

Plebe, Alessio, & De La Cruz Vivian M., (2016), *Neurosemantics: Neural Processes and the Construction of Linguistic Meaning* (Vol. 10), Springer.

La neurosemantica di cui trattano Alessio Plebe e Vivian De La Cruz è “il progetto di spiegare il significato delle parole in termini di computazioni del cervello” (p.1). I due autori presentano una ricerca originale, coraggiosa e metodologicamente consapevole. Il volume è composto da nove capitoli e diviso in due parti. La prima (capitoli 1-5) è dedicata alla discussione fondazionale del progetto, ai suoi concetti principali e alla sua localizzazione storica e logica in filosofia, psicologia e neuroscienze. La seconda parte introduce e illustra tre casi di applicazione elaborati dai due autori: un modello di neurosemantica per i nomi di oggetti (cap. 6), aggettivi di colore (cap. 7), termini morali (cap.8), e aggettivi numerali con considerazioni conclusive su possibili futuri sviluppi e ricerche in corso (cap.9). Ogni modello riproduce alcuni aspetti della competenza lessicale di queste categorie di parole in termini di interazione tra stimoli (ad esempio la percezione di una superficie verde e del suono “verde”) e attivazione di popolazioni di neuroni nell’area corticale, e comprende equazioni matematiche che simulano questa interazione. Presento qui brevemente il contenuto dei primi otto capitoli per poi concludere con un breve commento.

La centralità della matematica, definita “golden standard” per la semantica, è chiarita nell’Introduzione, nel capitolo 1, e ripresa nel capitolo 5. In estrema

sintesi, la tradizione di Leibniz, Boole e Frege aveva utilizzato il calcolo della logica per fare luce sul significato, poi questa tradizione ha mostrato difficoltà insormontabili che hanno via via portato alla “psicologizzazione” del significato delle parole, e quindi allo svincolamento del significato dei termini da quello degli enunciati o delle proposizioni, e alla ricerca sui concetti. Ma al paradigma cognitivista fino agli anni 2000, secondo Plebe e De La Cruz, è mancato un impianto matematico comparabile con quello della logica, nonché una migliore base empirica per studiare la natura del linguaggio, che ora è possibile grazie alla psicologia cognitiva, ma anche e soprattutto grazie allo studio del cervello.

Il progetto della nuova semantica modellizzabile matematicamente parte dunque dal neurone, e il capitolo 2 riprende in breve la storia delle scoperte che hanno portato alle conoscenze odierne sulla cellula fondamentale del sistema nervoso. Il neurone viene assunto come unità di computazione nel progetto della neurosemantica, anche se i due autori spiegano che l’accettazione del “dogma neuronale” ha ragioni più pragmatiche che teoriche, dato che nuovi candidati potrebbero emergere (i dendriti, ad esempio) studiando l’organizzazione del cervello. Il neurone, o meglio le popolazioni di neuroni, sono “plastiche”, cioè in grado di imparare dall’esperienza pressoché qualsiasi funzione. Questo, spiegano gli autori, rende l’idea stessa di unità computazionale applicabile al caso dei neuroni solo con la specificazione che si tratterà sempre di accettare un certo livello (alto) di semplificazione dei fenomeni. Il resto del capitolo 2 spiega come è strutturata la corteccia cerebrale e che cosa sappiamo sulle varie aree su cui è possibile, in una certa misura, mappare competenze di interesse semantico.

Il capitolo 3 affronta temi teorici propri del dibattito filosofico recente: che

cos'è una rappresentazione, una computazione, un meccanismo, e in che misura tali nozioni si possano applicare al caso dei neuroni. Plebe e De La Cruz riassumono le questioni sollevate da autori come Jerry Fodor, Fred Dretske e Ruth Millikan negli anni novanta del secolo scorso, in particolare il problema della normatività: se la rappresentazione (linguistica o cognitiva) di un oggetto è prodotta da questo come l'effetto di una causa, come possiamo distinguere le cause proprie da quelle improprie - ad esempio, affermare che la parola "cane" non si applica correttamente a un gatto visto di sfuggita, o a qualsiasi altro oggetto che ne causi l'occorrenza? La soluzione a cui erano arrivati gradualmente Dretske e Millikan è la cosiddetta *teleosemantica*: la funzione propria di una rappresentazione - ciò per cui è stata acquisita o selezionata dal sistema che la comprende - determina il suo significato, e quindi permette di determinare il *range* delle applicazioni scorrette. Plebe e De La Cruz adottano la teleosemantica per spiegare in che senso possiamo parlare di rappresentazioni neurali, ammettendo però che, in buona sostanza, essa ha il solo vantaggio di "essere la mossa più sicura per proteggere una teoria delle rappresentazioni neurali da una lunga lista di problemi (p. 44, traduzione mia). Successivamente, sempre nel capitolo 3, vengono discussi alcuni criteri per caratterizzare un processo come computazione, in particolare l'indipendenza dal mezzo (*medium independence*); la conclusione è che per parlare di computazione neurale questi criteri vanno ripensati, senza cadere d'altra parte in una forma di pancomputazionalismo, in cui ogni processo descrivibile matematicamente si definisce computazione: la computazione neurale è dunque «un'etichetta in cerca di teoria» (p. 47). Considerazioni analoghe nella forma valgono per la nozione di meccanismo, centrale oggi nel dibattito sulla spiegazione in filosofia della scienza. Se una

spiegazione meccanicista completa richiederebbe una corrispondenza uno a uno fra parti e operazioni del modello esplicativo e parti e operazioni del meccanismo reale, questo è impossibile di fatto - e forse per principio - nel caso dei neuroni, tenendo conto della mancanza di dettaglio della nostra conoscenza della corteccia cerebrale, ma anche della plasticità intrinseca al funzionamento del sistema. La spiegazione meccanicista è quindi un'«idea guida» per la neurosemantica, non un criterio di adeguatezza (p. 49). Sempre nel capitolo 3 Plebe e De La Cruz spiegano qual è il meccanismo fondamentale che permette la rappresentazione neurale: il riconoscimento di coincidenze (*coincidence detection*), cioè la capacità di rilevare segnali uguali e di reagire a tale rilevamento. I neuroni rappresentano oggetti e proprietà nel mondo imparando per associazione dal rilevamento di stimoli coincidenti. Gli autori difendono in una breve rassegna storica l'idea dell'apprendimento per associazione, che ha avuto poca fortuna nella scienza cognitiva classica. Il capitolo 3 contiene anche informazioni sull'organizzazione corticale e illustra le nozioni di mappa corticale e campo ricettivo. La prima parte del volume si chiude con il capitolo 4, in cui Plebe e De La Cruz propongono l'equazione matematica che computa l'attivazione di un'unità di rappresentazione neurale in una mappa corticale in risposta a uno stimolo, e che sarà alla base delle applicazioni nella seconda parte.

L'idea del capitolo 6 è la seguente: comprendere il nome di un oggetto, come "tavolo", consiste primariamente nel saper mettere insieme lo stimolo uditivo della parola "tavolo" con quello visivo dell'oggetto a cui si riferisce. Quindi la semantica dei nomi di oggetti, che sono i primi ad emergere nel lessico infantile, si può spiegare mediante un modello che simula la convergenza di processi visivi e uditivi, e descrive matematica-

mente l'attivazione delle aree corticali preposte. Gli autori spiegano che gli stimoli, così come le mappe topografiche corticali, sono molto semplificate rispetto alla realtà, ma forniscono ugualmente uno *sketch* del meccanismo del fenomeno da spiegare. Dopo una prima fase di sviluppo, il modello "impara", e la presentazione di nuovi stimoli visivi e uditivi conferma l'attivazione delle aree precedentemente associate a stimoli dello stesso tipo.

Il capitolo 7 fa un passo avanti in termini di complessità del fenomeno da spiegare: dai nomi agli aggettivi. Rispetto ai nomi di oggetti, gli aggettivi richiedono a chi apprende un cervello con abilità ulteriori rispetto al riconoscimento di stimoli cross-modalità con la stessa origine: la sensibilità sintattica all'ordine delle parole nella frase e la capacità di selezionare una dimensione percettiva (come forma o colore) rispetto alle altre – abilità che richiedono più memoria e maturità neurale. Gli autori presentano le equazioni descrittive di un meccanismo che connette input uditivi (suono della parola) e visivi (percezione del colore) con l'attivazione di popolazioni di neuroni a vari livelli. Molto interessante per chi scrive è la complicazione affrontata nel paragrafo 7.2.4., ovvero la questione del relativismo linguistico applicato ai termini di colore. In che misura i termini di colore a cui siamo esposti nella nostra comunità linguistica modificano il modo in cui percepiamo i colori? Com'è noto, l'universalismo linguistico e cognitivo di stampo chomskyano è stato recentemente messo alla prova da studi sperimentali in laboratorio e sul campo sulla varietà dei concetti di colore in popolazioni con lessici diversi. Plebe e De La Cruz, che hanno già affrontato questo tema in ricerche precedenti, ipotizzano che i fattori in base a cui varia la percezione dei colori siano non solo il lessico, ma anche la "dieta visuale", cioè il tipo di stimoli più frequentemente processati. L'ipotesi viene

testata controllando il modello di meccanismo in due condizioni diverse (con stimoli visivi tipici dei nostri climi, e dell'ambiente africano in cui vivono le popolazioni Berimno e Himba, popolazioni studiate dai linguisti per il loro diverso lessico dei colori) e in tre lingue diverse: inglese, Berimno e Himba, in un compito che mette in gioco la percezione categoriale (cioè la nostra capacità situare il discrimine tra due zone di stimolo contigue rispetto ad una dimensione, come le sfumature tra il blu e il verde). A parziale conferma dell'ipotesi, il modello mostra che i giudizi di similarità tra stimoli visivi di colore sono influenzati dalla lingua parlata, ma non altrettanto dalla dieta visiva.

Il capitolo 8, sulla neurosemantica dei termini morali, è il più pionieristico del volume: mentre per i nomi di oggetto e gli aggettivi di colore esiste già lavoro in neurosemantica, in questo campo la ricerca è agli inizi. L'idea guida dei due autori è l'espressivismo (termine che ci riporta alla filosofia analitica del linguaggio e in particolare all'etica come semantica di R. M. Hare): un termine morale esprime emozioni, e la morale è educazione emotiva. Su queste basi è possibile ipotizzare un modello in cui le aree corticali associate alle emozioni positive e negative vengono associate a stimoli visivi e uditivi di natura morale. Dopo una fase di training il modello sarà in grado di mostrare che ha imparato a giudicare "sbagliato" il furto, perché associa emozioni negative all'illustrazione che riproduce una scena di furto.

Come anticipavo all'inizio, le ricerche contenute in questo libro sono importanti e innovative, e spiccano per originalità nel campo della filosofia del linguaggio: costituiscono lo sforzo di rendere operativa e reale l'idea della matematizzazione del linguaggio, che è per alcuni coincidente con l'idea stessa di una teoria del significato. Anche solo per questo il volume di Plebe e De La Cruz me-

rita il massimo interesse. La mia competenza si ferma prima del livello tecnico richiesto a valutare nel merito le soluzioni proposte nei capitoli in cui il modello vero e proprio viene presentato nella struttura generale e nelle applicazioni specifiche, ed è auspicabile che gli scienziati cognitivi che lavorano in questo campo discutano queste proposte nel dettaglio.

Una delle domande la neurosemantica lascia aperta è di natura metodologica: se le nozioni filosofiche tradizionali – computazione, rappresentazione, meccanismo – sono tutte sostanzialmente inadatte a caratterizzare il progetto (come ci ricordano gli autori in varie parti del libro, e come ho notato sopra), allora abbiamo discusso per decenni di concetti che il progresso sulla conoscenza del cervello spazzerà via, e la neurosemantica ha bisogno di un altro paradigma, oppure la neurosemantica risponde a domande diverse da quelle che chi ha discusso di rappresentazione o spiegazione relativamente al linguaggio si era posto? L'impianto di questo libro sembra suggerire – per dirla con una frase idiomatica presa in prestito – che *the proof is in the pudding*: se la neurosemantica funzionerà e andrà avanti con modelli dell'uso del linguaggio sempre più adeguati esplicativamente e predittivamente, mostrerà che la sua ricetta filosofica è giusta, indipendentemente da quanto corrisponda con le nozioni tradizionali.

Elisabetta Lalumera

Università di Milano-Bicocca
elisabetta.lalumera@unimib.it