsi sempre più alla comunicazione naturale orale o, almeno, di assorbirne o riprodurne (via robotica) alcune istanze (gestualità, espressione del volto, etc). A questo punto, vale la pena ribadire che gli LLM sono rilevanti soprattutto nella misura in cui costituiscono un'estensione della cognizione umana, al di fuori di qualsiasi tentativo di simularla inutilmente. La presenza di indubbe simmetrie e asimmetrie connota le due forme di comunicazione ma il valore e le funzioni degli LLM, intesi come estensione della cognizione umana, sono evidenti anche in relazione alle asimmetrie. Le asimmetrie circoscrivono infatti le azioni dell'utente umano mediante forme di *prompting strategico*, utile per poi procedere soprattutto alla ricognizione delle informazioni, alla definizione organica delle stesse e all'integrazione delle conoscenze, accentuando così la natura intersoggettivo-sociale delle stesse anche in virtù della rimodulazione della capacità meta-rappresentazionale, narrativa e persuasiva che è un'istanza decisiva del *mindreading* (naturale) dell'utente (Adornetti & Ferretti, 2025). Se, quindi, il *mindreading* degli LLM è una metafora che sta per un comportamento che è da intendersi in tutta la sua complessità, la sua incidenza su quello dell'utente umano è carica di significato, anche in vista dell'esercizio della vigilanza epistemica che è forse l'esito più ragguardevole del *prompting strategico*. Chiarita la funzione positiva dell'asimmetria in relazione ai vantaggi cognitivi dell'utente umano, è utile ora una disamina complessiva del *continuum* tra esplicito e implicito, incentrandola – metodologicamente – sull'uso della metafora. Da questo punto di vista, la teoria della pertinenza può offrire strumenti di analisi significativi anche alla luce del cosiddetto "deflationary account of metaphors" (Sperber & Wilson, 2012, pp. 97-122). Una giustificazione deflazionistica non svaluta affatto la metafora intesa come risorsa cognitivo-linguistica: la funzione poetico-retorica, euristica ed estetica della metafora continua a suscitare interesse – nei parlanti, nei lettori e negli studiosi – nella misura in cui compensa gli usi letterali del linguaggio distanziandosene in termini decisi e palesi. Ciò avviene però in situazioni speciali, non normali. Il merito di una giustificazione deflazionistica consiste, invece, nella possibilità di valutare la metafora come un esempio normalissimo di quel *continuum* tra esplicito ed implicito che, come si è già visto, segna la comunicazione naturale ordinaria senza godere di alcuno statuto speciale. È rispetto a questo uso normale della metafora che si fa più interessante il confronto tra comunicazione ordinaria e quella artificiale tra utente umano e LLM.

5. L'indeterminatezza e gli LLM: la questione della metafora

Il *continuum* tra esplicito e implicito investe l'intera comunicazione naturale e, per questa ragione, non presenta alcuna unicità o specificità; nel caso dell'uso normale della metafora, il *continuum* comporta uno scorrimento cognitivo scandito da tre interpretazioni: la *letterale* (*literal*), l'*approssimativa* (*loose*) e l'*iperbolica* (*hyperbolic*) (Sperber & Wilson, 2012, p. 97). Così strutturato, il *continuum* determina una questione duplice: linguistica e cognitiva. Da un lato, esso investe l'interazione comunicativa *naturale* attivando il piano della comprensione nella misura in cui la possibile *implicatura* sollecita nell'interlocutore una riformulazione più esplicita che ne allinei la comprensione con l'intenzionalità del parlante (mittente) ma a partire da una nuova modulazione dell'enunciato/proposizione; dall'altro, *in foro interno* (del parlante e del suo interlocutore), l'eventuale rimodulazione implica un processo cognitivo di natura inferenziale che consiste proprio nell'esplorazione delle molteplici sfumature implicite in vista di quella più "pertinente" e, quindi, adatta ad un'*esplicatura* più risolutiva.

Il *continuum* tra esplicito e implicito è senza dubbio un elemento costitutivo dell'interazione linguistico-comunicativa naturale, intesa nella sua totalità, ma riguarda non solo le proiezioni metaforiche ma anche altre attività esposte all'*indeterminatezza* come l'esercizio della *traduzione* e i cosiddetti *atti linguistici indiretti* che potrebbero essere esaminate come altrettanti casi di studio sempre in vista di un'analisi contrastiva tra comunicazione naturale e quella artificiale via LLM (Ervas & Gola, 2016, pp. 76-78; Ruytenbeek, 2021; Garelli, 2024; Jucker, 2024).

Dalla seconda metà del Novecento gli studi e le ricerche sulla metafora si sono moltiplicati sensibilmente mettendone in evidenza la rilevanza sia presso i linguisti sia presso gli psicologi e i filosofi del linguaggio. Non più intesa come una risorsa retorico-stilistico-poetica, la metafora si configura come un dispositivo cognitivo-linguistico utile per una comprensione concettuale della realtà ormai del tutto svincolata dai criteri logici dell'astrazione da cui le stesse procedure della classificazione/categorizzazione dipendono (Ricoeur, 1975; Lakoff & Johnson, 1980). Anche l'accentuazione di questa funzione rientra però ancora in un modello teorico improntato all'unicità (*distinctiveness*) e alla completezza (*full-fledged theory*). Al contrario, l'uso normale della metafora si iscrive in situazioni in cui gli stessi usi letterali delle parole non sono affatto scontati: non è infatti così scontato che il letterale sia la controparte del metaforico o *ipso facto* la componente del linguaggio da compensare. Il letterale puro e codificato non è mai ravvisabile nella comunicazione naturale. In tal senso, va osservato che il significato letterale può evocare immagini o sensazioni tali da metterne in crisi la presunta struttura letterale a favore, invece, di significati più impliciti (Carston, 2010). Da questo punto di vista, il mito dell'univocità del significato letterale decade e diventa normale la possibilità di attivare interpretazioni e inferenze anche soltanto in relazione al solo significato letterale. La mente o intelligenza naturale tende costantemente a stabilire *concetti ad hoc* secondo diverse modalità (Garelli, 2024) in quanto l'implicito si annida, a vario titolo, in molteplici parole o enunciati, inducendo i parlanti e i loro interlocutori ad un continuo eserci-

zio di *costruzione del significato*, restringendo (*narrowing*) o estendendo (*broadening*) il significato letterale per poterne individuare la sfumatura più pertinente, vale a dire, il *concetto ad hoc*. Non si tratta dunque di concetti marginali o circoscritti agli usi retorici o poetici del linguaggio: nella loro dimensione ordinaria, il linguaggio e la comunicazione sono connotati da una costante *indeterminatezza* proprio a partire da quei significati letterali che il modello teorico di Grice preservava nel *codice* distanziandoli dai significati impliciti da ricondurre, mediante opportune inferenze, al livello dei primi (Grice, 1989). Secondo Sperber e Wilson, le interpretazioni letterali non sono affatto interpretazioni predefinite e automatiche e, quindi, non sono le prime interpretazioni che i parlanti prendono in considerazione rispetto alla presunta trasparenza del codice. Le interpretazioni letterali non sono affatto le più agevoli a disposizione dei parlanti (Sperber & Wilson, 2012, pp. 105-109). Sia rispetto a quelle sia rispetto alle interpretazioni metaforiche, il parlante ricorre all'ampia gamma degli "ostensive stimuli" con l'intenzione di meglio catturare l'attenzione del suo interlocutore per poi trasmettere il contenuto. Gli "ostensive stimuli" sono esterni al codice pur *accompagnandolo* significativamente. L'ipotesi di affidarsi al solo codice non è praticabile in quanto – va ribadito – le espressioni letterali, che vi fanno riferimento, non sono realmente esplicite. Da questo punto di vista, l'intera comunicazione naturale è connotata dall'indeterminatezza e dall'ambiguità che non sono da intendersi come un difetto della stessa bensì come la condizione strutturale della necessità da parte dei parlanti di far leva sulle proprie intenzioni rispetto alle molteplici istanze contestuali. In tal senso, il processo metaforico prende avvio da quello percettivo-corporeo (*embodiment*): la struttura implicita dei *concetti ad hoc* rinvia a quei tratti dei processi percettivo-corporei in grado di cogliere somiglianze tra diverse realtà senza più far riferimento ai rapporti determinati dall'uso letterale del linguaggio. L'individuazione dei significati è affidata esclusivamente ai processi inferenziali che si attivano congiuntamente con gli enunciati, rendendo così la comunicazione naturale un'attività intrinsecamente "context-sensitive" (sensibile al contesto) (Sperber & Wilson, 2012, pp. 101-105). È questo un processo di disambiguazione che si attiva costantemente in relazione sia al *continuum* tra esplicito e implicito sia al *continuum* tra le espressioni metaforiche e quelle (pseudo)letterali. La situazione si complica se il processo di disambiguazione viene preso in considerazione rispetto alla distinzione tra *metafore morte* e *vive* che si colloca all'interno di quella che è ancora una teoria della metafora a statuto speciale e, quindi, non aderente al modello della pertinenza. Lo statuto delle cosiddette *metafore morte* richiede un'analisi più articolata: sono metafore di natura ontologica in quanto scaturiscono da schemi o immagini legate fortemente all'*embodiment*; in tal senso, esse segnano i processi cognitivo-linguistici così fortemente da esporsi ad una vera e propria normalizzazione che si risolve nella loro lessicalizzazione e definizione convenzionale, riducendo così il ruolo di quelle espressioni letterali che *in primis* caratterizzano il dominio di partenza di una determinata metafora (Ervas & Gola, 2016, pp. 41-47). Sebbene strutturato metaforicamente, il dominio di arrivo non richiede nessun processo di disambiguazione: la sua configurazione è, al contempo, metaforica e lessicalizzata/convenzionale e quindi non sollecita – almeno in una prima fase della disamina – alcuna interpretazione. L'espressione "Sono fuori di me" è un ottimo esem-

pio di metafora morta: non richiede alcun processo di disambiguazione in quanto il riferimento allo *schema del contenitore* – funzionale allo statuto ontologico di questa tipologia di metafore – rimane invariato nel passaggio dal dominio di partenza a quello di arrivo. In questa prospettiva, la priorità cognitiva della metafora viene confermata e, soprattutto, se ne accentua l'allineamento con le normali espressioni letterali. Diversa è la situazione delle *metafore vive*: non sono convenzionali in quanto le implicature contenute sono nuove, alternative non solo alle espressioni letterali affini ma anche alle stesse *metafore morte*; le implicature sottostanti, contestuali all'evocazione di immagini nuove o inconsuete, espongono questo tipo di metafora ad interpretazioni che, rispetto al significato letterale del dominio di partenza, ne accentuano, da un lato, le potenzialità concettuali, dall'altro, lo espongono a quel processo di disambiguazione orientato verso il miglior grado di pertinenza. Le espressioni “il bosone della biologia” e “i quanti della coscienza” (esempi di metafora viva) non attivano collegamenti immediati con il *bosone* e con *i quanti della fisica* (dominio di partenza) e, quindi, sollecitano un processo di disambiguazione che richiede tempi di maturazione più dilatati in quanto, necessariamente e implicitamente, entrano in gioco immagini, informazioni di diversa natura, stimoli culturali che, paralleli alle espressioni letterali, ne garantiscono o risolvono la piena pertinenza.

Questa è la premessa per rendere la metafora ancora più funzionale al confronto tra comunicazione naturale e comunicazione artificiale tra utenti umani e LLM. Emerge un'asimmetria decisiva ai fini della presente disamina; un'asimmetria che integra quella inerente alle nozioni di *pertinenza* e *mindreading* affrontata nel precedente paragrafo.

Nella comunicazione naturale, le implicature inerenti alle metafore vive mettono il parlante e il suo interlocutore nella condizione di rilevarne la pertinenza mediante più inferenze di diversa natura e, soprattutto, mediante lo scorrimento del cosiddetto continuum tra implicito ed esplicito orientato verso l'individuazione dell'indizio più influente. Malgrado qualche pregiudizio, gli LLM danno prova di saper gestire le metafore: rispetto a determinati *prompting* dell'utente, gli LLM sono in grado di descrivere e spiegare le metafore, fornendo persino diversi esempi applicativi. Questa capacità si spiega in virtù del fatto che gli LLM attingono a fonti testuali che già prevedono determinate implicature, ricondotte *quantitativamente* a perifrasi letterali o a segmenti discreti, e disponibili per la combinazione delle stesse in diverse situazioni e/o in diversi enunciati. Nella comunicazione naturale, al contrario, le implicature presentano una diversa struttura e si caratterizzano soprattutto *qualitativamente*; non esiste, altresì, un campionario di implicature; il che comporta da parte dei parlanti e dei loro interlocutori la capacità di individuare la pertinenza in contesti non lineari e multiformi. L'asimmetria in questione è la seguente: *in primis* sembra che gli LLM possano gestire le sole metafore morte, la cui scomposizione è desumibile dalle fonti testuali (scrittura) a cui essi attingono in virtù di un sistema neurale artificiale che ne determina l'addestramento continuo; nella comunicazione naturale, invece, i parlanti attivano processi inferenziali in vista dell'individuazione di tratti impliciti pertinenti senza ricorrere né ad eventuali repertori o fonti né a criteri logico-linguistici coscientemente supervisionati ma poco aderenti alla dimensione qualitativa contestuale –

anche se, talvolta, inconscia o automatica – che orienta la ricognizione del *continuum* tra implicito ed esplicito. Questa è una sintesi ricorrente ma parzialmente efficace: è infatti noto che alcuni modelli LLM sono in grado di gestire metafore nuove (*metafore vive*) in contesti altrettanto nuovi, senza necessariamente far riferimento ad elementi semantici lessicalizzati; qualche difficoltà è tuttavia connessa a determinate configurazioni senso-motorie – non sempre coincidenti con i consueti schemi ontologici (contenitore, spostamento verticale o orizzontale etc.) – di alcune metafore che gli LLM non riconoscono e, quindi, non elaborano. Alcuni studi molto recenti, incentrati proprio sull’analisi sia delle metafore vive sia di quelle di matrice senso-motoria, riportano diversi esempi di metafore di questo tipo (Barattieri di San Pietro et al., 2023; Mangiaterra et al., 2025). Se è vero che gli LLM non si limitano alla gestione delle sole *metafore morte*, è altrettanto vero che essi possono applicarsi alle metafore nuove ma non a quelle che siano giustificabili in relazione a processi cognitivi del tutto radicati nell’*embodiment*. La sfera d’azione degli LLM non è quindi monolitica e agevole: la classica distinzione tra metafore morte e vive è ormai oggetto di diverse obiezioni nella misura in cui può essere ampiamente dimostrato che gli LLM possono applicarsi a metafore non ancora esposte alla lessicalizzazione, tranne a quelle strettamente connesse all’esperienza percettiva intesa nella sua totalità. Oltretutto, lo stesso quadro delle *metafore morte* non presenta più i contorni netti e forti del modello teorico tradizionale (à la Ricoeur): si tratta di metafore la cui lessicalizzazione consente la gestione degli LLM a condizione che questi ultimi siano attivati in relazione ad un’unica lingua; se, al contrario, l’obiettivo è quello di tradurre le metafore morte da una lingua ad un’altra, gli LLM mostrano diverse difficoltà di elaborazione: la lessicalizzazione si trasforma in un ostacolo, non consentendo quell’uso normale caratterizzato dalla possibilità di valorizzare il ruolo cognitivo-concettuale delle *metafore morte* mediante perifrasi e diverse rimodulazioni. Prende forma una vera e propria situazione limite: nei dispositivi di traduzione implementati negli LLM, le *metafore morte* possono “ritornare in vita” e rendere così meno praticabile il riferimento ai repertori, alle perifrasi e alle rimodulazioni letterali. D’altra parte, questo è un fenomeno ravvisabile nella stessa comunicazione naturale interlinguistica; il che stabilisce che la natura ontologica delle metafore non è sufficiente a garantire sempre la loro traducibilità. In tal senso, il processo di disambiguazione è duplice: da un lato, è connesso all’attivazione di più risorse cognitive, dall’altro, deve contemplare la possibilità di selezionare bene nella *lingua d’arrivo* immagini e/o perifrasi che siano aderenti al *dominio di arrivo* della metafora oggetto di traduzione. Ricorrendo allo strumento delle percentuali, soltanto parzialmente le metafore morte sono traducibili; l’eventuale evocazione di *immagini ad hoc*, vale a dire, di immagini non scontate ed implicite, potrebbe, da un lato, aggiungere un elemento di novità, non contemplato dall’uso lessicalizzato, dall’altro, richiedere uno scorrimento cognitivo tra le immagini più evocabili e quelle meno evocabili in vista dell’immagine più pertinente; si ricorre così a quel processo inferenziale sensibile agli aspetti contestuali che il parlante esercita costruendo o decostruendo il *continuum* tra esplicito ed implicito (Ervas & Gola, 2013; Kövecses, 2014; Ervas & Gola, 2016; Kövecses, 2020). È evidente che, al di fuori di repertori di metafore

morte, scomposte, descritte e rimodulate in termini di lessicalizzazione, diventa piuttosto arduo il tentativo di attribuire *in toto* alla comunicazione artificiale tra utenti umani e LLM la pratica di processi inferenziali sensibili al contesto, agli stimoli ostensivi, incluse le multiformi immagini evocabili, e agli eventuali significati letterali. Dal punto di vista strettamente semantico, può essere utile far osservare che gli LLM sono costruiti in modo tale da gestire istanze linguistiche (significati di parole o di espressioni, significati di metafore, etc) dotate di un tasso di frequenza elevato; in tal senso, essi non possono gestire o elaborare né il significato complessivo di alcune metafore vive né il valore di alcune parole polisemiche in quanto, in entrambe le situazioni, sono coinvolti “sensi” poco frequentemente rilevati (riferimento esclusivo alla lingua inglese) e, quindi, non progressivamente aderenti a quell’indeterminatezza che caratterizza sia le metafore vive sia le parole polisemiche (Meconi et al., 2025).

6. Conclusioni

Il motivo ispiratore del presente saggio è stata la possibilità di adottare la Teoria della Pertinenza di Sperber e Wilson nell’intento di descrivere l’interazione tra l’utente umano e gli LLM e di coglierne alcune analogie con la comunicazione naturale orale. Da questo punto di vista, è stato utile esaminare il dispositivo del *prompting* e verificarne l’aderenza delle funzioni al principio della pertinenza. Il *prompting* è il dispositivo che consente all’utente umano di gestire le *implicature* a partire dalla fase di ricognizione delle stesse fino alla definizione delle *esplicature*. Come si è potuto vedere, la ricognizione delle implicature è presente in entrambe le comunicazioni ma richiede modalità di realizzazione che non possono coincidere. Anche se non ancora supportata da una disamina *corpus-based*, da legare eventualmente anche alla valutazione di fattori di natura cognitiva e neuroscientifica, la questione della metafora mette lo studioso nella condizione non solo di individuare e giustificare alcune simmetrie – inerenti alla ricognizione della pertinenza mediante *prompting* – ma anche di prendere atto di alcune evidenti asimmetrie che però non concorrono minimamente a mettere in discussione né gli LLM né l’azione dell’utente. Per quanto lessicalizzate e ridotte ad un significato *quasi letterale*, le cosiddette *metafore morte* consentono all’utente di rendersi consapevole della presenza di molteplici implicature e di perfezionare il proprio capitale semantico e, quindi, di valorizzare le proprie competenze semantico-pragmatiche sia aprendosi ad una dimensione testuale intersoggettiva sia rimodulando continuamente le proprie capacità meta-rappresentazionali (*mindreading*). Si tratta di implicazioni positissime che però non azzerano o ridimensionano la specificità delle due forme di comunicazione. La comunicazione artificiale, quella che investe l’interazione tra gli utenti umani e gli LLM, non è dunque sovrapponibile alla comunicazione naturale in quanto non coordinata e motivata dall’attivazione di processi inferenziali sensibili ai multiformi fattori contestuali, agli stimoli ostensivi, alle molteplici immagini evocabili, congiunte sia alle metafore morte sia ai significati letterali.

Bibliografia

- Adornetti, I., & Ferretti, F. (2025). *Origins of Meaning*. Elsevier.
- Allott, N., & Textor, M. (2026). Literal and metaphorical meaning: in search of a lost distinction. *Inquiry: An Interdisciplinary Journal of Philosophy*, 69(2), 290-317.
- Barattieri di San Pietro, C., Frau, F., Mangiaterra, V., & Bambini, V. (2023). The pragmatic profile of ChatGPT: Assessing the communicative skills of a conversational agent. *Sistemi intelligenti*, 2, 379-400.
- Branda, F. (2026). When Artificial Intelligence Shapes the way we Think. *Philosophy & Technology*, 39(42). <https://doi.org/10.1007/s13347-026-01055-y>
- Carston, R. (2010). Metaphor: Ad Hoc Concepts, Literal Meaning and Mental Images. *Proceedings of the Aristotelian Society*, 110(3), 295-321.
- Chandaria, A. (2024, agosto). *Mindreading Models: Do LLMs Have Theory of Mind*. The Concept. <https://www.theconcept.ai/articles/mindreading-models-the-consequences-of-theory-of-mind-in-llms>
- De Caro, M., & Giovanola, B. (2025). *Intelligenze. Etica e politica dell'IA*. Il Mulino.
- Eragamreddy, N. (2025). The impact of AI on pragmatic competence. *Journal of Teaching English for Specific and Academic Purposes*, 169-189.
- Ervas, F., & Gola, E. (2013). Lessico e immaginazione nella traduzione delle metafore. *E.C. Rivista dell'Associazione italiana Studi semiotici*, 7(17), 91-6.
- Ervas, F., & Gola, E. (2016). *Che cos'è una metafora*. Carocci Editore.
- Forceville, C. (2025). Relevance theory as the foundation for an inclusive theory of communication. *Multimodal Communication*, 15(1), 5-20.
- Garello, S. (2024). *The Enigma of Metaphor. Philosophy, Pragmatics, Cognitive Science*. Springer.
- Grice, P. (1989). *Studies in the Way of Words*. Harvard University Press.
- Jucker, A. H. (2024). *Speech Acts. Discursive, Multimodal, Diachronic*. Cambridge University Press.
- Kövecses, Z. (2014). *Conceptual Metaphor Theory and the Nature of Difficulties in Metaphors Translation*. In D. R. Miller & E. Monti (a cura di), *Translating Figurative Language*, Ceslic.
- Kövecses, Z. (2020). *Extended Conceptual Metaphor Theory*. Cambridge University Press.
- Kravchenko, N., Kravets, O., Naumova, Y., Uriadova, V., & Shepelska, I. (2024). Analyzing creative metaphors in advertising: Integrating relevance theory and conceptual blending approaches. *Amazonia Investiga*, 13(82), 177-185.
- Lakoff, G., & Johnson, M. (1980). *Metaphors We Live By*. University of Chicago Press.
- Littlemore, J. (2019). *Metaphors in the Mind. Sources of Variation in Embodied Metaphor*. Cambridge University Press.
- Mangiaterra, V., Barattieri di San Pietro, C., Frau, F., Bambini, V., & Al-Azary, H. (2025). On choosing the vehicles of metaphors without a body: evidence from Large Language Models. In *Proceedings of the 2nd Workshop on Analogical Abstraction. Cognition, Perception, and Language (Analogy-Angle II)*, 37-44.

- Meconi, D., Stirpe, S., Martelli, F., Lavallo, L., & Navigli, R. (2025). Do Large Language Models Understand Word Senses? In *Proceedings of the 2025 Conference on Empirical Methods in Natural Language Processing*, 33897-33916.
- Ricoeur, P. (1975). *La métaphore vive*. Seuil.
- Roncaglia, G. (2023). *L'architetto e l'oracolo. Forme digitali del sapere da Wikipedia a ChatGPT*. Editori Laterza.
- Ruytenbeek, N. (2021). *Indirect Speech Acts*. Cambridge University Press.
- Sperber, D., & Wilson, D. (1995). *Relevance. Communication & Cognition*. Blackwell Publishing.
- Sperber, D., & Wilson, D. (2012). *Meaning and Relevance*. Cambridge University Press.
- Strachan, J. W. A., Albergo, D., Borghini, G., Pansardi, O., Scaliti, E., Gupta, S., Saxena, K., Rufo, A., Panzeri, S., Manzi, G., Graziano, M. S. A., & Becchio, C. (2024). Testing Theory of Mind in Large Language Models and humans. *Nature Human Behaviour*, 8(7), 1285-1295.