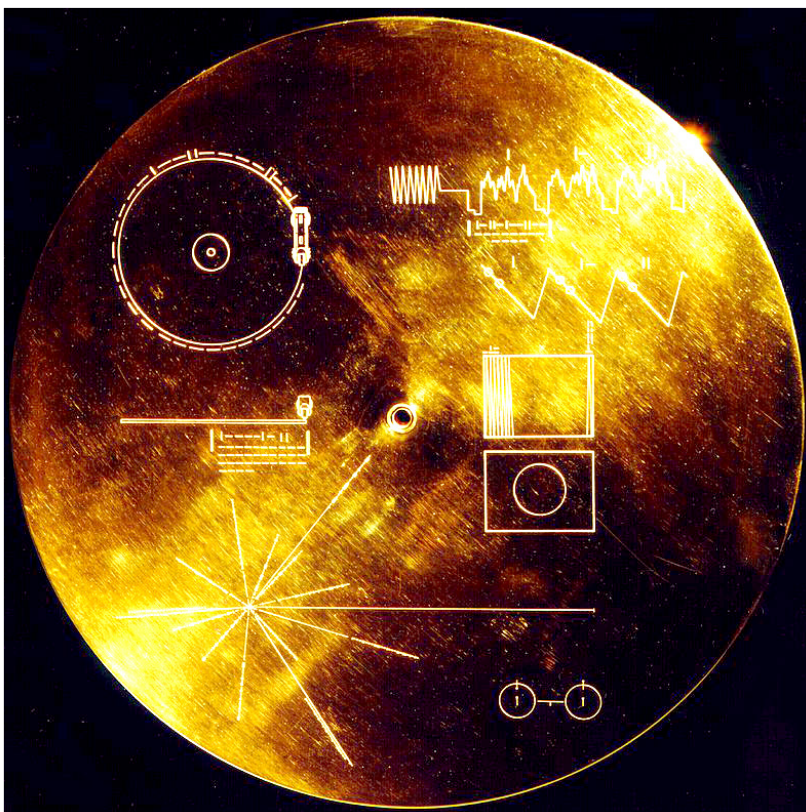




COMPTER PARLER SOIGNER

Tra linguistica e intelligenza artificiale

Atti

Pavia, 15-17 dicembre 2014

a cura di

Edoardo Maria Ponti – Marco Budassi

Pavia University Press
Scientifica

Atti



COMPTER PARLER SOIGNER

Tra linguistica e intelligenza artificiale

Atti

Pavia, Collegio Ghislieri, 15-17 dicembre 2014

a cura di

Edoardo Maria Ponti – Marco Budassi



Pavia University Press

Compter parler soigner : tra linguistica e intelligenza artificiale : atti : Pavia, Collegio Ghislieri, 15-17 dicembre 2014 / a cura di Edoardo Maria Ponti, Marco Budassi. – Pavia : Pavia University Press, 2016. – XIX, 119 p. ; 24 cm.

(Atti)

<http://archivio.paviauniversitypress.it/oa/9788869520389.pdf>

ISBN 9788869520372 (brossura)

ISBN 9788869520389 (e-book PDF)

In testa al front.: Ghislieri.

© 2016 Pavia University Press, Pavia

ISBN: 978-88-6952-038-9

Nella sezione *Scientifica* Pavia University Press pubblica esclusivamente testi scientifici valutati e approvati dal Comitato scientifico-editoriale.

I diritti di traduzione, di memorizzazione elettronica, di riproduzione e di adattamento anche parziale, con qualsiasi mezzo, sono riservati per tutti i paesi.

I curatori sono a disposizione degli aventi diritti con cui non abbiano potuto comunicare, per eventuali omissioni o inesattezze.

In copertina: *cover del disco Golden Record, a bordo della sonda Voyager lanciata nel 1977, contenente saluti in 55 lingue e un messaggio vocale di Jimmy Carter.*

Prima edizione: luglio 2016

Pavia University Press – Edizioni dell'Università degli Studi di Pavia
Via Luino, 12 – 27100 Pavia (PV) – Italia
www.paviauniversitypress.it – unipress@unipv.it

Printed in Italy

Sommario

Introduzione

Edoardo Maria Ponti e Marco Budassi	vii
---	-----

Calcolo di Lambek, logica lineare e linguistica

Claudia Casadio	1
1. Introduzione	1
2. La distinzione funzione–argomento	2
3. Calcolo di Lambek	4
4. Logica Lineare Non-commutativa	7
5. Calcolo di Lambek e logica lineare	11
6. Alcuni esempi linguistici	12
7. Conclusioni	16

Il processamento in tempo reale delle frasi complesse

Cristiano Chesi	21
1. Introduzione	21
2. Una rappresentazione ‘minimalista’ del problema	21
3. Alcune frasi difficili: frasi relative incassate e frasi scisse	24
4. Una rivisitazione <i>processing-friendly</i> delle Grammatiche Minimaliste	32
5. Conclusioni	36

Semantica distribuzionale. Un modello computazionale del significato

Alessandro Lenci	39
1. Introduzione	39
2. Principi e metodi di semantica distribuzionale	40
3. Valutare i modelli semantici distribuzionali	45
4. Le sfide per la semantica distribuzionale	46
5. Conclusioni e prospettive	48

ULISSE: una strategia di adattamento al dominio per l’annotazione sintattica automatica

Felice Dell’Orletta e Giulia Venturi	55
--	----

1. Introduzione	55
2. Adattamento al Dominio: definizione del problema	56
3. Cos'è un dominio?	61
4. Adattamento al Dominio: metodi	67

Alcuni principi linguistici per costruire risorse lessicali

Elisabetta Jezek	81
1. Introduzione	81
2. La risorsa	81
3. Preferenze di selezione e categorie ontologiche	84
4. <i>Mismatch</i> dei tipi semantici	86
5. Set lessicali instabili	87
6. Conclusioni	89

Alcune note sui Teoremi di Incompletezza di Gödel e la conoscenza delle macchine

Alessandro Aldini, Vincenzo Fano e Pierluigi Graziani	93
1. Gli argomenti gödeliani	93
2. Il punto di vista di Gödel	99
3. Macchine che conoscono se stesse	103
4. Conclusioni	111

<i>Abstracts in English</i>	115
---------------------------------------	-----

Introduzione

Edoardo Maria Ponti e Marco Budassi – Università degli Studi di Pavia e Istituto Universitario di Studi Superiori (I.U.S.S.) – {edoardomaria.ponti01, marco.budassi01}@universitadipavia.it

La scienza moderna nacque abbattendo una barriera: anche il mondo sublunare, fino ad allora concepito come il dominio degli accidenti e delle simpatie naturali, divenne ‘geometrizabile’ e meccanico al pari del mondo stellare (Koyré, 1957). Questa rivoluzione aprì un nuovo campo del sapere. Durante i nostri giorni, sta avendo luogo un processo forse più modesto nella portata storica, ma non troppo dissimile: i metodi quantitativi vengono applicati sempre più rigorosamente e pervasivamente a un oggetto tipicamente qualitativo, il linguaggio naturale.¹

Perché proprio ora? Questo processo di cambiamento è ascrivibile a due ordini di motivi. Il primo è lo sviluppo di calcolatori sempre più potenti e di algoritmi sempre più efficienti e accurati. Il secondo è l’accessibilità e la vastità dei dati su cui essi possono basarsi: basti pensare a quanto intenso sia il flusso dell’odierno scambio di informazioni. Tuttavia l’idea di accostare calcolatori e linguaggio è tutt’altro che nuova, e anzi appariva già nei lavori inaugurali di Intelligenza Artificiale, risalenti addirittura a prima che la disciplina assumesse questo nome.

Un esempio è il celebre lavoro di Turing (1950), che diede una risposta pratica al seguente interrogativo: le macchine possono pensare? Poiché la questione sembrava sul piano teorico irrisolvibile,² propose un *imitation game* (conosciuto anche come ‘test di Turing’) per valutare se le macchine possano comportarsi in modo indistinguibile dagli esseri umani. Da quali campi intellettivi sarebbe convenuto iniziare per far competere le macchine? Turing ne cita due: il gioco (nello specifico gli scacchi) e il linguaggio.

Alle macchine in questi due campi non arrisero tuttavia uguali fortune. Un calcolatore, Deep Blue della IBM, ha sconfitto Garry Kasparov da campione del mondo di scacchi in carica il 10 febbraio 1996 a Philadelphia (Campbell et al., 2002).³ Al contrario, ancora nessun calcolatore parla come un essere umano: chiunque può accorgersene dalle imperfezioni nelle traduzioni automatiche, dai dialoghi talora assurdi con gli assistenti vocali dei dispositivi, e dai motori di ricerca che funzionano per parole chiave piuttosto che per domande. In breve, il successo negli scacchi deriva dal fatto che le sue combinazioni, per quanto innumerevoli e calcolabili solo entro certi limiti, derivano da poche, semplici regole fisse. Il lin-

¹Resta comunque problematico se così facendo si vada almeno in parte a colmare lo iato fra ‘due culture’ (Snow, 1959), ovvero fra quella umanistica e quella scientifica.

²Per argomenti teorici in favore e contro l’equivalenza fra il funzionamento delle macchine e il pensiero umano si veda però il contributo di Aldini, Fano e Graziani in questo volume.

³Un altro recente successo è AlphaGo, un calcolatore addestrato tramite *deep learning* (cfr. *infra*). Esso ha sconfitto a Seoul il 9 marzo 2016 Lee Se-Dol, il campione del mondo di go, un gioco da tavolo orientale ancora più complesso degli scacchi.

guaggio umano, invece, è un sistema intrinsecamente ambiguo, che sa adattarsi alle più svariate situazioni comunicative e non crolla a fronte dei mutamenti (più o meno radicali) che si accumulano nel tempo in ciascuna delle differenti lingue, che restano nondimeno intertraducibili. Inoltre, da un lato esso è partizionabile in numerosi livelli interni, dall'altro interagisce con una serie di facoltà cognitive ad esso esterne.

Poiché dunque l'intuizione è antica ma i successi nelle applicazioni recenti, i già menzionati ordini di motivi che li hanno permessi meritano di essere approfonditi. Innanzitutto, partiamo dal primo aspetto: i dati. Nel tempo, si è passati da raccolte di attestazioni scarse e disomogenee alla sistematica creazione di risorse linguistiche. Le più basilari sono i corpora, cioè insiemi di testi rappresentativi di una certa varietà linguistica. Assicurare questa rappresentatività implica campionare testi autentici, in un formato 'leggibile' da un calcolatore, e selezionati attraverso criteri esterni, ovvero indipendenti dal contenuto dei testi (McEnery et al., 2006). Questo perché la campionatura assicura un bilanciamento dei dati. Inoltre, l'uso di criteri interni è evitato per non creare un circolo vizioso fra le ipotesi teoriche e i dati scelti per verificarle.

Ad oggi i corpora tengono traccia pressoché di ogni dimensione del linguaggio: ne esistono di statici e dinamici (con più varietà diacroniche), generalisti e specialistici (con gergo settoriale di un dato dominio), monolingue e plurilingue (paralleli o comparabili a seconda che offrano traduzioni dello stesso testo o meno). Inoltre, ognuno di questi corpora può essere annotato, ovvero arricchito con informazione extra-linguistica. La stessa annotazione però può riguardare a sua volta diversi livelli del linguaggio. Una loro lista non esaustiva include almeno la fonetica, le parti del discorso, i lemmi, la sintassi (costituenti e dipendenze), i ruoli semantici, l'accezione delle parole, la coreferenza, la struttura dell'informazione, gli atti linguistici e la stilistica.⁴

Questa attenzione ai corpora si è accompagnata, nel dibattito della disciplina della linguistica teorica, a una contrapposizione fra paradigmi: da una parte quello funzionalista e induttivo, che ne ha promosso l'ascesa, dall'altra quello formalista e deduttivo, che ha preferito fondare la validazione delle teorie soprattutto sull'introspezione (ovvero il giudizio del parlante nativo circa la grammaticalità di un enunciato).⁵ La pratica dell'introspezione è stata accusata di essere piuttosto fragile a livello epistemologico in quanto è soggettiva e inaccessibile all'osservazione: questo spiega in parte perché la validazione empirica su risorse linguistiche ad oggi è largamente predominante nella disciplina dell'Intelligenza Artificiale.

D'altro canto, tuttavia, anche l'approccio basato sui corpora ha subito critiche e su alcuni suoi aspetti discordano gli stessi suoi sostenitori. Infatti una sua versione più radicale, definita *corpus driven*, prevede di non usare alcuna etichetta o modello che precedano i dati, ovvero che non emergano direttamente da essi. A

⁴Per un esempio di risorsa linguistica si veda il contributo di Jezek in questo volume.

⁵Questa è un'approssimazione utile a fini espositivi ma non completamente corretta: la scuola formalista infatti si affida a varie altre fonti di evidenza, in primo luogo le neuroscienze e i dati comportamentali. A questo proposito, si veda il contributo di Chesi in questo volume.

questa versione viene però obiettato che, così facendo, si continuano a impiegare comunque strumenti teorici soggiacenti, solo reintrodotti surrettiziamente e in modo ingenuo perché non problematizzato. Una versione più moderata è quella esposta sopra e definita *corpus based*, perché richiede la campionatura e ammette l'annotazione. Al contrario della variante *corpus driven*, essa correrebbe secondo i detrattori il rischio di minare l'integrità dei dati: campionandoli e annotandoli potrebbe sofisticarli, inficiando così le generalizzazioni che si possono trarre da essi (la cosiddetta fallacia del *cherry picking*).⁶

Vale la pena rilevare anche un ultimo punto: non sempre, una volta prefissato un compito, si trova una risorsa linguistica adeguata ad esso. Infatti, la competenza e il tempo necessari ad annotare dei dati con informazione *gold* (ovvero esatta per definizione in quanto associata da valutatori umani) impediscono di creare risorse a sufficienza. Questa penuria viene spesso mitigata attingendo dati grezzi da fonti non appositamente create per il compito in questione. In genere la loro bassa qualità, che comporta una dose di cosiddetto rumore, è spesso controbilanciata dalla loro vasta mole. Al giorno d'oggi, infatti, abbiamo la fortuna di poter contare su una fonte inesauribile e variegata di dati: internet. In rete, o nella posta elettronica, riversiamo una quantità inimmaginabile di informazione non strutturata ma già in formato digitale. Accedervi, ripulirla e trarne dei corpora che possano essere rappresentativi di una varietà linguistica minimizzando il rumore è una delle maggiori sfide della disciplina (Kilgariff, Grefenstette, 2003).⁷ Altre volte, invece, una risorsa adeguata al compito prefissato esiste ma è troppo poco ampia: un rimedio è espanderla estraendo i dati più rilevanti da altre risorse.⁸

L'ampiezza dei dati porta il nostro discorso a un secondo punto fondamentale, quello della potenza computazionale necessaria per gestirli. Gordon Moore nel 1965 constatò sulla base del sessennio precedente che la complessità dei microcircuiti (misurata in termini di *transistor per chip*) aumentava in progressione geometrica nel tempo. Questa tendenza si è mantenuta costante fino ad oggi, ma sembra che ora rischi di rallentare per via del fatto che si stanno ormai raggiungendo dei limiti fisici nella dimensione dei *transistor*, perché a livello nanometrico il loro surriscaldamento e le interazioni subatomiche minacciano la loro resistenza e affidabilità (Waldrop, 2016). Tuttavia, sta al contempo prendendo piede la parallelizzazione massiva dei processori stessi e l'integrazione dei processori grafici: queste strategie fanno sì che la potenza di calcolo continui a salire.

Ovviamente, la potenza di calcolo necessaria dipende da ciò che il calcolatore deve eseguire e i fini di quest'operazione. Da un lato, i dati possono fornire evidenza per generalizzazioni teoriche riguardanti il linguaggio: questo è l'ambito proprio della Linguistica Computazionale (*Computational Linguistics*). Dall'altro,

⁶Per una più vasta bibliografia su teorie formaliste e funzionaliste, si vedano anche Van Valin (2001) e Graffi (2010).

⁷Risorse sviluppate in questo modo sono state presentate ad esempio da Baroni et al. (2009) e possono essere scaricate dal sito WaCky – The Web-As-Corpus Kool Yinitiative, URL: <<http://wacky.sslmit.unibo.it/doku.php?id=corpora>>.

⁸Un esempio di come questa estensione possa essere implementata è discusso da Dell'Orletta e Venturi in questo volume.

il calcolatore può imparare dai dati a imitare comportamenti e processi cognitivi umani riguardanti il linguaggio o svolgere compiti che richiedono la gestione, la manipolazione o il filtraggio di dati linguistici: questo è l'ambito del Trattamento Automatico del Linguaggio Naturale (*Natural Language Processing*). In quest'ambito, il calcolatore viene valutato attraverso protocolli altamente standardizzati, che spesso sono ospitati in veri e propri eventi (perlopiù annuali) dove i dati sono resi pubblici e chiunque può competere con il proprio algoritmo, i cosiddetti *shared tasks*. Alcuni esempi internazionali molto celebri sono CoNLL e SemEval, che propongono un ampio ventaglio di *tasks*. Non mancano anche simili eventi focalizzati su una singola lingua, come EvalIta per l'italiano.⁹

Ma come sono fatti questi algoritmi? Perlopiù, il calcolatore viene allenato per estrarre da alcuni dati la conoscenza essenziale da applicare poi su dati nuovi.¹⁰ Questa tecnica è definita *machine learning*, 'apprendimento automatico'. Essa consiste in un'inferenza probabilistica da esempi a valori:¹¹ una prima differenziazione è quella tra metodi generativi e metodi discriminativi. I primi imparano la distribuzione congiunta dei valori y e degli esempi x , ovvero in notazione $P(y, x)$. I secondi imparano la distribuzione dei valori condizionata dagli esempi, ovvero $P(y|x)$. In altri termini, i modelli discriminativi stabiliscono i confini tra i valori sulla base dei quali poi effettuare una predizione; i modelli generativi invece esplicitano le regolarità soggiacenti la generazione combinata di esempi e valori nei dati e producono potenzialmente altre coppie plausibili.¹² Un'altra rilevante dicotomia nell'apprendimento automatico passa tra metodi supervisionati e non supervisionati, che ora illustriamo.

Nell'apprendimento supervisionato, a ciascuno degli esempi che compongono il campione per l'allenamento della macchina (il *training set*) è associato un valore. In questo caso si dice che la macchina compie un'osservazione completa. Se il valore rientra in una serie di classi discrete si parla di classificazione. Al contrario, se il valore si trova nel continuo dei numeri reali, si parla di regressione. Inoltre l'algoritmo può mappare gli esempi su una matrice dei loro tratti oppure li può mappare direttamente nello spazio di calcolo: nel primo caso si parla di algoritmi *feature-based*, nel secondo di algoritmi *kernel-based*.¹³ Sulla base dell'osservazione la macchina crea dunque un modello, ovvero trova un parametro che, dato insieme agli esempi (o ai loro tratti), fornisca una distribuzione quanto più vicina ai dati originali. Infine, sempre sulla base del modello la macchina predice per ciascun esempio di dati nuovi il valore più probabile. Un algoritmo e i suoi parametri

⁹I dettagli riguardanti questi eventi si possono trovare in rete sui siti di CoNLL 2016, URL:<<http://www.conll.org/>>, SemEval-2016, URL:<<http://alt.qcri.org/semeval2016/>>, ed EvalIta, URL:<<http://www.evalita.it/2016/>>.

¹⁰Un esempio di compito è il *parsing*, ovvero l'analisi sintattica: dato un enunciato, bisogna assegnare ad ogni parola una funzione grammaticale e individuare le dipendenze fra parole. Ne parlano Dell'Orletta e Venturi, in questo volume.

¹¹Perché l'inferenza statistica sia valida, è richiesto che i valori delle variabili associate agli esempi e ai valori siano tra sé indipendenti e identicamente distribuiti nel campione osservato.

¹²Grazie alla legge di Bayes, le probabilità congiunte sono convertibili in probabilità condizionate.

¹³A certe condizioni i modelli creati da algoritmi *feature-based* e *kernel-based* sono equivalenti. Perciò spesso si preferisce l'uno o l'altro per semplici motivi di efficienza.

saranno tanto preferibili quanto più saranno in grado di fare questa predizione in modo accurato e rapido.

Ma come sapere in anticipo la qualità¹⁴ dell'algoritmo su dati nuovi? Fin dall'inizio si tiene separato dal *training set* (e perciò non viene impiegato per l'apprendimento) un insieme di dati di cui si conoscano i valori reali: questi verranno confrontati con i valori predetti dall'algoritmo una volta applicato su questi dati, che per questo motivo sono chiamati *test set*, 'campione di esame'. Un terzo gruppo di dati, a sua volta tenuto distinto, è il *validation set* o *development set*, 'campione di sviluppo' o 'campione di convalida'. Su questo prima della fase di valutazione sul *test set* vengono sperimentate varie configurazioni degli iperparametri, ovvero dei parametri dell'algoritmo che non sono appresi durante l'allenamento: quelli definitivi vengono decisi dallo sperimentatore in base alle prestazioni fornite sul *development set*. Quanto ai parametri da apprendere, trovare quelli globalmente ottimali è computazionalmente arduo, perciò si ricorre a delle euristiche per cercare almeno degli ottimi locali: lo spazio di questa ricerca viene definito da una *loss function*, 'funzione di perdita', che nell'apprendimento supervisionato rappresenta il costo di predire un certo valore dato il valore corretto. Questo spazio di ricerca viene comunemente 'esplorato' da strategie di ottimizzazione come ad esempio la discesa del gradiente.

Tuttavia, gli esempi di allenamento sono talvolta incompleti: a degli esempi (o ai loro tratti) non è associato alcun valore, o l'insieme è mancante di alcuni esempi. Nel caso una variabile interagisca col modello ma non siano mai osservabili i suoi valori, essa è definita variabile aleatoria (*random variable*). In queste eventualità, l'apprendimento è detto non supervisionato. Una delle applicazioni più rilevanti dell'apprendimento supervisionato in informatica e linguistica è quella del *word embedding*, 'immersione delle parole'. Questa tecnica permette di risolvere un problema notevole: come può un calcolatore imparare da stringhe di parole non conoscendone il significato? Bisogna tradurle in un linguaggio ad esso intelleggibile (i numeri), pur preservando le relazioni semantiche che intercorrono fra di esse, e senza poter far ricorso a dei valori annotati: nessuno saprebbe descrivere il significato di una parola senza ripetere la parola stessa.

Eppure il *word embedding* mappa ogni elemento di un vocabolario in un corrispettivo vettore di numeri reali all'interno di uno spazio multi-dimensionale continuo, dove parole semanticamente simili sono punti geometricamente vicini. Parole e numeri sembrano due domini davvero inconciliabili: ma come riesce? L'intuizione è che il significato di una parola è determinato dalle parole che la circondano nelle sue occorrenze nei testi. Questo ambito di studi, la Semantica Distribuzionale, fu reso popolare da Firth (1968) sottolineando l'importanza delle collocazioni, ovvero le coppie di parole che ricorrono assieme più frequentemente di quanto dovrebbe accadere per caso.¹⁵

¹⁴La qualità viene misurata attraverso una serie di metriche: molto diffuse sono ad esempio accuratezza, precisione, richiamo, *F1-score*, ecc.

¹⁵Varie tecniche di *word embedding* sono esposte da Lenci all'interno di questo volume. Un efficiente algoritmo di recente introduzione è Word2Vec (Mikolov et al., 2013).

Il *machine learning*, tuttavia, non è stato sempre l'unico protagonista dell'Intelligenza Artificiale, e a sua volta sta cambiando nelle proprie caratteristiche (Goodfellow et al., 2016). Prima del suo vasto utilizzo, i compiti linguistici erano svolti da algoritmi basati su regole esplicite e sviluppate 'artigianalmente' dagli sperimentatori. Poi, con l'avvento dell'apprendimento automatico, il contributo umano restò comunque presente nella selezione delle *features*, ma a determinare il loro peso e la loro relazione ai valori da predire provvedeva ora la macchina. Un'ulteriore sviluppo è consistito nell'aver fatto sì che la macchina selezionasse da sé anche le *features* dai dati grezzi, apprendendo da che evidenza apprendere. Questo è il *representation learning*, che prende il nome dal fatto che la rappresentazione dell'*input* è 'trasfigurata' in una più adatta all'apprendimento.

L'ultima frontiera di questa branca¹⁶ dell'Intelligenza Artificiale è il *deep learning*. I tratti immediatamente osservabili in ingresso sono manipolati 'distrucendo' tratti sempre più astratti e di alto livello: così, ad esempio, in un'immagine si passerà dalle linee, alle forme, agli oggetti. Il termine *deep*, 'profondo', sta proprio a indicare il numero elevato di queste trasformazioni. Da un punto di vista matematico, esse possono essere concepite come una suddivisione della funzione che mappa l'*input* all'*output* in funzioni più semplici. Da un punto di vista informatico, ogni stadio è lo stato della memoria del calcolatore dopo aver eseguito un insieme di istruzioni parallele.

A questo punto le macchine sono divenute capaci di pensare come esseri umani? Un nome alternativo del *deep learning* è quello di *artificial neural networks*, ovvero 'reti neurali artificiali', perché i lavori ai suoi albori erano ispirati all'apprendimento biologico (McCulloch, Pitts, 1943). Al giorno d'oggi, le neuroscienze continuano a ispirare le ricerche informatiche, pur senza aver mantenuto un ruolo predominante. Infatti, alcune architetture come le reti neurali convoluzionali imitano i neuroni del sistema visivo dei mammiferi (LeCun et al., 1998). D'altro canto, al contrario delle loro architetture, le modalità per allenare le reti neurali sono state sviluppate perlopiù in modo autonomo rispetto ai contributi delle neuroscienze (Hinton et al., 2006).

Il motivo di questo mancato apporto è il fatto che ad oggi non si è ancora raggiunta una conoscenza sufficiente riguardo al linguaggio nel nostro cervello, cioè a come questa funzione superiore emerga dagli elementi basilari, i neuroni e le sinapsi. Approcci *top-down* includono quelli comportamentali (come lo studio dei tempi di lettura¹⁷) e quelli di *neuro-imaging*. In particolare i potenziali evento-correlati hanno permesso di delineare modelli quadri-dimensionali (capaci cioè tanto di localizzare il processo quanto di descriverne le fasi temporali), come ad esempio quello di Friederici (2011). D'altro canto, l'approccio *bottom-up* si è spinto a dar conto delle proprietà computazionali della rete neurale del cervelletto (D'Angelo, Solinas, 2011). Tuttavia, la lacuna tra queste due direzioni di ricerca non è stata

¹⁶Bisogna infatti notare che esistono svariate altre branche, anzitutto l'approccio fondato sulle inferenze logiche e *database*, che fu a lungo un paradigma antagonista all'apprendimento automatico (Lenat, Guha, 1989).

¹⁷Un esempio è lo studio delle frasi complesse da parte di Chesi in questo volume.

ancora colmata.

Del resto, la fortuna del *deep learning* è nata dal giovamento che ne hanno tratto le applicazioni. Nel verso opposto, il *deep learning* e l'apprendimento automatico in genere sono sempre meno utili nell'aiutarci a capire i meccanismi sottostanti il linguaggio. Se da un lato essi forniscono una prospettiva olistica sulle funzioni cognitive superiori, integrando aree di ricerca finora separate (trattamento automatico del linguaggio naturale, visione, movimento e riconoscimento vocale), dall'altro frappongono l'ostacolo della totale intellegibilità dei modelli e dei tratti rilevanti per produrli. Steedman (2011) ha messo in risalto come questa barriera interpretativa si rifletta in una incomunicabilità anche fra le comunità degli studiosi che si occupano rispettivamente di linguistica teorica e computazionale.

Partendo dall'osservazione che a ogni livello del linguaggio i fenomeni hanno una distribuzione logaritmica, Steedman nota come i due tipi di studiosi si concentrino in realtà su porzioni differenti di questa stessa distribuzione. I computazionalisti cercano di catturare il maggior numero di fenomeni, limitandosi dunque ai più frequenti, mentre i teorici si dedicano alla *long tail*, la coda che contiene i fenomeni più rari. In breve, la soluzione proposta dall'autore è da un lato di raffinare i modelli probabilistici grazie alle scoperte teoriche, gettando così un ponte sulla *long tail*; dall'altro lato, di sviluppare formalismi con un'immediata spendibilità computazionale, ovvero limitati a combinazioni sintattiche in un dominio locale e risultanti in una struttura superficiale facilmente mappabile sul piano semantico. In questo modo, la teoria giova non solo alla componente dei dati, ma anche a quella degli algoritmi.

Questi requisiti vengono ad esempio soddisfatti da una famiglia di formalismi sviluppati in seno alla disciplina logica. Nello specifico, la logica non-commutativa mira a proporre analisi sensibili all'ordine degli elementi, cifra caratteristica dell'espressione linguistica. La logica non-commutativa si declina in un filone classico, ed in uno intuizionista. Proprio a questo secondo ramo è legato il *Syntactic Calculus* introdotto da Lambek e presentato da Casadio all'interno di questo volume. Le grammatiche logiche partono dal presupposto che i processi sottostanti la sintassi o la semantica di una lingua possano essere rappresentati come inferenze logiche; in questo modo, l'obiettivo è quello di ottenere un insieme di espressioni ben formate grazie a reti di dimostrazione. Oggi, i calcoli combinatori stanno vivendo una stagione di rinnovato interesse anche in campo applicativo anche grazie alla loro ibridazione con altri campi, come la semantica distribuzionale (Coecke et al., 2010).

La frammentazione ben illustrata da Steedman (2011), frutto delle opposte tensioni fra teoria e applicazioni, ha fatto sì che tra Linguistica (che si occupa di linguaggio, senza far necessariamente ricorso all'informatica) e Intelligenza Artificiale (che si occupa di agenti dai comportamenti razionali, non necessariamente il linguaggio) si è venuta ad addensare una galassia di branche con diversi gradi di affinità e autonomia; abbiamo già menzionato la Linguistica Computazionale e il Trattamento Automatico del Linguaggio Naturale, come pure accanto ad esse le Grammatiche Logiche e la Neurolinguistica. Ma vanno senz'altro aggiunte le *Digital Humanities*, che estendono l'informatica agli studi umanistici come arte e

storia, nonché il Riconoscimento Vocale (*Speech Recognition*), che tratta dati uditivi anziché testuali. In questa breve introduzione non possiamo rendere giustizia alla complessità e alle origini di tutte queste discipline, ma speriamo che riferimenti ad esse emergano almeno in parte dai contributi di vari autori che ci accingiamo a presentare.

I contributi rappresentano un resoconto il più possibile aggiornato di alcune delle prospettive fin qui menzionate, e propongono soluzioni ad alcuni dei loro problemi. Il contenuto dei contributi intreccia varie dicotomie: quella tra teoria (Casadio, Chesi, Graziani) e applicazioni (Lenci, Dell'Orletta), quella tra semantica (Lenci) e sintassi (Casadio, Dell'Orletta, Chesi). Fanno 'eccezione' il lavoro di Jezek perché esattamente a cavallo fra teoria, applicazioni, semantica e sintassi, e il lavoro di Aldini, Fano e Graziani perché costituito da una meta-riflessione. Tuttavia, l'ordine che abbiamo preferito seguire nel darne un breve riassunto qui di seguito e nell'inserirli all'interno del volume è ottenuto risalendo a ritroso lungo questa introduzione: partendo dai formalismi logici e passando attraverso le analisi della complessità delle frasi con esami psicolinguistici, la rappresentazione delle parole nella semantica distribuzionale, l'adattamento dell'analisi sintattica automatica a differenti domini, la creazione di risorse linguistiche, fino alla domanda originale: possono le macchine pensare?

*
* *

Il contributo di Casadio "Calcolo di Lambek, logica lineare e linguistica" analizza la comunicazione linguistica come un processo dinamico. Tale prospettiva permette una considerazione nuova di quelle grammatiche logiche per il linguaggio naturale, quali ad esempio la grammatica di Montague o le grammatiche categoriali, che assegnano un ruolo centrale al rapporto funzione-argomento nella classificazione e nell'organizzazione delle espressioni linguistiche. Questo calcolo logico si propone di ottenere un insieme di espressioni (e frasi) ben formate come risultato di appropriate dimostrazioni, rappresentando i processi sintattici e semantici di una lingua nei termini di inferenze logiche. In questo modo, è messo in luce il ruolo innovativo che può essere svolto dalla Logica Lineare Non-commutativa nello studio delle proprietà formali delle lingue naturali. Tale via è stata aperta da un decisivo articolo di Lambek (1997), nel quale è stata presentata la formulazione di un nuovo calcolo per l'analisi logica delle lingue naturali, da lui definito 'calcolo dei pregruppi' o *Pregroup Grammar*.

Procede su un'altra linea il contributo di Chesi "Il processamento in tempo reale delle frasi complesse". Esso si propone di trovare risposta a questa domanda: come mai risulta difficile capire certe frasi? Come prerequisito alla risposta, occorre comprendere come la complessità di un enunciato possa essere misurata in modo oggettivo. Dopo aver presentato il problema sullo sfondo teorico del 'Programma Minimalista' (Chomsky, 1995), Chesi passa a indagare il contributo in termini di complessità di differenti fattori strutturali, ovvero le frasi relative 'incassate' e le frasi scisse. La complessità di processazione è misurata nei termini di un rallentamento nella lettura, valutata con tecniche di *eye-tracking*. L'autore

rende poi conto dei risultati ottenuti empiricamente attraverso una teoria grammaticale capace di esplicitare le singole operazioni di costruzione della struttura frasale. Questa teoria grammaticale riprende e ribalta la prospettiva minimalista (passa dall'essere *bottom-up* all'essere *top-down*), così da prevedere quando si presenteranno problematicità nella processazione di un enunciato e quali saranno gli elementi a causarle.

Il contributo di Alessandro Lenci "Semantica distribuzionale: un modello computazionale del significato" si concentra su temi quali la natura, i fini e le sfide della semantica distribuzionale. Nella prospettiva della semantica distribuzionale, il significato delle parole è studiato attraverso l'analisi statistica dei contesti linguistici (estratti da corpora di testi) in cui tali parole ricorrono. Essa è basata sulla cosiddetta Ipotesi Distribuzionale, secondo la quale la similarità semantica tra due lessemi correla con la similarità delle distribuzioni di tali lessemi nei contesti linguistici. Un tale punto di vista permette di fornire definizioni nuove a concetti quali quello di 'similarità semantica' dei lessemi (ovvero, la similarità delle distribuzioni contestuali di due o più lessemi) o quello di 'significato' (ovvero, una proprietà che emerge dall'uso di una data parola in relazione al contesto linguistico). Nel suo contributo, Lenci descrive il metodo standard per costruire modelli semantici distribuzionali e come questi possano essere valutati in maniera efficace. Un punto cruciale per la ricerca futura sarà determinare con successo quale tipologia di informazione semantica possa essere catturata sulla base di proprietà contestuali e quali parti del significato rimangano invece al di fuori delle possibilità dei modelli distribuzionali.

L'argomento di cui discute il contributo di Felice Dell'Orletta e Giulia Venturi "ULISSE: una Strategia di Adattamento al Dominio per l'Annotazione Sintattica Automatica" è la *Domain Adaptation* ('Adattamento al Dominio') di algoritmi destinati all'annotazione sintattica attraverso l'addestramento automatico. Le principali difficoltà di tali algoritmi sono per lo più legate ad aspetti del linguaggio naturale quali, per fare alcuni esempi, l'ambiguità, i mutamenti diacronici o appunto la variazione della lingua rispetto al dominio di conoscenza. La questione del *Domain Adaptation* può essere così posta: può un corpus annotato rappresentativo di una certa varietà linguistica essere utilizzato per l'analisi sintattica di un altro corpus rappresentativo di un'altra varietà? Nel loro intervento, Dell'Orletta e Venturi affrontano la *Domain Adaptation* presentando un algoritmo sviluppato dall'Italia NLP Lab 1 chiamato ULISSE (*Unsupervised Linguistically-driven Selection of dEpendency parses*).

L'articolo di Jezek "Alcuni principi linguistici per costruire risorse lessicali" si occupa di analizzare alcuni problemi di tipo semantico o (più generalmente) linguistico che possono emergere quando si progettano risorse lessicali destinate al Trattamento Automatico del Linguaggio Naturale. In particolare, Jezek si focalizza su una risorsa (T-PAS) che mira a creare un inventario di strutture predicatorgo per i verbi italiani tipizzate e derivate da un corpus attraverso un *clustering* manuale di informazioni di tipo distribuzionale. Queste strutture si presentano sotto la forma di *pattern* verbali, dove per ogni *slot* argomentale è esplicitato quale sia il tipo semantico atteso. Una problematicità messa in evidenza per questa

operazione è legata al fatto che, in specifiche posizioni argomentali del verbo, alcune delle parole-*filler* non combaciano perfettamente con il tipo semantico che ci si aspetterebbe per un certo *slot*. Da ultimo, viene rilevato come i membri di ciascun set lessicale mostrino talora una debole coesione semantica. Dopo aver affrontato una a una queste questioni, Jezek valuta se le informazioni legate a ciascuna di esse debbano o meno essere inserite in una risorsa linguistica come T-PAS, e pone l'obiettivo futuro di tecniche per acquisire strutture predicato-argomento e set lessicali in maniera semi-automatica.

Da ultimo, l'articolo di Aldini, Fano e Graziani "Alcune note sui Teoremi di Incompletezza di Gödel e la conoscenza delle macchine" passa in rassegna alcuni dei principali argomenti in favore e contro la tesi dell'irriducibilità della mente umana a una macchina. Centro della discussione degli autori è comprendere quali sia la relazione tra i Teoremi di Gödel e la conoscenza che gli uomini e le macchine possono avere di se stessi. Il celebre lavoro di Turing (1950) ha dato il via a un vero e proprio progetto chiamato 'meccanicismo', il cui obiettivo primario è la costruzione di macchine in grado di riprodurre o se non altro simulare i comportamenti intelligenti dell'uomo. Tuttavia, i teoremi di Gödel sono sempre stati visti come uno strumento efficace per confutare la tesi meccanicista sia in prospettiva 'estensionale' (le procedure e i risultati della mente umana sono meccanizzabili), sia in prospettiva 'intensionale' (la mente umana è una particolare macchina). Questo contributo apporta alla *querelle* importanti chiarificazioni, che vanno dall'analisi del concetto di dimostrabilità intuitiva all'indagine del nesso tra conoscenza/non-conoscenza da parte di una macchina della propria fattività e del proprio codice, dando così rigore alle idealizzazioni in gioco.

*
* *

Questo volume raccoglie gli atti di un convegno tenutosi presso il Collegio Ghislieri a Pavia dal 15 al 17 dicembre 2014, supportato dal Dipartimento di Studi Umanistici, Sezione di Linguistica Teorica e Applicata dell'Università degli Studi di Pavia e dallo IUSS Center for Neurocognition, Epistemology and theoretical Syntax (NEtS) dell'Istituto Universitario di Studi Superiori di Pavia. Il convegno è stato organizzato da un gruppo di alunni provenienti dal Collegio Ghislieri o dall'Almo Collegio Borromeo, e i singoli interventi si sono tenuti in lingua italiana. Perciò, abbiamo deciso di conservarne lo spirito mantenendo la lingua originaria e occupandoci della curatela noi stessi alunni. Le motivazioni di queste scelte sono però più profonde.

La scelta della lingua italiana vuole soprattutto rendere omaggio al contributo degli studiosi italiani alla ricerca internazionale nel campo del linguaggio, della logica e della computazione. Menzionarli tutti sarebbe impossibile, perciò a titolo di esempio ci si può volgere ai momenti cruciali della nascita della disciplina. In particolare, alcuni di questi studiosi hanno detenuto il ruolo di protagonisti negli elementi che abbiamo ricordato come trainanti del progresso tecnologico attuale: gli algoritmi e i dati. Trasferendoci nell'Inghilterra agli albori dell'età vittoriana, Ada Lovelace, figlia del celebre poeta Byron, concepisce l'archetipo dei programmi informatici. E lo fa commentando la minuta descrizione del primo calcolatore

meccanico, l'*Analytical Engine*, inventato da un altro britannico, Charles Babbage. Paradossalmente, a fare da intermediario ai due connazionali è un italiano: Lovelace annota fra le righe un articolo su una conferenza tenuta a Torino da Babbage e scritto dal giovane ingegnere (e futuro primo ministro del Regno d'Italia) Luigi Federico Menabrea (Menabrea, King countess of Lovelace, 1843). Dopo circa un secolo, nasce il primo esempio di corpus digitalizzato (su schede perforate) per merito di un gesuita italiano, Roberto Busa: si tratta dell'*Index Thomisticus*, che raccoglie l'*opera omnia* di Tommaso d'Aquino (Busa, 1974-1980), costruito in un laboratorio a Gallarate con il supporto dell'IBM.

D'altro canto, crediamo che la scelta di lasciare a noi stessi studenti la cura-tela di questi atti, a prezzo di una minore esperienza e competenza, possa essere giustificata dalla possibilità di affrontare tematiche che di primo acchito potrebbero risultare ostiche con l'occhio di chi ricorda ancora gli sforzi e le modalità per iniziare a padroneggiarle. Così gli obiettivi principali di questo volume sono offrire a coloro che si avvicinano alla disciplina qualche punto di riferimento per iniziare a orientarsi e, se possibile, far appassionare ulteriormente chi già vi è familiare approfondendone e scoprendone aspetti nuovi. Infatti, sempre più persone sono chiamate ad apportare un contributo alla comunità scientifica che sta attorno a questa disciplina. Una congiuntura ha voluto che qualche giorno prima di questo convegno si tenesse la prima conferenza nazionale di linguistica computazionale in Italia (CLiC-it) a Pisa, presso l'Istituto di Linguistica Computazionale "Antonio Zampolli" del CNR. Questa iniziativa è stata replicata nel 2015 a Trento presso la Fondazione Bruno Kessler, facendo contestualmente nascere anche un'associazione dedicata, l'Associazione Italiana di Linguistica Computazionale (AILC).

Il titolo che abbiamo dato agli atti, l'*explicit* della *Petite cosmogonie portative* di Raymond Queneau, vuole essere questo invito, ricordando che non si tratta solo di occuparsi di calcolo (*compter*) e linguaggio (*parler*), ma anche di prendersi cura (*soigner*) dello sviluppo scientifico della disciplina e dell'impatto che queste nuove tecnologie avranno su economia e società, tanto pervasivo quanto può esserlo il nostro uso del linguaggio.

Ringraziamenti

L'istituzione che ha reso possibile questo evento è il Collegio Ghislieri: i nostri più sinceri ringraziamenti vanno anzitutto al rettore Andrea Belvedere, che non solo ha concesso i finanziamenti e le strutture necessarie, ma ci ha incentivati e sostenuti durante varie fasi dell'organizzazione. La nostra gratitudine va anche all'Associazione Alunni e al suo presidente Emilio Girino per aver accettato di sostenere i costi della pubblicazione cartacea di questo volume.

Ringraziamo, oltre agli autori che compaiono nel volume, anche gli altri partecipanti al convegno, Andrea Moro dell'Istituto Universitario di Studi Superiori, Egidio D'Angelo dell'Università degli Studi di Pavia e Bernardo Magnini della Fondazione Bruno Kessler a Trento.

Durante l'organizzazione, si è dimostrato insostituibile anche l'apporto di due nostri compagni, Camilla Balbi e Guglielmo Inglese. Inoltre, ringraziamo la professoressa Elisabetta Jezek per la sua attività di supervisione. Non possiamo inoltre non menzionare Amar Hadzahasanovic per un suo prezioso consiglio.

Ringraziamo, infine, i docenti del Dipartimento di Studi Umanistici, Sezione di Linguistica Teorica e Applicata dell'Università degli Studi di Pavia e dell'Istituto Universitario di Studi Superiori, per la disponibilità nella moderazione degli interventi, in particolare le professoresse Silvia Luraghi, Maria Freddi e Irina Prodanof.

Bibliografia

- Baroni M., S. Bernardini, A. Ferraresi, E. Zanchetta (2009). "The WaCky wide web: a collection of very large linguistically processed web-crawled corpora". *Language resources and evaluation*, 43.3, pp. 209–226.
- Busa R. (1974-1980). *Index Thomisticus*. Stuttgart-Bad Cannstatt: Frommann-Holzboog.
- Campbell M., A. J. Hoane, F. Hsu (2002). "Deep blue". *Artificial intelligence*, 134.1, pp. 57–83.
- Chomsky N. (1995). *The minimalist program*. Cambridge (Massachusetts): The MIT Press.
- Coecke B., M. Sadrzadeh, S. Clark (2010). "Mathematical foundations for a compositional distributional model of meaning". *arXiv preprint arXiv:1003.4394*.
- D'Angelo E., S. Solinas (2011). "Realistic modeling of large-scale networks: spatio-temporal dynamics and long-term synaptic plasticity in the cerebellum", *Advances in Computational Intelligence*, pp. 547–553.
- Firth J. R. (1968). *Selected papers of JR Firth, 1952-59*. Bloomberg: Indiana University Press.
- Friederici A. D. (2011). "The brain basis of language processing: from structure to function". *Physiological reviews*, 91.4, pp. 1357–1392.
- Goodfellow I., Y. Bengio, A. Courville (2016). "Deep Learning". Book in preparation for MIT Press. URL: <http://www.deeplearningbook.org>.
- Graffi G. (2010). *Due secoli di pensiero linguistico: dai primi dell'Ottocento a oggi*. Milano: Carocci.
- Hinton G., S. Osindero, Y.-W. Teh (2006). "A fast learning algorithm for deep belief nets". *Neural computation*, 18.7, pp. 1527–1554.
- Kilgarrriff A., G. Grefenstette (2003). "Introduction to the special issue on the web as corpus". *Computational linguistics*, 29.3, pp. 333–347.
- Koyré A. (1957). *From the closed world to the infinite universe*. Baltimore: Johns Hopkins University Press.
- Lambek J. (1997). "Type grammar revisited", in Alain Lecomte, Francois Lamarche, Guy Perrier (cur.), *Logical aspects of computational linguistics*. Berlin: Springer, pp. 1–27.
- LeCun Y., L. Bottou, Y. Bengio, P. Haffner (1998). "Gradient-based learning applied to document recognition". *Proceedings of the IEEE*, 86.11, pp. 2278–2324.

- Lenat D. B., R. V. Guha (1989). *Building large knowledge-based systems; representation and inference in the Cyc project*. Reading: Addison-Wesley Longman.
- McCulloch W. S., W. Pitts (1943). "A logical calculus of the ideas immanent in nervous activity". *The bulletin of mathematical biophysics*, 5.4, pp. 115–133.
- McEnery T., R. Xiao, Y. Tono (2006). *Corpus-based language studies: An advanced resource book*. London: Taylor & Francis.
- Menabrea L. F., A. King countess of Lovelace (1843). *Sketch of the analytical engine invented by Charles Babbage, Esq.* London: Taylor.
- Mikolov T., I. Sutskever, K. Chen, G. S. Corrado, J. Dean (2013). "Distributed representations of words and phrases and their compositionality". *NIPS Proceedings: Advances in neural information processing systems*, pp. 3111–3119.
- Snow C. P. (1959). "Two cultures". *Science*, 130.3373, pp. 419–419.
- Steedman M. (2011). "Romantics and revolutionaries". *Linguistic Issues in Language Technology*, 6.11, [senza numerazione].
- Turing A. (1950). "Computing machinery and intelligence". *Mind*, 59.236, pp. 433–460.
- Van Valin R. D. (2001). *An introduction to syntax*. Cambridge: Cambridge University Press.
- Waldrop M. M. (2016). "The chips are down for Moore's law". *Nature News*, 530.7589, pp. 144–147.

Calcolo di Lambek, logica lineare e linguistica

Claudia Casadio – Università “Gabriele D’Annunzio” Chieti e Pescara
claudia.casadio@unich.it

1. Introduzione

La logica lineare (Girard, 1987) rappresenta un’area di ricerca nuova ed interessante, sia dal punto di vista teorico, sia dal punto di vista delle sue possibili applicazioni, in particolare nell’ambito delle discipline computazionali e linguistiche. Per quanto riguarda la dimensione del linguaggio, si tratta di un territorio per molti aspetti inesplorato, come è stato sottolineato da J. Y. Girard (1995), che si presta allo scambio interdisciplinare e alla produzione di nuovi modelli interpretativi. La versione della logica lineare chiamata logica non commutativa (Abrusci, 1991; Abrusci, Ruet, 1999) nasce dall’esigenza di sviluppare una logica sensibile non solo al numero e all’uso delle risorse, ma anche al loro ordine, proprietà distintiva delle espressioni e dei contesti linguistici; questa logica ammette sia una formulazione classica, che una formulazione intuizionista. È stato lo studio delle proprietà della formulazione intuizionista che ha portato alla scoperta delle sue relazioni con il calcolo per il linguaggio naturale o *Syntactic Calculus*, introdotto nel 1958 dal matematico canadese J. Lambek (1958): il Calcolo di Lambek è stato, ed è tuttora, oggetto di studio nell’ambito della logica, dell’informatica e della linguistica teorica.¹⁸

L’aspetto forse più promettente degli studi recenti in questo campo è rappresentato dalla diretta corrispondenza individuata tra i due connettivi del Calcolo di Lambek, *left slash*: “\” e *right slash*: “/”, che presiedono alla generazione dei tipi sintattici, e le due implicazioni: post-implicazione “—o” e retro-implicazione “o—” della logica lineare non-commutativa, a loro volta formulabili nei termini dei connettivi *par*: “⌘”, *times*: “⊗” della logica lineare moltiplicativa.¹⁹ In particolare, è stato dimostrato che il Calcolo di Lambek corrisponde al frammento moltiplicativo della logica lineare non-commutativa intuizionista senza esponenziali (Abrusci, 1991; Casadio, Lambek, 2001; Abrusci, 2002).

¹⁸Per un quadro degli studi sul Calcolo di Lambek e le sue recenti applicazioni a livello logico e linguistico si richiamano Van Benthem (1988), Van Benthem (1991), Morrill (1994), Morrill (2010) e Moortgat (1997).

¹⁹Precisamente, nel calcolo PNCL (*Pure Non-commutative Classical Linear propositional logic* (Abrusci, 1991)) la post-implicazione lineare “—o” corrisponde all’operazione “\” del Calcolo di Lambek, e la retro-implicazione lineare “o—” corrisponde alla operazione “/”. Sulla base della definizione delle due implicazioni in PNCL, valgono le seguenti equivalenze: $A \multimap B = A^\perp \wp B$ e $B \multimap A = B \wp A^\perp$, dove $()^\perp$ (post-negazione) e $^\perp()$ (retro-negazione) sono le due negazioni ammesse dalla logica lineare non-commutativa.

La logica lineare non-commutativa introduce una prospettiva dinamica in base a cui la comunicazione linguistica è analizzata come un processo che può svilupparsi ed essere afferrato da diversi punti di vista. Questa prospettiva dinamica consente di considerare sotto una nuova luce le grammatiche logiche per il linguaggio naturale, come la grammatica di Montague o le grammatiche categoriali che, richiamandosi al principio fregeano di composizionalità del significato, assegnano un ruolo basilare alla relazione funzione-argomento nella classificazione ed organizzazione delle risorse linguistiche (Van Benthem, 1991; Moortgat, 1997; Morrill, 2010). Le grammatiche dei tipi logici, assumendo l'impostazione di teoria delle dimostrazioni in base a cui i processi sintattici (e semantici) di una lingua possano essere rappresentati nei termini di inferenze logiche, si propongono di ottenere l'insieme delle espressioni ben formate, ed in particolare delle frasi, come risultato di appropriate dimostrazioni. In questo quadro, ha particolare rilievo il metodo delle reti di dimostrazione (*proof nets*), introdotti nell'ambito della logica lineare per fornire una rappresentazione geometrica dei processi logici messi in atto da una dimostrazione. L'impiego dei *proof nets* e la concezione ad essi associata di una geometria dell'interazione hanno rappresentato una scelta metodologica importante che consente di afferrare e comprendere in modo immediato le proprietà logiche delle dimostrazioni, fornendo inoltre uno strumento proficuo per l'analisi formale delle lingue naturali.

2. La distinzione funzione-argomento

Le grammatiche dei tipi logici si ispirano a due fondamentali orientamenti della analisi semantica del linguaggio naturale: la 'teoria delle categorie di significato' formulata all'inizio del secolo da Edmund Husserl e la 'teoria composizionale del significato' di Gottlob Frege.²⁰ Le *Bedeutungskategorien* di Husserl richiamano la concezione aristotelica delle parti del discorso, ma in un nuovo quadro teorico che rimanda alle ricerche sul significato di Bolzano e Frege. Nell'analisi di Husserl il discorso si struttura in parti che esibiscono due fondamentali aspetti relativamente al modo in cui si costituisce il loro significato: alcune parti, come nomi o enunciati, sono significative di per sé stesse, ovvero sono dotate di un significato indipendente. Altre parti, al contrario, non sono significative di per sé stesse, ma hanno bisogno di essere completate per dare luogo ad una espressione dotata di significato; si tratta delle parole sincategorematiche, come *e*, *o*, *è*, la cui funzione linguistica dipende dal significato dell'intero enunciato (o contesto) che contribuiscono a determinare. La connessione dei significati, che è alla base dell'articolazione linguistica, è il risultato della combinazione di parti incomplete con parti adeguate a completarle.

²⁰Il tema delle categorie di significato è analizzato nelle *Ricerche Logiche* di Husserl, in particolare nella Terza e nella Quarta Ricerca (Husserl, 1913); per una analisi della teoria del significato di Husserl si veda Casadio (2011). Per la teoria del significato di Frege si vedano Frege (1892) e Frege (1893/1903) e per la corrispondenza con Husserl Gabriel (1976).

La distinzione tra significati indipendenti o completi e significati non indipendenti o incompleti è molto vicina al modo in cui Frege caratterizza la nozione di funzione come espressione 'insatura' che richiede di essere completata da argomenti appropriati per dare luogo ad un certo valore. Frege riconosce, inoltre, una relazione diretta tra espressioni complete (nomi propri o enunciati) e funzioni, dal momento che una funzione, insatura, può sempre essere ricavata da una espressione completa, mediante l'individuazione di una sua parte come variabile (Frege, 1893/1903). Raccogliendo in un unico quadro teorico questi risultati, come avviene nell'ambito delle ricerche sviluppate dalla Scuola Logica Polacca,²¹ otteniamo alcuni principi che svolgono un ruolo basilare nell'analisi semantica del linguaggio naturale: (i) vi sono parti del discorso che sono incomplete; (ii) tali parti del discorso corrispondono a funzioni che possono essere rappresentate, mediante appropriate notazioni algebriche o logiche, come categorie funtoriali; (iii) l'articolazione linguistica è il risultato di processi di saturazione delle categorie delle funzioni con gli argomenti appropriati, dando luogo a valori che sono entità complete, in particolare le proposizioni.

Poiché una funzione è identificata dalle categorie dei suoi argomenti e dalla categoria del valore che risulta per quegli argomenti, la notazione algebrica introdotta da Kazimierz Ajdukiewicz²² si è dimostrata particolarmente adeguata nell'ambito dei modelli orientati ad utilizzare funzioni per analizzare formalmente il linguaggio naturale. Tale notazione rappresenta i simboli di funzione (i funtori) mediante indici frazionari. L'indice frazionario associato ad un funtore prende al denominatore gli indici degli argomenti della funzione corrispondente e al numeratore l'indice del valore assunto dalla funzione. Si ottiene così la notazione standard di una grammatica categoriale in cui A/B è il simbolo di una funzione che per un argomento di tipo B assume un valore di tipo A :

$$\frac{\text{funtore} : A/B \quad \text{argomento} : B}{A}$$

corrispondente all'operazione di applicazione funzionale che in una grammatica a costituenti immediati è rappresentata dalla regola di riscrittura (regola di cancellazione o eliminazione):

$$A/B \quad B \rightarrow A$$

Ad esempio, un verbo transitivo come *dare*, è analizzato come un'espressione incompleta corrispondente ad una funzione che richiede due argomenti: l'oggetto diretto: SN_1 e l'oggetto indiretto: SN_2 . L'applicazione della categoria che corrisponde a tale funzione, alle categorie degli argomenti, genera come valore la categoria dell'espressione risultante, il predicato verbale:

²¹L'approfondimento e la formalizzazione della teoria del significato di Husserl ha caratterizzato la Scuola Logica Polacca in un ambito di ricerche dedicate ai fondamenti della matematica ed all'analisi formale del linguaggio. Si veda in proposito McCall (1967).

²²Si veda Ajdukiewicz (1935), il noto saggio dedicato alla connessione sintattica, comparso sul primo numero della rivista *Studia Philosophica*.

dare: $(SV/SN_2)/SN_1$, un libro: SN_1 , a Maria: $SN_2 \rightarrow$
 dare un libro a Maria: SV .

In base alla distinzione tra espressioni complete ed incomplete, Frege divide i segni linguistici in due categorie: (i) la categoria dei nomi propri, a cui appartengono le espressioni complete (che non hanno posti per argomento); (ii) la categoria dei nomi di funzione, a cui appartengono le espressioni incomplete (che hanno uno o più posti per argomento). La prima classe include i nomi di oggetti (nomi propri, descrizioni definite), gli enunciati (come nomi dei particolari oggetti che sono i valori di verità) e, in generale, i nomi dei corsi di valori delle funzioni.

La nozione di funzione può essere estesa a qualsiasi segno linguistico in cui un costituente è lasciato indeterminato: una funzione è ottenuta rimuovendo una o più occorrenze di un'espressione completa (un nome proprio) da un'espressione completa. Si hanno così gli ingredienti essenziali per definire un insieme di tipi logici per il linguaggio naturale a partire da una base finita costituita dai due tipi N (nomi) ed S (enunciati), come nel calcolo formulato da Geach (1972), in cui, rimuovendo una espressione di tipo $:N$ da un enunciato di tipo $:S$, si ottiene il tipo $:S/N$ di una funzione di primo livello corrispondente ad un predicato:

Socrate $:N$ sogna $:S/N \rightarrow$ Socrate sogna $:S$

mentre, rimuovendo una espressione di tipo $:S/N$ da un enunciato, si ottiene una funzione di secondo livello di tipo $:S/(S/N)$, come nel caso dei sintagmi quantificazionali:

Ogni uomo $:S/(S/N)$ sogna $:S/N \rightarrow$ Ogni uomo sogna $:S$

Queste sono le assunzioni su cui sono state fondate le grammatiche categoriali di Ajdukiewicz (Ajdukiewicz, 1935), Bar-Hillel (Bar-Hillel, 1953; Bar-Hillel et al., 1960) e il calcolo sintattico di Lambek (Lambek, 1958).

3. Calcolo di Lambek

Lambek presenta il suo calcolo sintattico in due noti articoli: "The mathematics of sentence structure" del 1958 e "On the calculus of syntactic types" del 1961. Il calcolo fornisce una procedura effettiva per determinare le frasi di una lingua (naturale o formale) e lo studio delle sue proprietà logiche e matematiche ha dimostrato che può essere applicato in modo proficuo a diversi ambiti di indagine: dalla teoria linguistica, alla analisi computazionale alla traduzione automatica.

Il Calcolo di Lambek (L) è un sistema deduttivo che definisce un insieme di tipi logici o sintattici ed è chiuso rispetto alle tre operazioni binarie $A \bullet B$ (prodotto), C/B (C over B), $A \backslash C$ (A under C), in cui M è un sistema moltiplicativo (monoide) con identità I (per ogni $A, B, C \in M$):

- a. $A \bullet B = \{x.y \in M \mid x \in A \wedge y \in B\}$
- b. $C/B = \{x \in M \mid \forall y \in B, x.y \in C\}$
- c. $A \backslash C = \{y \in M \mid \forall x \in A, x.y \in C\}$

I teoremi elencati in (3) sono dimostrabili in L sulla base degli assiomi in (2) e delle regole di inferenza in (1) (per ogni $A, B, C \in M$):

- (1) a. $A \bullet B \rightarrow C \leftrightarrow A \rightarrow C/B$ (*implicazione o divisione a destra*)
 b. $A \bullet B \rightarrow C \leftrightarrow B \rightarrow A \backslash C$ (*implicazione o divisione a sinistra*)
 c. $A \rightarrow B, B \rightarrow C \Rightarrow A \rightarrow C$ (*transitività*)
- (2) a. $X \rightarrow X$ (*identità*)
 b. $(A \bullet B) \bullet C \leftrightarrow A \bullet (B \bullet C)$ (*associativa*)
- (3) a. $(A/B) \bullet B \rightarrow A$ (*cancellazione a destra*)
 b. $A \bullet (A \backslash B) \rightarrow B$ (*cancellazione a sinistra*)
 c. $B \rightarrow (A/B) \backslash A$ (*type raising*)
 d. $B \rightarrow A/(B \backslash A)$ (*type raising*)
 e. $(A \backslash B)/C \leftrightarrow A \backslash (B/C)$ (*associatività*)
 f. $(A/B)/C \leftrightarrow A/(C \bullet B)$ (*def. $\bullet$*)
 g. $(A/B) \bullet (B/C) \rightarrow (A/C)$ (*composizione*)
 h. $A/B \rightarrow (A/C)/(B/C)$ (*Geach law*)

Il sistema moltiplicativo M non è necessariamente associativo: se si ammette la proprietà associativa (2b) come in Lambek (1958), M è un semigruppato che consente di formare liberamente costituenti a destra o a sinistra (*unbracketed strings*); se non viene ammessa la proprietà associativa si ottiene il sistema più restrittivo formulato in Lambek (1961) e studiato in numerose applicazioni linguistiche (si vedano in proposito Moortgat, 1997; Morrill, 2010; Morrill, 1994).

Il calcolo sintattico L (generato dal sistema moltiplicativo M) può essere pensato come un sistema universale di regole che si applica ad un dizionario specificato per una certa lingua (naturale o formale)²³ definendo appropriati insiemi di parole (o stringhe di parole) chiamati tipi sintattici, in cui S è il tipo degli enunciati dichiarativi o frasi della lingua, N è il tipo dei nomi come *Mary, John*. Partendo da questi due tipi di base, applicando le regole del calcolo L , possono essere derivati i tipi dei vari costituenti di un enunciato (elencati nel Lessico B), come NP (sintagma nominale), VP (sintagma verbale), PP (sintagma preposizionale), ecc., nei cui termini si può dimostrare se una certa sequenza di tipi sintattici, assegnata ad una stringa di parole della lingua, è una frase grammaticale della lingua.

²³Facciamo qui riferimento alla lingua inglese, daremo quindi alcuni esempi relativi alla lingua italiana.

- a. $B = [N, NP, S]$
- b. if $\text{Mary} \rightarrow NP$ and $\text{Mary works} \rightarrow S$,
then $\text{works} \rightarrow NP \backslash S$;
- c. if $\text{apple} \rightarrow N$ and $\text{an apple} \rightarrow NP$,
then $\text{an} \rightarrow NP / N$;
- d. if $\text{Mary ate an apple} \rightarrow S$ and $\text{Mary} \rightarrow NP$,
then $\text{ate an apple} \rightarrow NP \backslash S$;
if $\text{an apple} \rightarrow NP$, then $\text{ate} \rightarrow (NP \backslash S) / NP$.

Di seguito presentiamo la derivazione di alcune frasi della lingua italiana, usando la procedura inferenziale in base a cui ogni applicazione di una regola del calcolo L è indicata dal simbolo scritto a fianco della linea orizzontale.

CALCOLO SINTATTICO L: *Applicazione a destra / e a sinistra *

- a. *Maria mette una tovaglia nuova sul tavolo.*
- b. **mette (una tovaglia nuova) (sul tavolo)**

$$\frac{\frac{(VP/PP)/NP \quad NP}{VP/PP} / \quad PP}{VP} /$$
- c. **Maria (mette una tovaglia nuova sul tavolo)**

$$\frac{NP \quad VP = NP \backslash S}{S} \backslash$$
- a. *Gianni corre velocemente sotto la pioggia.*
- b. **corre velocemente (sotto la pioggia)**

$$\frac{VP \quad VP \backslash VP}{VP} \backslash \quad \frac{VP \backslash VP}{VP} \backslash$$
- c. **Gianni (corre velocemente sotto la pioggia)**

$$\frac{NP \quad VP = NP \backslash S}{S} \backslash$$

CALCOLO SINTATTICO L: *Type Raising o Expansion: EXP*

- a. *Molte ragazze partono domani.*
- b. **(Molte ragazze) (partono domani)**

$$\frac{\frac{NP}{S/(NP \backslash S)} \quad EXP \quad NP \backslash S}{S} /$$

4. Logica Lineare Non-commutativa

Il Calcolo di Lambek L ha una traduzione naturale nella logica lineare non-commutativa (*Non-commutative Linear Logic*: NLL), un sistema della logica lineare che si ottiene escludendo le regole strutturali di indebolimento (*weakening*): $A \vdash A, A$; contrazione (*contraction*): $A, A \vdash A$; scambio (*exchange*): $A, B \vdash B, A$ (vedi Abrusci, 1991; Abrusci, 2002; Abrusci, 2014). Precisamente, il Calcolo di Lambek corrisponde al frammento moltiplicativo MNLL e può ricevere sia una formulazione intuizionista che una formulazione classica, come è dimostrato in Abrusci (1991). Il sistema completo NLL ha il seguente alfabeto (Abrusci, 1991; Abrusci, Ruet, 1999):

Variabili Proposizionali :	$P, Q, \dots$
Costanti Logiche (<i>bottom, one, true, zero</i>):	$\perp, 1, T, 0$
Connettivi binari :	moltiplicativi (<i>times, par</i>): $\otimes, \wp$ additivi (<i>with, plus</i>): $\&, \oplus$
Connettivi unari (<i>of course, why not</i>):	$!, ?$
Simbolo di sequente :	$\vdash$
Simboli ausiliari:	$(,), \dots$

NLL e il suo frammento moltiplicativo MNLL sono sistemi classici perché i sequenti possono avere più di una conclusione e valgono leggi logiche classiche come le leggi *De Morgan* e una particolare formulazione delle leggi della *Negazione*. Si ha la seguente definizione di sequente nel calcolo ad una parte (*one-sided sequent calculus*):

$\vdash \Gamma$ è un sequente sse Γ è una sequenza finita di formule di NLL (MNLL) .

Una proprietà distintiva della logica lineare non-commutativa è la presenza di due negazioni:

- *post*-negazione di $A = A^\perp$
- *retro*-negazione di $A = {}^\perp A$

Esse sono introdotte mediante una definizione metalinguistica (Abrusci, 1991; Abrusci, 2002). Su questa base, l'insieme delle formule di NLL è definito induttivamente (per la corrispondenza con il Calcolo di Lambek, limitiamo la definizione al frammento moltiplicativo MNLL):

- ciascuna variabile proposizionale è una formula di MNLL;
- se A è una formula di MNLL, $A^\perp, {}^\perp A$ sono formule di MNLL;
- se A e B sono formule di MNLL, $A \otimes B$ è una formula di MNLL;
- se A e B sono formule di MNLL, $A \wp B$ è una formula di MNLL;
- nient'altro è una formula di MNLL.

Formule che contengono le due implicazioni “ $\multimap$ ” and “ $\multimap$ ” (corrispondenti alle operazioni “ $\backslash$ ” , “ $/$ ” del Calcolo di Lambek) possono essere tradotte in formule di MNLL contenenti il connettivo $\wp$ e una delle due negazioni. Se A e B sono formule di MNLL,

- $A \multimap B$ è tradotta dalla formula $A^\perp \wp B$ di MNLL ;
- $B \multimap A$ è tradotta dalla formula $B \wp A^\perp$ di MNLL .

Queste traduzioni richiamano la definizione del connettivo di *implicazione* della logica proposizionale nei termini della *disgiunzione* e della *negazione*, e svolgono un ruolo fondamentale nella costruzione delle dimostrazioni assicurando la transizione dai tipi logici implicazionali del Calcolo di Lambek a tipi classici definiti nei termini dei connettivi duali *par*: $\wp$ e *times*: $\otimes$, e delle due negazioni lineari. Una proprietà caratteristica dei sistemi NLL, MNLL è l’architettura simmetrica in base a cui le leggi classiche di De Morgan valgono per insiemi duali di connettivi ed in particolare per i connettivi $\wp$ e $\otimes$, che qui riportiamo :

Leggi De Morgan

$$\begin{aligned} (A \wp B)^\perp &= B^\perp \otimes A^\perp , \\ {}^\perp(A \wp B) &= {}^\perp B \otimes {}^\perp A , \\ (A \otimes B)^\perp &= B^\perp \wp A^\perp , \\ {}^\perp(A \otimes B) &= {}^\perp B \wp {}^\perp A . \end{aligned}$$

La distribuzione della negazione sulle due sottoformule ha l’effetto di cambiare il connettivo nel suo duale invertendo l’ordine delle formule:

$$\begin{aligned} (A^\perp \wp B)^\perp &= B^\perp \otimes A^{\perp\perp} & (B \wp {}^\perp A)^\perp &= A \otimes B^\perp , \\ {}^\perp(A^\perp \wp B) &= {}^\perp B \otimes A & {}^\perp(B \wp {}^\perp A) &= {}^\perp\perp A \otimes {}^\perp B , \\ (A^\perp \otimes B)^\perp &= B^\perp \wp A^{\perp\perp} & (B \otimes {}^\perp A)^\perp &= A \wp B^\perp , \\ {}^\perp(A^\perp \otimes B) &= {}^\perp B \wp A & {}^\perp(B \otimes {}^\perp A) &= {}^\perp\perp A \wp {}^\perp B . \end{aligned}$$

La negazione lineare è una operazione involutiva: nel calcolo dei sequenti a due parti (*two-sided sequent calculus*), la retro (post)-negazione di una formula di MNLL che occorre in un sequente, ha l’effetto di spostare tale formula sul lato opposto del simbolo di sequente. Se la formula è una premessa, la sua retro (post)-negazione sarà una conclusione, e viceversa; questo scambio di ruoli, che esprime la dualità premesse-conclusioni, rispetta l’ordine mostrato qui sotto:

$$\begin{aligned} \frac{A, \Gamma \vdash \Delta}{\Gamma \vdash A^\perp, \Delta} ()^\perp & \quad \frac{\Gamma \vdash A, \Delta}{{}^\perp A, \Gamma \vdash \Delta} {}^\perp() \\ \frac{\Gamma, B \vdash \Delta}{\Gamma \vdash \Delta, {}^\perp B} {}^\perp() & \quad \frac{\Gamma \vdash \Delta, B}{\Gamma, B^\perp \vdash \Delta} ()^\perp \end{aligned}$$

Dalla definizione delle due negazioni in MNLL segue che la legge logica della doppia negazione vale solo per le negazioni simmetriche: ${}^\perp(A^\perp)$, $({}^\perp A)^\perp$, mentre le formule negate due volte sullo stesso lato non possono essere ridotte a formule positive:

Leggi della Doppia Negazione

$$\perp(A^\perp) = (\perp A)^\perp = A ,$$

$$\perp\perp(A) \neq (A)^{\perp\perp} \neq A .$$

Il seguente teorema è dimostrabile nel calcolo dei sequenti a una parte di MNLL:

Se il sequente $\vdash \Gamma, A$ è dimostrabile in MNLL, anche i sequenti $\vdash A^{\perp\perp}, \Gamma$ e $\vdash \Gamma, \perp\perp A$ sono dimostrabili in MNLL .

In base a questo teorema nel sistema MNLL si danno due leggi cicliche della negazione (*cyclic negation laws*) che consentono a una formula negata due volte sullo stesso lato di essere spostata, indietro: $()^{\perp\perp}$ o in avanti: $\perp\perp()$, in un sequente:

Leggi Cicliche della Negazione

$$\frac{\vdash \Gamma, A}{\vdash A^{\perp\perp}, \Gamma} ()^{\perp\perp} \quad \frac{\vdash A, \Gamma}{\vdash \Gamma, \perp\perp A} \perp\perp()$$

Il calcolo dei sequenti ad una parte per MNLL include una regola *Assioma*, due regole *Cut*, due regole per la congiunzione moltiplicativa $\otimes$, e una regola per la disgiunzione moltiplicativa: $\wp$.²⁴

ASSIOMA

$$\frac{}{\vdash A^\perp, A} \text{id}$$

REGOLE CUT

$$\frac{\vdash \Gamma, A, \Gamma' \vdash A^\perp, \Delta}{\vdash \Gamma, \Delta, \Gamma'} \text{cut1} \quad \frac{\vdash \Gamma, A \vdash \Delta, A^\perp, \Delta'}{\vdash \Delta, \Gamma, \Delta'} \text{cut2}$$

REGOLE $\otimes$

$$\frac{\vdash \Gamma, A, \Gamma' \vdash B, \Delta}{\vdash \Gamma, A \otimes B, \Delta, \Gamma'} \otimes 1 \quad \frac{\vdash \Gamma, A \vdash \Delta, B, \Delta'}{\vdash \Delta, \Gamma, A \otimes B, \Delta'} \otimes 2$$

REGOLA $\wp$

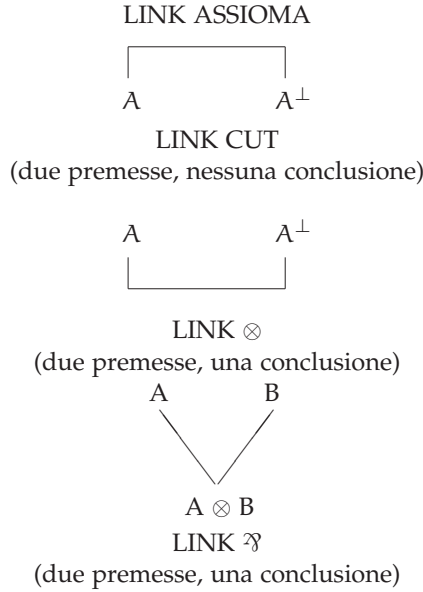
$$\frac{\vdash \Gamma, A, B}{\vdash \Gamma, A \wp B} \wp$$

²⁴Qui sono presentate le regole per le formule $A^\perp$, le regole corrispondenti per le formule $\perp A$ sono analoghe. Le due regole per *Cut* e $\otimes$ servono per rendere conto del contesto.

Un sequente $\vdash \Gamma$ è dimostrabile in MNLL se e solo se c'è una dimostrazione di $\vdash \Gamma$ in MNLL. La nozione di *proof net* svolge un ruolo fondamentale in logica lineare (Girard, 1987): i *proof nets* sono grafi associati a dimostrazioni.²⁵

Dal momento che una dimostrazione di un sequente corrisponde ad una deduzione da premesse $\Gamma = A_1, \dots, A_n$ a conclusioni $\Delta = B_1, \dots, B_n$, essa può essere rappresentata mediante un grafo orientato π di occorrenze di formule (i *vertici* di π), in cui ciascun vertice in π è la *premessa* di al massimo un link in π e la *conclusione* di esattamente un link in π . Girard (1987) definisce una classe di grafi con queste proprietà come la classe delle *proof structures* della logica lineare moltiplicativa (MLL) e mostra che ad ogni dimostrazione D di un sequente $\vdash \Gamma$ in MLL può essere associata una *proof structure* (commutativa) $\Phi(D)$ le cui conclusioni sono esattamente le formule in Γ . I *proof nets* per MLL sono quindi individuati come un particolare sottoinsieme delle *proof structures* di MLL, quelle che non presentano *short trips*.²⁶

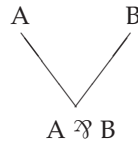
Nella logica lineare non-commutativa MNLL, che risulta dalla esclusione della regola strutturale dello scambio (*exchange*), i *proof nets* non-commutativi hanno una buona proprietà geometrica: sono grafi planari, in cui i links non si intersecano. I *proof nets* in MNLL sono costruiti per mezzo dei seguenti *links*²⁷:



²⁵*Proof nets* per il Calcolo di Lambek sono analizzati in Abrusci (1991) e Moortgat (1997). Le proprietà dei *proof nets* non-commutativi sono studiate in particolare in Abrusci, Ruet (1999) e Abrusci (2014).

²⁶Per la spiegazione di queste nozioni e le relative dimostrazioni si rimanda a Girard (1987).

²⁷Di nuovo i links sono dati per formule $A^\perp$; i corrispondenti links con formule $^\perp A$ seguono per dualità e sono analoghi.



Come si vede dalle regole per i connettivi di MNLL, il *link- $\otimes$* connette due formule che appartengono a contesti separati, mentre il *link- $\bowtie$* connette due formule che appartengono allo stesso contesto.

5. Calcolo di Lambek e logica lineare

I tipi sintattici del Calcolo di Lambek L possono essere tradotti in MNLL mediante una procedura che sostituisce ciascuna formula con $/$ o $\backslash$ in una formula con una appropriata sequenza di $\bowtie$ e negazioni $()^\perp$, ${}^\perp()$, iniziando dalle formule semplici $A\backslash B$ e A/B :

$$\begin{aligned} A\backslash B &\text{ è tradotta in } A^\perp \bowtie B, \\ B/A &\text{ è tradotta in } B \bowtie {}^\perp A. \end{aligned}$$

Su questa base si possono ottenere tipi più complessi:

Verbo Intransitivo:	IV	$NP\backslash S$	$NP^\perp \bowtie S$
Verbo Transitivo:	TV	$(NP\backslash S)/NP$	$NP^\perp \bowtie S \bowtie {}^\perp NP$
Verbo Ditransitivo:	DTV	$((NP\backslash S)/PP)/NP$	$NP^\perp \bowtie S \bowtie {}^\perp PP \bowtie {}^\perp NP$
Sintagma Quantificazionale:	Q_1	$S/(NP\backslash S)$	$S \bowtie ({}^\perp S \otimes NP)$
Sintagma Quantificazionale:	Q_2	$(S/NP)\backslash S$	$(NP \otimes S^\perp) \bowtie S$
Avverbio Intransitivo:	IVA	$(NP\backslash S)\backslash(NP\backslash S)$	$(S^\perp \otimes NP^\perp) \bowtie NP^\perp \bowtie S$

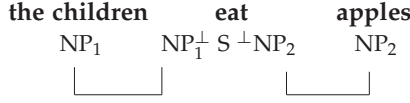
In L, il tipo IV dei verbi intransitivi e il tipo TV dei verbi transitivi sono introdotti nel calcolo dei sequenti a due parti (a sinistra) o a una parte (a destra), mediante le due implicazioni $/$ e $\backslash$:

$$\begin{array}{lll} \text{IV: } run & NP\backslash S \vdash NP\backslash S & run \vdash NP\backslash S \\ \text{TV: } eats & (NP_1\backslash S)/NP_2 \vdash (NP_1\backslash S)/NP_2 & eats \vdash (NP_1\backslash S)/NP_2 \end{array}$$

In MNLL questi tipi sono introdotti come sequenti del calcolo ad una parte, in cui la virgola corrisponde al “ $\bowtie$ ”:

$$\begin{aligned} \text{IV: } &\vdash run^\perp, NP^\perp, S \\ \text{TV: } &\vdash eat^\perp, NP^\perp, S, {}^\perp NP \end{aligned}$$

La definizione dei tipi del Calcolo di Lambek nei termini di $\bowtie$ e delle due negazioni $()^\perp$, ${}^\perp()$, ha l’effetto di rappresentare i posti per argomento di una categoria funzionale (come IV o TV) secondo un ordine lineare, che permette di collegarli simultaneamente con gli argomenti appropriati; le modalità di tali connessioni possono essere adeguatamente rappresentate mediante i *proof nets* non-commutativi di MNLL, come mostrato in questo esempio:



Questo grafo planare ha la buona proprietà di rappresentare simultaneamente le due derivazioni, diverse ma equivalenti, in base a cui la stringa corrispondente al verbo transitivo TV: *eats* viene combinata con i suoi due argomenti: il soggetto NP_1 e il complemento oggetto NP_2 :

$$\begin{array}{c}
 \frac{\frac{\vdash \text{the} - \text{children}^\perp, NP \quad \vdash \text{eat}^\perp, NP^\perp, S, \perp NP}{\vdash \text{eat}^\perp, \text{the} - \text{children}^\perp, S, \perp NP} \text{ cut} \quad \vdash \text{apples}^\perp, NP}{\vdash \text{apples}^\perp, \text{eat}^\perp, \text{the} - \text{children}^\perp, S} \text{ cut} \\
 \hline
 \text{the} - \text{children}, \text{eat}, \text{apples} \vdash S
 \end{array}$$

$$\begin{array}{c}
 \frac{\vdash \text{the} - \text{children}^\perp, NP \quad \frac{\frac{\vdash \text{eat}^\perp, NP^\perp, S, \perp NP \quad \vdash \text{apples}^\perp, NP}{\vdash \text{apples}^\perp, \text{eat}^\perp, NP^\perp, S} \text{ cut}}{\vdash \text{apples}^\perp, \text{eat}^\perp, \text{the} - \text{children}^\perp, S} \text{ cut} \\
 \hline
 \text{the} - \text{children}, \text{eat}, \text{apples} \vdash S
 \end{array}$$

Figura 1. MNLL: *Derivazioni di TV*

Le due derivazioni equivalenti sono ottenute ciascuna mediante due *cuts* operati in ordine inverso, prima il verbo transitivo è applicato al soggetto e poi all'oggetto, e viceversa. Poiché in MNLL vale il teorema di *cut-elimination*, come provato in Abrusci (1991), derivazioni *cut-free* possono essere date per ogni espressione linguistica che si ottiene con *cuts* semplici, o *cuts* complessi che possono essere così ridotti. In questo senso il calcolo dei sequenti ad una parte per MNLL rappresenta una procedura efficiente per ottenere dimostrazioni mediante *cuts*.

6. Alcuni esempi linguistici

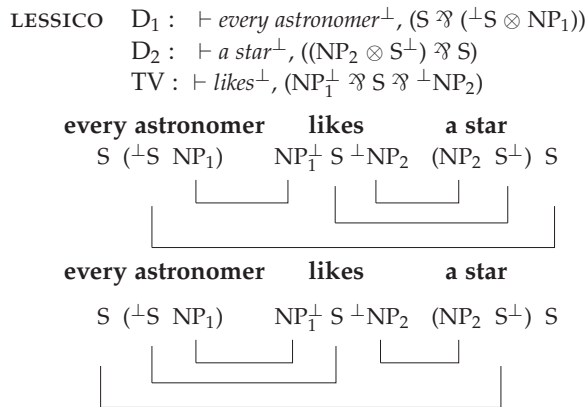
Consideriamo ora alcune applicazioni linguistiche che risultano dalla traduzione del Calcolo di Lambek nel frammento moltiplicativo della logica lineare non-commutativa MNLL. In primo luogo, prendiamo in esame alcuni esempi relativi al campo d'azione (*scope*) dei quantificatori (Montague, 1976; Benthem, Ter Meulen, 1996), mostrando come si possa rappresentarlo mediante *proof net* non-commutativi. Quindi viene proposta l'analisi, sempre avvalendosi della geometria dei *proof net*, di alcuni casi di contesti subordinati, conosciuti nella letteratura come dipendenze illimitate (*unbounded dependencies*) (Van Benthem, 1991; Moortgat, 1997; Morrill, 2010; Morrill, 1994).

6.1. Quantificatori

Poiché i *proof nets* in MNLL hanno la buona proprietà di rappresentare simultaneamente tutte le derivazioni equivalenti di una data stringa di parole, un *proof*

net può essere associato all'insieme di tutte le dimostrazioni equivalenti di una data stringa di parole (espressione o frase). Le ambiguità reali corrisponderanno quindi a *proof net* differenti, come nel caso delle ambiguità di campo d'azione esemplificate da enunciati con quantificatori come '*Every astronomer likes a star*' ('Ogni astronomo ama una stella'), che ricevono due rappresentazioni diverse corrispondenti al *narrow scope* (campo d'azione ridotto) vs. *wide scope* (campo d'azione ampio) di ciascun quantificatore (vedi Montague, 1976).

Nel seguente esempio, i sintagmi quantificazionali ricevono il tipo D_1 e D_2 per le rispettive occorrenze pre-verbali e post-verbali e i due *proof net* sono grafi planari in forma normale, con *links* paralleli:



La lettura *narrow scope* del quantificatore subordinato D_2 , associata alla lettura *wide scope* del quantificatore principale D_1 : '*every astronomer is such that he likes a star*' ('ogni astronomo è tale che egli ama una stella'), è analizzata in MNLL come il processo che inizia nel contesto del predicato verbale e termina nel quantificatore principale, in cui occorre il tipo S non connesso da alcun link ad altri tipi.

D'altro lato, la lettura *wide scope* del quantificatore subordinato D_2 , associata alla lettura *narrow scope* del quantificatore principale D_1 , *there is a star such that every astronomer likes it* (c'è una stella tale che ogni astronomo la ama), è analizzata in MNLL come il processo che va dal soggetto principale, il quantificatore D_1 , al complemento verbale con il quantificatore D_2 , in cui occorre il tipo S non connesso ad altri tipi.

Ai due *proof net* corrispondono le dimostrazioni in Figura 2., nel calcolo dei sequenti a una parte per MNLL. Vediamo dunque che i *proof net* della logica non-commutativa rappresentano un metodo efficiente per inferire le proprietà di campo d'azione dei quantificatori, a partire dall'informazione fornita dalle assegnazioni di tipi sintattici agli elementi del lessico e dalle operazioni associate ai connettivi di MNLL (Casadio, Lambek, 2002).

6.2. Dipendenze frasali

Le dipendenze illimitate sono configurazioni di parole che presentano una relazione a distanza tra due costituenti: un argomento che occorre in genere all'inizio

$$\begin{array}{c}
\frac{\frac{\frac{\vdash \text{likes}^\perp, \text{NP}_1^\perp, S, \perp \text{NP}_2 \quad \vdash \text{astar}^\perp, \text{NP}_2 \otimes S^\perp, S}{\vdash \text{astar}^\perp, \text{likes}^\perp, \text{NP}_1^\perp, S} \text{ cut}}{\vdash \text{astar}^\perp, \text{likes}^\perp, \text{NP}_1^\perp \wp S} \text{ } \\
\frac{\vdash \text{every astronomer}^\perp, S, \perp S \otimes \text{NP}_1 \quad \vdash \text{astar}^\perp, \text{likes}^\perp, \text{NP}_1^\perp \wp S}{\vdash \text{astar}^\perp, \text{likes}^\perp, \text{every astronomer}^\perp, S} \text{ cut} \\
\frac{\vdash \text{astar}^\perp, \text{likes}^\perp, \text{every astronomer}^\perp, S}{\text{every astronomer, likes, astar} \vdash S} \perp(i)
\end{array}$$

$$\begin{array}{c}
\frac{\vdash \text{every astronomer}^\perp, S, \perp S \otimes \text{NP}_1 \quad \vdash \text{likes}^\perp, \text{NP}_1^\perp, S, \perp \text{NP}_2}{\vdash \text{likes}^\perp, \text{every astronomer}^\perp, S, \perp \text{NP}_2} \text{ cut} \\
\frac{\vdash \text{likes}^\perp, \text{every astronomer}^\perp, S, \perp \text{NP}_2 \quad \vdash \text{astar}^\perp, \text{NP}_2 \otimes S^\perp}{\vdash \text{astar}^\perp, \text{likes}^\perp, \text{every astronomer}^\perp, S} \text{ } \\
\frac{\vdash \text{astar}^\perp, \text{likes}^\perp, \text{every astronomer}^\perp, S}{\text{every astronomer, likes, astar} \vdash S} \perp(i)
\end{array}$$

Figura 2. MNLL: *Quantificatori*

della frase (*filler*) e una posizione originale di base, detta *gap* o posto per argomento. Dato che questi due elementi non sono adiacenti e possono essere messi in comunicazione attraverso un contesto relativamente ampio, hanno rappresentato un problema per le teorie linguistiche e in particolare per le grammatiche dei tipi logici (Ades, Steedman, 1982).

Le proprietà della logica lineare non-commutativa e la sua architettura simmetrica consentono di delineare un nuovo modo per analizzare le dipendenze illimitate, facendo ricorso alle proprietà della dualità ed al cambio di prospettiva che ne risulta. I seguenti sono alcuni esempi di interrogative con dipendenze frasali che prenderemo in esame:²⁸

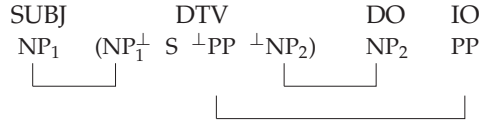
- a) *Elisa wrote a letter to John* (Elisa ha scritto una lettera a Giovanni)
- b) *Whom did Elisa write a letter to?* (A chi ha scritto Elisa una lettera?)
- c) *To whom did Elisa write a letter?* (A chi ha scritto Elisa una lettera?)
- d) *What did Elisa write?* (Che cosa ha scritto Elisa?)

Nell'analisi che segue, assegnamo il tipo S alle frasi che sono enunciati, il tipo Q alle frasi che sono domande, il tipo NP ai sintagmi nominali, il tipo DTV ai verbi transitivi con un oggetto diretto e un oggetto indiretto, il tipo AUX all'ausiliare; inoltre, consideriamo solo i pronomi interrogativi al caso accusativo *whom*, *what*, che rimandano ad una posizione di complemento diretto o indiretto nella frase:

LESSICO	NP :	$\vdash \text{Elisa}^\perp, \text{NP}_1$
	DTV :	$\vdash \text{write}^\perp, (\text{NP}_1^\perp \wp S \wp \perp \text{PP} \wp \perp \text{NP}_2)$
	AUX :	$\vdash \text{did}^\perp, (Q \wp \perp \text{INF} \wp \perp \text{NP}_1)$
	WH :	$\vdash \text{whom}^\perp, (\bar{Q} \wp (\perp \perp \text{NP}_i \otimes \perp Q))$
	P :	$\vdash \text{to}^\perp, (\text{PP} \wp \perp \text{NP}_3)$

²⁸ Anche in questo caso ci avvaliamo di esempi della lingua inglese, in cui le dipendenze frasali presentano una struttura tipicamente introdotta dai sintagmi-*wh* come *who*, *whom*, *what*, ecc.; si vedano Ades, Steedman (1982) e Chomsky (1995).

Su questa base, i verbi transitivi DTV, come *give* o *write*, che prendono un oggetto diretto DO e un oggetto indiretto IO, hanno la seguente struttura argomentale:

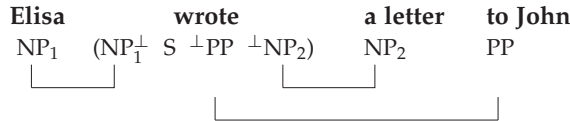


Essa è esemplificata dalla seguente dimostrazione nel calcolo dei sequenti ad una parte:

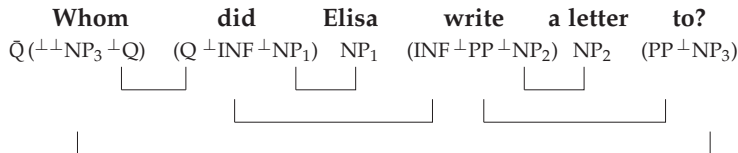
$$\begin{array}{c}
 \frac{\frac{\frac{\vdash Elisa^\perp, NP_1 \quad \vdash wrote^\perp, NP_1^\perp, S, ^\perp PP, ^\perp NP_2}{\vdash wrote^\perp, Elisa^\perp, S, ^\perp PP, ^\perp NP_2} \text{ cut } \quad \vdash a \text{ letter}^\perp, NP_2}{\vdash a \text{ letter}^\perp, wrote^\perp, Elisa^\perp, S, ^\perp PP} \text{ cut } \quad \vdash to \ John^\perp, PP}{\vdash to \ John^\perp, a \text{ letter}^\perp, wrote^\perp, Elisa^\perp, S} \text{ cut } \\
 \hline
 Elisa, wrote, a \text{ letter}, to \ John \vdash S \quad (i)^\perp
 \end{array}$$

Figura 3. MNLL: *Dipendenza frasale*

La dimostrazione e tutte quelle ad essa equivalenti corrispondono a questo semplice *proof net*:



La domanda relativa all'oggetto diretto DO si ottiene assegnando al pronome-*wh* il tipo: $(\bar{Q} \ \mathfrak{A} \ (^{\perp\perp} NP_i \otimes ^\perp Q))$; in base a questo assegnamento, la formula $(^{\perp\perp} NP_3 \otimes ^\perp Q)$ è duale della formula $(Q \ \mathfrak{A} \ ^\perp NP_3)$ che risulta dalla composizione dei costituenti della clausola con cui il sintagma-*wh* è in relazione, per NP_3 argomento della preposizione *to*. La relazione di dipendenza si ottiene eseguendo il *cut* tra il $\mathfrak{A}$ -link nella formula $(Q \ \mathfrak{A} \ ^\perp NP_3)$ della clausola principale e il $\otimes$ -link nella formula $(^{\perp\perp} NP_3 \otimes ^\perp Q)$ del sintagma-*wh*. Nel *proof net* questo *cut* risulta nel lungo link che connette il tipo a sinistra $^{\perp\perp} NP_3$ al tipo $^\perp NP_3$ nella parte finale della dipendenza frasale:²⁹



²⁹La formula con due negazioni nel tipo del sintagma-*wh* è collegata alla posizione definita da Chomsky come 'traccia' del movimento che l'argomento ha compiuto per raggiungere la parte iniziale della frase interrogativa, assumendo la forma di sintagma quantificazionale; vedi Chomsky (1995).

Per dimostrare il sequente *Whom, did, Elisa, write, a-letter, to* $\vdash \bar{Q}$, deriviamo il sequente che include l'argomento $\perp NP_2$ a cui il sintagma nominale *a-letter* si applica:

$$\frac{\frac{\vdash did^\perp, Q, \perp INF, \perp NP_1 \quad \vdash Elisa^\perp, NP_1}{\vdash Elisa^\perp, did^\perp, Q, \perp INF} \text{ cut} \quad \vdash write^\perp, INF, \perp PP, \perp NP_2}{\vdash write^\perp, Elisa^\perp, did^\perp, Q, \perp PP, \perp NP_2} \text{ cut}$$

Poi applichiamo gli argomenti NP_2 e PP :

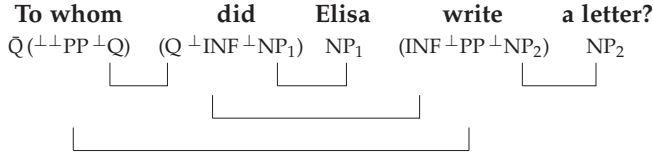
$$\frac{\frac{\vdash write^\perp, Elisa^\perp, did^\perp, Q, \perp PP, \perp NP_2 \quad \vdash aletter^\perp, NP_2}{\vdash aletter^\perp, write^\perp, Elisa^\perp, did^\perp, Q, \perp PP} \text{ cut} \quad \vdash to^\perp, PP, \perp NP_3}{\vdash to^\perp, aletter^\perp, write^\perp, Elisa^\perp, did^\perp, Q, \perp NP_3} \text{ cut}$$

Infine applichiamo il pronome relativo *Whom* alla formula risultante:

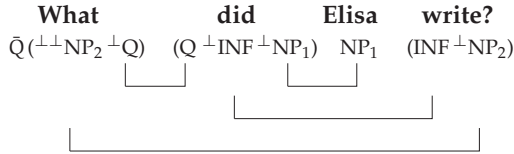
$$\frac{\vdash Whom^\perp, \bar{Q}, (\perp\perp NP_3 \otimes \perp Q) \quad \vdash to^\perp, aletter^\perp, write^\perp, Elisa^\perp, did^\perp, (Q \wp \perp NP_3)}{\vdash to^\perp, aletter^\perp, write^\perp, Elisa^\perp, did^\perp, Whom^\perp, \bar{Q}} \text{ cut} \quad (i)^\perp$$

$Whom, did, Elisa, write, a letter, to \vdash \bar{Q}$

Il caso in cui la preposizione *to* occorre all'inizio della frase insieme al pronome-*wh* segue lo stesso schema: la differenza è che il tipo $(\bar{Q} \wp (\perp\perp PP \otimes \perp Q))$ è assegnato al composto lessicale *to-whom* e il $\wp$ -link collega il posto per argomento di tipo $\perp PP$ dell'oggetto indiretto nel DTV e il tipo $\perp\perp PP$ nel sintagma-*wh*:



La forma interrogativa dell'oggetto diretto DO si ottiene assegnando il tipo $(\bar{Q} \wp (\perp\perp NP_2 \otimes \perp Q))$ al pronome-*wh*, in cui la formula $(\perp\perp NP_2 \otimes \perp Q)$ risulta duale della formula $(Q \wp \perp NP_2)$, che corrisponde alle clausola in cui manca l'oggetto diretto cui si riferisce il pronome relativo:



7. Conclusioni

In questo lavoro abbiamo cercato di mettere in luce, con l'aiuto di alcuni esempi, il ruolo innovativo che può essere assolto dalla logica lineare non commutativa nello studio delle proprietà formali delle lingue naturali. In particolare, abbiamo cercato di mostrare che le grammatiche dei tipi logici possono essere riformulate e considerate sotto una nuova luce nei termini delle rappresentazioni geometriche

offerte dalle reti di dimostrazione (*proof nets*) della logica lineare non commutativa. Queste strutture, come risultato dell'interazione delle due negazioni e dei connettivi duali $\wp$, $\otimes$, costituiscono un contesto dinamico per la rappresentazione del linguaggio naturale, nel quale una sequenza di formule (un enunciato) può essere considerato da diversi punti di vista.

In un articolo del 1999 Lambek (1997) ha presentato la formulazione di un nuovo calcolo per l'analisi logica delle lingue naturali, da lui definito calcolo dei pregruppi o *Pregroup Grammar*. Non entreremo qui nei dettagli di questo calcolo che rappresenta un'alternativa innovativa del *Syntactic Calculus* del 1958; per gli aspetti teorici e le applicazioni linguistiche si vedano in particolare: Lambek (2008), Casadio, Lambek (2001), Casadio, Lambek (2002) e Casadio, Lambek (2008). Come il calcolo sintattico era basato sulla operazione di *residuation* e su monoidi residuati, il calcolo dei pregruppi si basa sulle nozioni di *compact closed category* e aggiunto, riformulando in questo contesto le proprietà della logica lineare non-commutativa moltiplicativa (Casadio, Lambek, 2002; Lambek, 2004; Lambek, 2001; Lambek, 2008).

Un aspetto caratteristico del calcolo dei pregruppi è la definizione di categoria che ammette due aggiunti, un aggiunto di destra (*right adjoint*) e un aggiunto di sinistra (*left adjoint*) in modo analogo alle due negazioni della logica lineare non-commutativa, precisamente:

$$\begin{array}{ll} \text{post-negazione di } A = A^\perp & \text{vs. aggiunto di destra di } A = A^r \\ \text{retro-negazione di } A = {}^\perp A & \text{vs. aggiunto di sinistra di } A = A^l \end{array}$$

Per via della semplificazione ottenuta nei processi di computazione, il calcolo dei pregruppi è stato applicato in breve tempo a numerose lingue naturali: inglese, italiano, francese, tedesco, polacco, persiano, giapponese e molte altre (Casadio, Lambek, 2008; Lambek, 2010; Sadrzadeh, 2008). Le proprietà del calcolo sono attualmente oggetto di studio, ma è importante sottolineare la caratteristica fondamentale di questi sistemi: l'attenzione verso l'ordine delle risorse che vengono prese in considerazione. Tale sensibilità all'ordine, che risulta dalla soppressione della regola strutturale dello scambio, si riflette nella presenza di due negazioni (due aggiunti) e nella conseguente architettura simmetrica e planare delle reti di dimostrazione (*proof net*).

Bibliografia

- Abrusci V. M. (1991). "Phase semantics and sequent calculus for pure noncommutative classical linear propositional logic". *The Journal of Symbolic Logic*, 56.4, pp. 1403–1451.
- Abrusci V. M. (2002). "Classical conservative extensions of Lambek calculus". *Studia Logica*, 71.3, pp. 277–314.
- Abrusci V. M. (2014). "On residuation", in C. Casadio, B. Coecke, M. Moortgat, P. Scott (cur.), *Categories and types in logic, language, and physics*. Berlin: Springer, pp. 14–27.
- Abrusci V. M., P. Ruet (1999). "Non-commutative logic I: the multiplicative fragment". *Annals of pure and applied logic*, 101.1, pp. 29–64.

- Ades A. E., M. J. Steedman (1982). "On the order of words". *Linguistics and philosophy*, 4.4, pp. 517–558.
- Ajdukiewicz K. (1935). "Die syntaktische Konnexität". *Studia philosophica*, 1, pp. 1–27.
- Bar-Hillel Y., C. Gaifman, E. Shamir (1960). "On categorial and phrase structure grammars". *Bulletin of the research council of Israel*, 9F, pp. 1–16.
- Bar-Hillel Y. (1953). "A quasi-arithmetical notation for syntactic description". *Language*, 29, pp. 47–58.
- Benthem J. van, A. Ter Meulen (1996). *Handbook of logic and language*. Elsevier.
- Casadio C., J. Lambek (2008). *Recent computational algebraic approaches to morphology and syntax*. Milano: Polimetrica.
- Casadio C. (2011). "Espressione e significato nelle Ricerche Logiche di Husserl", in Domenico Bosco, Roberto Garaventa, Luigi Gentile, Claudio Tuozzolo (cur.), *Logica, ontologia ed etica*. Milano: Franco Angeli.
- Casadio C., J. Lambek (2001). "An algebraic analysis of clitic pronouns in Italian", in P. De Groote, G. Morrill, C. Retoré (cur.), *Logical aspects of computational linguistics*. Berlin: Springer, pp. 110–124.
- Casadio C., J. Lambek (2002). "A tale of four grammars". *Studia Logica*, 71.3, pp. 315–329.
- Chomsky N. (1995). *The minimalist program*. Cambridge (Massachusetts): The MIT Press.
- Frege G. (1892). "Über Sinn und Bedeutung". *Zeitschrift für Philosophie und philosophische Kritik*, 100, pp. 25–50.
- Frege G. (1893/1903). *Die Grundgesetze der Arithmetik*. Vol. 1. Jena: Hermann Pohle.
- Gabriel G. (1976). "Frege an Husserl, Husserl an Frege", *Wissenschaftlicher Briefwechsel*. Leipzig: Meiner.
- Geach P. T. (1972). "A program for syntax", in D. Davidson, G. Harman (cur.), *Semantics of natural language*. Dordrecht: Reidel, pp. 483–497.
- Girard J.-Y. (1987). "Linear logic". *Theoretical computer science*. Vol. 50. 1, pp. 1–102.
- Girard J.-Y. (1995). "Linear logic: its syntax and semantics", in Jean-Yves Girard, Yves Lafont, Laurent Regnier (cur.), *Advances in linear logic*. Cambridge: Cambridge University Press, pp. 1–42.
- Husserl E. (1913). *Logische Untersuchungen*. Halle.
- Lambek J. (1958). "The mathematics of sentence structure". *American mathematical monthly*, 65, pp. 154–170.
- Lambek J. (1961). "On the calculus of syntactic types", in R. Jacobson (cur.), *Structure of language and its mathematical aspects*. Providence: American Mathematical Society.
- Lambek J. (1997). "Type grammar revisited", in Alain Lecomte, Francois Lamarche, Guy Perrier (cur.), *Logical aspects of computational linguistics*. Berlin: Springer, pp. 1–27.
- Lambek J. (2001). "Type grammars as pregroups". *Grammars*, 4.1, pp. 21–39.
- Lambek J. (2004). "A computational algebraic approach to English grammar". *Syntax*, 7.2, pp. 128–147.

- Lambek J. (2008). *From word to sentence: a computational algebraic approach to grammar*. Milano: Polimetrica.
- Lambek J. (2010). "Exploring feature agreement in French with parallel pregroup computations". *Journal of logic, language and information*, 19.1, pp. 75–88.
- Storrs McCall (cur.), cur. (1967). *Polish logic, 1920-1939*. Oxford: Clarendon Press.
- Montague R. (1976). *Formal philosophy: selected papers*. New Haven: Yale University Press.
- Moortgat M. (1997). "Categorical type logics", in J. Van Benthem, A. Ter Meulen (cur.), *Handbook of logic and language*. Amsterdam: Elsevier, pp. 93–177.
- Morrill G. (1994). *Type logical grammar categorical logic of signs*. Dordrecht: Kluwer.
- Morrill G. (2010). *Categorical grammar: logical syntax, semantics, and processing*. Oxford: Oxford University Press.
- Sadrzadeh M. (2008). "Pregroup analysis of Persian sentences", in C. Casadio, J. Lambek (cur.), *Recent computational algebraic approaches to morphology and syntax*. Milano: Polimetrica.
- Van Benthem J. (1988). "The lambek calculus", in R. T. Oehrle (cur.), *Categorical grammars and natural language structures*. Springer, pp. 35–68.
- Van Benthem J. (1991). "Language in action". *Journal of philosophical logic*, 20.3, pp. 225–263.

Il processamento in tempo reale delle frasi complesse

Cristiano Chesi – Istituto Universitario di Studi Superiori di Pavia –
cristiano.chesi@iusspavia.it

1. Introduzione

Questa riflessione parte da una domanda apparentemente molto semplice: come mai risulta difficile capire certe frasi? La risposta a questa domanda purtroppo non potrà essere altrettanto semplice: si dovrà prima di tutto tentare di specificare in modo chiaro come la complessità di un enunciato possa essere misurata in modo oggettivo, quindi come questa misura possa essere correlata con una difficoltà quantificabile (ad esempio un rallentamento nei tempi di lettura) legata al processamento di queste frasi.

Il tema della complessità è sempre stato di estremo interesse per la ricerca in linguistica formale, guadagnando recentemente una posizione centrale all'interno del cosiddetto 'Programma Minimalista' (Chomsky, 1995). Tale programma ha avuto come obiettivo primario quello di semplificare l'inventario di elementi da usare per spiegare in modo preciso come funziona una grammatica adatta a descrivere un linguaggio umano.

In questo lavoro cercherò prima di tutto di presentare alcuni aspetti cruciali del Programma Minimalista, quindi illustrerò brevemente alcuni risultati di esperimenti che mostrano come il processamento in tempo reale di certe frasi che produciamo e cerchiamo di comprendere tenda ad essere più o meno complesso, infine tenterò di riconciliare l'aspetto formale e quello empirico dello studio di queste frasi in un'unica teoria grammaticale volta a spiegare come mai certi aspetti di queste strutture ci risultino così ostici.

2. Una rappresentazione 'minimalista' del problema

Varie idee, talvolta addirittura contraddittorie, si sono susseguite nel corso degli ultimi venti anni all'interno del 'Programma Minimalista' (Chomsky, 1995). Alla luce delle recenti evoluzioni di questo approccio, possiamo delineare un minimo comun denominatore all'interno del Programma, identificabile con la necessità di avere un lessico e almeno un'operazione ricorsiva di base che permetta di combinare produttivamente gli elementi lessicali in strutture frasali complete.

Il lessico è una componente di base della grammatica che custodisce le singole parole della nostra lingua. Assumiamo per semplicità di non considerare nessuna analisi morfologica o distinzione tra radici e suffissi all'interno di un lessico minimale: in pratica sia *mangia* che *mangiare* saranno due entrate lessicali distinte del

nostro lessico.³⁰ Ad esse devono essere associate, sempre nel lessico, informazioni sulla pronuncia delle parole e sul loro significato. Alle singole parole devono inoltre essere associati quei tratti che permettono di combinare tra loro gli elementi lessicali in sintagmi e frasi. Questi tratti possono essere considerati come i bordi sagomati nei tasselli di un puzzle: due parole che stanno bene in sequenza hanno 'sagomature' compatibili (ad esempio 'Gianni mangia') mentre altre producono sequenze agrammaticali (ad esempio '*Gianni mela'³¹).

La principale operazione che permette di combinare questi elementi lessicali in sequenze sfruttando i loro tratti si chiama *Merge* e, semplicemente, permette di costruire l'insieme di due elementi lessicali (*una* e *mela* diventa [una mela]). L'operazione di *Merge* si dice ricorsiva, poiché può applicarsi ricorsivamente ai risultati di altre operazioni di *Merge* precedentemente avvenute:

- (1.) a. [mangia [una mela]]
b. [Gianni [mangia [una mela]]]

In ogni seria teoria grammaticale c'è bisogno di almeno un meccanismo ricorsivo che permetta di fare un uso infinito (frasi sempre nuove) di mezzi finiti (elementi lessicali). Tale necessità è giustificata dal fatto che qualunque lingua naturale annovera un numero infinito di frasi possibili e solo un numero finito di elementi lessicali.³²

Vista in termini più precisi (formali), l'operazione di *Merge* è legittima se e solo se un tratto di selezione (categoriale) viene soddisfatto. Adottando l'approccio di Stabler (1997) alla formalizzazione delle idee minimaliste, possiamo assumere che ogni elemento lessicale ha dei tratti categoriali (parte del discorso, in inglese *Part of Speech*, POS), quali determinante (D), nome (N), verbo (V):

- (2.) [D Gianni], [N mela], [D una], [V mangia]

Questi tratti 'di base' sono a loro volta selezionabili: il tratto =N indica che l'elemento che lo possiede deve combinarsi con un elemento di tipo N affinché il risultato del *Merge* sia soddisfacente. I tratti vengono cancellati dopo essere stati combinati via *Merge*, onde evitare che nuove operazioni di *Merge* utilizzino gli stessi tratti:

- (3.) Lessico: [D Gianni], [N mela], [=ND una], [=D=DV mangia]

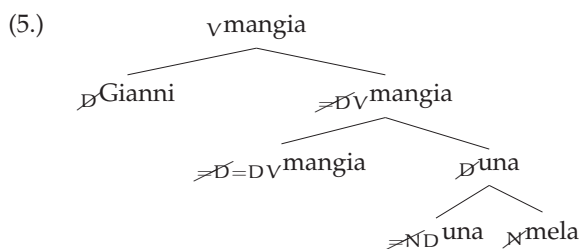
³⁰Questa idea, nota come 'lista completa' (in inglese '*full listing hypothesis*' Jurafsky, Martin, 2008, p. 35), non è psicolinguisticamente molto sostenibile, ma è computazionalmente corretta e semplifica di molto il ragionamento che vogliamo fare in questa sede. Assumere una derivazione morfologica parallela all'archiviazione lessicale non invalida la logica dell'argomentazione qui riportata.

³¹L'asterisco ad inizio frase viene solitamente utilizzato per indicare una frase agrammaticale come appunto '*Gianni mela' oppure '*Gianni mangiano'. Diversi gradi di grammaticalità/accettabilità sono espressi solitamente utilizzando altri diacritici quali il punto interrogativo: '??quali esercizi non sai chi ha risolto?'.
³²Sebbene esista la possibilità di coniare neologismi, un'aggiunta di un elemento ad una lista finita, restituisce comunque una lista finita.

- (4.) I Merge: [_D [_{≠ND} una] [_N mela]]
 II Merge: [_V [_{≠D=DV} mangia] [_D [_{≠ND} una] [_N mela]]]
 III Merge: [_V [_D Gianni] [_V [_{≠D=DV} mangia] [_D [_{≠ND} una] [_N mela]]]]

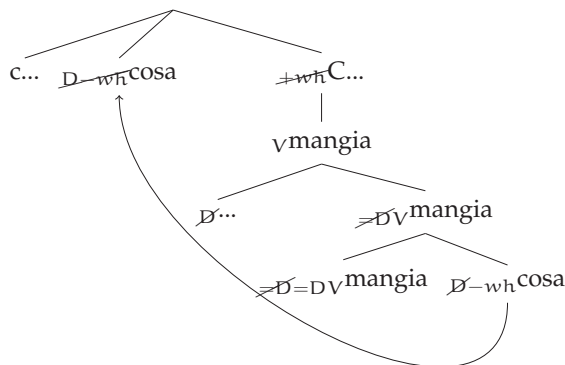
In questo modo, un verbo che assegna due ruoli tematici (paziente: ‘una mela’ e agente ‘Gianni’) avrà due tratti di tipo ‘D’ (determinante) e li cancellerà nell’ordine prestabilito, assegnando due distinti ruoli tematici a due distinti sintagmi nominali (indicati in questo formalismo come elementi di tipo ‘D’).

La relazione tra gli elementi lessicali combinati via *Merge* può essere espressa tanto in termini insiemistici, quanto in termini di dominanza e costituenza se si assume che l’elemento che seleziona domina (‘proietta sopra’, in termini tecnici) l’elemento selezionato; in questo caso la derivazione in 4. corrisponderà all’albero sintattico in 5. (che altro non è se non la storia della derivazione narrata dal basso verso l’altro, *bottom-up*, attraverso le varie operazioni di *Merge*). Le foglie dell’albero (le etichette più in basso, sotto cui non esistono rami), lette da sinistra a destra restituiscono invece la frase generata (appunto: ‘Gianni mangia una mela’).



All’interno del Programma Minimalista si assume che esista una variante ‘non-locale’ dell’operazione di *Merge*, una variante cioè che permetta di far entrare in relazione sintagmi o elementi lessicali in posizioni non adiacenti. Facciamo un esempio: ‘Cosa mangia?’ è una domanda il cui sintagma nominale [_D cosa] ha lo stesso ruolo del sintagma nominale [_D una mela] della frase in 5.: entrambe sono il paziente del predicato *mangiare*. Per render conto di questa simmetria, si assume che il sintagma [_D cosa] sia interpretato come paziente (e riceva dal verbo il ruolo di oggetto diretto) esattamente come [_D una mela]. La posizione periferica sinistra di [D cosa] è determinata da un altro tratto (un tratto interrogativo ‘-wh’ associato all’elemento lessicale *cosa*: [_D -wh cosa]) che viene richiamato più in alto nella struttura da un altro elemento lessicale chiamato complementatore (C). Tale complementatore può essere foneticamente nullo, così come il soggetto in lingue che l’italiano. Per ottenere il corretto ordine degli elementi all’interno di una frase, si deve assumere che il complementatore nelle frasi interrogative sia dotato di un tratto “+wh” che induce l’operazione responsabile del riordinamento degli elementi: come nel *Merge* locale, il *Merge* non-locale (detto *Internal Merge* o *Move*) recupera un elemento dotato di tratto ‘-wh’ già presente nella struttura fin lì costruita e lo solleva fino ad una posizione adiacente a C in cui verrà effettivamente pronunciato. Come risultato dell’operazione, i tratti ‘+wh’ e ‘-wh’ vengono

cancellati per evitare di essere utilizzati per nuove operazioni di spostamento. La parte rilevante della derivazione è descritta in 6.:³³



(6.)

Da un punto di vista puramente formale questa visione della grammatica ci permette di avere una teoria descrittivamente adeguata:³⁴ le frasi generabili sono effettivamente frasi che riteniamo ben formate in italiano e la struttura ad esse assegnata è coerente con l'interpretazione che un parlante nativo associa alla frase in esame. Da un punto di vista psicolinguistico, invece, potremmo non essere pienamente soddisfatti: uno scollamento sostanziale è presente quando la teoria grammaticale prevede che l'ordine in cui le parole sono introdotte nella derivazione sintattica sia 'mela, una, mangia, Gianni' mentre l'ordine con cui noi parlanti nativi quelle parole effettivamente le processiamo in una frase (sia generandole che comprendendole) è 'Gianni mangia una mela'.

Ricordiamoci che il nostro obiettivo era quello di catturare l'oggettiva difficoltà nel processamento di frasi particolarmente ostiche per i parlanti nativi attraverso la nostra grammatica. Vediamo quindi adesso come potremmo ragionare in questo senso e quali indizi di difficoltà di processamento potremmo individuare data una derivazione sintattica come quella fin qua descritta.

3. Alcune frasi difficili: frasi relative incassate e frasi scisse

Alcune frasi possono risultare difficili da comprendere per varie ragioni: prima di tutto potrebbero esserci parole che non si conoscono o che sono infrequenti o non familiari. Ad esempio 'Gianni ha mangiato una mela ieri' è più facile per

³³La copia non pronunciata dell'elemento mosso (cioè quella più in basso) viene indicata tra parentesi angolate (cioè <cosa>).

³⁴Il lettore interessato a questioni formali ed empiriche sulla reale adeguatezza di questo modello è rinvio alla lettura di altri saggi necessariamente più tecnici come Chesi (2015).

molti parlanti italiani di 'Gianni ha mangiato un litchi ieri'; questa difficoltà è evidente guardando i tempi di lettura.³⁵ Un rallentamento nella regione di [un litchi] contro [una mela] è legato alla più bassa frequenza di *litchi* (dato facilmente recuperabile ispezionando corpora dell'italiano) e alla bassa familiarità (poche persone conoscono il frutto cinese chiamato litchi), che rendono questo elemento difficilmente accessibile nel lessico mentale (se non addirittura assente). Tra gli altri fattori di rallentamento procedurale è inoltre da considerare la difficoltà articolatoria di questa parola (*litchi* si pronuncia /'littʃi/, non /'litki/ come le regole di lettura dell'italiano porterebbero a pensare).

Un ulteriore livello di complessità è dato dall'ambiguità delle parole di una frase e dalle molteplici interpretazioni che a quella frase si possono dare in virtù dell'ambiguità delle parole stesse: 'la vecchia legge la regola' è più difficile di 'la vecchia rilegge la regola' poiché nel primo caso sussiste un'ambiguità lessicale in ogni elemento della frase che conduce a rianalizzare la frase per scovare la seconda delle due distinte interpretazioni ('una signora anziana legge una normativa' oppure 'una vecchia normativa regola qualcosa') mentre nella seconda frase questa possibilità non sussiste e solo un'interpretazione è possibile, in virtù della non ambiguità del verbo *rileggere*. Nel primo caso, utilizzando strumenti di analisi dei movimenti oculari (*eye-tracking*) si possono osservare regressioni: il lettore torna indietro e rilegge più volte vari segmenti della frase mostrata per intero sullo schermo prima di informare lo sperimentatore che ha compreso i due sensi della frase.

In questo lavoro non prenderemo in considerazione questi fattori, che di sicuro contribuiscono alla difficoltà percepita di un enunciato. Quello che ci interessa andare ad investigare è il contributo in termini di complessità di precisi e misurabili fattori strutturali, che, tenuti costanti tutti gli altri fattori (quali la frequenza e familiarità delle parole, così come l'ambiguità delle parole e delle strutture) rendono più o meno facile il recupero del significato di un enunciato.

Le strutture maggiormente utilizzate in questo senso per valutare la difficoltà percepita da un parlante nativo sono le frasi relative del tipo 'il signore che ha mangiato la mela è Gianni' e le frasi scisse: 'è Gianni che ha mangiato la mela'. Questo perché in esse vengono utilizzate sia operazioni di costruzione della struttura locale (*Merge*) che operazioni di movimento, e quindi di ricostruzione di una dipendenza non-locale.

È in effetti interessante notare che in una grammatica formale, come quella minimalista, se è possibile applicare un'operazione una volta, dovrebbe essere le-

³⁵Un esempio di esperimento di rilevazione dei tempi di lettura è il *self-paced reading* (lettura auto-regolata): al soggetto viene chiesto di leggere un pezzetto di frase su uno schermo, quindi di premere un tasto per leggerne un altro pezzo; ad esempio:

i.	Gianni	_____	_____	
ii.	_____	ha mangiato	_____	
iii.	_____	_____	un litchi	_____ oppure _____ una banana _____
iv.	_____	_____	_____	ieri.

Un rallentamento nel segmento iii. nella condizione [un litchi] (contro [una banana]) mostra che la diversa parola ha richiesto un tempo di processamento più lungo.

gittimo poterla applicare sempre se il contesto rimane lo stesso. In particolare se abbiamo un sintagma nominale [un signore] e possiamo modificarlo con una frase relativa 'che un amico aveva salutato' ottenendo '[un signore [che un amico aveva salutato]] è caduto dalla bici' allora dovrebbe esser sempre possibile modificare un sintagma nominale del tipo 'un signore' con una frase relativa. Ciò è apparentemente vero in certi contesti:

- (7.) E venne [il cane [che morse [il gatto [che rincorse [il topo [che mio padre comprò]]]]]]

ma non in altri:

- (8.) ??[Un signore [che un amico [che mio padre conosce] aveva salutato] è caduto dalla bici].

La seconda frase è assolutamente incomprensibile per la maggior parte dei parlanti dell'italiano. Eppure l'aggiunta di una frase relativa, potenzialmente ottenuta attraverso un'operazione di *Merge*, risulta possibile nel primo caso e 'difficile' nel secondo. Bever (1970) e altri hanno notato che se in questa struttura si cambiano alcune parole la difficoltà del processamento si riduce:

- (9.) [Un signore [che **qualcuno** [che **io** conosco], aveva salutato]] è caduto dalla bici].

In pratica più i sintagmi nominali ([un signore], [qualcuno] e [io]) sono diversi, migliore è la comprensione della frase. Bisogna notare che la struttura 'profonda' della frase in 9. è la seguente:³⁶

- (10.) [Un signore; [che **qualcuno**_i [che **io** conosco _{-i}], aveva salutato _{-j}]] è caduto dalla bici].

In pratica, l'oggetto diretto di *conosco* è [qualcuno] mentre l'oggetto diretto di *salutato* è [un signore]. Queste dipendenze scavalcano rispettivamente il pronome soggetto di *conosco* ([io]) e il soggetto di *salutato* ([qualcuno]). Poiché solo [un amico] e [mio padre] nella frase in 8. sono stati rimpiazzati da [io] e [qualcuno] in 9. il miglioramento nella comprensione delle frasi deve essere legato alla differenza tra queste coppie di elementi.

Per facilitare la comprensione degli aspetti più tecnici, si utilizzeranno alcuni termini specifici che ci permetteranno di far riferimento alle posizioni coinvolte in questo tipo di strutture: in 9. [un signore] è la testa della relativa e [qualcuno] il soggetto della relativa. La posizione vuota dopo *salutato* è la posizione base dell'oggetto della relativa, che subisce un movimento (dipendenza non-locale). Infine, le parentesi quadre prima di *che* e dopo la posizione vuota sono i confini della relativa. Vediamo ora nel dettaglio l'evidenza empirica che mostra una differenza di difficoltà nella percezione di queste frasi.

³⁶In pratica, la lineetta bassa indica la posizione in cui un elemento non pronunciato deve essere interpretato come tema del predicato; il pedice indica il sintagma nominale a cui si riferisce.

3.1. Evidenza empirica

Gordon et al. (2001), utilizzando la tecnica del *self-paced reading* (vedere nota 35), mostrano come certi tipi di relative siano più complesse di altre. In particolare le relative del tipo 11.a (relative soggetto, poiché la ‘testa della relativa’, [the banker], occupa la posizione di soggetto del predicato *praised* nella frase relativa) vengono lette più rapidamente delle relative del tipo 11.b (relative oggetto, poiché [the banker], occupa la posizione di oggetto del predicato *praised*):

- (11.) a *The banker* [that _ *praised the barber*] *climbed the mountain.*
 il banchiere che _ lodò il barbiere scalò la montagna
- b *The banker* [that *the barber praised* _] *climbed the mountain.*
 il banchiere che il barbiere lodò _ scalò la montagna

Questo rallentamento è evidente nelle parole critiche *praised* in 11.a contro *barber* 11.b e *praised* in 11.b contro *barber* 11.a (entrambi i segmenti critici vengono letti 300 ms più lentamente nella condizione oggetto rispetto a quella soggetto). Anche nell’accuratezza delle risposte alle domande di comprensione delle suddette frasi si nota una significativa differenza, con le relative soggetto più semplici da capire (93% delle risposte alle domande che i lettori forniscono sono corrette) delle relative oggetto (87% di accuratezza).

L’interpretazione di questa asimmetria è stata principalmente ricondotta al fatto che solo nella relative oggetto il soggetto della frase relativa viene ‘scavalcato’ dalla dipendenza non-locale come mostrato di seguito:

- (12.) The banker [that the barber praised _] climbed the mountain
- 

Ipotizzando che questa dipendenza non-locale sia stata generata utilizzando l’operazione di movimento e quindi dei tratti particolari per ‘spostare’ la testa della relativa nella sua posizione periferica, possiamo intuire che questi tratti debbano giocare un ruolo cruciale nello spiegare la facilitazione descritta dall’esempio in 9..

In effetti, in un secondo esperimento, Gordon e colleghi mostrano come un pronome di seconda persona (*you*) in posizione di soggetto della relativa effettivamente riduca la complessità della relativa oggetto, in pratica rendendola indistinguibile dalla relativa soggetto:

- (13.) a *The banker* [that *you praised* _] *climbed the mountain.*
 il banchiere che tu lodasti _ scalò la montagna
- b *The banker* [that _ *praised you*] *climbed the mountain.*
 il banchiere che _ lodò te scalò la montagna

Il segmento verbale della relativa *praised* è letto molto più rapidamente nella condizione 13.a (quando il soggetto della relativa oggetto è un pronome) rispetto alla condizione 11.b (quando il soggetto è una descrizione definita), mentre i tempi di lettura delle parole critiche *you* e *praised* non differiscono significativamente nelle condizioni (13.a e 13.b, cioè quando la struttura è una relativa soggetto e il pronome è l'oggetto diretto del predicato *praised*; in questo caso in effetti non si osserva nessuno scavalco).

Risultati del tutto paragonabili sono stati ottenuti da Gordon e colleghi utilizzando un nome proprio al posto del pronome, sempre in posizione di soggetto della frase relativa:

- (14.) a *The **banker** [that **Ben** praised _] climbed the mountain.*
 il banchiere che Ben lodò _ scalò la montagna
- b *The **banker** [that _ praised **Ben**] climbed the mountain.*
 il banchiere che _ lodò Ben scalò la montagna

Per capire se fondamentale è solo il ruolo del soggetto oppure anche il ruolo della testa della relativa a rendere più o meno semplice la comprensione di queste frasi, Gordon e colleghi utilizzano una tipologia di frasi simili alle relative fin qui mostrate, le 'frasi scisse' (*cleft sentences*, in inglese):

- (15.)
- a It was the **banker** that the **lawyer** saw _ in the parking.
 b It was the **banker** that **Bill** saw _ in the parking.
 c It was **John** that the **lawyer** saw _ in the parking.
 d It was **John** that **Bill** saw _ in the parking.

Le frasi scisse permettono di utilizzare come testa della struttura relativa sia descrizioni definite ([il banchiere], [l'avvocato]) che nomi ([John], [Bill]) e potenzialmente pronomi (vedere esperimento successivo), permettendo tutte le possibili combinazioni di elementi. In effetti, i nomi propri e i pronomi possono essere teste nelle frasi scisse ma non nelle relative restrittive.³⁷

A questo punto è interessante osservare che di nuovo la differenza tra relativa soggetto e oggetto ricompare: in particolare, si osserva un significativo rallentamento (di circa 250 ms) del verbo della relativa (*saw*) nelle strutture relative oggetto in cui la testa e il soggetto sono entrambi descrizioni definite (13.a) o entrambi nomi propri (13.c). Il rallentamento nel processamento del verbo della relativa è

³⁷'Gianni che tutti conoscono' non è una frase relativa restrittiva (cioè non significa esattamente 'tra tutti gli individui di nome Gianni, quello a cui mi riferiscono è quello che tutti conoscono') ma una relativa appositiva (cioè 'esiste un individuo chiamato Gianni e tutti lo conoscono'). Per una descrizione precisa della tipologia di frasi relative e per le varie ipotesi di struttura, fare riferimento a Bianchi (2002).

evidente anche nei casi di *mismatch* (13.b e 13.c), ma la differenza tra relativa soggetto e oggetto è espressa in appena 100 ms di rallentamento e, generalmente, in questi secondi casi, il verbo è letto 150 ms più velocemente che nei casi di *matching* descritti precedentemente.

In conclusione, il ruolo della testa della relativa sembra sia importante, ma lo è in dipendenza dalla tipologia del soggetto della relativa: due sintagmi nominali dello stesso tipo producono gli effetti di rallentamento più forti in conseguenza di una difficoltà maggiore percepita dai lettori.

Un’ulteriore conferma di questo risultato è dato dall’esperimento di Warren, Gibson (2005), l’ultimo che descriveremo, che mette a confronto tutte le possibili combinazioni di sintagmi nominali nelle due posizioni: descrizioni definite, nomi propri e pronomi.

(16.)

- a. It was the **banker** that the **lawyer** avoided _ at the party.
- b. It was the **banker** that **Dan** avoided _ at the party.
- c. It was the **banker** that **we** avoided _ at the party.
- d. It was **Patricia** that the **lawyer** avoided _ at the party.
- e. It was **Patricia** that **Dan** avoided _ at the party.
- f. It was **Patricia** that **we** avoided _ at the party.
- g. It was **you** that the **lawyer** avoided _ at the party.
- h. It was **you** that **Dan** avoided _ at the party.
- i. It was **you** that **we** avoided _ at the party.

I risultati grezzi dei tempi di lettura del verbo *avoided* sono riportati in Tabella 1..³⁸

Condizione	D-D	D-N	D-P	N-D	N-N	N-P	P-D	P-N	P-P
Tempi lettura (ms)	365	319	306	348	347	291	348	311	291
Err. Stand. (ms)	-19	-12	-14	-18	-21	-14	-18	-15	-13

Tabella 1. Tempi di lettura e errore standard in millisecondi (D = descrizioni definite, N = nomi propri, P = pronomi; ad esempio D-D = 16.a; D-N = 16.b; P-P = 16.i etc.).

3.2. Interpretazione dei risultati dei tempi di lettura

Molte teorie psicolinguistiche si sono susseguite, tentando di spiegare asimmetrie simili a quelle dei tempi di lettura riportati nella Tabella 1.. In pratica, una teoria psicolinguisticamente interessante dovrebbe riuscire a spiegare come mai il verbo viene letto più lentamente (e la frase viene compresa con maggiore difficoltà), ad esempio, nei casi D-D (‘Era il banchiere che l’avvocato ha evitato alla festa’, 16.a)

³⁸Grazie a Tessa Warren per aver condiviso questi risultati.

e N-N ('Era Patricia che Dan ha evitato alla festa', 16.e), mentre i tempi di lettura e la difficoltà percepita sono molto minori nel caso P-P ('Eri tu che noi abbiamo evitato alla festa', 16.i).

Due sono le principali interpretazioni dei dati descritti nel capitolo precedente, entrambe facenti cruciale riferimento ad una limitazione d'uso della 'memoria di lavoro': da un lato, si propone l'idea di un costo di integrazione di nuovo materiale referenziale (*Linguistic Integration Cost*, Gibson, 1998, pp. 12-13) legato all'interpretazione di nuovi referenti nominali che, secondo una gerarchia di 'accessibilità', produrrebbero uno sforzo, per quanto riguarda l'aggiornamento in memoria di descrizioni definite, maggiore di quello richiesto dai nomi propri, che è maggiore, a sua volta, di quello richiesto dai pronomi referenziali (*io* e *tu*). Questa idea spiegherebbe come mai le condizioni in cui si trovano una descrizione definita e un nome proprio siano più semplici di quella in cui si presentano due descrizioni definite, ma non spiega come mai anche con due nomi propri il costo di processamento sia paragonabile a quello che si incontra con due descrizioni definite.

Gordon e colleghi, invece, riadattando un vecchio modello di interferenza nella memoria di lavoro (Crowder, 1976) suggeriscono che quel che conta è la similarità tra i due sintagmi coinvolti nella dipendenza, spiegando così che nel caso di *matching* (nome con nome o descrizione definita con descrizione definita) si ottengono i risultati peggiori in termini di complessità (difficoltà percepita maggiore), mentre quando i due sintagmi nominali sono differenti, i tempi di lettura nella regione critica del verbo della relativa migliorano.

Un raffinamento dell'idea di Gordon e colleghi utilizza l'operazione di movimento descritta nella Sezione 2.: tale operazione era responsabile della creazione di dipendenze non-locali come quella tra la testa di una frase scissa e la posizione oggetto a questa collegata: Friedmann et al. (2009), sfruttando la teoria della località, formulata inizialmente in Rizzi (1990) e riadattata al caso del processamento linguistico degli afasici da Grillo (2008), vede da un punto di vista puramente descrittivo la relazione tra una posizione di base (Y) e la posizione in cui un elemento viene mosso (X). Tale relazione è tanto più difficile quanto più un elemento Z, se è presente, è strutturalmente compatibile per soddisfare la relazione che l'elemento X è chiamato a soddisfare. Nell'esempio 12., *banker* è X, la posizione di oggetto vuota è Y, mentre *barber* è Z.

In pratica, l'elemento X [the banker] si muove perché marcato durante l'introduzione nella derivazione da un tratto $-R$ ($[-_R$ the banker]) che verrà richiamato dal complementatore con il fine di spostare l'elemento nella sua posizione più alta (a sinistra nella frase pronunciata). Ricordiamoci che in una derivazione minimalista standard, prima si crea, via *Merge*, il costituente strutturalmente più basso, cioè il predicato con l'oggetto diretto ([saw [the banker]]); quindi si aggiunge il soggetto ([the lawyer [saw [the banker]]]); infine, sempre via *Merge*, il complementatore, dotato di un tratto che richiama l'oggetto diretto in posizione di testa della relativa, scavalcando il soggetto con l'operazione di movimento descritta graficamente in 12. ($[-_R$ the banker] $[+_R$ that] [the lawyer [saw $[-_R$ the banker]]]).

L'idea originale di Friedmann e colleghi è di utilizzare questi tratti che scate-

nano il movimento per prevedere il grado di difficoltà della relazione non-locale da stabilire. In pratica, si possono presentare quattro casi determinati dai tratti associati agli elementi X, Y e Z in gioco:

- **Identità:** i tratti in Y e in Z sono gli stessi. Z è strutturalmente più vicino a X e X ricerca esattamente i tratti presenti sia in Z che in Y per il movimento; Z deve quindi soddisfare la dipendenza, pena l'agrammaticalità della frase (è questo in soldoni il principio di località, Rizzi, 1990) (ad es. *"*Cosa credi chi ha visto?"*).
- **Disgiunzione:** i tratti in Y e in Z sono diversi. X richiede i tratti di Y per soddisfare la relazione di movimento e Z può quindi tranquillamente essere scavalcato senza produrre una frase agrammaticale (ad es. *"Cosa credi che Gianni abbia visto?"*)
- **Inclusione:** Y è più ricco di Z e i tratti che X richiede sono parzialmente soddisfatti da Z e completamente soddisfatti da Y; in questo caso Y e X possono stabilire una relazione di movimento, ma tale relazione sarà più complessa da processare della relazione precedente (caso di disgiunzione) (ad es. *"È Gianni che il maestro ha visto?"*).
- **Inclusione inversa:** Z è più ricco di Y e i tratti che X richiede sono parzialmente soddisfatti da Y e completamente soddisfatti da Z; in questo caso Z e X devono stabilire una relazione di movimento; se X e Y provassero a stabilire una relazione, la frase sarebbe agrammaticale come nel caso dell'identità (ad esempio *"Cosa credi quale studente abbia visto?"*).

Associando i tratti appropriati alla testa e al soggetto della relativa potremmo essere in grado di ottenere una prima metrica di complessità: seguendo Friedmann e colleghi, ipotizziamo che i sintagmi nominali costituiti da una descrizione definita (come [il banchiere]) abbiano un tratto determinante (D) che è selezionato dal predicato e un tratto di restrizione lessicale N (cioè [_{DN} il banchiere]). Ipotizziamo poi che i sintagmi nominali espressi attraverso un nome proprio abbiano lo stesso numero di tratti, ma che la restrizione lessicale N sia particolare in quanto permetta la legittimazione attraverso un solo elemento (un nome proprio appunto) di un tratto di determinazione e la specificazione della restrizione lessicale (cioè [_{DN_{pr}} Gianni]); infine assumiamo che i pronomi siano dotati di tratto determinante, ma non di restrizione lessicale (cioè [_D noi]).

Poiché la testa della relativa è mossa in virtù di un tratto '-R' associato all'oggetto diretto, questa operazione legittima lo scavalco del soggetto della relativa senza incorrere nel caso di identità. Ricadendo però in un caso di inclusione, siamo in grado di prevedere, secondo Friedmann e colleghi, che i casi in cui i tratti condivisi sono di più presenteranno le maggiori difficoltà (in pratica, i casi D-D e N-N), mentre nei casi in cui i tratti condivisi tra la testa e il soggetto sono meno, la frase risulterà più semplice (ogni volta che si presenta un pronome). Riassumendo le previsioni che la teoria della località (L) riesce a fare, abbiamo:

tipo	D-D	D-N	D-P	N-D	N-N	N-P	P-D	P-N	P-P
ms	365	319	306	348	347	291	348	311	291
L	difficile	?	facile	?	difficile	facile	facile	facile	facile

Tabella 2. Previsioni di complessità utilizzando la teoria della località basata su tratti (L).

Per facilitare la lettura della tabella sono stati inseriti dei colori che facilitano l'individuazione dei tempi di lettura più lenti (grigio scuro) e di quelli più rapidi (bianco). Come si riesce facilmente a notare, la teoria qui descritta riesce a prevedere correttamente alcune difficoltà e alcuni casi 'semplici', ma senza fornire nessuna apprezzabile gradualità della difficoltà percepita, né, soprattutto, del perché il rallentamento dei tempi di lettura si riscontra esattamente nel segmento del verbo della frase relativa.

4. Una rivisitazione *processing-friendly* delle Grammatiche Minimaliste

Finora abbiamo visto che esiste una teoria grammaticale semplice per spiegare la costruzione di una frase, la teoria minimalista, in grado di render conto sia della costruzione dei sintagmi adiacenti, che della creazione di dipendenze non-locali, quali quella che si richiede nelle frasi scisse, in cui un elemento, la testa della relativa, è 'mosso' dalla posizione tematica di base (subito dopo il verbo, nel caso dell'oggetto diretto) scavalcando la posizione del soggetto. Questo approccio ci fornisce una teoria elegante per capire come l'ordine delle parole a la loro interpretazione compositiva viene a formarsi. Abbinato ad una teoria della località basata sui tratti, riusciamo anche ad intuire come mai alcune frasi possano risultare più complesse di altre. Rimane adesso da capire se un piccolo raffinamento di questa teoria potrebbe permetterci di ottenere una maggiore gradualità nel giudizio di complessità, ovvero potrebbe dirci precisamente quanto una frase è più o meno difficile di un'altra e dove questa particolare difficoltà si presenterà. Un passo in questa direzione viene fatto da quelle grammatiche che 'rovesciano' la costruzione della struttura: in pratica, si assume che le parole all'interno di frase siano generate all'incirca nell'ordine in cui poi vengono lette e pronunciate, cioè da sinistra a destra, anziché da destra a sinistra come previsto dall'approccio Minimalista standard.³⁹ Questa idea parrebbe l'uovo di Colombo, ma una proposta in tal senso si ritrova per la prima volta nella tesi Colin Phillips del 1996 (Phillips, 1996) ed è stata largamente ignorata dal filone di ricerca principale. Il lavoro di

³⁹Questo approccio viene solitamente chiamato *Top-Down*, dall'alto al basso (Chesi, 2012), per distinguerlo dall'approccio standard che invece costruisce la struttura dal basso all'alto (e, conseguentemente, da destra a sinistra in lingue come l'italiano o l'inglese).

Phillips, tra gli altri, ha ispirato il modello grammaticale descritto in Chesi (2004), Chesi (2012) e Chesi (2015) che qui rapidamente illustreremo.

Del modello Minimalista standard, viene mantenuta l'idea del lessico (una lista di parole con associati tratti), le operazioni di *Merge* e *Move* e il fatto che tali operazioni si basino sui tratti associati alle parole nel lessico. Come nel *Merge* minimalista, il *Merge* nel nuovo modello usa il tratto $=X$ per esprimere un requisito di selezione che solo un costituente portatore del tratto X potrà soddisfare; la differenza rispetto al modello standard è che prima (in senso strettamente temporale) si dovrà valutare l'elemento portatore del tratto $=X$, quindi si cercherà di soddisfare questa 'aspettativa' che è venuta a crearsi con un elemento di tipo X , dove X può essere uno solo o anche un insieme di tratti. Le aspettative possono in effetti essere abbastanza complesse: un predicato verbale potrà selezionare un sintagma nominale ($=DP$) con i tratti D (determinazione) e N (restrizione lessicale) espressi, mentre un nome potrà richiedere una specificazione di una qualche restrizione (frase relativa, $=R$) che ha una struttura frasale completa (ovvero un adeguato complementatore, C , un soggetto, S , e un predicato, V , con specificazione temporale esplicita, T , cioè, stando alla grammatica dell'italiano, questa lista ordinata di tratti: $R = \langle C S T V \rangle$).

Il movimento, invece, in questo nuovo modello, utilizza esplicitamente l'idea della memoria di lavoro e vi inserisce i sintagmi nominali che non hanno ancora ricevuto un ruolo tematico; se incontriamo quindi un soggetto in posizione preverbale (ad esempio 'Gianni corre'), l'ordine di processamento degli elementi porterà a concludere che prima *Gianni* viene inserito nella derivazione per soddisfare qualche tratto particolare (ad esempio un requisito che richiede che un soggetto sia espresso per ogni frase di senso compiuto), quindi viene inserito nella memoria di lavoro finché *corre* non viene processato e quindi cancellato dalla memoria e re-integrato nella struttura come agente del verbo *correre*. Si assume inoltre che la memoria funzioni come una pila di oggetti, cioè se più elementi vengono inseriti in memoria, il primo da riutilizzare sarà l'ultimo che vi è stato inserito. Prendendo la parte rilevante della struttura delle frasi in questione, ad esempio la frase scissa 'Era Gianni che tu hai visto' possiamo derivare la struttura partendo dal lessico e dai tratti ad esso associati come indicato sotto:⁴⁰

- (17.) Lessico: $[_{DNpr=R} \text{Gianni}] [D \text{ tu}] [C \text{ che}] [T \text{ hai}] [_{cop=DP(no-tem-rol)} \text{era}]$
 $[_{V=DP=DP} \text{visto}]$

Derivazione:

1. $[_{cop=DP(no-tem-rol)} \text{era } DP[DN \dots]]$

La copula viene inserita nella derivazione e l'aspettativa per un DP a cui non assegna ruolo tematico viene proiettata.

⁴⁰La derivazione indicata è estremamente semplificata per i fini di questo articolo. Si ricorda che $=R$ corrisponde all'espansione dei seguenti tratti $\langle C S T V \rangle$ e che 'S' è un tratto che richiede l'inserimento di un sintagma nominale (potenzialmente anche foneticamente nullo). Il lettore interessato è invitato ad approfondire i vari aspetti tecnici in Chesi (2015).

2. $[_{cop} \text{ era } DP[_{DN_{pr}=R} \text{ Gianni}]] M = ([_{DN_{pr}} \text{ Gianni}])$
Gianni soddisfa i requisiti di selezione, ma non avendo il ruolo tematico assegnato, viene inserito nella memoria di lavoro (M), suggerendo che una dipendenza non-locale dovrà essere stabilita con un predicato appropriato.
3. $[_{cop} \text{ era } DP[_{DN_{pr} \not=R} \text{ Gianni } R[_{CSTV} \dots]]]$
 La restrizione viene espansa con la sequenza di tratti C(omplementatore), S(oggetto), T(empo) e V(erbo) lessicale.
4. $[_{cop} \text{ era } DP[_{DN_{pr}} \text{ Gianni } R[_{C} \text{ che } [_{S} [_{D} \text{ tu}]] \text{ T hai } [_{V=DP=DP} \text{ visto}]]]] M = ([_{DN_{pr}} \text{ Gianni}], [_{D} \text{ tu}])$
 In ordine, $[_{C} \text{ che}]$, $[_{D} \text{ tu}]$ $[_{T} \text{ hai}]$ e $[_{V=DP=DP} \text{ visto}]$ vengono inseriti nella derivazione; *tu* può soddisfare il requisito di soggetto del predicato, ma in questa posizione, prima cioè di aver inserito *visto* non riceve nessun ruolo tematico, quindi viene inserito come *Gianni* nella memoria di lavoro.
5. $[_{cop} \text{ era } DP[_{DN_{pr}} \text{ Gianni } R[_{C} \text{ che } [_{S} [_{D} \text{ tu}]] \text{ T hai } [_{V=DP=DP} \text{ visto } DP[\dots] [_{VDP}[\dots]]]]]] M = ([_{DN_{pr}} \text{ Gianni}], [_{D} \text{ tu}])$
 I due ruoli tematici indicati dalla sequenza di tratti $=DP =DP$ sul verbo lessicale vengono proiettati nella struttura.
6. $[_{cop} \text{ era } DP[_{DN_{pr}} \text{ Gianni } R[_{C} \text{ che } [_{S} [_{D} \text{ tu}]] \text{ T hai } [_{V} \text{ visto } DP[_{D} \text{ tu}]] [_{V} DP[_{DN_{pr}} \text{ Gianni}]]]]]]]$
 Secondo l'operazione di movimento, per soddisfare i requisiti tematici del verbo si vanno a pescare, in ordine inverso rispetto a quello con cui vi sono stati inseriti, i sintagmi nominali nella memoria di lavoro; quindi prima *tu*, poi *Gianni*, soddisfacendo così correttamente i ruoli di agente e paziente assegnati dal verbo proprio come previsto nella teoria standard.

I vantaggi di questo approccio sono notevoli: da un lato si ha una derivazione che corrisponde esattamente all'ordine in cui le parole vengono lette e pronunciate; dall'altro si fa esplicito riferimento alla memoria di lavoro e si può notare (passi 4–6 in 17.) che tale memoria è popolata contemporaneamente dai due elementi che quindi potranno in qualche modo interferire generando la difficoltà che vorremmo spiegare.

In effetti sappiamo che l'interferenza nella memoria di lavoro è una delle principali limitazioni all'accesso delle informazioni (Anderson, Neely, 1996; Crowder, 1976).⁴¹ Inoltre sappiamo che il momento più critico è rappresentato dal recupero dell'informazione, con pochissimi effetti in fase di archiviazione e mantenimento dell'informazione nella memoria di lavoro (Dillon, Bittner, 1975; Gardiner et al., 1972; Tehan, Humphreys, 1996). Volendo quindi attribuire al verbo il ruolo principale di elemento scatenante la fase del recupero, e volendo utilizzare i tratti che questo porta come indizi per selezionare un elemento in memoria piuttosto che

⁴¹Vedere Nairne (2002) per una panoramica sul tema.

un altro (seguendo l'intuizione del modello *cue-based*, Van Dyke, McElree, 2006) possiamo formulare una precisa metrica basata sui seguenti tratti a disposizione degli elementi lessicali in gioco:

- Descrizioni definite: $\langle D, \text{num}, N \rangle$
 (18.) Nomi propri: $\langle D, \text{num}, N_{pr} \rangle$
 Pronomi personali: $\langle D, \text{case}, \text{pers}, \text{num} \rangle$

In pratica i tratti di numero (*num*, cioè plurale o singolare) sono specificati in tutti e tre le tipologie di elementi, quelli di caso solo sui pronomi (*io* contro *me*), così come i tratti marcati di persona (prima o seconda; la terza persona viene solitamente considerata come una forma di default); infine la restrizione lessicale *N* è diversa per quanto concerne nomi comuni (*N* appunto) e nomi propri (*N_{pr}*) questo perché i due elementi si comportano in modo diverso (in particolare i nomi propri possono costituire un sintagma nominale completo senza la necessità di essere introdotti da un determinante). La restrizione lessicale è invece assente nei pronomi personali.

Sapendo di richiedere al lettore un piccolo sforzo, presenterò adesso una esplicita metrica di complessità (*Feature Retrieval Cost*) che si basa appunto sul recupero degli elementi in memoria in base ai tratti a disposizione: in parole semplici, il recupero di un elemento sarà più facile quando maggiori saranno i tratti che lo contraddistinguono e che sono indicati dal verbo; più difficile sarà il recupero quando in memoria ci saranno più elementi tutti ugualmente compatibili con i tratti che il verbo richiede.

$$(19.) \quad C_{FRC}(x) = \prod_{i=1}^{m_x} \frac{(1 + nF_i)^{m_i}}{1 + dF_i}$$

In termini più precisi, il costo del recupero di un elemento in memoria, giunti alla parola *x*, sarà proporzionale al prodotto dei costi del recupero dei singoli elementi in memoria e questi costi cresceranno in modo esponenziale in base al numero degli elementi in memoria (*m_i*) e al numero dei tratti non distinti che dovranno essere recuperati (*nF_i*); tale costo decrescerà in base al numero di tratti distinti che il verbo esplicitamente richiede (*dF_i*).

In pratica quindi, nel caso D-D ('It was [*D, num_sing, N* the lawyer] who [*D, num_sing, N* the businessman] avoided...') il costo sarà particolarmente alto: $C_{FRC}(\text{avoided}) = 16$. Ovvero $16 \cdot 1$, poiché 16 sarà il costo del recupero del sintagma soggetto ([*+D, +num_sing, N* the businessman] visto che *nF* = 3, cioè $\langle D, \text{num_sing}, N \rangle$, *m* = 2, poiché due elementi sono in memoria in quel momento, e *dF* = 0 poiché nessun tratto espresso dal verbo distingue [the lawyer] da [the businessman]); mentre il costo per il recupero dell'oggetto diretto [the lawyer] è 1, poiché *nF* = 0, *m* = 1 e *dF* = 0.

Stesso esatto costo nel caso N-N ('It was [*D, num_sing, N_{prop}* Patricia] who [*D, num_sing, N_{prop}* Dan] avoided...') visto che entrambe le restrizioni lessicali saranno di tipo *N_{pr}* (ovvero indistinguibili l'una dall'altra).

Nel caso dei pronomi, invece, ('It was [D, num_sing, per_II, case_nom, you] who [D, num_plur, per_I, case_nom we] avoided...') il costo sarà molto minore: $C_{FRC}(\text{avoided}) = 2$. Cioè $2 \cdot 1$, poiché 2 sarà il costo del recupero del pronome soggetto [D, num_plur, per_I, case_nom we] (visto che $nF = 1$, cioè $\langle D \rangle$, $m = 2$, poiché due elementi sono in memoria in quel momento, e $dF = 1$ poiché il caso nominativo di we è selezionato dal verbo); il costo per il recupero dell'oggetto diretto sarà sempre 1 (poiché $nF = 0$, $m = 1$ e $dF = 0$).

Così facendo possiamo assegnare un valore ad ogni condizione, ad esempio D-N ('It was [D, num_sing, N the lawyer] who [D, num_sing, Nprop Dan] avoided...') ottenendo costi intermedi: $C_{FRC}(\text{avoided}) = 12,25$. Ovvero $12,25 \cdot 1$, poiché $12,25$ sarà il costo del recupero del sintagma soggetto [$+D$, $+num_sing$, Nprop Dan] (visto che $nF = 2,5$, cioè 2 per $\langle D, num_sing \rangle + 0,5$ per la differenza tra $\langle N \rangle$ e $\langle Nprop \rangle$, $m = 2$, poiché due elementi sono in memoria in quel momento, e $dF = 0$ poiché nessun tratto espresso dal verbo distingue [the lawyer] da [Dan]).

Oppure P-D ('It was [D, num_sing, per_II, case_nom, you] who [D, num_sing, N the lawyer] avoided...'): $C_{FRC}(\text{avoided}) = 9$. Con 9 che sarà il costo del recupero del sintagma soggetto [D, num_sing, N the lawyer], visto che $nF = 2$, cioè 2 per $\langle D, num_sing \rangle$, $m = 2$, poiché due elementi sono in memoria in quel momento, e $dF = 0$ poiché nessun tratto espresso dal verbo distingue [the lawyer] da [you]).

In conclusione, a tutte le combinazioni possiamo adesso dare un costo preciso che può essere valutato proporzionalmente al rallentamento dei tempi di lettura (costi più alti dovrebbero corrispondere a tempi più alti di lettura):

tipo	D-D	D-N	D-P	N-D	N-N	N-P	P-D	P-N	P-P
ms	365	319	306	348	347	291	348	311	291
L	difficile	?	facile	?	difficile	facile	facile	facile	facile
FRC	16	12,25	4,5	12,25	16	4,5	9	9	1

Tabella 3. Comparazione fra previsioni di complessità utilizzando la teoria della località basata su tratti (L) e la nuova metrica basata sul recupero in memoria di lavoro (FRC, *Feature Retrieval Cost*).

Nella maggior parte dei casi, queste previsioni sembrano andare nella direzione giusta, migliorando in modo cruciale le aspettative di tutte le teorie fin qui discusse.

5. Conclusioni

In questo breve articolo ho cercato di illustrare il processamento di alcune frasi in cui minime variazioni dei sintagmi nominali producono apprezzabili differenze in termini di difficoltà percepita. È stato mostrato come per poter render conto di queste differenze occorra una teoria grammaticale esplicita e una derivazione precisa della struttura frasale che permetta di prevedere il corretto ordine delle parole e l'impatto, in termini di processamento in tempo reale, delle singole operazioni di costruzione della struttura: la teoria minimalista e l'applicazione

dell'idea della località in termini di tratti coinvolti nelle dipendenze non-locali ha permesso di inquadrare precisamente l'ipotesi della 'similarità' proposta da Gordon e colleghi. Alcune varianti della teoria standard si sono comunque rese necessarie per migliorare le previsioni: si è dovuto invertire la direzionalità della costruzione della struttura frasale, partendo dall'alto (complementatore), anziché dal basso (oggetto diretto). Così facendo, siamo stati in grado di prevedere dove le problematicità di un enunciato si presenteranno (cioè sul verbo della struttura relativa) e quali elementi causeranno maggiore difficoltà (cioè sintagmi nominali coinvolti nelle dipendenze non-locali che presentano una forte similarità di tratti). Sebbene l'approccio *Top-down* non sia ad oggi la versione più accreditata della teoria minimalista, i vantaggi che dalla sua adozione si traggono sono decisamente soddisfacenti.

Bibliografia

- Anderson M. C., J. H. Neely (1996). "Interference and inhibition in memory retrieval". *Memory*, 22, p. 586.
- Bever T. G. (1970). "The cognitive basis for linguistic structures", in R Hayes (cur.), *Cognition and language development*. New York: Wiley & Sons, pp. 279–362.
- Bianchi V. (2002). "Headed relative clauses in generative syntax. Part I". *Glott International*, 6.7, pp. 197–204.
- Chesi C. (2004). "Phases and cartography in linguistic computation toward a cognitively motivated computational model of linguistic competence". Tesi di dott. Università di Siena.
- Chesi C. (2012). *Competence and computation: toward a processing friendly Minimalist Grammar*. Padova: Padova Unipress.
- Chesi C. (2015). "On directionality of phrase structure building". *Journal of psycholinguistic research*, 44.1, pp. 65–89.
- Chomsky N. (1995). *The minimalist program*. Cambridge (Massachusetts): The MIT Press.
- Crowder R. G. (1976). *Principles of learning and memory*. Hillsdale: Lawrence Erlbaum.
- Dillon R. F., L. A. Bittner (1975). "Analysis of retrieval cues and release from proactive inhibition". *Journal of verbal learning and verbal behavior*, 14.6, pp. 616–622.
- Friedmann N., A. Belletti, L. Rizzi (2009). "Relativized relatives: Types of intervention in the acquisition of A-bar dependencies". *Lingua*, 119.1, pp. 67–88.
- Gardiner J. M., F. I. M. Craik, J. Birtwistle (1972). "Retrieval cues and release from proactive inhibition". *Journal of Verbal Learning and Verbal Behavior*, 11.6, pp. 778–783.
- Gibson E. (1998). "Linguistic complexity: locality of syntactic dependencies". *Cognition*, 68.1, pp. 1–76.
- Gordon P. C., R. Hendrick, M. Johnson (2001). "Memory interference during language processing". *Journal of experimental psychology: learning, memory, and cognition*, 27.6, p. 1411.

- Grillo N. (2008). *Generalized minimality: syntactic underspecification in Broca's aphasia*. Utrecht: LOT.
- Jurafsky D., J. Martin (2008). *Speech and language processing*. 2^a ed. Upper Saddle River: Prentice Hall.
- Nairne J. S. (2002). "The myth of the encoding-retrieval match". *Memory*, 10.5-6, pp. 389-395.
- Phillips C. (1996). "Order and structure". Tesi di dott. Massachusetts Institute of Technology.
- Rizzi L. (1990). *Relativized minimality*. Cambridge (Massachusetts): The MIT Press.
- Stabler E. (1997). "Derivational minimalism", in Christian Retoré (cur.), *Logical aspects of computational linguistics*. Berlin: Springer, pp. 68-95.
- Tehan G., M. S. Humphreys (1996). "Cuing effects in short-term recall". *Memory & Cognition*, 24.6, pp. 719-732.
- Van Dyke J. A., B. McElree (2006). "Retrieval interference in sentence comprehension". *Journal of Memory and Language*, 55.2, pp. 157-166.
- Warren T., E. Gibson (2005). "Effects of NP type in reading cleft sentences in English". *Language and cognitive processes*, 20.6, pp. 751-767.

Semantica distribuzionale. Un modello computazionale del significato

Alessandro Lenci – Università di Pisa – alessandro.lenci@unipi.it

1. Introduzione

I modelli semantici distribuzionali rappresentano il significato delle parole attraverso l'analisi statistica dei contesti linguistici in cui ricorrono. La similarità semantica dei lessemi è misurata sulla base della similarità delle loro distribuzioni contestuali costruite con informazioni estratte da corpora testuali. Nella semantica distribuzionale il significato è dunque una proprietà emergente dall'uso delle parole nei contesti linguistici. L'approccio distribuzionale all'analisi semantica è stato oggetto di intensa ricerca negli ultimi anni, sia in Linguistica Computazionale che nelle Scienze Cognitive. L'uso di informazioni distribuzionali estratte da corpora per misurare la similarità semantica tra termini lessicali è diventato un paradigma centrale nel Trattamento Automatico della Lingua (TAL). Lo sviluppo della semantica distribuzionale ha beneficiato della disponibilità sempre crescente di dati testuali e di una maggiore potenza di calcolo, che hanno consentito lo sviluppo di nuovi metodi di analisi dei contesti linguistici e la loro applicazione su larga scala, per sviluppare modelli realistici del lessico semantico utilizzabili sia in contesti applicativi, sia per modellazioni linguistiche e (neuro)cognitive.

Dopo un lungo percorso iniziato più di mezzo secolo fa, con periodi alterni di oblio e popolarità dovuti al succedersi di diversi paradigmi di ricerca dominanti nel TAL e in linguistica, la semantica distribuzionale ha raggiunto oggi una sua maturità. Molti algoritmi per la costruzione di spazi semantici distribuzionali e dati per la loro valutazione sono ora disponibili. Numerosi articoli di rassegna hanno contribuito a consolidare e approfondire gli aspetti metodologici dei modelli distribuzionali (Turney, Pantel, 2010; Baroni, Lenci, 2010). Negli ultimi anni, la ricerca ha concentrato molti dei suoi sforzi sullo sviluppo di metodi di ottimizzazione dei modelli per la loro applicazione a corpora di grandi dimensioni e sullo studio e valutazione dei parametri che determinano effetti significativi sulla qualità e sulla natura delle rappresentazioni semantiche costruite dai modelli distribuzionali. Tuttavia, questioni importanti rimangono ancora aperte, sia per giungere a una migliore comprensione del tipo di informazioni che questi metodi permettono di estrarre dai dati linguistici, sia per estendere la loro applicazione alla modellazione di nuovi fenomeni semantici.

Lo scopo di questo contributo è fornire una breve rassegna dello stato dell'arte della semantica distribuzionale nell'ambito dei modelli computazionali dell'analisi linguistica. Dopo una presentazione dei meccanismi principali alla base dei modelli distribuzionali e della loro valutazione sperimentale, illustreremo alcune sfide che questa metodologia di analisi semantica deve affrontare.

2. Principi e metodi di semantica distribuzionale

2.1. L'Ipotesi Distribuzionale

Il fondamento epistemologico della semantica distribuzionale è l'Ipotesi Distribuzionale: la similarità semantica si correla con la similarità delle distribuzioni nei contesti linguistici. Zellig S. Harris è normalmente ritenuto il pioniere dell'Ipotesi Distribuzionale (Harris, 1954). Infatti egli considerava la metodologia distribuzionale come l'unico approccio scientifico possibile allo studio del significato linguistico. Nei suoi ultimi lavori, Harris propose un metodo per classificare le parole sulla base dei loro contesti attraverso la raccolta e l'analisi delle relazioni di dipendenza sintattica, espresse in termini di operatori e argomenti (Harris, 1991).

A partire dagli anni Sessanta molte implementazioni dell'Ipotesi Distribuzionale sono state realizzate per la costruzione automatica di *thesauri* (G. Grefenstette, 1994). Un contributo cruciale allo sviluppo della semantica distribuzionale è infatti giunto dal *vector space model* nell'*Information Retrieval* (Salton et al., 1975), che ha permesso di apportare miglioramenti fondamentali alla metodologia originale di Harris rispetto sia alla natura dei dati che alla loro formalizzazione matematica, accelerando così la diffusione del paradigma distribuzionale in linguistica computazionale. Negli ultimi venti anni, la possibilità di applicare tale metodologia su ampia scala a corpora di grandi dimensioni ha fatto sì che l'approccio distribuzionale sia diventato il paradigma semantico di riferimento nel TAL.

2.2. La struttura dei modelli semantici distribuzionali

Nei modelli distribuzionali, le parole sono rappresentate come vettori costruiti a partire dalla loro distribuzione nei contesti linguistici, e la similarità tra le parole è approssimata attraverso la misura della distanza geometrica tra vettori. Il metodo standard per costruire modelli semantici distribuzionali è tipicamente formato da quattro fasi principali (Turney, Pantel, 2010):

1. per ciascuna parola target, vengono raccolti e contati i contesti generando così una matrice di co-occorrenza;
2. le frequenze sono poi trasformate in pesi statistici in grado di riflettere meglio l'importanza dei contesti;
3. la matrice che si ottiene è molto grande e sparsa, ovvero la maggior parte delle sue entrate è zero. Per tale motivo, vengono applicate tecniche matematiche per ridurre il numero delle sue dimensioni;
4. la similarità semantica delle parole target viene misurata attraverso la similarità dei corrispondenti vettori riga nella matrice.

I vettori costituiscono il principale strumento di rappresentazione matematica del contenuto lessicale nella semantica distribuzionale. Un vettore è una lista ordinata di numeri reali $(v_1, \dots, v_n)$, nella quale ciascun valore v_i è l' i -esima componente del vettore. I vettori hanno interpretazioni geometriche. Se si assume un sistema di

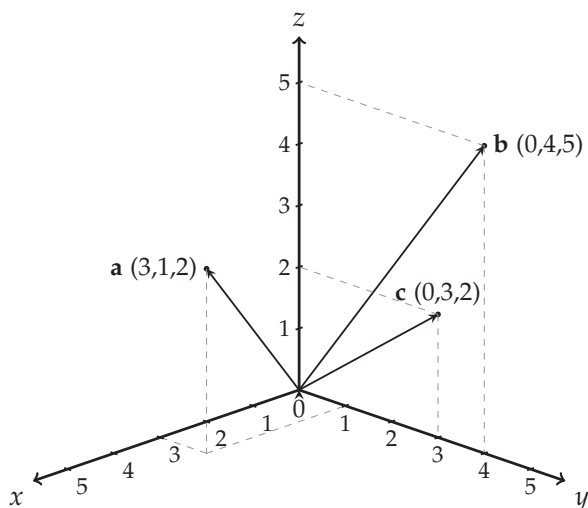


Figura 4. Vettori in uno spazio tridimensionale

assi cartesiani come quello nella Figura 4., i vettori **a**, **b**, e **c** identificano punti nello spazio e le loro componenti corrispondono alle coordinate dei punti sugli assi. Gli stessi vettori possono essere anche rappresentati da frecce che congiungono un punto origine a un punto finale. L'origine è fissata all'incrocio degli assi cartesiani, e le coordinate del punto finale corrispondono alle componenti del vettore. I modelli distribuzionali forniscono, dunque, una rappresentazione geometrica del lessico come uno spazio vettoriale semantico.

Ad esempio, supponiamo di aver contato in un corpus quante volte i nomi *auto*, *gatto*, *cane* e *camion* co-occorrono con i verbi *mangiare*, *guidare* e *correre*, ottenendo la seguente distribuzione di frequenza:

<i>auto</i>	<i>guidare</i>	3	<i>cane</i>	<i>mangiare</i>	3
<i>auto</i>	<i>correre</i>	2	<i>cane</i>	<i>correre</i>	4
<i>gatto</i>	<i>mangiare</i>	4	<i>camion</i>	<i>guidare</i>	2
<i>gatto</i>	<i>correre</i>	3	<i>camion</i>	<i>correre</i>	3

Rappresentiamo dunque *auto*, *gatto*, *cane* e *camion* con i seguenti vettori, indicando il vettore distribuzionale di un lessema con il lessema stesso in grassetto:

auto = (0, 3, 2)
gatto = (4, 0, 3)
cane = (3, 0, 4)
camion = (0, 2, 3)

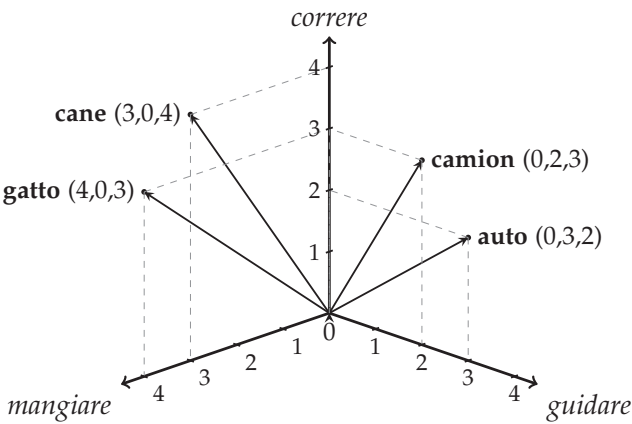


Figura 5. Vettori distribuzionali

La prima componente dei vettori è la frequenza di co-occorrenza con *mangiare*, la seconda con *guidare* e la terza con *correre* (zero indica il caso in cui una parola non ricorra mai in un dato contesto). Questi vettori corrispondono a punti (o frecce) nello spazio tridimensionale illustrato nella Figura 5.. Gli assi dello spazio sono etichettati con contesti linguistici (in questo caso specifico le parole *mangiare*, *guidare* e *correre*), e la posizione delle parole target nello spazio è determinata dalla loro frequenza di co-occorrenza. Come si è detto, invece delle frequenze è possibile usare vari tipi di pesi statistici (cfr. *infra*, Sezione 2.2.1.).

Una delle misure più comuni di similarità distribuzionale è il coseno dell’angolo θ tra due vettori:

$$\text{sim}_{\cos}(\mathbf{u}, \mathbf{v}) = \cos(\theta) = \frac{\mathbf{u} \cdot \mathbf{v}}{\|\mathbf{u}\| \|\mathbf{v}\|} = \frac{\sum_{i=1}^n u_i v_i}{\sqrt{\sum_{i=1}^n u_i^2} \sqrt{\sum_{i=1}^n v_i^2}}$$

I coseni dei vettori delle quattro parole target sono riportati nella tabella seguente:

<i>auto</i>	1			
<i>gatto</i>	0.33	1		
<i>cane</i>	0.44	0.96	1	
<i>camion</i>	0.92	0.50	0.66	1
	<i>auto</i>	<i>gatto</i>	<i>cane</i>	<i>camion</i>

Se due vettori sono geometricamente allineati sulla stessa linea e puntano nella medesima direzione, l’angolo tra di loro misura 0 gradi e il coseno è 1 (massima similarità); viceversa, se i due vettori sono indipendenti (ortogonali), il loro angolo è vicino a 90 gradi e il coseno è uguale a 0 (assenza di similarità). Il coseno

rappresenta dunque la similarità in termini geometrici e misura la similarità tra vettori con la loro vicinanza nello spazio. Nella Figura 5., **cane** e **gatto** sono molto vicini e puntano nella stessa direzione, perché hanno valori simili nelle stesse componenti. L'angolo che essi formano è dunque più piccolo di quello con **auto** e **camion**: più piccolo l'angolo, maggiore il coseno e la similarità tra i vettori.

I modelli distribuzionali hanno molteplici opzioni e possibilità di realizzazione, conseguenti a vari parametri in ciascuna fase del processo di costruzione. I valori di tali parametri possono modificare in maniera anche molto significativa la struttura dello spazio semantico.

2.2.1. Parametri

Un modello semantico distribuzionale riflette il comportamento delle parole nell'uso linguistico ed è quindi, per definizione, fortemente dipendente dal tipo di corpus che viene analizzato. Mentre negli anni Novanta venivano usati corpora specializzati di medie dimensioni per l'acquisizione e costruzione di *thesauri* distribuzionali (G. Grefenstette, 1994; Nazarenko et al., 2001), più di recente si è passati all'uso di corpora di grandi dimensioni, trasversali rispetto ai generi testuali e ai domini tematici. Sono stati usati in semantica distribuzionale corpora giornalistici o formati da articoli di enciclopedia (Peirsman, Geeraerts, 2009), corpora di riferimento bilanciati come il British National Corpus (Sadrzadeh, E. Grefenstette, 2011), grandi corpora ottenuti dal Web (Agirre et al., 2009), o combinazioni di corpora di varie tipologie (Baroni, Lenci, 2010). La tendenza a usare corpora di grandi dimensioni è principalmente motivata dalla necessità di aumentare la copertura delle risorse lessicali distribuzionali riducendo al contempo la sparsità dei dati, che notoriamente può avere effetti negativi sulla qualità degli spazi semantici.

Un altro parametro cruciale nell'implementazione dei modelli semantici distribuzionali è la definizione dei contesti. Tre tipi di contesti linguistici sono tipicamente usati: nei modelli basati su documenti (*document models*), come la *Latent Semantic Analysis* (LSA) (Landauer, Dumais, 1997), le parole sono simili se appaiono negli stessi documenti o negli stessi paragrafi; nei modelli basati sulle parole (*word models*), viene invece considerata una finestra di parole che ricorrono intorno ai lessemi target (Lund, Burgess, 1997; Sahlgren, 2008; Ferret, 2013); i modelli sintattici (*syntactic models*) sono vicini all'approccio originale di Harris, poiché usano come contesti le relazioni di dipendenza sintattica delle parole target (Curran, 2004; Padó, Lapata, 2007; Baroni, Lenci, 2010). I modelli basati sulle parole hanno un parametro aggiuntivo rappresentato dalla dimensione della finestra per la selezione delle parole contesto, finestra che può andare da poche parole a un intero paragrafo, mentre i modelli sintattici hanno bisogno di specificare le relazioni di dipendenza che sono selezionate per definire i contesti di co-occorrenza (Baroni, Lenci, 2010; Peirsman, Heylen et al., 2007). Alcuni esperimenti suggeriscono che i modelli sintattici tendano a identificare vicini distribuzionali che sono tassonomicamente correlati al target, principalmente co-ponimi, mentre i modelli basati sulle parole sono più orientati verso l'identificazione di relazioni associative (Van de Cruys, 2008; Peirsman, Heylen et al., 2007; O. Levy, Goldberg, 2014). La questione se i contesti sintattici forniscano veramente un vantaggio reale rispetto ai

modelli basati su finestre di parole è comunque ancora aperta. Una differenza molto più sostanziale esiste invece nei confronti dei modelli basati sui documenti, che sono fortemente orientati verso l'identificazione di vicini distribuzionali appartenenti a domini tematici ad ampio spettro (Sahlgren, 2006).

Altri parametri dei modelli distribuzionali hanno ricevuto particolare attenzione: i pesi statistici dei contesti e le misure di similarità dei vettori. Un'ampia gamma di scelte esiste per entrambi (Curran, 2004; Bullinaria, J. P. Levy, 2007), ma oggi la pratica più comune consiste nell'usare la *Positive Pointwise Mutual Information* come peso statistico per misurare la salienza dei contesti e il coseno come misura di similarità semantica. Questi infatti sono in grado di offrire le migliori prestazioni in vari tipi di esperimenti semantici (Turney, Pantel, 2010).

I vettori nella matrice di co-occorrenza forniscono una rappresentazione esplicita della distribuzione nei contesti (O. Levy, Goldberg, 2014). Ciascuna dimensione del vettore, infatti, corrisponde a un contesto specifico nel quale la parola target è stata osservata. I vettori di concorrenza espliciti hanno molte dimensioni (tipicamente nell'ordine delle centinaia di migliaia o anche più) e sono sparsi. Vari tipi di tecniche sono perciò usate per ridurre le dimensioni dei vettori e limitare così la complessità computazionale. L'approccio più comune consiste nel proiettare l'originale matrice sparsa in una matrice densa a ridotta dimensionalità usando metodi come *Singular Value Decomposition* (Landauer, Dumais, 1997), *Non-Negative Matrix Factorization* (Van de Cruys, 2010), e *Latent Dirichlet Allocation* (Blei et al., 2003). Crucialmente, le dimensioni dei vettori ridotti non corrispondono più a contesti espliciti, ma piuttosto a dimensioni semantiche 'nascoste', ovvero implicite, nell'originale distribuzione dei dati. Le tecniche di riduzione di matrici permettono di rimuovere il rumore nei dati estratti dai corpora e di sfruttare le ridondanze e correlazioni tra i contesti linguistici, migliorando così la qualità degli spazi semantici (Turney, Pantel, 2010). Un'alternativa molto popolare ed efficace è il *Random Indexing* (Sahlgren, 2006): invece di ridurre una matrice costruita in precedenza, rappresentazioni vettoriali con un numero ridotto di dimensioni sono costruite incrementalmente assegnando a ciascuna parola un vettore casuale che viene poi sommato ai vettori delle parole che co-occorrono con essa.

Molte ricerche sono state dedicate a comprendere l'impatto di questi parametri sulle performance dei modelli distribuzionali. Gli studi più recenti e comprensivi sono quelli di Lapesa, Evert (2014) e Kiela, Clark (2014), che analizzano un'ampia gamma di parametri quali il tipo di corpus, l'uso di metodologie di lemmatizzazione o *stemming* del testo, la tipologia dei contesti (dipendenze sintattiche contro co-occorrenze, direzione e dimensione della finestra, ecc.), i pesi statistici dei contesti, le misure di similarità e le tecniche di riduzione delle dimensioni dei vettori. Questi esperimenti permettono di individuare le configurazioni migliori dei parametri dei modelli distribuzionali in vari compiti semantici.

2.2.2. Contare o predire

I modelli che abbiamo descritto sopra usano un approccio alla costruzione delle rappresentazioni distribuzionali basato sul conteggio delle frequenze delle parole nei testi: le co-occorrenze estratte da corpora sono contate, poi pesate e infine

opzionalmente ridotte per costruire dei vettori densi. Per tale motivo sono anche definiti *count models*. Recentemente è apparsa una nuova famiglia di modelli distribuzionali basata su una tecnica di predizione: algoritmi neurali creano direttamente rappresentazioni distribuzionali dense e a bassa dimensionalità, imparando a predire in maniera ottimale i contesti di una parola target (Mikolov et al., 2013). Per tale motivo sono anche definiti *prediction models* e le rappresentazioni costruite da essi *embedding*, perché le parole sono *embedded* ('incassate') dentro uno spazio lineare a bassa dimensionalità formato da *feature* latenti. Vari tipi di regolarità linguistiche sono stati identificati negli spazi semantici formati dagli *embedding* neurali: ad esempio, il fatto che *re* e *regina* hanno la stessa relazione di genere di *uomo* e *donna* viene rappresentato attraverso le relazioni tra i rispettivi vettori nello spazio. In tal modo, il vettore di una parola (es. **regina**) può essere ricostruito dalle rappresentazioni delle altre parole attraverso una semplice aritmetica vettoriale (es., **re** – **uomo** + **donna**). Alcuni esperimenti hanno anche mostrato che i *prediction model* sono in grado di superare i *count models* in vari *task* (Baroni, Dinu et al., 2014).

Nonostante la loro crescente popolarità, la questione se gli *embedding* neurali siano una reale innovazione rispetto ai modelli più tradizionali è ancora lontana dall'essere risolta. Per esempio, le stesse regolarità catturate dai *prediction model* sono anche catturate dai modelli basati su conteggi di frequenza espliciti (O. Levy, Goldberg, 2014). Quando i parametri di questi ultimi sono accuratamente raffinati, non vengono osservate differenze significative nelle prestazioni tra le due tipologie di modelli (O. Levy, Goldberg, Dagan, 2015). È possibile che la ricerca futura dimostrerà se gli *embedding* offrano dei chiari vantaggi, ma per ora i due approcci non differiscono in maniera sostanziale per gli aspetti del significato che sono in grado di catturare. Sono semplicemente modi alternativi di costruire rappresentazioni distribuzionali.

3. Valutare i modelli semantici distribuzionali

La dicotomia classica tra metodi intrinseci e metodi estrinseci di valutazione nel TAL si applica anche alla semantica distribuzionale. Le valutazioni intrinseche hanno come obiettivo quello di misurare la qualità della risorsa stessa, confrontandola con i giudizi di valutatori umani o con simili risorse semantiche che possono essere utilizzate come standard di riferimento (*gold standard*). Le valutazioni estrinseche misurano invece il contributo specifico della risorsa per migliorare le performance di un sistema in cui è inserita.

La valutazione intrinseca dei modelli semantici distribuzionali viene realizzata attraverso la comparazione con risorse lessicali come i sinonimi del TOEFL (Landauer, Dumais, 1997), *thesauri* specializzati (G. Grefenstette, 1994), Wordnet (Curran, Moens, 2002; Padró et al., 2014; Anguiano, Denis, 2011), dizionari di sinonimi (Van der Plas, Tiedemann, Manguin, 2011), ecc. La valutazione intrinseca dei modelli di semantica distribuzionale è una questione complessa per molteplici ragioni. Prima di tutto, i modelli distribuzionali catturano una nozione molto ampia di vicinanza semantica (cfr. *infra*, Sezione 4.). Esiste dunque uno slittamento

inevitabile tra i risultati prodotti dai modelli computazionali e le risorse che invece contengono relazioni lessicali di tipo classico, quali dizionari dei sinonimi, *thesauri* e Wordnet. Un secondo tipo di problema è dovuto al fatto che i modelli semantici distribuzionali riflettono le specificità del corpus da cui sono costruiti, e possono dunque identificare relazioni che mancano in risorse lessicali di tipo generale. È infatti difficile, forse impossibile, verificare la validità di una relazione semantica fuori contesto (Muller et al., 2014). Allo scopo di superare tali limiti, sono state sviluppate risorse specificatamente orientate verso la valutazione dei modelli semantici distribuzionali. Uno degli standard di riferimento più usati è WordSim-353 (Finkelstein et al., 2002), con 353 coppie di parole inglesi associate a un punteggio di similarità semantica espresso da valutatori umani. Recentemente è stata anche realizzata una versione multilingue di questo *dataset* (Leviant, Reichart, 2015).

Per quanto riguarda la valutazione estrinseca, l'uso di *feature* distribuzionali è utile ogni volta in cui è necessario misurare la similarità tra parole o tra porzioni di testo più ampie. Vari esperimenti sono stati dedicati all'uso delle risorse distribuzionali nell'*Information Retrieval* per misurare la similarità tra le *query* (Alfonseca et al., 2009; Claveau, Kijak, 2015). Nei compiti di sostituzione lessicale (McCarthy, Navigli, 2007), i modelli distribuzionali sono usati per identificare potenziali sostituti di una parola prima del processo di disambiguazione (Fabre et al., 2014). La similarità distribuzionale è anche utilizzata per determinare il senso di una parola in un corpus (McCarthy, Koeling et al., 2007). Modelli di semantica distribuzionale si sono dimostrati efficaci in applicazioni di TAL più complesse come il *Textual Entailment* e la generazione automatica di riassunti (Cheung, Penn, 2013). Inoltre, gli *embedding* neurali sono stati usati con successo per migliorare sistemi di *Semantic Role Labeling* e *Named Entity Recognition* (Collobert, Weston, 2008).

4. Le sfide per la semantica distribuzionale

I modelli semantici distribuzionali sono stati oggetto di molte critiche, anche all'interno della comunità della linguistica computazionale. L'approccio esclusivamente induttivo al significato della semantica distribuzionale è molto utile dal punto di vista del TAL, ma rimane una questione aperta se le statistiche di co-occorrenza da sole siano sufficienti per affrontare aspetti semantici più profondi, oppure siano soltanto in grado di fornire un'approssimazione molto superficiale del significato lessicale (Sahlgren, 2008; Lenci, 2008; Koller, 2015).

L'Ipotesi Distribuzionale è di fatto una definizione della similarità semantica in termini di prossimità in uno spazio. La similarità semantica è però essa stessa una nozione molto vaga, che va dalla similarità tra parole alla similarità tra relazioni (Turney, 2006; Baroni, Lenci, 2010; Turney, 2013). La similarità semantica in senso stretto, come relazione tra parole che condividono simili tratti semantici (es. *auto* e *camion*), deve inoltre essere distinta dall'associazione semantica tra parole, come *auto* e *ruota* (Budanitsky, Hirst, 2006; Agirre et al., 2009). Questi due tipi di relazioni semantiche hanno proprietà molto differenti e tuttavia sono discriminate a fatica dai modelli semantici distribuzionali. Anche *gold standard*

come WordSim-353 contengono in realtà molte coppie di parole associate che non sono semanticamente simili in senso stretto (Agirre et al., 2009). Per affrontare questo problema, è stato recentemente sviluppato il *dataset* SimLex-999 allo scopo di valutare specificatamente la capacità di modelli semantici distribuzionali di catturare la similarità semantica piuttosto che relazioni di associazione semantica (Hill et al., 2015).

Un problema aggiuntivo è dato dal fatto che sia la similarità semantica che l'associazione semantica sono in realtà termini che coprono tipi molto diversi di relazioni lessicali. Per esempio, i sinonimi, i co-iponimi e perfino gli antonimi possono essere detti semanticamente simili perché condividono un elevato numero di tratti. Le associazioni semantiche dall'altro lato includono la metonimia, relazioni locative, o la semplice appartenenza al medesimo campo semantico (Morris, Hirst, 2004). Questa nozione ampia e graduata di associazione è al tempo stesso utile e problematica per le applicazioni del TAL, perché è molto difficile tracciare un limite chiaro tra le relazioni associative rilevanti e quelle non rilevanti (Sahlgren, 2008; Ferret, 2013). In generale, i vicini distribuzionali identificati dei modelli semantici distribuzionali hanno relazioni molto diverse con la parola target, suggerendo che i modelli semantici distribuzionali forniscono in realtà una rappresentazione 'a grana grossa' del significato lessicale. Recentemente i *dataset* BLESS (Baroni, Lenci, 2011) e EVALution (Santus, Yung et al., 2015) sono stati progettati proprio per valutare l'abilità dei modelli distribuzionali di discriminare differenti tipi di relazioni lessicali. Questa rappresenta un'importante area di ricerca nella semantica distribuzionale (Van der Plas, Tiedemann, 2006; Lenci, Benotto, 2012; Santus, Lenci et al., 2014; The Pham et al., 2015).

Una questione fondamentale è determinare quale tipologia di informazione semantica può essere catturata sulla base delle proprietà contestuali e quali parti del significato delle parole rimangono invece al di là delle possibilità dei modelli distribuzionali, a meno di non integrare le co-occorrenze contestuali con altre informazioni. Vari lavori recenti si focalizzano su questo problema. Gupta et al. (2015) mostrano che alcuni aspetti dell'informazione referenziale sono accessibili dal punto di vista distribuzionale, mentre Herbelot, Ganesalingam (2013) suggeriscono che l'informatività dei lessemi (ovvero la distinzione tra parole con maggiore o minore peso informativo) è difficile da ricavare su base distribuzionale. Zarcone et al. (2015) mostrano che non solo la prototipicità semantica di un argomento rispetto al predicato (il cosiddetto *thematic fit*), ma anche i vincoli determinati dal tipo semantico richiesto dal predicato possono essere catturati da modelli distribuzionali per rappresentare fenomeni di *coercion* e metonimia logica. In una prospettiva simile, lavori recenti propongono di creare un ponte tra semantica formale e distribuzionale (Guevara, 2011; E. Grefenstette, 2013), in maniera da combinare le rappresentazioni distribuzionali dei significati lessicali con le capacità inferenziali dei sistemi formali (Erk, 2013; Boleda, Erk, 2015).

4.1. Oltre le parole

La maggior parte dei lavori sulla semantica distribuzionale è dedicata all'analisi delle parole in isolamento, ma negli ultimi anni la ricerca si è anche focalizzata

sull'estensione di questi modelli per rappresentare unità semantiche più ampie, come sintagmi e frasi, seguendo due approcci principali. Il primo consiste nel considerare i sintagmi in aggiunta alle parole come elementi target a cui viene associata una rappresentazione vettoriale unitaria (es. Baldwin et al., 2003). Tuttavia questo rimane un approccio largamente minoritario a causa della scarsità di dati per ricostruire la distribuzione di unità complesse. La seconda opzione consiste nel modellare la composizionalità semantica all'interno del paradigma distribuzionale, sulla base dell'assunzione che l'informazione semantica sui sintagmi può essere computata combinando informazioni sulle rappresentazioni vettoriali dei loro componenti. Mitchell, Lapata (2010) propongono un modello generale per la semantica distribuzionale composizionale e analizzano diverse funzioni di combinazione vettoriale. Più di recente, è stato anche proposto un *task* di valutazione di aspetti legati alla semantica composizionale nell'ambito delle campagne SemEval (Marelli et al., 2014). Diversi tipi di unità sintagmatiche sono state esaminate, quali le coppie aggettivo-nome (Baroni, Zamparelli, 2010), verbo-nome (Mitchell, Lapata, 2010), e le frasi. Un altro settore di ricerca estremamente promettente riguarda l'uso di *feature* extralinguistiche per integrare i dati distribuzionali con altri tipi di informazione multimodale (Bruni et al., 2014).

5. Conclusioni e prospettive

La semantica distribuzionale è un paradigma scientifico giovane, ma nonostante la sua breve storia è stata in grado di guadagnare un grande successo nella comunità del TAL, con interessi crescenti nella ricerca linguistica e nelle scienze cognitive. Come è stato illustrato nelle sezioni precedenti, la varietà di modelli semantici distribuzionali sta aumentando rapidamente. Inoltre si è ottenuta una comprensione molto più profonda degli effetti e dei ruoli dei vari tipi di parametri rilevanti per le rappresentazioni distribuzionali. Certamente la semantica distribuzionale deve affrontare ancora molte sfide. Sotto molti punti di vista, i modelli semantici distribuzionali tuttora forniscono una rappresentazione molto grezza del significato, e i loro limiti (così come le loro potenzialità) devono essere tuttora esplorate. Al tempo stesso, il numero dei compiti semantici che sono affrontati da questi modelli è costantemente in espansione, andando ben al di là delle originali applicazioni dell'Ipotesi Distribuzionale per l'identificazione dei sinonimi. Tutto ciò rende questo ambito di ricerca estremamente vitale e promettente per ottenere una maggiore comprensione del significato linguistico.

Bibliografia

- Agirre E., E. Alfonseca, K. Hall, J. Kravalova, M. Paşca, A. Soroa (2009). "A study on similarity and relatedness using distributional and WordNet-based approaches". *Proceedings of Human Language Technologies: the 2009 Annual Conference of the North American Chapter of the Association for Computational Linguistics*, pp. 19–27.

- Alfonseca E., K. Hall, S. Hartmann (2009). "Large-scale computation of distributional similarities for queries". *Proceedings of Human Language Technologies: The 2009 Annual Conference of the North American Chapter of the Association for Computational Linguistics, Companion Volume: Short Papers*, pp. 29–32.
- Anguiano E. H., P. Denis (2011). "FreDist: Automatic construction of distributional thesauri for French". *18ème conférence sur le traitement automatique des langues naturelles (TALN 2011)*, pp. 119–124.
- Baldwin T., C. Bannard, T. Tanaka, D. Widdows (2003). "An empirical model of multiword expression decomposability". *Proceedings of the ACL 2003 workshop on multiword expressions: analysis, acquisition and treatment*. Vol. 18, pp. 89–96.
- Baroni M., G. Dinu, G. Kruszewski (2014). "Don't count, predict! A systematic comparison of context-counting vs. context-predicting semantic vectors". *Proceedings of the 52nd Annual Meeting of the Association for Computational Linguistics*, pp. 238–247.
- Baroni M., A. Lenci (2010). "Distributional memory: a general framework for corpus-based semantics". *Computational Linguistics*, 36.4, pp. 673–721.
- Baroni M., A. Lenci (2011). "How we BLESSed distributional semantic evaluation". *Proceedings of the GEMS 2011 Workshop on GEometrical Models of Natural Language Semantics*, pp. 1–10.
- Baroni M., R. Zamparelli (2010). "Nouns are vectors, adjectives are matrices: representing adjective-noun constructions in semantic space". *Proceedings of the 2010 Conference on Empirical Methods in Natural Language Processing*, pp. 1183–1193.
- Blei D. M., A. Y. Ng, M. I. Jordan (2003). "Latent Dirichlet Allocation". *Journal of machine learning research*, 3, pp. 993–1022.
- Boleda G., K. Erk (2015). "Distributional semantic features as semantic primitives – or not". Intervento all'AAAI Spring Symposium on Knowledge Representation and Reasoning.
- Bruni E., N.-K. Tran, M. Baroni (2014). "Multimodal distributional semantics". *Journal of artificial intelligence research (JAIR)*, 49, pp. 1–47.
- Budanitsky A., G. Hirst (2006). "Evaluating Wordnet-based measures of lexical semantic relatedness". *Computational linguistics*, 32.1, pp. 13–47.
- Bullinaria J., J. P. Levy (2007). "Extracting semantic representations from word co-occurrence statistics: A computational study". *Behavior research methods*, 39 (3), pp. 510–526.
- Cheung J. C. K., G. Penn (2013). "Probabilistic domain modelling with contextualized distributional semantic vectors". *Proceedings of ACL 2013*, pp. 392–401.
- Claveau V., E. Kijak (2015). "Thésaurus distributionnels pour la recherche d'information et vice-versa". *Actes de la 13ème Conférence en Recherche d'Information et Applications (CORIA 2015)*.
- Collobert R., J. Weston (2008). "A unified architecture for Natural Language Processing: Deep Neural Networks with Multitask Learning". *Proceedings of the 25th International Conference on Machine Learning*, pp. 160–167.

- Curran J. R., M. Moens (2002). "Improvements in automatic thesaurus extraction". *Proceedings of the ACL-02 workshop on Unsupervised lexical acquisition*. Vol. 9, pp. 59–66.
- Curran J. R. (2004). "From distributional to semantic similarity". Tesi di dott. University of Edinburgh.
- Erk K. (2013). "Towards a semantics for distributional representations". *Proceedings, 10th International Conference on Computational Semantics (IWCS-2013)*.
- Fabre C., N. Hathout, L.-M. Ho-Dac, F. Morlane-Hondée, P. Muller, F. Sajous, L. Tanguy, T. Van de Cruys (2014). "Présentation de l'atelier SemDis 2014: sémantique distributionnelle pour la substitution lexicale et l'exploration de corpus spécialisés". *Actes de la conférence Traitement Automatique du Langage Naturel (TALN 2014)*, pp. 196–205.
- Ferret O. (2013). "Identifying bad semantic neighbors for improving distributional thesauri". *51st Annual Meeting of the Association for Computational Linguistics (ACL 2013)*, pp. 561–571.
- Finkelstein L., E. Gabrilovich, Y. Matias, E. Rivlin, Z. Solan, G. Wolfman, E. Ruppín (2002). "Placing Search in Context: The Concept Revisited". *ACM Transactions on Information Systems*, 20.1, pp. 116–131.
- Grefenstette E. (2013). "Towards a formal distributional semantics: simulating logical calculi with tensors". *Proceedings of the Second Joint Conference on Lexical and Computational Semantics*. Pp. 1–10.
- Grefenstette G. (1994). *Explorations in automatic thesaurus discovery*. Norwell: Kluwer Academic Publishers.
- Guevara E. (2011). "Computing semantic compositionality in distributional semantics". *Proceedings of the Ninth International Conference on Computational Semantics*, pp. 135–144.
- Gupta A., G. Boleda, M. Baroni, S. Padó (2015). "Mapping conceptual features to referential properties". Intervento alla Conference on Empirical Methods in Natural Language Processing (EMNLP 2015).
- Harris Z. S. (1954). "Distributional structure". *Word*, 10.2-3, pp. 146–162.
- Harris Z. S. (1991). *A theory of language and information: a mathematical approach*. Oxford: Clarendon Press.
- Herbelot A., M. Ganesalingam (2013). "Measuring semantic content in distributional vectors". *Proceedings of the 51st Annual Meeting of the Association for Computational Linguistics*. Vol. 2, pp. 440–445.
- Hill F., R. Reichart, A. Korhonen (2015). "SimLex-999: Evaluating Semantic Models with (Genuine) Similarity Estimation". *Computational Linguistics*, 41 (4), pp. 1–32.
- Kiela D., S. Clark (2014). "A systematic study of semantic vector space model parameters". *Proceedings of the 2nd Workshop on Continuous Vector Space Models and their Compositionality (CVSC) at EACL*, pp. 21–30.
- Koller A. (2015). "Top-down questions for distributional semantics". Intervento al Workshop on formal and distributional semantics.

- Landauer T. K., S. T. Dumais (1997). "A solution to Plato's problem: The latent semantic analysis theory of acquisition, induction, and representation of knowledge". *Psychological review*, 104.2, p. 211.
- Lapesa G., S. Evert (2014). "A large scale evaluation of distributional semantic models: Parameters, interactions and model selection". *Transactions of the Association for Computational Linguistics*, 2, pp. 531–545.
- Lenci A. (2008). "Distributional semantics in linguistic and cognitive research". *Italian journal of linguistics*, 20.1: *From context to meaning: distributional models of the lexicon in linguistics and cognitive science*, pp. 1–31.
- Lenci A., G. Benotto (2012). "Identifying hypernyms in distributional semantic spaces". **SEM 2012: The First Joint Conference on Lexical and Computational Semantics (SemEval 2012)*, pp. 75–79.
- Leviant I., R. Reichart (2015). "Judgment language matters: multilingual vector space models for judgment language aware lexical semantics". *ArXiv e-prints*. arXiv: 1508.00106 [cs.CL].
- Levy O., Y. Goldberg (2014). "Linguistic regularities in sparse and explicit word representations". *Proceedings of the Eighteenth Conference on Computational Language Learning (CoNLL)*, pp. 171–180.
- Levy O., Y. Goldberg, I. Dagan (2015). "Improving distributional similarity with lessons learned from word embeddings". *Transactions of the ACL*, 3, pp. 211–225.
- Lund C., K. Burgess (1997). "Modelling parsing constraints with high-dimensional context space". *Language and cognitive processes*, 12.2-3, pp. 177–210.
- Marelli M., L. Bentivogli, M. Baroni, R. Bernardi, S. Menini, R. Zamparelli (2014). "Semeval-2014 task 1: evaluation of compositional distributional semantic models on full sentences through semantic relatedness and textual entailment". *Proceedings of the 8th International Workshop on Semantic Evaluation (SemEval 2014)*, pp. 1–8.
- McCarthy D., R. Koeling, J. Weeds, J. Carroll (2007). "Unsupervised acquisition of predominant word senses". *Computational Linguistics*, 33.4, pp. 553–590.
- McCarthy D., R. Navigli (2007). "Semeval-2007 task 10: English lexical substitution task". *Proceedings of the 4th International Workshop on Semantic Evaluations*, pp. 48–53.
- Mikolov T., K. Chen, G. Corrado, J. Dean (2013). "Efficient estimation of word representations in vector space". *Proceedings of Workshop at ICLR 2013*, pp. 1–12.
- Mitchell J., M. Lapata (2010). "Composition in distributional models of semantics". *Cognitive Science*, 34.8, pp. 1388–1439.
- Morris J., G. Hirst (2004). "Non-classical lexical semantic relations". *Proceedings of the HLT-NAACL Workshop on Computational Lexical Semantics*, pp. 46–51.
- Muller P., C. Fabre, C. Adam (2014). "Predicting the relevance of distributional semantic similarity with contextual information". *52nd Annual Meeting of the Association for Computational Linguistics (ACL 2014)*, pp. 479–488.

- Nazarenko A., P. Zweigenbaum, B. Habert, J. Bouaud (2001). "Corpus-based extension of a terminological semantic lexicon". *Recent Advances in Computational Terminology*, pp. 327–351.
- Padó S., M. Lapata (2007). "Dependency-based construction of semantic space models". *Computational Linguistics*, 33.2, pp. 161–199.
- Padró M., M. Idiart, C. Ramisch, A. Villavicencio (2014). "Nothing like good old frequency: studying context filters for distributional thesauri". *Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pp. 419–424.
- Peirsman Y., D. Geeraerts (2009). "Predicting strong associations on the basis of corpus data". *Proceedings of the 12th Conference of the European Chapter of the Association for Computational Linguistics*, pp. 648–656.
- Peirsman Y., K. Heylen, D. Speelman (2007). "Finding semantically related words in Dutch. Cooccurrences versus syntactic contexts". *Proceedings of the CoSMO workshop at CONTEXT-07*, pp. 9–16.
- Sadrzadeh M., E. Grefenstette (2011). "A compositional distributional semantics, two concrete constructions, and some experimental evaluations", *Quantum Interaction*, pp. 35–47.
- Sahlgren M. (2006). "The word-space model". Tesi di dott. University of Stockholm.
- Sahlgren M. (2008). "The distributional hypothesis". *Italian Journal of Linguistics*, 20.1, pp. 33–54.
- Salton G., A. Wong, C.-S. Yang (1975). "A vector space model for automatic indexing". *Communications of the ACM*, 18.11, pp. 613–620.
- Santus E., A. Lenci, Q. Lu, S. Schulte im Walde (2014). "Chasing hypernyms in vector spaces with entropy". *Proceedings of the 14th Conference of the European Chapter of the Association for Computational Linguistics*. Vol. 2, pp. 38–42.
- Santus E., F. Yung, A. Lenci, C.-R. Huang (2015). "EVALution 1.0: an evolving semantic dataset for training and evaluation of distributional semantic models". *Proceedings of the 4th Workshop on Linked Data in Linguistics: Resources and Applications*, pp. 64–69.
- The Pham N., A. Lazaridou, M. Baroni (2015). "A multitask objective to inject lexical contrast into distributional semantics". Intervento al 53rd Annual Meeting of the Association for Computational Linguistics (ACL 2015).
- Turney P. D. (2006). "Similarity of semantic relations". *Computational Linguistics*, 32.3, pp. 379–416.
- Turney P. D. (2013). "Distributional semantics beyond words: supervised learning of analogy and paraphrase". *Transactions of the Association for Computational Linguistics (TACL)*, pp. 353–366.
- Turney P. D., P. Pantel et al. (2010). "From frequency to meaning: vector space models of semantics". *Journal of artificial intelligence research*, 37.1, pp. 141–188.
- Van de Cruys T. (2008). "A comparison of bag of words and syntax-based approaches for word categorization". *Proceedings of the ESSLLI Workshop on Distributional Lexical Semantics*, pp. 47–54.

- Van de Cruys T. (2010). "A non-negative tensor factorization model for selectional preference induction". *Natural Language Engineering*, 16.04, pp. 417–437.
- Van der Plas L., J. Tiedemann (2006). "Finding synonyms using automatic word alignment and measures of distributional similarity". *Proceedings of the COLING/ACL, Main conference poster sessions*, pp. 866–873.
- Van der Plas L., J. Tiedemann, J.-L. Manguin (2011). "Synonym acquisition across domains and languages", *Advances in distributed agent-based retrieval tools*. Berlin: Springer, pp. 41–57.
- Zarcone A., S. Padó, A. Lenci (2015). "Same same but different: type and typicality in a distributional model of complement coercion". *NetWordS 2015 word knowledge and word usage*, pp. 91–94.

ULISSE: una strategia di adattamento al dominio per l'annotazione sintattica automatica

Felice Dell'Orletta e Giulia Venturi – Istituto di Linguistica Computazionale “A. Zampolli”, Consiglio Nazionale delle Ricerche (ILC-CNR) –
felice.dellorletta@ilc.cnr.it giulia.venturi@ilc.cnr.it

1. Introduzione

Fino alla metà degli anni Ottanta gli algoritmi utilizzati per sviluppare sistemi di annotazione linguistica automatica sfruttavano modelli della lingua formati da insiemi complessi di regole scritte a mano e definite a priori a partire dalla conoscenza del sistema linguistico da formalizzare (detti anche ‘modelli simbolici’). Il compito degli algoritmi era quello di identificare le regole da applicare in un determinato contesto e di eseguire la specifica azione. Le principali difficoltà di questa tipologia di algoritmi sono per lo più legate ad aspetti intrinseci del linguaggio naturale, quali l'ambiguità linguistica, la variazione della lingua rispetto al dominio di conoscenza o al genere testuale, le evoluzioni lungo l'asse temporale, ecc.

Basandosi sulla formalizzazione a priori del linguaggio, i modelli simbolici a regole si scontrano pertanto con la difficoltà intrinseca di modellare l'intera conoscenza di una lingua attraverso un insieme finito di regole. L'ostacolo è rappresentato dalla necessità di costruire un modello che sia adattabile alle variazioni linguistiche o all'analisi di *input* ‘non standard’ (es. la lingua dei *social media*, trascrizioni di parlato, ecc.).

Dalla metà degli anni Ottanta grazie al crescente sviluppo di risorse linguistiche, quali ad esempio lessici, *thesauri*, corpora linguisticamente annotati, e al progresso degli studi nel campo dell'Intelligenza Artificiale, iniziano ad essere sviluppati algoritmi basati sull'apprendimento automatico (*Machine Learning*). Tra questi particolare successo hanno avuto gli algoritmi basati su apprendimento supervisionato che risultano particolarmente efficienti ed accurati nella risoluzione di compiti di classificazione. Per poter sfruttare questi algoritmi all'interno dei processi di analisi linguistica, i compiti di annotazione vengono trasformati in compiti di classificazione probabilistica condotta utilizzando un modello statistico estratto da un corpus ‘di addestramento’ nel quale ogni elemento linguistico è associato ad una specifica classe di appartenenza. L'algoritmo è in grado, in fase di addestramento (*training*), di apprendere le varie caratteristiche proprie di un determinato elemento linguistico appartenente ad una classe e di valutare per ognuna di queste caratteristiche il contributo nel determinare la corrispondente classe di appartenenza. Terminata la fase di addestramento, in fase di analisi, l'algoritmo è in grado di estrarre da un nuovo testo preso in esame (non presente nel corpus di addestramento) le varie caratteristiche che contribuiscono a determinare

la classe di appartenenza dell'elemento linguistico. Questi algoritmi sono quindi in grado di estrarre automaticamente dai corpora annotati i modelli linguistici che rappresentano la conoscenza della lingua necessaria per risolvere un determinato compito di analisi linguistica (come ad esempio l'annotazione morfo-sintattica, sintattica, ecc.).

Nonostante questi algoritmi siano più efficienti, accurati e robusti, tuttavia anch'essi soffrono del problema legato alla variazione della lingua rispetto al dominio di conoscenza, al genere testuale, ecc. In questo caso il problema principale è legato alla necessità di adattare il corpus 'di addestramento' alla varietà linguistica in esame. È questo il tema di cui discuteremo in questo contributo finalizzato a definire il problema noto come *Domain Adaptation* e a descrivere in particolare un recente approccio, basato su di un algoritmo chiamato ULISSE (Dell'Orletta et al., 2011), sviluppato dall'ItalianNLP Lab⁴² dell'Istituto di Linguistica Computazionale "Antonio Zampolli" del CNR di Pisa. Questo algoritmo permette di ottenere ottimi risultati nell'adattamento dei sistemi di annotazione linguistica a diverse varietà linguistiche o domini di conoscenza. Nello specifico, affronteremo il caso dell'Adattamento al Dominio (*Domain Adaptation*) di algoritmi di annotazione sintattica automatica dal momento che questa annotazione rappresenta uno dei livelli di analisi linguistica più complessi e più utilizzati in numerosi compiti di Trattamento Automatico del Linguaggio Naturale, come ad esempio estrazione di relazioni, semplificazione automatica, traduzione automatica, ecc. per i quali l'analisi sintattica costituisce il punto di partenza.

2. Adattamento al Dominio: definizione del problema

"Può un corpus annotato [rappresentativo di una certa varietà linguistica] essere utilizzato per l'analisi sintattica di un altro corpus [rappresentativo di un'altra varietà]?"⁴³ La domanda posta nel 2001 da Daniel Gildea ci introduce ad un problema cruciale con cui si devono confrontare i sistemi automatici di analisi sintattica (*syntactic parser* o *parser*) basati su metodi di apprendimento supervisionato: la loro precisione di analisi diminuisce in modo considerevole quando i sistemi vengono applicati su corpora rappresentativi di un dominio diverso rispetto a quello su cui sono stati addestrati. Come riporta Gildea (2001) stesso, il *Collins' parser* (Collins, 1999) addestrato sugli articoli del Wall Street Journal diminuisce la propria accuratezza di analisi di circa 5 punti percentuali quando testato sui testi del Brown Corpus benchè rappresentativi dello stesso genere testuale (giornalistico). O ancora, lo stesso *parser* addestrato sugli articoli della Penn Treebank (Marcus et al., 1993) e testato sul GENIA corpus (Tateisi et al., 2005), un corpus di *abstract* di articoli biomedici, diminuisce di circa 7 punti percentuali (Clegg, Shepherd, 2005). Da qui la necessità di mettere a punto metodi di *Domain Adaptation*, che permettano di contenere la diminuzione dell'accuratezza nell'analisi di testi di nuovi domini.

⁴²Italian Natural Language Processing Laboratory, URL:<www.italianlp.it>

⁴³"Can training data from one corpus be applied to parsing another?" (Gildea, 2001).

La definizione di una strategia per adattare un *parser* statistico ad un nuovo dominio si inserisce nel contesto più ampio delle sfide poste dall'utilizzo di strumenti di Trattamento Automatico del Linguaggio per l'annotazione linguistica automatica di corpora testuali di dominio. A partire dai primi studi condotti negli anni Ottanta, l'attenzione si è in questo senso focalizzata sulla necessità di accordare le varie fasi di elaborazione automatica del testo alle specificità di un determinato linguaggio specialistico.⁴⁴ Dal momento che gli strumenti di analisi automatica del testo sono tipicamente sviluppati per l'elaborazione della lingua usata nei giornali, considerata rappresentativa della lingua comune, l'obiettivo è stato sin d'allora quello di adattare gli strumenti al riconoscimento della struttura linguistica specifica di testi di dominio.

Gli approcci seguiti negli anni Ottanta si basavano sulla possibilità di individuare le peculiarità di un determinato linguaggio specialistico a partire dalle differenze linguistiche esistenti tra lingua standard e di dominio. Tale approccio aveva il suo fondamento nella *Theory of Sublanguages* di Zellig Harris.⁴⁵ Secondo Harris ogni linguaggio specialistico è descrivibile come un 'sottoinsieme' dell'insieme rappresentato dal linguaggio naturale. In analogia con la teoria matematica degli insiemi, esso può dunque essere operativamente denominato *sublanguage*, inteso come "un sottoinsieme del linguaggio comune che si comporta essenzialmente come l'intero linguaggio, essendo però circoscritto a uno specifico dominio" (Grishman, Kittredge, 1986).⁴⁶ Ogni linguaggio specialistico è pertanto studiabile nei termini di restrizione o ampliamento, intersezione o deviazione rispetto alle caratteristiche proprie del linguaggio comune.

Lo studio delle somiglianze e differenze tra un *sublanguage* e la lingua comune aveva un duplice intento, teorico e applicativo. Da un lato, l'analisi dei *sublanguages* ne rendeva interessante lo studio sul piano teorico in quanto "microcosmi dell'intero linguaggio"⁴⁷ in grado di fornire informazioni sul linguaggio naturale stesso (Grishman, Kittredge, 1986). Dall'altro, tale interesse era finalizzato a indagare come le caratteristiche proprie di un *sublanguage* potessero ripercuotersi sull'accuratezza dell'analisi linguistica automatica di testi di dominio.

I domini specialistici oggetto di maggiore interesse erano quelli caratterizzati da un linguaggio altamente specialistico come quello biomedico. È stato infatti questo il dominio di conoscenza al centro delle principali attività di ricerca condotte in quegli anni. Il progetto considerato pioniere in questo ambito, il *Linguistic String Project*,⁴⁸ avviato nel 1965, era espressamente finalizzato a definire una strategia di annotazione sintattica di testi di argomento biomedico, primo passo di un processo completo di trattamento automatico dell'informazione contenuta in testi

⁴⁴Vedi tra gli altri gli studi di Kittredge (1982), Grishman, Nhan et al. (1984), Grishman, Kittredge (1986) e Lehrberger (1986).

⁴⁵Per una descrizione completa della teoria di Harris si rimanda a Harris (1968).

⁴⁶"A subsystem of language that behaves essentially like the whole language, while being limited in reference to a specific subject domain" (Grishman, Kittredge, 1986).

⁴⁷"Microcosms of the whole language" (Grishman, Kittredge, 1986).

⁴⁸Linguistic String Project, URL:<<http://cs.nyu.edu/cs/projects/lsp/index.html>>

di letteratura scientifica (Sager et al., 1987).

Già nell'ambito di questi primi lavori lo studio in particolare delle peculiarità sintattiche rivestiva un interesse specifico. Come ricordato da Bonzi (1990), infatti, "scoprire le regolarità e le principali differenze nell'uso di strutture sintattiche tra diverse discipline e tipologie di testi può aiutare a costruire *parser* più efficienti, può aiutare a scoprire approcci migliori per rintracciare automaticamente i termini che meglio descrivono un documento [...]".⁴⁹ La centralità dell'annotazione sintattica del testo è ricordata più recentemente da Nivre: i risultati dei sistemi di analisi sintattica automatica sono alla base di numerosi compiti di analisi semantica del testo.⁵⁰ Proprio l'importanza dell'analisi sintattica come punto di partenza di analisi più complesse finalizzate all'accesso al contenuto del testo impone che un *parser* per poter essere davvero utile a tal fine debba soddisfare quattro requisiti fondamentali: quelli di *robustness*, *disambiguation*, *accuracy* e *efficiency* (Nivre, 2006).

Oggi sviluppare sistemi che soddisfino tali requisiti è reso ancora più complesso dal fatto che gli strumenti si trovano a dover analizzare non solo testi rappresentativi di diversi domini ma anche 'linguaggi non canonici',⁵¹ come ad esempio trascrizioni del parlato, testi prodotti nei social media (blog, Twitter, Facebook, ecc.), sms, testi di apprendenti e più in generale testi caratterizzati da una sintassi 'non standard'. La difficoltà di sviluppare sistemi di analisi sintattica accurati e robusti ha imposto pertanto la definizione di nuove strategie di adattamento al dominio.

Queste strategie di adattamento variano a seconda del tipo di paradigma che sta alla base del particolare sistema di analisi sintattica automatica utilizzato. Due sono i principali paradigmi: quelli basati su grammatiche (*grammar-driven approach*) e quelli basati su approcci indotti automaticamente dai dati (*data-driven approach*). Secondo la definizione offerta da Nivre (2006), il primo "dipende da un'approssimazione più o meno soddisfacente del linguaggio", approssimazione definita in modo deduttivo dal linguista sulla base delle sue intuizioni; il secondo "dipende da un'inferenza induttiva condotta a partire da un campione più o meno rappresentativo del linguaggio".⁵² Nel primo caso, l'annotazione sintattica del testo è condotta usando grammatiche formali (come ad esempio le 'grammatiche libere da contesto') in grado di definire un linguaggio *L*, dove tale linguaggio è l'insieme delle frasi derivabili nella grammatica data. In questo tipo di approccio analizzare una frase vuol dire derivare tutte (o alcune) le possibili analisi della

⁴⁹"Uncovering the regularities and major differences in the use of syntactic patterns among various disciplines and text types may help to build more efficient parsers for natural language input, to uncover better ways of automatically finding terms which best describe a document [...]" (Bonzi, 1990).

⁵⁰"If we move to applications that require some kind of semantic analysis of individual sentences, the role of parsing becomes more evident" (Nivre, 2006).

⁵¹Per una definizione dei cosiddetti '*Non-Canonical Languages*' si consulti il sito del First Workshop on Syntactic Analysis of Non-Canonical Language (SANCL) tenutosi nel 2012, URL:<<https://sites.google.com/site/sancl2012/>>.

⁵²"Depends on a more or less satisfactory language approximation [...] on inductive inference from a more or less representative language sample" (Nivre, 2006).

frase data la grammatica G . Il compito della disambiguazione (cioè della scelta di un'unica rappresentazione sintattica della frase date tutte le possibili interpretazioni ammesse dalla grammatica) viene condotto attraverso l'utilizzo di regole o restrizioni per permettere al sistema di scegliere una interpretazione piuttosto che un'altra.

Nel secondo caso, l'annotazione sintattica consiste nel generare l'analisi corretta per una frase attraverso un processo di inferenza induttiva a partire da un cosiddetto *training corpus* di riferimento, "un corpus di riferimento nel quale un esperto ha assegnato l'analisi corretta ad ogni segmento rilevante del testo" (Nivre, 2006).⁵³ È centrale pertanto per questa famiglia di strumenti d'analisi la presenza di un corpus annotato manualmente dal quale gli strumenti 'apprendono' ad associare al testo l'informazione linguistica corretta grazie ad un costante processo inferenziale di classificazione probabilistica. Da questo corpus gli strumenti, utilizzando algoritmi di apprendimento automatico, ricavano un modello matematico probabilistico da applicare all'annotazione linguistica di un corpus 'sconosciuto'. In questo caso, il problema della disambiguazione è intrinsecamente affrontato durante la creazione del modello statistico. Il *parser* estrae direttamente dal corpus le statistiche che gli permettono di risolvere il problema della disambiguazione attraverso un processo di assegnamento di un punteggio di probabilità di analisi ad ogni interpretazione possibile della frase.

I due diversi approcci all'analisi sintattica implicano diversi comportamenti dei sistemi in fase di analisi di testi caratterizzati da una sintassi 'non standard'. Entrano in gioco modi diversi di soddisfare due dei requisiti fondamentali che deve avere un *parser* nell'analizzare un testo: quello 1) di essere 'robusto'⁵⁴ nel trattare *input* mal formato o diverso dal linguaggio per l'analisi del quale è stato sviluppato e 2) di essere in grado di 'disambiguare'⁵⁵ tra possibili analisi diverse. Nei *grammar-driven parser*, ogni frase da analizzare deve far parte del linguaggio definito dalla grammatica; questo non accade nei casi di frasi appartenenti a domini diversi rispetto a quelli per i quali la grammatica è stata definita o quando la frase è appunto caratterizzata da una sintassi non standard. Per questo motivo i *parser* basati sulle grammatiche tendono a soffrire di una scarsa robustezza che può essere migliorata attraverso l'introduzione di nuove regole e/o il 'rilassamento' di quelle esistenti. In ogni caso aumentare la copertura e la robustezza di un sistema a regole, oltre ad essere un processo molto complesso, vuol dire aumentare il grado di possibili interpretazioni sintattiche di una frase, rendendo così il pro-

⁵³"I.e., a reference corpus of texts, where each relevant text segment has been assigned its correct analysis by a human expert" (Nivre, 2006).

⁵⁴Il requisito di 'robustezza' (*robustness*) di uno strumento di *parsing* sintattico è così definito: "A system P for parsing texts in language L satisfies the requirement of robustness if and only if, for any text $T = (x_1, \dots, x_n)$ in L , P assigns at least one analysis to every text sentence $x_i \in T$ " (Nivre, 2006, p. 41).

⁵⁵Il requisito di 'disambiguazione' (*disambiguation*) di uno strumento di *parsing* sintattico è così definito: "A system P for parsing texts in language L satisfies the requirement of disambiguation if and only if, for any text $T = (x_1, \dots, x_n)$ in L , P assigns at most one analysis to every text sentence $x_i \in T$ " (Nivre, 2006, p. 42).

blema della disambiguazione un problema molto difficile e computazionalmente oneroso.

Differentemente, i sistemi *data-driven* sono capaci di indurre i vincoli per guidare il processo di analisi delle frasi senza creare un linguaggio entro il quale restringere le possibili frasi da analizzare. Questa caratteristica permette a questi *parser* di analizzare anche frasi non ben fondate, cioè che non rispettano una struttura sintattica definita da una grammatica di riferimento o normativa. Nivre ricorda che, anche a discapito della correttezza d'analisi, l'approccio *data-driven* consente di assegnare sempre comunque almeno un'analisi alla frase in *input*. Questo fa sì che al cambiare di dominio anche testi caratterizzati da un linguaggio 'diverso' da quello del *training* corpus di riferimento siano sempre comunque analizzati. Questo permette di superare il problema della robustezza tipico degli approcci *grammar-driven*, che per riuscire ad analizzare testi di un dominio specialistico (rappresentativi di un linguaggio 'non standard') richiedono un considerevole impegno umano per estendere la grammatica.

Ponendosi in una prospettiva di *Domain Adaptation*, nel caso dei sistemi *grammar-driven* il problema da affrontare è quello della mancanza di 'copertura' della grammatica, la capacità del sistema di tenere conto dei fenomeni sintattici propri del linguaggio che appartiene alla grammatica. Il metodo seguito dai primi studi basati su questo approccio era dunque quello di tracciare "un profilo del sublanguage in grado di descrivere le relative frequenze distribuzionali di frasi e strutture testuali diverse"⁵⁶ rispetto al profilo della lingua comune (Kittredge, 1982). L'attenzione per le caratteristiche di un *sublanguage* era finalizzata alla creazione di una nuova grammatica formale, costruita a partire dagli usi sintattici tipici, diversi da quelli propri della lingua comune. La direzione di ricerca è esposta chiaramente da Grishman, Nhan et al. (1984) che dichiarano l'intenzione di trovare "una procedura euristica in grado di individuare le informazioni di dominio a partire da testi presi come campioni di un sublanguage", una strategia che permettesse di "adattare una grammatica generica alla sintassi di un particolare sublanguage".⁵⁷

Per quanto riguarda gli approcci *data-driven*, data la loro capacità di produrre una rappresentazione sintattica indipendentemente dalla frase da analizzare, il problema dell'adattamento a nuovi domini si trasforma da un problema di robustezza a uno di accuratezza. Per un *parser data-driven* il miglior metodo per ridurre questo problema consiste nella costruzione di una risorsa da usare in fase di addestramento sintatticamente annotata (*treebank*) in modo manuale e composta da testi rappresentativi della particolare varietà testuale o linguistica da analizzare. Questo approccio per l'adattamento, che possiamo definire 'supervisionato', comporta notevoli costi, sia in termini di tempo sia di competenze necessarie per

⁵⁶"a refined sublanguage profile stating the relative frequencies of different sentence and text structures" (Kittredge, 1982).

⁵⁷"a discovery procedure [...] – a procedure which can determine the domain dependent information from sample texts in the sublanguage [...] adapt a broad-coverage grammar to the syntax of a particular sublanguage" (Grishman, Nhan et al., 1984).

l'annotazione manuale. Per questo motivo sono stati sviluppati negli ultimi anni numerosi algoritmi di *Domain Adaptation* definiti 'semi-supervisionati'. L'obiettivo di questi metodi è quello di definire delle strategie di scelta di frasi che, all'interno di un ampio corpus rappresentativo del dominio che si vuole analizzare, contengono caratteristiche tipiche della varietà linguistica del dominio. Dopo essere state selezionate, in base al tipo di algoritmo semi-supervisionato utilizzato, queste frasi vengono corrette manualmente (processo di *Active Learning*) o analizzate automaticamente (processo di *Self-Training*) e vengono aggiunte al *training corpus* originale per introdurre nuova conoscenza nel sistema di analisi.

In questi ultimi anni la maggior parte dei lavori finalizzati allo sviluppo di sistemi per l'analisi sintattica si sono dedicati allo studio di approcci *data-driven*, grazie soprattutto alla loro maggiore robustezza. Saranno infatti questi i sistemi di annotazione a cui faremo riferimento in questo lavoro. In quanto segue, dopo aver mostrato come sia possibile misurare la 'distanza' tra un corpus di addestramento ed il testo da analizzare, quantificandone le differenze rispetto a diversi livelli di descrizione linguistica (Sezione 3.), illustreremo i principali algoritmi di adattamento 'semi-supervisionato' di *parsers data-driven* (Sezione 4.) focalizzandoci in particolare sul recente algoritmo chiamato ULISSE (Dell'Orletta et al., 2011) e mostRANDONE i risultati ottenuti su domini e lingue diverse. Sebbene i metodi che verranno descritti possano essere applicati, con piccoli accorgimenti, sia a sistemi di analisi che producono rappresentazioni sintattiche a costituenti sia a dipendenze, in questo lavoro ci focalizzeremo, al solo scopo esplicativo, su questi ultimi.

3. Cos'è un dominio?

Come osservato da Gildea (2001), lo studio delle variazioni linguistiche in corpora rappresentativi di varietà linguistiche diverse è di fondamentale importanza nella progettazione e sviluppo di sistemi statistici per l'analisi sintattica. Sottolineando l'importanza degli studi di Douglas Biber (Biber, 1993; Biber et al., 1998) sulle variazioni tra registri, Gildea chiarisce che "le frequenze di occorrenza di varie strutture nel corpus di addestramento si riflettono nel modello probabilistico indotto dal *parser*".⁵⁸ Ponendosi in un'ottica di adattamento al dominio, l'analisi delle caratteristiche linguistiche proprie del corpus di addestramento è dunque di centrale importanza soprattutto allo scopo di metterne in luce le differenze rispetto alle strutture specifiche dei testi che si intende analizzare. Sono infatti proprio queste differenze a causare possibili difficoltà di analisi poiché, non viste in fase di addestramento, non sono presenti nel modello statistico indotto dal *parser*.

Partendo proprio da queste considerazioni, in quanto segue illustriamo una metodologia di analisi finalizzata a descrivere in modo comparativo la distribuzione di alcune caratteristiche lessicali, morfo-sintattiche e sintattiche all'interno di varietà testuali diverse. La metodologia messa a punto ci permette di offrire una definizione operativa di dominio inteso come "una classe di testi definita da

⁵⁸"The frequencies of various structures in training data are reflected in a statistical parser's probability model" (Gildea, 2001).

una distribuzione interna di caratteristiche linguistiche diversa rispetto ad un altro insieme di testi”.

Tale metodologia si ispira alla prospettiva di indagine di Douglas Biber, finalizzata allo studio delle specificità linguistiche proprie di una data varietà linguistica (o registro). Secondo Biber, è un'analisi comparativa tra registro e lingua standard che consente di arrivare a definire cosa significa comune o raro, un'analisi condotta alla ricerca di similitudini e differenze. In quest'ottica, una descrizione completa spesso implica un'analisi composita durante la quale caratteristiche linguistiche acquisite da più livelli di descrizione linguistica svolgono tutte un ruolo importante nella ricostruzione delle dimensioni di variazione sottese tra le varietà (Biber, 1993).

In quanto segue illustreremo e discuteremo un esempio della metodologia di analisi da noi messa a punto. Abbiamo preso in considerazione due domini, quello giuridico e quello biomedico, e due lingue, italiano e inglese. Le caratteristiche rintracciate nei testi rappresentativi di questi due domini sono state messe a confronto con quelle di testi giornalistici considerati rappresentativi della lingua comune (standard). I corpora qui analizzati sono stati precedentemente utilizzati per lo sviluppo di algoritmi di adattamento al dominio di *parser* statistici durante il 'First Shared Task on Dependency Parsing of Legal Texts', tenutosi nell'ambito del 'Workshop on Semantic Processing of Legal Texts' (SPLeT 2012),⁵⁹ e il 'Workshop on Biomedical Natural Language Processing' (BioNLP 2013).⁶⁰ In entrambe le occasioni, i corpora giornalistici sono stati usati per addestrare i *parser* usati, mentre i testi giuridici e biomedici (annotati in modo manuale) sono stati impiegati per verificarne l'accuratezza in uno scenario di adattamento al dominio. Le differenze tra i corpora di lingua standard e quelli 'di dominio' discusse in quanto segue hanno pertanto una duplice lettura: da un lato, ci permettono di definire empiricamente cos'è un 'dominio' rispetto a una varietà di lingua che consideriamo standard; dall'altro, ci permettono di mettere in luce la necessità di sviluppare strategie di adattamento al dominio che colmino la distanza tra il corpus di addestramento e il corpus da analizzare, distanza intesa come diversa distribuzione di eventi linguistici (caratteristiche) che minano la precisione dei *parser*.

Già a partire da caratteristiche di base del testo come la lunghezza della frase (calcolata come il numero medio di *tokens* presenti in una frase) esistono differenze significative tra i testi giuridici e i testi di letteratura biomedica, da una parte, e gli articoli di giornale, dall'altra. Come mostrano infatti le Figure 6. e 7., gli articoli di giornale contengono frasi mediamente più corte di quelle contenute nei testi di dominio. Più in particolare, la Figura 6. mette a confronto corpora in lingua inglese: 1) la porzione di 447.000 *token* della Penn Treebank (PTB) (Marcus et al., 1993) contenente articoli del Wall Street Journal (WSJ) usata in fase di addestramento (*WSJ_gold*), 2) una collezione più ampia (39.285.425 *token*) di WSJ contenente articoli linguisticamente annotati in modo automatico (dunque senza

⁵⁹SPLeT 2012, URL:<<https://sites.google.com/site/splet2012workshop/shared-task>>

⁶⁰BioNLP 2013, URL:<<http://compbio.ucdenver.edu/BioNLP2013/index.shtml>>

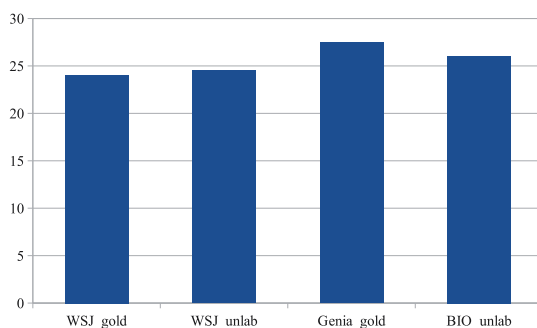


Figura 6. Lunghezza media delle frasi nei testi biomedici e negli articoli di giornale in inglese.

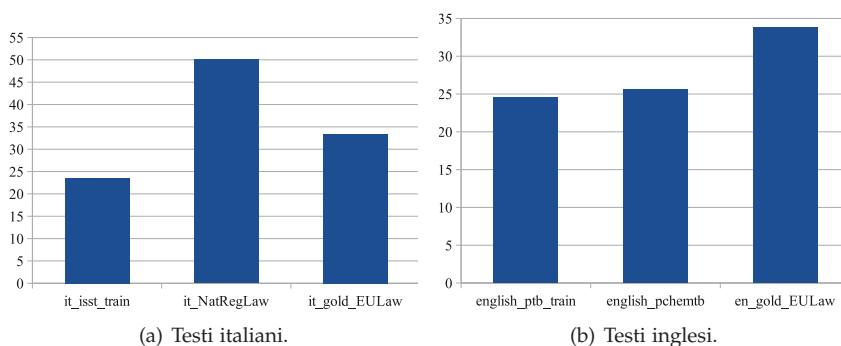


Figura 7. Lunghezza media delle frasi nei testi giuridici italiani e inglesi a confronto con gli articoli di giornale nelle due lingue.

revisione manuale) (*WSJ_unlab*), 3) GENIA treebank (Tateisi et al., 2005) composta da *abstract* di articoli biomedici in inglese (circa 493.000 *token*) (*Genia_gold*), 4) un corpus più ampio di articoli di argomento biomedico annotati automaticamente (9.776.890 *token*) (*BIO_unlab*). Come si può notare, le frasi contenute negli articoli di giornale sia annotati in modo manuale (*WSJ_gold*) sia annotati automaticamente (*WSJ_unlab*) sono mediamente più corte di quelle contenute nei testi biomedici (*Genia_gold* e *BIO_unlab*).

Tendenze simili si registrano nell'analisi di testi di dominio giuridico, sia per l'italiano che per l'inglese. La Figura 7.a mostra i confronti tra la lunghezza delle frasi di 1) articoli di giornale contenuti nella porzione della Italian Syntactic-Semantic Treebank (ISST-TANL) (Dell'Orletta, Marchi, Montemagni, Venturi et al., 2013) (71.568 *token*) usata in fase di addestramento dai partecipanti al 'First Shared Task on Dependency Parsing of Legal Texts' (*it_issst_train*), 2) una collezione di te-

sti giuridici emanati dallo stato italiano e dalle regioni (5.194 token) (*it_NatRegLaw*) e dalla Commissione Europea (5.662 token) (*it_gold_EULaw*), tutti annotati in modo manuale. Anche in questo caso, le frasi contenute negli articoli di giornale sono mediamente più corte di quelle contenute nei testi giuridici e in particolare quelle dei testi giuridici statali e regionali risultano essere le più lunghe.

Interessante inoltre come i testi giuridici si caratterizzino per frasi mediamente più lunghe sia dei testi giornalistici sia dei testi biomedici. Questo emerge dall'analisi della Figura 7.b dove sono messi a confronto corpora giornalistici, di testi di argomento biomedico e testi giuridici. Sono in particolare riportati i valori delle lunghezze medie delle frasi contenute 1) nella porzione di articoli giornalistici della PTB usata come corpus di addestramento (446.573 token) (*english_ptb_train*), 2) in una collezione di *abstract* biomedici inglesi annotati in modo manuale (5.001 token) (*english_pchemtb*), 3) in una collezione di testi giuridici emanati dalla Commissione Europea (5.621 token) (*en_gold_EULaw*) anch'essi annotati in modo manuale.

Se ne può dunque concludere che la lunghezza media della frase rappresenta già un primo livello di differenza tra lingua standard e di dominio. È questa una caratteristica che, seppure di base, è strettamente legata ad aspetti di complessità di analisi: se i testi da analizzare contengono frasi più lunghe di quelle contenute nel corpus di addestramento, essi avranno di conseguenza distribuzioni diverse di categorie morfo-sintattiche e caratteristiche sintattiche tipiche di frasi mediamente più lunghe (come ad esempio alberi sintattici più profondi, relazioni di dipendenza sintattica più lunghe, ecc.).

Per quanto riguarda le caratteristiche morfo-sintattiche, i testi del dominio biomedico e giuridico si caratterizzano per una percentuale maggiore di preposizioni e sostantivi e per una minore presenza di verbi (vedi Figura 8.). Come si può notare, tuttavia, vi sono delle differenze significative tra i due domini, ad esempio i testi biomedici si differenziano dai testi giuridici per una percentuale decisamente maggiore di sostantivi. È quello che emerge da un'analisi più dettagliata condotta calcolando la distribuzione delle dieci triple più frequenti di categorie morfo-sintattiche legate da una relazione di dipendenza sintattica⁶¹ nei corpora di testi biomedici (vedi Tabella 4.). Gli *abstract* biomedici (*Genia_gold* e *BIO_unlab* nella Tabella) risultano essere nettamente caratterizzati da una maggiore frequenza di triple di tipo nominale (es. sequenze di due nomi e una preposizione, NN-NN-IN; nome-preposizione-nome, NN-IN-NN; aggettivo-nome-preposizione, JJ-NN-IN) rispetto a quelle di tipo verbale (es. sequenze di determinante-nome-verb, DT-NN-VBZ) rispetto agli articoli di giornale (*WSJ_gold* e *WSJ_unlab* nella Tabella).

Anche a livello delle caratteristiche sintattiche, gli articoli di giornale si differenziano decisamente rispetto ai testi di dominio. Ad esempio sia i testi giuridici che quelli biomedici si caratterizzano per sequenze mediamente più lunghe di complementi preposizionali che modificano a cascata un sostantivo (Figura 9.). Questo tratto risulta specifico in particolare dei testi di dominio giuridico e

⁶¹Ogni tripla è formata da una sequenza di tre categorie morfo-sintattiche legate all'interno di un albero sintattico da relazioni di dipendenza di tipo 'figlio-padre-nonno'.

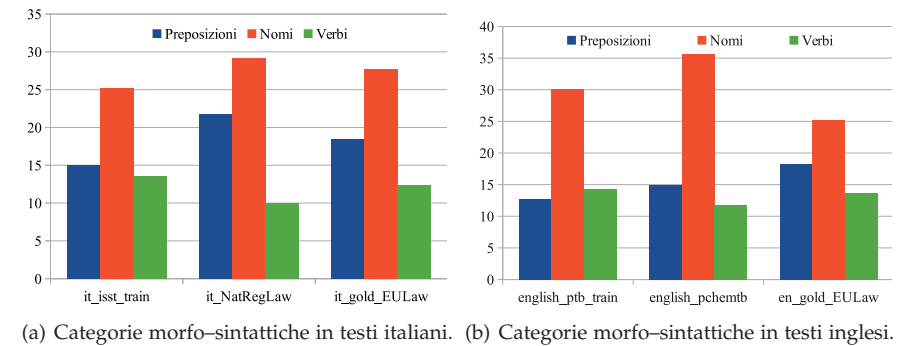


Figura 8. Distribuzione di categorie morfo-sintattiche nei testi giuridici italiani e inglesi a confronto con gli articoli di giornale nelle due lingue.

WSJ_gold		WSJ_unlab		Genia_gold		BIO_unlab	
Triple	% Freq	Triple	% Freq	Triple	% Freq	Triple	% Freq
DT-NN-IN	2.03	DT-NN-IN	1.72	NN-NN-IN	3.66	DT-NN-IN	2.87
.-VBD-ROOT	1.61	.-VBD-ROOT	1.30	NN-IN-NN	2.93	NN-IN-NN	2.39
NN-IN-NN	1.11	JJ-NN-IN	0.99	DT-NN-IN	2.48	JJ-NN-IN	2.08
JJ-NN-IN	1.10	NN-IN-NN	0.97	JJ-NN-IN	1.96	NN-NN-IN	1.73
.-VBZ-ROOT	1.09	NNP-NNP-IN	0.87	NN-NNS-IN	1.88	IN-NN-IN	1.72
NNP-NNP-IN	0.95	DT-NN-VBD	0.85	JJ-NNS-IN	1.77	JJ-NNS-IN	1.36
DT-NN-VBZ	0.89	NN-VBD-ROOT	0.80	IN-NN-IN	1.65	.-VBD-ROOT	1.33
DT-NN-VBD	0.87	JJ-NNS-IN	0.79	NN-CC-IN	1.64	NNS-IN-NN	1.13
JJ-NNS-IN	0.87	NNP-NNP-VBD	0.78	NNS-IN-NN	1.56	NNP-NN-IN	1.03
IN-NN-IN	0.87	.-VBZ-ROOT	0.75	NN-NN-CC	1.47	NN-IN-VBN	0.93

Tabella 4. Distribuzione percentuale delle dieci triple più frequenti di categorie morfo-sintattiche legate da una relazione di dipendenza sintattica nei testi biomedici e negli articoli di giornale.

soprattutto dei testi giuridici italiani regionali e statali.

Altri due tratti caratterizzanti i domini presi in esame, in particolare i testi giuridici italiani statali e regionali, sono la lunghezza media delle relazioni di dipendenza sintattica (calcolata come la distanza in *token* tra la testa sintattica e il proprio dipendente) (Figura 10.) e la profondità dell'albero sintattico (calcolata come il numero massimo di relazioni di dipendenza consecutive) (Figura 11.). Come messo in evidenza da McDonald e Nivre (McDonald, Nivre, 2007), l'analisi di dipendenze lunghe è una delle principali cause di errore dei sistemi statistici di analisi sintattica. Più in generale, relazioni di dipendenza lunghe e alberi sintattici profondi sono considerati due dei principali elementi di complessità sintattica di una frase (Gibson, 1998; Yngve, 1960; Frazier, 1985). Rispetto a queste osservazioni, i testi dei domini analizzati risultano dunque caratterizzati non solo da eventi linguistici 'distanti' da quelli della lingua comune ma si contraddistinguono anche per una generale maggiore complessità linguistica.

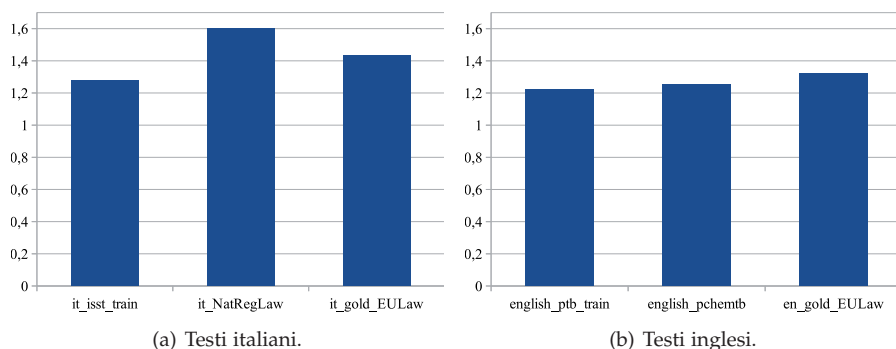


Figura 9. Lunghezza media di complementi preposizionali che modificano un sostantivo nei testi giuridici e biomedici italiani e inglesi a confronto con gli articoli di giornale nelle due lingue.

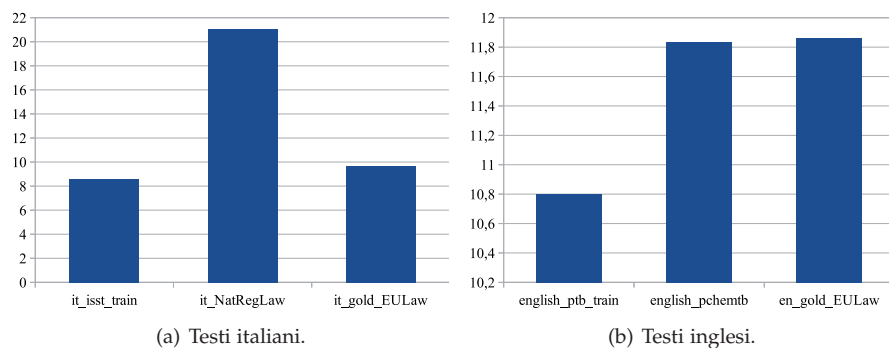


Figura 10. Lunghezza media delle relazioni di dipendenza sintattica nei testi giuridici e biomedici italiani e inglesi a confronto con gli articoli di giornale nelle due lingue.

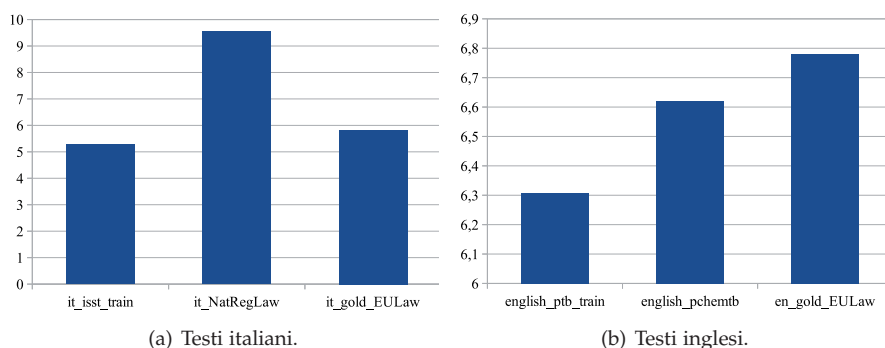


Figura 11. Profondità media dell'albero sintattico nei testi giuridici e biomedici italiani e inglesi a confronto con gli articoli di giornale nelle due lingue.

Veniamo infine alla distribuzione del lessico, un elemento centrale nella definizione di un dominio. Entrambi i domini presi in considerazione contengono una ridotta percentuale dei lemmi presenti nei testi giornalistici usati in fase di addestramento, con alcune differenze significative tra i due domini.⁶² Il corpus di articoli biomedici in inglese e quello di testi giuridici europei in inglese ha una percentuale pari rispettivamente al 60% e all'88% di lemmi presenti nel corpus di testi giornalistici di riferimento così come la collezione di testi giuridici italiani statali e regionali e quella di testi giuridici europei in italiano ha rispettivamente l'81% e l'88% di lemmi del corpus di addestramento. Ne emerge dunque che i due domini si differenziano in modo significativo rispetto al lessico: il dominio giuridico, essendo caratterizzato da un linguaggio che "è un sottinsieme, distinto ma non separato dal linguaggio generale" (Cassese, 1992), si contraddistingue per una maggiore percentuale di lessico giornalistico (cioè comune), caratteristica dei testi giuridici sia italiani sia inglesi. Al contrario, il dominio biomedico è nettamente più distante a livello lessicale dalla lingua giornalistica.

In conclusione, è interessante far notare come anche le differenze nella distribuzione del lessico hanno una diretta ripercussione sull'accuratezza dell'analisi sintattica di testi di dominio. I lemmi sconosciuti, non visti cioè in fase di addestramento, sono quelli per cui il sistema di analisi commetterà il maggior numero di errori, a livello di riconoscimento sia delle categorie morfo-sintattiche sia delle relazioni di dipendenza sintattica nelle quali sono coinvolti.

4. Adattamento al Dominio: metodi

A partire dal 2007 sono state organizzate una serie di conferenze finalizzate a fare il punto sullo stato dell'arte in materia di adattamento al dominio di strumenti di Trattamento Automatico del Linguaggio. Il 'Domain Adptation Track' dello 'Shared Task on Dependency Parsing' organizzato durante la conferenza 'Conference on Computation Natural Language Learning' (CoNLL) è stata la prima occasione in cui si è svolta una competizione tra approcci diversi finalizzati ad adattare un *parser* a domini diversi rispetto a quelli di addestramento, nello specifico testi biomedici, chimici e trascrizioni di parlato infantile (Nivre et al., 2007). In seguito, sono state organizzate numerose altre iniziative tra le quali, nel 2010, il 'Workshop on Domain Adaptation for Natural Language Processing' (DANLP) della 'Conference of the Association of Computational Linguistics' che ha permesso di fare il punto sulle tipologie di approcci di *Domain Adaptation* per diversi compiti di Trattamento Automatico del Linguaggio,⁶³ non solo per sistemi di annotazione sintattica automatica. Più recentemente, nel 2012 in occasione del 'First Workshop on Syntactic Analysis of Non-Canonical Language' (SANCL) e dello 'Shared Task' collegato,⁶⁴ il concetto di dominio è stato allargato alle cosiddette

⁶²In questo caso abbiamo riportato i risultati acquisiti da sezioni dei corpora di dimensioni simili dal momento che la distribuzione del lessico dipende strettamente dalla lunghezza del testo.

⁶³DANLP 2010, URL:<<https://sites.google.com/site/danlp2010/>>

⁶⁴SANCL 2012, URL:<<https://sites.google.com/site/sancl2012/>>

‘varietà non canoniche della lingua’.

Tutte queste iniziative hanno affrontato il tema del *Domain Adaptation* in relazione alla lingua inglese. Poche sono state le occasioni in cui il tema è stato analizzato in rapporto a lingue diverse. Un'eccezione è stato il 'Domain Adaptation Track' (Dell'Orletta, Marchi, Montemagni, Venturi et al., 2013) organizzato in occasione della campagna di valutazione 'Natural Language Processing and Speech Tools for Italian' (EVALITA) nel 2011 e finalizzato a mettere a confronto metodi di adattamento per la lingua italiana. Inoltre, in occasione del 'First Shared Task on Dependency Parsing of Legal Texts' (Dell'Orletta, Marchi, Montemagni, Plank et al., 2012) è stato possibile mettere a confronto metodi di adattamento al dominio per la lingua italiana e inglese.

Come ricordato nella Sezione 2. di questo contributo, la situazione ideale per l'adattamento di un sistema di annotazione sintattica *data-driven* consiste nel poter disporre di una collezione di testi del dominio target annotati in modo manuale. Tuttavia, il principale ostacolo è rappresentato dal fatto che costruire una risorsa sintatticamente annotata in modo manuale è costoso, sia in termini di tempo sia di competenze necessarie. È questo il motivo per cui poche sono le *treebank* oggi esistenti per un dominio diverso da quello giornalistico. I due principali esempi sono: il GENIA corpus (Tateisi et al., 2005), contenente testi biomedici⁶⁵ per un totale di circa 18 mila frasi, e la Google Web Treebank (Petrov, McDonald, 2012) costruita in occasione dello 'Shared Task' del 'First Workshop on Syntactic Analysis of Non-Canonical Language' (SANCL-2012) e contenente diverse tipologie di testi del web (Yahoo answers, e-mail, newsgroup, blog e commenti), per un totale di circa 20 mila frasi.

Data la difficoltà di reperire testi di dominio annotati manualmente, la maggior parte degli studi si è invece concentrata sullo sviluppo di strategie di adattamento semi supervisionato raggruppabili in due classi: quelle che si basano su metodi di *Active Learning* e quelle basate su metodi di *Self-Training*. Entrambe partono da un presupposto comune: il fatto che siano disponibili grandi collezioni di testi di dominio da annotare in modo automatico. L'obiettivo comune è quello di selezionare all'interno di queste grandi collezioni di testi di dominio esempi, cioè frasi 'informative' del dominio target, che possono essere utilizzati in fase di addestramento in aggiunta al corpus di addestramento di partenza.

Diversa è tuttavia la strategia di selezione degli esempi adottata e il modo in cui la loro 'informatività' viene calcolata. I metodi di *Active Learning* si basano su funzioni statistiche in grado di selezionare quegli esempi che, corretti a mano e reinseriti nel corpus di addestramento come nuovo *training*, permettono al sistema di migliorare le proprie accuratèzze. È il caso ad esempio del lavoro di Attardi et al. (2013) che, come criterio di selezione delle frasi da rivedere a mano, hanno utilizzato la probabilità restituita dal classificatore durante la selezione delle azioni del *parser* nella fase di analisi. L'intuizione è stata quella di selezionare come frasi più informative quelle per le quali il *parser* aveva 'più dubbi'

⁶⁵GENIA corpus, URL:<<http://www.nactem.ac.uk/genia/genia-corpus>>

nella scelta dell'interpretazione sintattica. Quelle frasi molto probabilmente erano quelle che contenevano caratteristiche sconosciute rispetto a quelle viste in fase di addestramento.

A differenza dei metodi di adattamento al dominio basati su *Active Learning*, che hanno l'obiettivo di ridurre i tempi e l'impegno di annotazione manuale limitandola ai soli esempi considerati più informativi, i metodi di *Self-Training* non prevedono una fase di revisione manuale. Come nel caso dell'*Active Learning*, questi metodi si differenziano tra loro rispetto alla strategia di selezione delle frasi analizzate automaticamente da inserire nel corpus di addestramento come nuovo *training*. Le frasi selezionate vengono direttamente inserite, senza revisione manuale, come nuovi dati di addestramento in aggiunta a quelli originari. In questo modo il *parser* si 'autoaddestra' su di una collezione nella quale sono uniti esempi annotati a mano (*gold*) ed esempi annotati in modo automatico. In questo caso la difficoltà sta nel definire una funzione che stabilisca sia il livello di 'affidabilità' dell'annotazione automatica sia l' 'informatività' degli esempi da selezionare. I lavori di adattamento al dominio basati su metodi di *Self-Training* (McClosky et al., 2006; Sagae, Tsujii, 2007; McClosky, Charniak, 2008; Sagae, 2010; McClosky et al., 2010 solo per citare alcuni dei più rilevanti) si differenziano pertanto tra loro rispetto alla strategia di selezione automatica dell'affidabilità e dell'informatività degli esempi annotati automaticamente. In quanto segue introdurremo un innovativo metodo di selezione di affidabilità e informatività di frasi annotate in modo automatico (Sezione 4.1.) e ne mostreremo il funzionamento in un sistema di adattamento al dominio basato su *Self-Training* (Sezione 4.2.).

4.1. ULISSE: selezionare alberi sintattici corretti e informativi

Il punto di partenza fondamentale di ogni processo di adattamento al dominio basato su metodi di *Self-Training* è la definizione di una strategia di selezione di alberi sintattici all'interno di un corpus di frasi annotate automaticamente. Per questo motivo si sono diffusi negli ultimi anni diversi approcci finalizzati a stabilire il grado di affidabilità dei risultati dell'annotazione automatica di una frase. Rispetto al metodo utilizzato questi approcci si possono classificare in tre grandi tipi: 1) quelli supervisionati che utilizzano un classificatore esterno per stabilire se le frasi sono state annotate in modo più o meno corretto (Kawahara, Uchimoto, 2008); 2) quelli che combinano i risultati di più *parser* scegliendo le frasi per i quali la maggior parte dei *parser* ha attribuito la stessa analisi (Sagae, Tsujii, 2007); 3) quelli non supervisionati che per la selezione degli alberi sintattici utilizzano solo statistiche acquisite dalle frasi analizzate dal *parser* che si vuole adattare. In questo contributo ci siamo focalizzati su un algoritmo non supervisionato evitando di essere influenzati dalla selezione del corpus di addestramento (metodi supervisionati) e dall'accuratezza dei *parser* utilizzati (metodi che combinano più *parser*). In quanto segue illustreremo un innovativo algoritmo e la sua applicazione in uno scenario di *Self-Training*. Questo algoritmo, chiamato ULISSE, ci ha permesso di migliorare le prestazioni dei *parser* quando applicati su testi di dominio distanti rispetto al loro corpus di addestramento.

ULISSE (*Unsupervised Linguistically-driven Selection of dEpendency parses*) è in grado, data una collezione di frasi automaticamente annotate a livello sintattico a dipendenze, di assegnare ad ogni frase annotata (di qui in avanti riferita come *dependency parse*) un punteggio che quantifica il grado di correttezza dell'analisi automatica. Descriviamo ora il funzionamento dell'algoritmo, mentre per una presentazione più dettagliata rimandiamo al lavoro di Dell'Orletta e colleghe (Dell'Orletta et al., 2011).

La strategia di analisi di ULISSE può essere divisa in due fasi. In una prima fase, ULISSE crea un modello statistico utilizzando un insieme di caratteristiche 'linguisticamente motivate' estratte da un ampio corpus annotato automaticamente. Per la costruzione del modello, ULISSE estrae caratteristiche linguisticamente motivate dai *dependency parse* e da loro sotto-alberi definiti sulla base di criteri 'linguisticamente motivati'. Ad ognuna di esse assegna la probabilità di apparire nel corpus analizzato date alcune 'proprietà' delle frasi all'interno delle quali esse appaiono. Proprietà prese in considerazione sono ad esempio la lunghezza della frase, la distribuzione delle categorie morfo-sintattiche, le parole che le compongono e il loro ordinamento. Vengono quindi acquisite in questa fase le scelte di analisi che il *parser* ha fatto in determinati contesti linguistici. Nella fase successiva, ULISSE assegna un punteggio di accuratezza (di affidabilità dell'analisi) a frasi analizzate automaticamente usando il modello statistico creato nella prima fase. Questo punteggio viene ottenuto calcolando le caratteristiche e i sotto-alberi definiti nella prima fase e presenti nel *dependency parse* preso in considerazione. A queste caratteristiche vengono associati i pesi statistici contenuti nel modello. Una funzione di assegnamento del punteggio combina i pesi per ottenere il risultato finale. L'intuizione su cui si basa ULISSE è che le scelte prese con più frequenza dal *parser* in contesti linguistici simili possono essere considerate più affidabili.

Il punteggio assegnato da ULISSE può essere utilizzato per creare un ordinamento di affidabilità all'interno di un ampio corpus di frasi analizzate. Questo ordinamento può essere considerato il punto di partenza di un processo di *Self-Training*. Oltre alla capacità di creare un ordinamento di affidabilità tra frasi analizzate, altro elemento peculiare di ULISSE è la sua abilità di selezionare le frasi caratterizzanti i fenomeni linguistici tipici del dominio di partenza. Questo è possibile grazie al fatto di fare affidamento su caratteristiche linguisticamente motivate estratte dai *dependency parse*.

La scelta di quali caratteristiche usare per selezionare le frasi corrette e informative è un elemento chiave all'interno dell'intero processo. Tutte le caratteristiche su cui fa affidamento ULISSE sono scelte a partire dalla struttura sintattica a dipendenze e mirano a catturare aspetti dell'annotazione sintattica che caratterizzano in modo particolare il dominio in esame. Questo, unito alla capacità di associare un punteggio di affidabilità all'analisi automatica, rende ULISSE un algoritmo particolarmente adatto per essere usato in un contesto di *Domain Adaptation*.

Possiamo classificare le caratteristiche usate da ULISSE in due tipi: 'globali', calcolate cioè per ogni *dependency parse*, e 'locali', calcolate per ogni singola relazione di dipendenza sintattica. A partire dall'intero albero sintattico vengono calcolate le seguenti caratteristiche che, come abbiamo visto nella Sezione 3.,

rappresentano elementi di variazione tra domini:

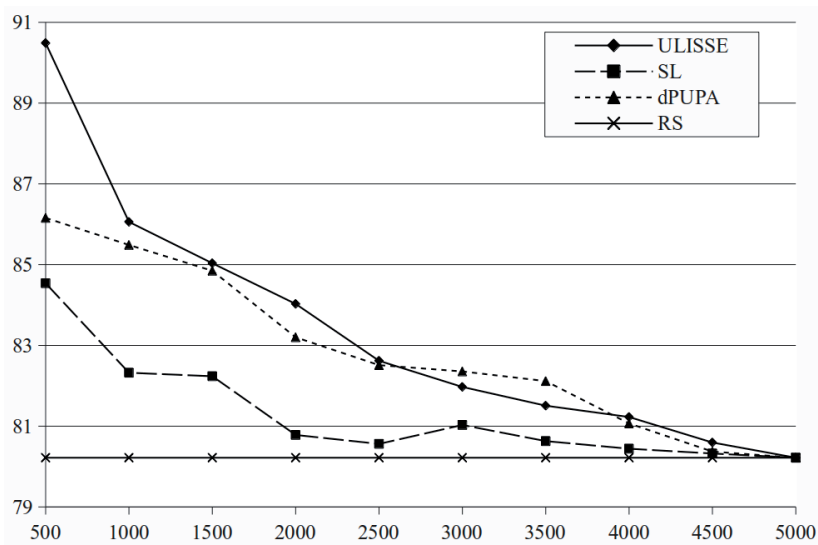
- **profondità media dell'albero sintattico**, calcolata come il numero massimo di relazioni di dipendenza sintattica consecutive;
- **lunghezza media di complementi preposizionali**, calcolata come il numero medio di complementi preposizionali consecutivi che modificano un nome;
- **valenza verbale**, calcolata come il numero di dipendenze sintattiche effettive di un verbo;
- **radici sintattiche verbali**, calcolata come il numero di radici sintattiche di tipo verbale di una frase rispetto al numero totale di radici sintattiche;
- **subordinate contro frasi principali**, calcolata come 1) il rapporto tra frasi subordinate e principali e 2) posizione delle subordinate rispetto alla principale all'interno di una frase;
- **lunghezza delle relazioni di dipendenza sintattica**, calcolata come la distanza tra la testa e il dipendente (in *token*).

Queste caratteristiche, oltre a delineare i tratti linguistici peculiari del dominio analizzato, sono fortemente connesse con la misura della complessità sintattica di una frase. Ad esempio, la profondità dell'albero sintattico di una frase può essere considerata un indicatore di crescente complessità sintattica (Yngve, 1960; Frazier, 1985; Gibson, 1998), così come la posizione della subordinata rispetto alla principale: frasi che contengono subordinate in posizione post-verbale sono più semplici di frasi in cui la subordinata precede la frase principale (Miller, Weinert, 1998). McDonald, Nivre (2007) mostrano che i *parsers* statistici hanno una diminuzione di accuratezza quando si trovano a dover analizzare frasi che contengono relazioni di dipendenza molto lunghe.

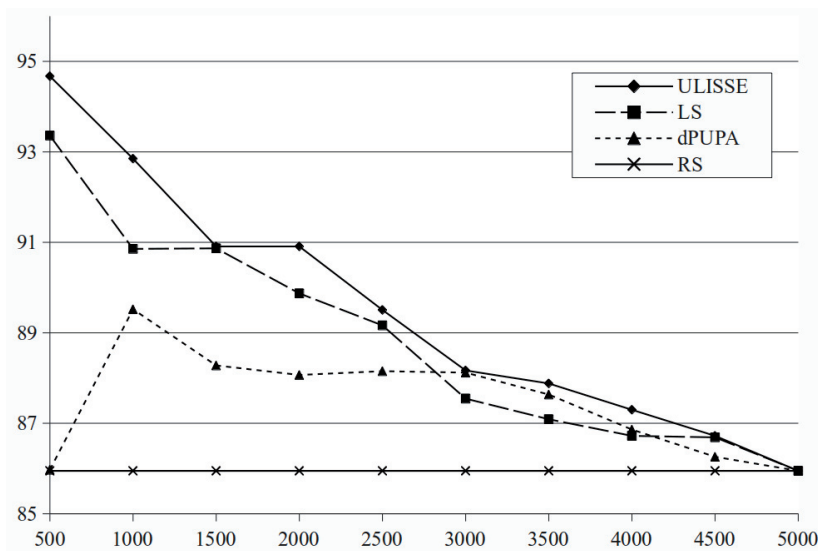
La caratteristica 'locale' di ULISSE associa ad ogni relazione di dipendenza un punteggio di plausibilità. Tale punteggio viene calcolato tenendo in considerazione per ogni relazione: la categoria morfo-sintattica (*part-of-speech*) del dipendente, della testa sintattica, e del padre della testa sintattica, e la relazione sintattica che li lega. Ad ogni *dependency parse* viene associato il peso della relazione sintattica meno plausibile.

Per dimostrare l'efficacia di ULISSE nel predire l'accuratezza dell'analisi sintattica fatta da un *parser*, riportiamo gli esperimenti descritti in Dell'Orletta et al. (2011) dove ULISSE è stato testato su corpora di testi italiani e inglesi annotati automaticamente utilizzando il *parser* DeSR (Attardi, 2006). In tutti gli esperimenti ULISSE si è dimostrato più efficace rispetto agli altri algoritmi di selezione presi come riferimento (*baseline*).⁶⁶

⁶⁶In Dell'Orletta et al. (2011) è possibile trovare altri esperimenti realizzati con ULISSE su diversi domini e utilizzando diversi *parser* che dimostrano l'indipendenza della sua efficacia rispetto alle lingue, ai domini e alle diverse strategie di analisi sintattica utilizzate.



(a) Testi italiani.



(b) Testi inglesi.

Figura 12. Accuratezza di analisi (LAS) di gruppi di frasi annotate in modo automatico e ordinate da ULISSE e dalle *baselines*.

La Figura 12. mostra i risultati degli esperimenti. In particolare, la Figura 12.a riporta i risultati degli esperimenti condotti su un corpus di testi italiani e la Figura 12.b quelli condotti su un corpus di testi inglesi. Per la lingua italiana il *parser* DeSR è stato addestrato sulla ISST-TANL, una *treebank* di articoli giornalistici e periodici di 3.109 frasi (71.285 *token*) usata nella campagna di valutazione di strumenti di Trattamento Automatico del Linguaggio per l'italiano Evalita'09⁶⁷ (Bosco et al., 2009). Come ampio corpus di riferimento da cui ULISSE ha estratto il modello statistico è stata usata una porzione del corpus di articoli giornalistici CLIC-ILC (Marinelli et al., 2003) pari a 1.104.237 frasi (22.830.739 *token*). Il *test set* usato per valutare l'accuratezza di DeSR è quello distribuito in occasione di Evalita'09 composto da 260 frasi (5.011 *token*). Per gli esperimenti condotti sull'inglese, DeSR è stato addestrato sulle sezioni 2–11 della parte della Penn Treebank (Marcus et al., 1993) usata in occasione del 'CoNLL 2007 Shared Task on Dependency Parsing' (Nivre et al., 2007) che comprende 447.000 *token* e circa 18.600 frasi. Come corpus di riferimento è stato usato l'ampio corpus distribuito in occasione del 'CoNLL 2007 Shared Task on Dependency Parsing', composto da 1.625.606 frasi (39.285.425 *token*), e come *test set* la sezione 23 della parte della Penn Treebank pari a 214 frasi (5.003 *token*).

In ogni figura, l'asse delle ordinate riporta l'accuratezza di analisi del *parser* espressa in *Labelled Attachment Score* (LAS), che misura la percentuale di *token* ai quali il *parser* ha assegnato correttamente sia la testa sintattica sia il tipo di relazione di dipendenza. Sull'asse delle ascisse sono riportate le porzioni crescenti dei primi k *token* della lista di *dependency parse* ordinati da ULISSE e le *baseline*. Il valore di k progredisce per intervalli di 500 *token* e va dai primi 500 *token* all'intera porzione del *test set* (dunque $k=500$, $k=1000$, $k=1500$, ecc.). La linea etichettata con RS rappresenta l'accuratezza di DeSR sull'intero *test set*, quella etichettata come ULISSE rappresenta l'ordinamento prodotto da ULISSE, la linea SL rappresenta l'accuratezza ottenuta utilizzando la lunghezza della frase come criterio per ordinare i *dependency parse*. In base a questa *baseline* le frasi annotate sono state ordinate dalle più corte alle più lunghe, assumendo che le più corte fossero le più corrette perché, come dimostrato da McDonald, Nivre (2007), frasi più lunghe sono più difficili da analizzare. La lunghezza della frase come criterio di selezione dei *dependency parse* più affidabili costituisce pertanto una *baseline* piuttosto efficace. Le performance di ULISSE sono inoltre state confrontate con dPUPA, che rappresenta la versione adattata per l'analisi sintattica a dipendenze di un algoritmo di ordinamento sviluppato per l'analisi sintattica a costituenti (Reichart, Rappoport, 2009). PUPA rappresenta lo stato dell'arte degli approcci di selezione non supervisionati.

Esaminando la Figura 12. possiamo notare come, nonostante la bontà delle *baseline* scelte, per entrambe le lingue ULISSE risulti il criterio di ordinamento migliore. Sebbene la selezione dei *dependency parse* guidata dalla lunghezza della frase risulti un buon criterio di ordinamento, ULISSE è stato in grado di selezionare frasi più lunghe che hanno un livello di accuratezza maggiore rispetto a frasi

⁶⁷Evalita'09, URL:<<http://evalita.fbk.eu/index.html>>

molto brevi. Ad esempio, se osserviamo l'esperimento sull'italiano, mentre le frasi contenute nella prima porzione di 500 *token* ordinate con *SL* hanno una lunghezza media di 7,88 *token* (e una accuratezza di 84,54% di LAS), ULISSE è stato in grado di ordinare nella stessa porzione dei primi 500 *token* frasi con una lunghezza media di 9,5 *token* (e un'accuratezza di 90,48%). Nella stessa porzione di testo *PUPA* ha selezionato frasi con una lunghezza media di 12 *token* ma con un'accuratezza minore di quella ottenuta da ULISSE (86,16%), dimostrandosi tuttavia un ottimo criterio di selezione.

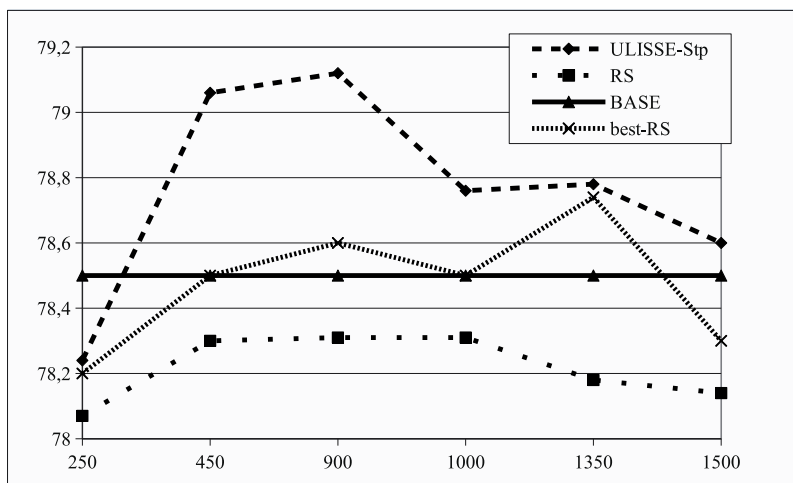
Abbiamo quindi visto come ULISSE sia in grado di selezionare e ordinare per affidabilità decrescente di analisi le frasi all'interno di un ampio corpus facendo unicamente riferimento alle caratteristiche specifiche del corpus stesso. Gli esperimenti mostrati sono stati condotti in uno scenario che possiamo definire *in-domain*, dal momento che il dominio dei testi usati in fase di addestramento del *parser* è lo stesso di quello dell'ampio corpus di documenti su cui è stato testato ULISSE. Sia per l'italiano sia per l'inglese gli esperimenti sono stati condotti su testi giornalistici. I buoni risultati raggiunti ci hanno suggerito come l'algoritmo di selezione potesse essere utilizzato con successo anche nel caso in cui il corpus di riferimento usato da ULISSE fosse diverso da quello su cui il *parser* era stato originariamente addestrato, scenario che possiamo definire *out-domain*. Questo, unito alla sua capacità di estrarre frasi rappresentative del dominio analizzato (grazie all'uso di caratteristiche linguisticamente motivate), fanno di ULISSE un algoritmo particolarmente adatto per essere usato in un contesto di *Self-Training per Domain Adaptation*.

4.2. ULISSE per l'Adattamento al Dominio

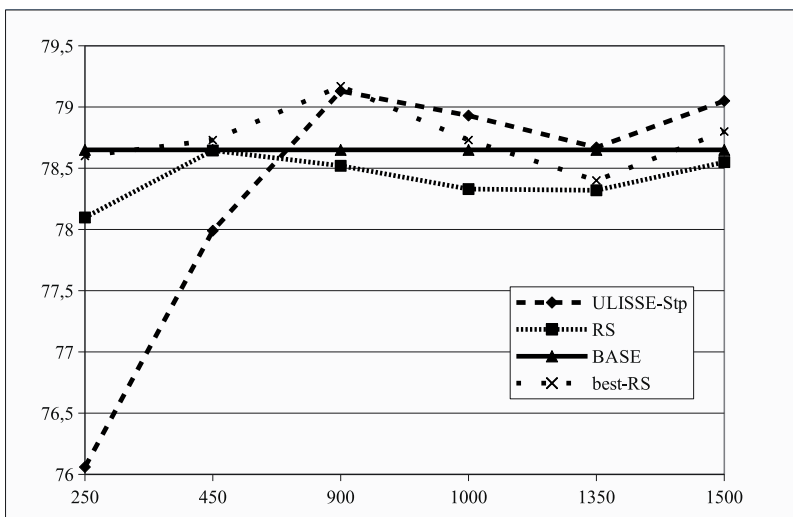
In quanto segue illustreremo alcuni esperimenti condotti usando ULISSE in uno scenario *out-domain*. Gli esperimenti, originariamente condotti da Dell'Orletta et al. (2013) hanno dimostrato come l'algoritmo sia in grado di selezionare e ordinare frasi annotate automaticamente sufficientemente affidabili e informative per essere direttamente inserite nel corpus di addestramento originario affinché il *parser* si 'autoaddestri' adattandosi ad un nuovo dominio. In questo paragrafo ci soffermeremo sugli esperimenti condotti su testi scritti in lingua inglese.

Le Figure 13. e 14. riportano i risultati degli esperimenti di adattamento al dominio condotti con l'obiettivo di adattare il *parser* DeSR all'analisi di *abstract* di argomento biomedico e chimico in inglese. Il *parser* è originariamente addestrato sulle sezioni 2–11 della parte della Penn Treebank (Marcus et al., 1993). Sono stati usati come corpora di riferimento, da cui ULISSE ha potuto estrarre il modello statistico, le collezioni distribuite in occasione del 'CoNLL 2007 Shared Task on Domain Adaptation' (Nivre et al., 2007). Si tratta di una collezione di *abstract* di argomento chimico (di qui in avanti riferiti come *CHEM*) composta da 396.128 frasi (10.482.247 *token*) e di una collezione di *abstract* di argomento biomedico (di qui in avanti riferiti come *BIO*) composta da 375.421 frasi (9.776.890 *token*). Gli esperimenti condotti su *CHEM* sono stati testati su di una porzione di 195 frasi (5.001 *token*) e quelli condotti su *BIO* su una collezione di 200 frasi (5.017 *token*). Le performance di DeSR 'autoaddestrato' usando i *dependency parse* selezionati da

ULISSE sono state messe a confronto con i risultati ottenuti 1) da DeSR addestrato solo sulla Penn Treebank (*BASE* nelle figure), 2) usando in fase di addestramento frasi annotate automaticamente e selezionate in modo casuale nell'ampio corpus di dominio (*RS*) sulla base della media di un processo di *10-fold cross validation*, 3) usando il miglior risultato ottenuto nel processo di *10-fold cross validation* (*best-RS*).



(a) *Abstract* di argomento chimico.



(b) *Abstract* di argomento biomedico.

Figura 13. Accuratezza di analisi (LAS) in uno scenario di *Self-Training* usando le frasi ordinate da ULISSE e dalle *baselines*.

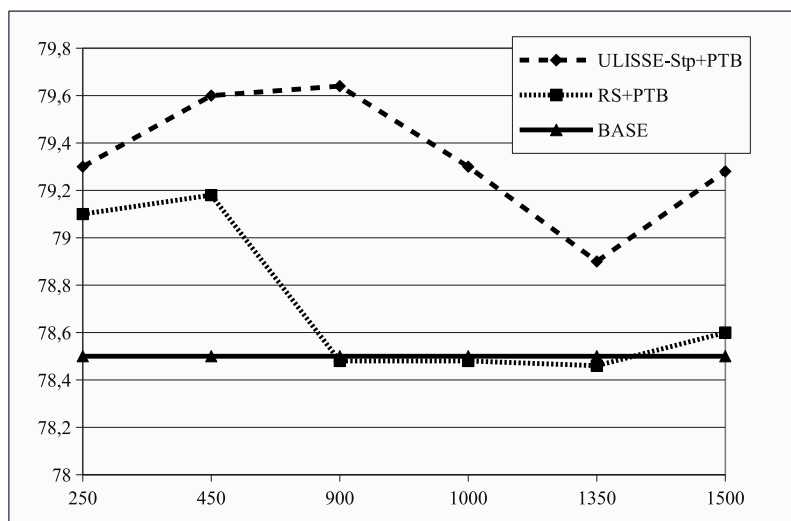
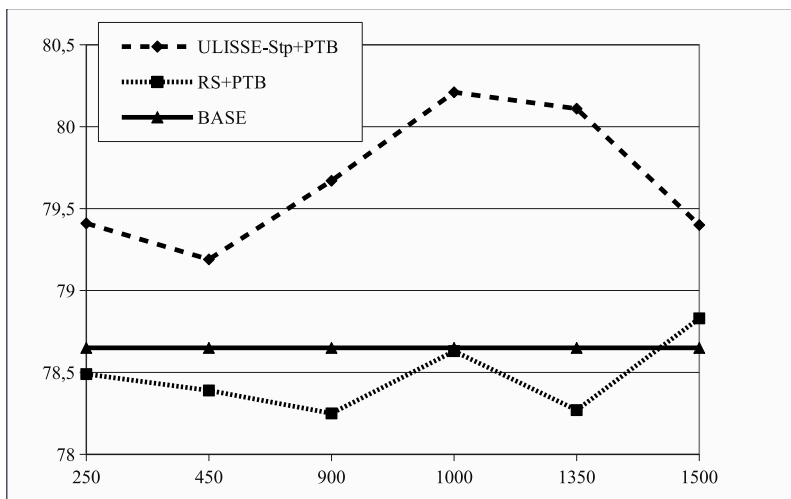
(a) *Abstract di argomento chimico.*(b) *Abstract di argomento biomedico.*

Figura 14. Accuratezza di analisi (LAS) in scenari di *Self-Training* usando le frasi ordinate da ULISSE e dalle *baselines* in aggiunta alle frasi della Penn Treebank.

Sono stati condotti due tipi di esperimenti: un primo nel quale l'accuratezza di DeSR è stata testata su *CHEM* e *BIO* usando in fase di addestramento solo le frasi selezionate da ULISSE e dalle tre *baselines* (si tratta dunque di un corpus

di addestramento costruito in modo completamente automatico); un secondo nel quale a queste frasi sono state aggiunte quelle del corpus di addestramento di partenza di DeSR (in questo caso il corpus di addestramento è composto da una parte annotata in modo manuale e da una in modo automatico). In ogni figura, l'asse delle ordinate riporta l'accuratezza di analisi del *parser* espressa in *Labelled Attachment Score* (LAS), mentre sull'asse delle ascisse sono riportate le dimensioni in *token* delle frasi selezionate da ULISSE e dalle *baseline* (250 migliaia di *token*, 450, 900, ecc.) progressivamente aggiunte in fase di addestramento.

Nel primo tipo di esperimenti (Figure 13.a e 13.b), ULISSE è risultato il miglior algoritmo di *Self-Training*, sebbene nell'esperimento di adattamento al dominio degli *abstract* di argomento biomedico solo aggiungendo 900 mila *token* selezionati da ULISSE l'accuratezza di DeSR migliora rispetto ai risultati raggiunti aggiungendo frasi selezionate dalle *baseline*. In generale, si può comunque notare che 900 mila frasi di addestramento è la quantità massima dopo la quale l'accuratezza del *parser* inizia a decrescere. Ciò significa che dopo un certo numero di nuove frasi ULISSE non riesce più a selezionare dall'ampio corpus annotato automaticamente di partenza *dependency parse* che abbiano un buon livello di affidabilità e siano sufficientemente informative. È interessante notare come DeSR addestrato sui soli *dependency parse* selezionati da ULISSE abbia un'accuratezza maggiore di quella ottenuta usando i testi giornalistici (annotati in modo manuale) della Penn Treebank.

Dai risultati ottenuti nel secondo tipo di esperimenti (Figure 14.a e 14.b), emerge che le frasi selezionate da ULISSE, unite a quelle del corpus di addestramento originale, sono quelle che permettono a DeSR di avere le accuratezze maggiori. Questo risultato dimostra come l'algoritmo di selezione sia non solo in grado di scegliere i *dependency parse* più affidabili ma anche quelli che contengono le specificità del dominio. Sono queste le frasi che, essendo rappresentative degli eventi linguistici tipici del dominio da analizzare, aggiungono nuova informazione al corpus di addestramento originario permettendo a DeSR di adattarsi.

Bibliografia

- Attardi G., M. Simi, A. Zanelli (2013). "Domain Adaptation by Active Learning". *Evaluation of natural language and speech tool for Italian*. B. Magnini, F. Cutugno, M. Falcone, E. Pianta (cur.) Heidelberg: Springer-Verlag, pp. 77–85.
- Attardi G. (2006). "Experiments with a multilanguage non-projective dependency parser". *Proceedings of the Tenth Conference on Computational Natural Language Learning (CoNLL-X '06)*, pp. 166–170.
- Biber D. (1993). "Using register-diversified corpora for general language studies". *Computational linguistics journal*. Vol. 19(2), pp. 219–241.
- Biber D., S. Conrad, R. Reppen (1998). *Corpus linguistics. Investigating language structure and use*. Cambridge: Cambridge University Press.
- Bonzi S. (1990). "Syntactic patterns in scientific sublanguages: a study of four disciplines". *Journal of the American Society for Information Science*, 41 (2), pp. 121–131.

- Bosco C., S. Montemagni, A. Mazzei, V. Lombardo, F. Dell'Orletta, A. Lenci (2009). "Parsing Task: comparing dependency parsers and treebanks." *Proceedings of Evalita'09 (Evaluation of NLP and Speech Tools for Italian)*, [senza numerazione].
- Cassese S. (1992). "Introduzione allo studio della normazione". *Rivista trimestrale di diritto pubblico*, 2, pp. 307–330.
- Clegg A. B., A. J. Shepherd (2005). "Evaluating and integrating treebank parsers on a biomedical corpus". *Proceedings of the Workshop on Software*, pp. 14–33.
- Collins M. (1999). "Head-driven statistical models for natural language parsing". Tesi di dott. University of Pennsylvania.
- Dell'Orletta F., S. Marchi, S. Montemagni, B. Plank, G. Venturi (2012). "The SPLeT-2012 shared task on dependency parsing of legal texts". *Proceedings of the Fourth Workshop on Semantic Processing of Legal Texts (SPLeT 2012)*, pp. 42–51.
- Dell'Orletta F., S. Marchi, S. Montemagni, G. Venturi, T. Agnoloni, E. Francesconi (2013). "Domain adaptation for dependency parsing at Evalita 2011". *Evaluation of natural language and speech tool for Italian*. B. Magnini, F. Cutugno, M. Falcone, E. Pianta (cur.) Heidelberg: Springer-Verlag, pp. 58–69.
- Dell'Orletta F., G. Venturi, S. Montemagni (2011). "ULISSE: an unsupervised algorithm for detecting reliable dependency parses". *Proceedings of CoNLL 2011*, pp. 115–124.
- Dell'Orletta F., G. Venturi, S. Montemagni (2013). "Unsupervised linguistically-driven reliable dependency parses detection and self-training for adaptation to the biomedical domain". *Proceedings of the 12th workshop on Biomedical Natural Language Processing (BioNLP-2013)*, pp. 45–53.
- Frazier L. (1985). "Syntactic complexity", in L. Karttunen D.R. Dowty, A.M. Zwicky (cur.), *Natural language parsing*. Cambridge: Cambridge University Press.
- Gibson E. (1998). "Linguistic complexity: locality of syntactic dependencies". *Cognition*. Vol. 68, pp. 1–76.
- Gildea D. (2001). "Corpus variation and parser performance". *Proceedings of Empirical Methods in Natural Language Processing (EMNLP 2001)*, pp. 167–202.
- R. Grishman, R. Kittredge (cur.), cur. (1986). *Analyzing language in restricted domains: sublanguage description and processing*. Hillsdale: Lawrence Erlbaum.
- Grishman R., N. T. Nhan, E. Marsh, L. Hirschman (1984). "Automated determination of sublanguage syntactic usage". *Proceedings of the 10th International Conference on Computational Linguistics*, pp. 96–100.
- Harris Z. (1968). *Mathematical Structures of Language*. Hoboken: Wiley.
- Kawahara D., K. Uchimoto (2008). "Learning reliability of parses for domain adaptation of dependency parsing". *Proceedings of IJCNLP 2008*, pp. 709–714.
- Kittredge R. (1982). "Variation and homogeneity of sublanguages", in R. Kittredge, J. Lehrberger (cur.), *Sublanguage: studies of language in restricted semantic domains*. Berlin: De Gruyter, pp. 107–137.
- Lehrberger J. (1986). "Sublanguage analysis". *Analyzing language in restricted domains: sublanguage description and processing*. R. Grishman, R. Kittredge (cur.) New York: Lawrence Erlbaum, pp. 19–38.

- Marcus M. P., M. A. Marcinkiewicz, B. Santorini (1993). "Building a large annotated corpus of English: The Penn Treebank". *Computational linguistics*, 19.2, pp. 313–330.
- Marinelli R., L. Biagini, R. Bindi, S. Goggi, M. Monachini, P. Orsolini, E. Picchi, S. Rossi, N. Calzolari, A. Zampolli (2003). "The Italian PAROLE corpus: an overview", in A. Zampolli (cur.), *Linguistica computazionale*. Pisa: Giardini, pp. 401–421.
- McClosky D., E. Charniak (2008). "Self-training for biomedical parsing". *Proceedings of the 46th Annual Meeting of the Association for Computational Linguistics on Human Language Technologies*, pp. 101–104.
- McClosky D., E. Charniak, M. Johnson (2010). "Automatic domain adaptation for parsing". *Proceedings of the HLT-NAACL'2010*, pp. 28–36.
- McClosky D., E. Charniak, M. Johnson (2006). "Reranking and self-training for parser adaptation". *Proceedings of ACL 2006*, pp. 337–344.
- McDonald R., J. Nivre (2007). "Characterizing the errors of data-driven dependency parsing models". *Proceedings of the the EMNLP-CoNLL*, pp. 122–131.
- Jim Miller, Regina Weinert (cur.), cur. (1998). *Spontaneous spoken language. Syntax and discourse*. Oxford: Clarendon Press.
- Nivre J. (2006). *Inductive dependency parsing*. Berlin: Springer.
- Nivre J., J. Hall, S. Kubler, R. McDonald, J. Nilsson, S. Riedel, D. Yuret (2007). "The CoNLL 2007 Shared Task on Dependency Parsing". *Proceedings of the the EMNLP-CoNLL*, pp. 915–932.
- Petrov S., R. McDonald (2012). "Overview of the 2012 shared task on parsing the web". *Proceedings of the First Workshop on Syntactic Analysis of Non-Canonical Language (SANCL 2012)*, pp. 1–8.
- Reichart R., A. Rappoport (2009). "Automatic selection of high quality parses created by a fully unsupervised parser". *Proceedings of CoNLL 2009*, pp. 156–164.
- Sagae K. (2010). "Self-training without reranking for parser domain adaptation and its impact on semantic role labeling". *Proceedings of the 2010 Workshop on Domain Adaptation for Natural Language Processing (DANLP 2010)*, pp. 37–44.
- Sagae K., J. Tsujii (2007). "Dependency parsing and domain adaptation with LR models and parser ensemble". *Proceedings of EMNLP-CoNLL 2007*, pp. 1044–1050.
- N. Sager, C. Friedman, M. Lyman (cur.), cur. (1987). *Medial language processing*. Reading: Addison-Wesley Publishing Company.
- Tateisi Y., A. Yakushiji, T. Ohta, J. Tsujii (2005). "Syntax annotation for the GENIA corpus". *Companion volume to Proceedings of IJCNLP'05*, pp. 222–227.
- Yngve V. H. (1960). "A model and an hypothesis for language structure". *Proceedings of the American Philosophical Society*, 104.5, pp. 444–466.

Alcuni principi linguistici per costruire risorse lessicali

Elisabetta Jezek – Università di Pavia – jezek@unipv.it

1. Introduzione

Risorse linguistiche estese e dettagliate che forniscono informazioni relazionali riguardo ai predicati e ai loro argomenti sono strumenti indispensabili per un ampio numero di applicazioni di *Natural Language Processing* (NLP), che per procedere all'analisi e alla comprensione del testo richiedono di identificare in primo luogo i partecipanti all'evento espresso da un predicato. Per questa ragione, la costruzione di risorse contenenti strutture predicato-argomenti dei verbi è molto diffusa. In particolare, corpora annotati manualmente con questo tipo di informazioni costituiscono spesso la base per lo sviluppo di modelli probabilistici che identificano automaticamente le relazioni semantiche veicolate dai costituenti frasali nel testo, come nel caso del *Semantic Role Labeling* (Gildea, Jurafsky, 2002), oggi implementato anche con analisi distribuzionali (Roth, Woodsend, 2014).

Tuttavia, la codifica di tali informazioni nelle risorse spesso non è uniforme e aderente a uno standard, nonostante gli sforzi effettuati nel quadro di iniziative come ISO (International Organization for Standardization). Per questa ragione, la valutazione delle *performance* computazionali di tali risorse risulta difficile. In questo contributo, affronterò il problema della corretta identificazione e codifica delle strutture appropriate per i verbi, che ritengo costituisca una delle ragioni primarie di molte delle difficoltà incontrate in questo ambito. Per fare ciò, attraverso un *case study* incentrato su un repertorio *corpus-based* di strutture predicato-argomento per i verbi italiani, esaminerò i problemi linguistici che emergono nella sua costruzione e commenterò le scelte operate relative alla codifica delle informazioni semantiche, alla luce del complesso *trade-off* tra esigenze di modellizzazione linguistica e necessità applicative di tipo computazionale. La struttura del contributo è la seguente: nella Sezione 2 introdurrò la risorsa oggetto d'analisi e la metodologia impiegata; nelle Sezioni 3, 4 e 5 mi focalizzerò su tre differenti fenomeni semantici che si presentano nella sua costruzione e sulle scelte di codifica adottate; nella Sezione 6 trarrò le conclusioni di tale indagine.

2. La risorsa

La risorsa oggetto dell'analisi è T-PAS, un archivio di *Typed Predicate Argument Structures* (strutture argomentali 'tipate') per l'italiano acquisite da corpora attraverso il *clustering* manuale di informazioni distribuzionali sui verbi italiani.⁶⁸ I

⁶⁸T-PAS è sviluppato presso il Dipartimento di Studi Umanistici dell'Università di Pavia, in collaborazione con il gruppo *Human Language Technology* della Fondazione Bruno Kessler (FBK) di Trento

T-PAS sono *pattern* verbali estratti da corpora che includono la specificazione del tipo semantico atteso per ogni posizione argomentale, come per esempio [[Human]] *guida* [[Vehicle]].⁶⁹ La risorsa è concepita per condurre analisi linguistica e per applicazioni volte alla processazione del linguaggio naturale e, nello specifico, all'identificazione di differenze di significato per i verbi basate su *pattern* testuali. T-PAS è la prima risorsa per l'italiano in cui le proprietà di selezione semantica e le distinzioni di senso dei predicati verbali sono descritte su base esclusivamente empirica. I *pattern* verbali più salienti sono individuati utilizzando una procedura lessicografica chiamata *Corpus Pattern Analysis* (CPA, Hanks, 2004; Hanks, 2013), la quale è basata sull'analisi della co-occorrenza statistica degli elementi lessicali degli *slot* sintattici dei verbi in esempi estratti da corpora.

Punti di riferimento per il progetto T-PAS sono FrameNet (Ruppenhofer et al., 2010) e VerbNet (Kipper-Schuler, 2005). FrameNet e VerbNet sono però differenti da T-PAS perché le strutture che essi identificano non sono acquisite da corpora secondo una procedura sistematica (si veda, tuttavia, Bonial et al., 2013). Un'altra risorsa da menzionare è PDEV (Hanks, Pustejovsky, 2005), un dizionario di *pattern* dei verbi inglesi, il quale costituisce il risultato principale dell'applicazione della procedura CPA all'inglese. Quanto all'italiano, un progetto complementare è LexIt (Lenci et al., 2012), una risorsa che fornisce informazioni distribuzionali acquisite automaticamente su verbi, aggettivi e nomi. A differenza di T-PAS, LexIt non produce un inventario di *pattern* e le categorie utilizzate per classificare la semantica degli argomenti non sono derivate da corpora (si veda la Sezione 3.). Repertori di significati come MultiWordNet (Pianta et al., 2002) e Senso Comune (Oltramari et al., 2013) sono risorse alle quali T-PAS può essere collegato con facilità, nell'ottica di popolare le prime con distinzioni di senso per i verbi basate su *pattern* estratti da corpora.

2.1. Panoramica della risorsa

La prima versione di T-PAS contiene l'analisi di 1000 verbi a polisemia media, selezionati attraverso un'estrazione random dal set totale dei lemmi verbali ad alta disponibilità del Sabatini, Coletti (2007), secondo le seguenti proporzioni: 10% di verbi a due significati, 60% di verbi da tre a cinque significati, 30% di verbi da sei a undici significati. La risorsa è formata da tre componenti (cfr. Figura 15.):

- a. un repertorio di T-PAS estratti da corpora e connessi a unità lessicali (verbi);
- b. un elenco di circa 230 tipi semantici per i nomi (ovvero, [[Human]], [[Location]], [[Event]], [[Artifact]]), ottenuti attraverso la generalizzazione di set di elementi lessicali trovati in posizione argomentale all'interno del corpus;

e con il supporto tecnico della Facoltà di Informatica della Masaryk University a Brno (Repubblica Ceca); la risorsa è disponibile con licenza Creative Common Attribution 3.0 e scaricabile sul sito di T-PAS, URL:<<http://tpas.fbk.eu>>.

⁶⁹I nomi dei tipi semantici sono convenzionalmente annotati tra doppie parentesi quadre e con la lettera maiuscola. Inoltre, il metalinguaggio è inglese, come per le altre componenti della risorsa (per esempio le relazioni grammaticali).

- c. un corpus di frasi che istanziano T-PAS, annotate per unità lessicale (verbo) e numero di *pattern*.⁷⁰

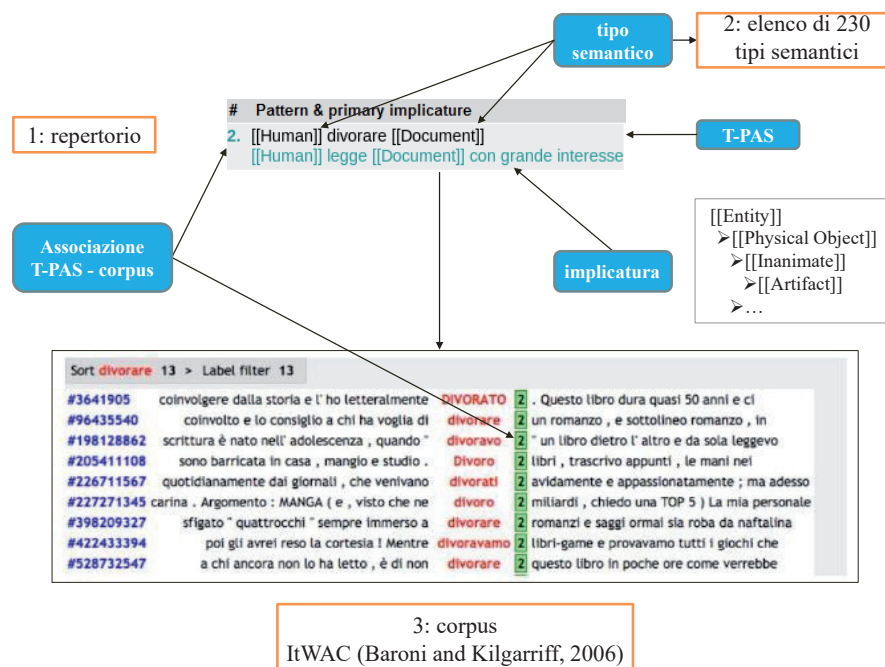


Figura 15. La risorsa T-PAS

Al momento, la metodologia di acquisizione di T-PAS e il *tagging* dei tipi semantici avvengono interamente in maniera manuale. In particolare, il lessicografo: 1) sceglie un verbo target e crea un campione di 250 concordanze nel corpus; 2) esaminando le righe del corpus, identifica i vari tipi di strutture sintagmatiche che corrispondono ai contesti minimi in cui tutte le parole, incluso il verbo, sono disambiguate; 3) identifica le preferenze di tipo semantico di ogni posizione argomentale della struttura, controllando il set lessicale dei *filler* e assegnando un tipo semantico atteso a ogni posizione argomentale; 4) registra i *pattern* identificati nella risorsa usando il *Pattern Editor*; a ogni *pattern* corrisponde un numero identificativo; 5) assegna tutti i 250 esempi del campione ai *pattern* corrispondenti annotando ogni riga del corpus con il numero del *pattern* ad essa associato; 6) propone una descrizione del significato associato al *pattern*, espresso sotto forma

⁷⁰Il corpus di riferimento è una versione ridotta di ItWac (Baroni, Kilgarriff, 2006), messo a punto presso il laboratorio di Natural Language Processing della Masaryk University (Repubblica Ceca) da J. Pomikalek rimuovendo i *7-grams duplicates* da ItWac. I *corpus tools* impiegati sono stati Manatee, Bonito e Sketch Engine (Kilgarriff et al., 2004).

di implicatura ancorata alle preferenze di selezione del *pattern*; per esempio, il T-PAS nella Figura 15. ha l'implicatura: *[[Human]] legge [[Document]] con grande interesse.*

Tutti i passaggi appena menzionati sono spiegati nel dettaglio nelle linee guida fornite al lessicografo. La procedura sopra descritta solleva una serie di questioni a livello teorico. Nello specifico: a) definire quale set di tipi semantici è necessario e appropriato per identificare le proprietà di selezione argomentale dei verbi e organizzarlo in una struttura tassonomica che consenta di identificare il livello di selezione; b) rappresentare il fatto che gli elementi lessicali che popolano specifiche posizioni argomentali nell'ambito di uno stesso significato verbale spesso non corrispondono al tipo semantico atteso in quella posizione, e c) rappresentare il fatto che i membri dei set lessicali di singoli tipi (per esempio *[[Garment]]*) cambiano da verbo a verbo in modo analogo alle 'somiglianze di famiglia' individuate da Wittgenstein, così che non sembra possibile effettuare solide generalizzazioni riguardo a quale sia il set di elementi lessicali rappresentativo di un tipo specifico. Nella rimanente parte del contributo, esaminerò nell'ordine ciascuno di questi tre punti.

3. Preferenze di selezione e categorie ontologiche

Come già menzionato nella Sezione 2.1., applicando la procedura CPA all'analisi delle concordanze per circa 2000 verbi inglesi e italiani è stata ottenuta una lista di circa 230 tipi semantici attraverso il *clustering* manuale a partire dal set di elementi lessicali trovati nelle posizioni argomentali dei *pattern* all'interno del corpus. Sebbene i tipi identificati assomiglino molto a categorie concettuali/ontologiche per i nomi, devono in realtà essere considerati oggetti linguistici, ovvero classi semantiche, dato che sono indotti dall'analisi delle proprietà di selezione dei verbi di cui i nomi costituiscono argomento. Essi riflettono il modo in cui attraverso la lingua predichiamo proprietà o relazioni tra entità reali o immaginarie e concorrono a determinare il significato acquisito dal verbo nel contesto.

Per esempio, *[[Furniture]]*, come in *[[Human]] arredare [[Room|Building]]* (con *[[Furniture]]*), non è una categoria astratta ma una classe semantica identificata generalizzando la lista dei collocati statisticamente rilevanti che riempiono la posizione argomentale. Un esempio è il seguente:

(20.) *arredare*

[[Human]] arreda [[Room|Building]] (con *[[Furniture]]*)

Set lessicale *[[Furniture]]* = {mobile, mobilio, mobili, tavolo, letto, tavolino, divano, sedia, statua, quadro...}

I tipi semantici differiscono dalle categorie di entità definite sulla base di assiomi ontologici, come quelle di DOLCE (*Descriptive Ontology for Linguistic and Cognitive Engineering*). Queste, sebbene aspirino a catturare le categorie concet-

tuali sottostanti al linguaggio naturale,⁷¹ non basano le distinzioni categoriali su osservazioni sistematiche e sul *clustering* di dati linguistici.

Nel progetto, la lista dei tipi identificati attraverso l'analisi *bottom-up* illustrata nella Sezione 2. è organizzata in una gerarchia al fine di individuare il livello di selezione dei singoli verbi. La relazione principale della struttura tassonomica è la relazione 'IS_A' (*subsumption*), come nel caso di [[Plane]] che è un tipo di [[Vehicle]], che è un tipo di [[Artifact]] (vedi Tabella 5.), e così via.

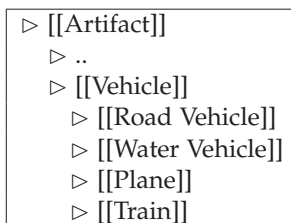


Tabella 5. Frammento della tassonomia.

Allo stato attuale, l'organizzazione tassonomica dei tipi è prevalentemente basata su scelte che riflettono le intuizioni del lessicografo riguardo al significato dei termini trovati nel corpus e sulla comparazione manuale di set lessicali appartenenti a verbi che entrano in una relazione 'IS_A'.

Una delle proprietà chiave della tassonomia è che non mostra la stessa granularità per tutto lo spazio semantico. Al contrario, risulta essere asimmetrica e irregolare, e rivela maggiore attenzione alle attività in cui partecipano esseri umani rispetto, per esempio, all'astronomia e alla scienza in generale (il che conferma che la lingua di tutti i giorni è più antropocentrica rispetto alla scienza). In altre parole, l'antropocentricità è una caratteristica primaria della tassonomia dei tipi indotta dall'analisi di corpora. Per fare un esempio, la lista dei tipi differenzia un gran numero di strumenti e altri artefatti come tipi semantici, che sono creati e usati per scopi diversi dagli esseri umani. L'intero universo al di fuori del pianeta terra, invece, è rappresentato solo da pochi elementi lessicali che rientrano sotto il tipo [[Location]]. Inoltre, la lista dei tipi individua [[Horses]] e [[Dogs]] come tipi semantici, ma non ha interesse nel concetto scientifico di <chordate>. Questo avviene perché gli esseri umani *cavalcano* un cavallo, *ferrano* un cavallo, *sellano* un cavallo; i cavalli *si imbibizzarriscono*, i cavalli *galoppiano* e i cavalli *trottano* o *vanno al trotto*. Per tutti questi verbi, il concetto <horse> serve a distinguere un particolare tipo di evento da eventi che coinvolgono altri animali (Hanks, 2000). Dall'altro lato, invece, non esistono eventi codificati nel linguaggio naturale in cui la classe <chordate> assuma valore distintivo.

Nei *task* semantici di disambiguazione (come ad esempio *Word Sense Disambiguation*) è indispensabile un meccanismo di base che raggruppi parole e che distin-

⁷¹"[A]iming at capturing the ontological categories underlying natural language and human common sense" (Gangemi et al., 2002).

gua ‘un cavallo che galoppa’ da ‘un’inflazione che galoppa’ o ‘una disoccupazione che galoppa’. La tassonomia dei tipi è fondata su tale meccanismo, consente di distinguere un *pattern* di utilizzo del verbo da tutti gli altri, e predice il senso che il verbo assume nei diversi contesti.

In conclusione, l’informazione sulle preferenze di selezione indotte dall’analisi di corpora è centrale in una risorsa lessicale e va distinta dall’informazione di natura ontologica, in quanto la prima riguarda il piano della lingua, la seconda quello concettuale. La distinzione tra tipi semantici indotto da corpora e categorie ontologiche è spesso sottovalutata nella costruzione di risorse lessicali; per la disambiguazione sono spesso impiegate categorie definite a prescindere dall’uso linguistico, con risultati limitati indotti dalla metodologia stessa.

4. Mismatch dei tipi semantici

Come descritto nella Sezione 2., per acquisire i *pattern* dai dati del corpus, l’annotatore di T-PAS recupera una ad una le concordanze che presentano il verbo analizzato (limitatamente al campione analizzato), individua le strutture rilevanti, analizza il set lessicale per ogni relazione grammaticale e assegna a ogni posizione argomentale un tipo semantico atteso. Un problema ricorrente che sorge in questa fase è l’assenza di corrispondenza del tipo semantico associato ai *filler* argomentali nell’ambito dello stesso senso del verbo con il tipo semantico atteso per quella posizione; per esempio, generalmente raggiungiamo una *[[Location]]* (‘raggiungere l’isola’), ma nel corpus troviamo molti contesti in cui si raggiungono *[[Artifacts]]*, come in ‘raggiungere il semaforo’, ‘raggiungere la macchina’, eccetera.

La mancanza di corrispondenza tra il tipo semantico assegnato dal verbo e il tipo semantico associato al *filler* argomentale all’interno della stessa relazione grammaticale può essere definito come un caso di *mismatch* (dei tipi semantici).⁷² I *mismatch* possono essere classificati secondo vari parametri. Di seguito sono riportati tre esempi di *mismatch* classificati secondo i criteri seguenti: 1) classe verbale; 2) relazione grammaticale target (in corsivo); tipo di *shift* (negli esempi, *[[Human]] as [[Event]]*, *[[Artifact]] as [[Human]]*, *[[Event]] as [[Human]]*, *[[Artifact]] as [[Location]]*, *[[Event]] as [[Location]]*, eccetera).⁷³

(21.) Verbi aspettuali

[[Human]-subj interrompe [[Event]-obj]

‘Arriva Mirko e interrompe *la conversazione*.’ (matching)

‘Il presidente interrompe *l’oratore*.’ (mismatch: *[[Human]] as [[Event]]*)

(22.) Verbi di comunicazione

⁷²In letteratura, l’espressione *semantic type coercion* è spesso usata per rendere conto del fenomeno sottostante ai casi di *mismatch* dei tipi semantici (si veda Pustejovsky, Jezek (2008) per un’analisi *corpus-based* di casi di *coercion*).

⁷³Negli esempi, le occorrenze sono associate ai tipi semantici derivati da uno studio CPA di questi verbi. Questo studio è stato usato come base per costruire il *dataset* per il SemEval-2010 *shared task* sul fenomeno della *coercion*, descritto in Pustejovsky, Rumshisky et al. (2010).

[[Human]-subj] *annuncia* [[Event]-obj]

'Lo *speaker* annuncia la partenza.' (matching)

'L'*altoparlante* annunciava l'arrivo del treno.' (mismatch: [[Artifact]] as [[Human]])

'Una *telefonata anonima* avvisa la polizia.' (mismatch: [[Event]] as [[Human]])

(23.) Verbi di movimento telici

[[Human]-subj] *raggiunge* [[Location]-obj]

'Abbiamo raggiunto l'*isola* alle 5.' (matching)

'Ho raggiunto il *semaforo* e ho svoltato a destra.' (mismatch: [[Artifact]] as [[Location]])

'Gli invitati arrivano al *concerto* in ritardo.' (mismatch: [[Event]] as [[Location]])

Le informazioni relative ai *mismatch* sono importanti, tanto a livello linguistico quanto a livello computazionale, in quanto permettono di individuare le metonimie, molto diffuse nei testi (Markert, Nissim, 2007). Nella risorsa T-PAS, tali informazioni sono presenti secondo due codifiche. Se due tipi semantici alternano regolarmente all'interno della stessa posizione argomentale, la codifica è l'alternanza di tipo a livello di singolo T-PAS. Alternanze di tipo comuni in posizione di soggetto sono per esempio [[Human|Institution]] e [[Human|Body Part]], come in: [[Human|Body Part]] *sanguina*. Elementi lessicali non canonici che superano una determinata soglia statistica sono invece etichettati come membri 'onorari' di un tipo semantico in contesti specifici. Per esempio, *semaforo* e *concerto* sono inseriti come membri del set lessicale del tipo [[Location]] nel contesto di *raggiungere*, ma non nel contesto di altri verbi. I membri onorari sono annotati come 'a' = argomenti anomali. In questo modo, tutti i casi di *mismatch* possono essere facilmente estratti dalla risorsa e possono essere utilizzati come dati di addestramento in applicazioni rivolte alla risoluzione della metonimia (Pustejovsky, Rumshisky et al., 2010).

5. Set lessicali instabili

Un'altra caratteristica saliente dei set lessicali, oltre al fatto che non combaciano perfettamente con i tipi semantici, è che i membri di ciascun set possiedono una debole coesione semantica. I set lessicali che popolano un tipo semantico si mostrano *shimmering* (lett. 'cangianti') (Jezek, Hanks, 2010), ovvero i loro membri cambiano da verbo a verbo: alcune parole entrano mentre altre escono, secondo un principio che ricorda le 'somiglianze di famiglia' di Wittgenstein. Verbi diversi selezionano differenti membri prototipici di un tipo semantico anche se il resto del set rimane uguale. Per esempio, i verbi *lavare* e *amputare* tipicamente selezionano entrambi [[Body Part]] come oggetto diretto. Una persona può lavare un(a) {gamba | braccio | piede}, oppure amputare un(a) {gamba | braccio | piede}. Prototipicamente, tuttavia, una persona lava {viso | mani | capelli}, ma non amputa {viso | mani | capelli}:

- (24.) *lavare* [[Body Part]-obj]
 Set lessicale [[Body Part]] = {denti, mano, piede, viso, faccia, schiena, testa, orecchio, volto}
- (25.) *amputare* [[Body Part]-obj]
 Set lessicale [[Body Part]] = {arto, gamba, braccio, dito, orecchio, mano, piede}

Un altro esempio è fornito da verbi che esprimono azioni tipicamente svolte con membri del tipo [[Document]] (come *leggere, pubblicare, spedire* o *tradurre*):

- (26.) *leggere* = {opinione, commento, libro, avvertenza, recensione, giornale, cronaca, intervista, articolo, blog, messaggio, poesia, notizia, resoconto, racconto, fumetto, romanzo}
- pubblicare* = {articolo, intervista, nota, avviso, bando, studio, libro, sondaggio, notizia, volume, saggio, testo, giornale, comunicato, monografia, rivista, documento, inesattezza, annuncio}
- spedire* = {cartolina, pacco, e-mail, fax, messaggio, lettera, raccomandata, foto, telegramma, copia, sms, merce, curriculum, pacchetto, invito, vaglia, posta, file, libro, coupon, modulo}
- tradurre* = {testo, frase, bibbia, parola, brano, libro, poesia, vocabolo, concetto, opera, versetto, termine, idea, nome, spiegazione, romanzo, canzone, pensiero, vangelo, espressione}

I dati in 26. mostrano che, sebbene da un punto di vista intuitivo [[Document]] sia un tipo facilmente definibile, i suoi membri a livello linguistico variano a seconda del contesto verbale (la sottolineatura individua i membri presenti in più set). Ciò riflette il fatto che nel mondo reale non compiamo esattamente gli stessi tipi di operazioni con gli oggetti denotati dai nomi che rappresentano questo tipo. Per esempio, il tradurre è un'attività che tipicamente viene fatta con [[Document]] come i *libri*, le *poesie* e i *romanzi*, mentre ci sono altri [[Document]] come i *giornali* che non vengono in genere tradotti – se anche in linea di principio sarebbe possibile farlo – ma con i quali svolgiamo altre attività. Un *giornale* tipicamente viene letto, un *messaggio* tipicamente viene mandato, e così via (si veda Jezek, Lenci, 2007, per una discussione approfondita del comportamento distribuzionale delle parole che denotano [[Document]]).

Questi dati sui set lessicali permettono di capire come sia complessa la relazione tra il sistema concettuale e le restrizioni di selezione che i verbi impongono ai loro argomenti. Il comportamento selettivo dei verbi è strettamente connesso con la struttura del nostro sistema concettuale, ma le restrizioni di selezione non combaciano perfettamente con i tipi semantici. In altre parole, quando parliamo di una classe di oggetti, isoliamo tale classe dal nostro sistema concettuale in un modo che è parzialmente indipendente dal suo stato categoriale. Ne consegue che l'assunto secondo cui le restrizioni di selezione possano essere catturate in termini di tipi semantici è troppo forte.

L'analisi empirica dei set lessicali che popolano specifiche posizioni argomentali in relazione ai verbi target può migliorare la nostra comprensione del fenomeno della selezione degli argomenti e offrire degli indizi più facili da riconoscere per poter efficacemente operare la disambiguazione dei significati lessicali in contesto. Per tali ragioni, questo tipo di informazione è rilevante a livello di risorsa lessicale. Allo stato attuale, tale informazione non è presente in T-PAS. In termini puramente statistici, tale informazione può tuttavia essere aggiunta in modo automatico annotando, per ogni membro canonico del set, informazioni riguardo al contesto nel modo seguente:⁷⁴

(27.) [[Document]]:

```
{...libro <leggere_7429/284411; scrivere_4435/510270;
pubblicare_2207/188905, recensire_271/8861...>
...romanzo <scrivere_1137/510270; pubblicare_587/188905;
leggere_700/284411; ambientare_89/10418...>
...articolo <pubblicare_4024/188905; apparire_1303/154907;
scrivere_1352/510270; leggere_482/284411...>
...lettera <scrivere_3484/510270; inviare_2276/315547;
spedire_395/23552; ricevere_1426/164275...>}
```

In tale sistema, un nodo nella tassonomia dei tipi (ovvero un tipo semantico) non è concepito come il target di tutti e i soli elementi lessicali che appartengono a quel nodo, ma piuttosto come il target degli elementi lessicali che tipicamente appartengono a quel nodo. La tassonomia dei tipi non è concepita nei termini di una struttura rigida sì/no, ma come una struttura di preferenze collocazionali basata su dati statistici, chiamata *shimmering lexical sets* e popolata con parole che entrano e escono dai tipi semantici a seconda del contesto, e la cui frequenza relativa nei vari contesti può essere misurata (e comparata) nei corpora. Una struttura *shimmering* di questo tipo preserva, seppure in una forma più debole, i vantaggi predittivi dell'organizzazione gerarchica concettuale e allo stesso tempo mantiene la validità empirica della descrizione del linguaggio naturale.

6. Conclusioni

In questo contributo ho esaminato un repertorio di strutture predicato-argomento per i verbi italiani, tipizzate e derivate dall'analisi di corpora. Obiettivo di questo lavoro era utilizzare questa risorsa come *case study* per evidenziare la complessità del *trade-off* tra modellizzazione linguistica ed esigenze computazionali nella costruzione di risorse lessicali per scopi di NLP. Le conclusioni principali che possono essere tratte da questa analisi sono le seguenti. I tipi semantici indotti dall'indagine dei set lessicali degli argomenti dei verbi nei corpora non devono essere

⁷⁴In 27. il primo numero (7429) è il numero totale delle occorrenze di *libro* con *leggere*, mentre il secondo (284411) è il numero totale delle occorrenze di *leggere* nel corpus di riferimento (una versione ridotta di ItWac – si veda la nota 70).

confusi con le categorie ontologiche; è opportuno tenere i due livelli concettualmente e strutturalmente separati. In alcuni casi, il comportamento linguistico può essere spiegato sulla base di restrizioni ontologiche, mentre in altri è necessario rendere conto di fenomeni linguistici che riflettono l'ontologia solo parzialmente. Per scopi di NLP, i tipi semantici *corpus-driven* sono rilevanti in quanto riflettono fenomeni linguistici, ovvero le preferenze di selezione dei verbi che sono predittive dei loro significati. L'annotazione di un corpus con tali informazioni ha inoltre come prodotto collaterale la costruzione di un repertorio *corpus-based* di spostamenti metonimici, dato che il *mismatch* tra set lessicali e tipi semantici è perlopiù dovuto a interpretazioni metonimiche indotte dal contesto. Anche tale fenomeno linguistico è di grande interesse per le applicazioni computazionali. Recenti lavori propongono di ricavare l'informazione sui set lessicali, anziché attraverso l'annotazione manuale, in modo automatico utilizzando procedure di *clustering* di rappresentazioni vettoriali (Ponti et al., 2016) o attraverso risorse esterne come per esempio WordNet (Feltracco et al., 2016). Ricerche future potranno condurre all'elaborazione di metodologie semi-automatiche per l'individuazione di T-PAS e per la strutturazione della tassonomia dei tipi.

Ringraziamenti

Desidero ringraziare Patrick Hanks per il suo contributo e supporto scientifico al progetto T-PAS e James Pustejovsky, Laure Vieu, Alessandro Lenci, Chiara Melloni, Francesca Frontini e Valeria Quochi per la discussione dei fenomeni di *coercion*. Inoltre, ringrazio Bernardo Magnini, Anna Feltracco e Edoardo Maria Ponti per la collaborazione in merito all'analisi dei set lessicali, e gli organizzatori del convegno *Compter parler soigner. Tra linguistica e intelligenza artificiale*, Collegio Ghislieri, Pavia, 15–17 dicembre 2014, dove il contenuto di questo contributo è stato presentato, per i loro utili commenti e suggerimenti.

Bibliografia

- Baroni M., A. Kilgarriff (2006). "Large linguistically-processed web corpora for multiple languages". *Proceedings of the Eleventh Conference of the European Chapter of the Association for Computational Linguistics: Posters & Demonstrations*, pp. 87–90.
- Bonial C., O. Hargraves, M. Palmer (2013). "Expanding VerbNet with Sketch Engine". *Proceedings of the 6th International Conference on Generative Approaches to the Lexicon*, pp. 44–54.
- Feltracco A., L. Gatti, S. Magnolini, B. Magnini, E. Jezek (2016). "Using WordNet to build lexical sets for Italian verbs". *Proceedings of the 8th International Global WordNet Conference 2016*, pp. 100–104.
- Gangemi A., N. Guarino, C. Masolo, A. Oltramari, L. Schneider (2002). "Sweetening ontologies with DOLCE", in A. Gómez-Pérez, V.R. Benjamins (cur.), *Knowledge engineering and knowledge management: ontologies and the semantic Web*. Berlin: Springer, pp. 166–181.

- Gildea D., D. Jurafsky (2002). "Automatic labeling of semantic roles". *Computational linguistics*, 28.3, pp. 245–288.
- Hanks P. (2000). "Contributions of lexicography and corpus linguistics to a theory of language performance". *Proceedings of the Ninth EURALEX International Congress*, pp. 3–13.
- Hanks P. (2004). "Corpus pattern analysis". *Proceedings of the XI Euralex International Congress*. Vol. 1, pp. 87–98.
- Hanks P. (2013). *Lexical analysis: norms and exploitations*. Cambridge (Massachusetts): MIT Press.
- Hanks P., J. Pustejovsky (2005). "A pattern dictionary for natural language processing". *Revue Française de linguistique appliquée*, 10.2, pp. 63–82.
- Jezek E., P. Hanks (2010). "What lexical sets tell us about conceptual categories". *Lexis*, 4, pp. 7–22.
- Jezek E., A. Lenci (2007). "When GL meets the corpus: a data-driven investigation of semantic types and coercion phenomena". *Proceedings of GL*, pp. 10–11.
- Kilgarriff A., P. Rychly, P. Smrz, D. Tugwell (2004). "The Sketch Engine". *Proceedings of the XI Euralex International Congress*, pp. 105–111.
- Kipper-Schuler K. (2005). "VerbNet: a broad-coverage, comprehensive verb lexicon". Tesi di dott. University of Pennsylvania.
- Lenci A., G. Lapesa, G. Bonansinga (2012). "LexIt: a computational resource on Italian argument structure". *Proceedings of the Eighth International Conference on Language Resources and Evaluation (LREC'12)*, pp. 3712–3718.
- Markert K., M. Nissim (2007). "SemEval-2007 task 8: metonymy resolution". *Proceedings of the Fourth International Workshop on Semantic Evaluations (SemEval-2007)*, [senza numerazione].
- Oltramari A., G. Vetere, I. Chiari, E. Jezek, F. M. Zanzotto, M. Nissim, A. Gangemi (2013). "Senso Comune: a collaborative knowledge resource for Italian", *The People's Web Meets NLP: Collaboratively Constructed Language Resources*. Berlin: Springer, pp. 45–67.
- Pianta E., L. Bentivogli, C. Girardi (2002). "Multiwordnet: developing an aligned multilingual database". *Proceedings of the First International Conference on Global WordNet*, pp. 293–302.
- Ponti E. M., E. Jezek, B. Magnini (2016). "Grounding the Lexical Sets of Anti-Causative Pairs on a Vector Model". *DSALT: Distributional Semantics And Linguistic Theory*.
- Pustejovsky J., E. Jezek (2008). "Semantic coercion in language: Beyond distributional analysis". *Italian journal of linguistics*, 20.1, pp. 175–208.
- Pustejovsky J., A. Rumshisky, A. Plotnick, E. Jezek, O. Batiukova, V. Quochi (2010). "SemEval-2010 task 7: argument selection and coercion". *Proceedings of the 5th international workshop on semantic evaluation*, pp. 27–32.
- Roth M., K. Woodsend (2014). "Composition of word representations improves semantic role labelling". *Proceedings of EMNLP'14*, pp. 407–413.
- Ruppenhofer J., M. Ellsworth, M. R. Petruck, C. Johnson, J. Scheffczyk (2010). *FrameNet II: Extended theory and practice*. Manuscript. International Computer Science Institute, University of Berkeley.

Sabatini F., V. Coletti (2007). *Il Sabatini-Coletti. Dizionario della lingua italiana 2008*. Milano: Rizzoli-Larousse.

Alcune note sui Teoremi di Incompletezza di Gödel e la conoscenza delle macchine

Alessandro Aldini*, Vincenzo Fano[‡] e Pierluigi Graziani[‡] – *[‡]Università di Urbino, [‡]Università di Chieti-Pescara – {alessandro.aldini,vincenzo.fano}@uniurb.it pierluigi.graziani@unich.it

1. Gli argomenti gödeliani

Un'idea abbastanza diffusa è quella secondo cui sia possibile rappresentare l'intera soggettività umana mediante algoritmi (ad esempio: D. Dennett, J. Fodor, P. Churchland). Non ci occuperemo di questo punto di vista fortemente metafisico, che incontra difficoltà in relazione ai celebri esperimenti mentali proposti da T. Nagel, J. Searle e F. Jackson. Per contro, ci interesseremo al progetto, più limitato, di riprodurre o simulare meccanicamente i comportamenti intelligenti dell'uomo, lanciato dal celebre articolo di Turing del 1950, progetto usualmente chiamato 'meccanicismo'. Tale progetto è stato ampiamente discusso e, se da un lato ha fornito strumenti ai suoi sostenitori, dall'altro ha indotto gli anti-meccanicisti a costruire sue confutazioni.⁷⁵ Ciò che qui ci interessa è analizzare come i teoremi di Gödel siano stati quasi subito visti come uno strumento per confutare la tesi meccanicista, sia essa intesa in modo estensionale (le procedure e i risultati della mente umana sono meccanizzabili), sia intensionale (la mente umana è una particolare macchina). Turing stesso comprese tale implicazione dei teoremi di Gödel (Turing, 1950); inoltre P. Rosenbloom, G. Kemeny, E. Nagel e J. R. Newman negli anni Cinquanta dello scorso secolo elaborarono argomentazioni aventi al loro centro l'idea che i teoremi di Gödel potessero offrire uno strumento logico per confutare la tesi filosofica del meccanicismo.⁷⁶ Nonostante questa ampia tradizione, l'argomento gödeliano anti-meccanicista è legato al nome del filosofo inglese John Randolph Lucas, che nel 1961 (Lucas, 1961) elaborò un'argomentazione avente lo scopo di dimostrare, basandosi sui teoremi di Gödel, che non è possibile rappresentare l'intelligenza umana con una Macchina di Turing (Turing, 1936). L'argomento elaborato da Lucas può essere così schematizzato:⁷⁷

⁷⁵Testi interessanti sono, ad esempio, Marciszewski, Murawski (1995) e Webb (1980).

⁷⁶Vedi Bruni (2004) e Fano, Graziani (2013).

⁷⁷Ricordiamo l'enunciato del cosiddetto primo Teorema di Gödel: se U è un sistema formale coerente, in grado di esprimere una certa porzione di aritmetica, allora esiste un enunciato G , formulato nel linguaggio L del sistema, tale che G è indecidibile in U , ossia né dimostrabile, né refutabile. L'incompletezza sintattica di U ha la conseguenza che vi è una formula ben formata chiusa vera nel modello standard del sistema che non è dimostrabile nel sistema. Il secondo Teorema di Gödel è di fatto un corollario del primo e asserisce che: se U è un sistema formale coerente in grado di esprimere una certa porzione di aritmetica, allora U non può provare la propria coerenza.

- L1. Supponiamo per assurdo che esista una Macchina di Turing (MT) che abbia esattamente le stesse capacità intellettuali dell'uomo.
- L2. Essa deve essere in grado di produrre tutti i teoremi di un qualche sistema formale U che contiene l'aritmetica formale (AF).⁷⁸
- L3. Tuttavia essa non è in grado di produrre come vera la formula G (di Gödel) di U .
- L4. Per contro l'uomo ha la capacità intellettuale di vedere che G è vera.
- L5. Dunque MT non riproduce tutte le capacità intellettuali dell'uomo.

L'enunciato indecidibile G di AF viene deciso attraverso un'argomentazione meta-matematica: G dice, tramite 'gödelizzazione', di non essere dimostrabile; se fosse dimostrabile sarebbe falso, ma poiché U non dimostra falsità, allora G non deve esservi dimostrabile e, dunque, G è vero. Questo argomento non può essere formalizzato in AF, poiché richiederebbe di definirvi la nozione di verità che, per un famoso teorema dimostrato dal logico Alfred Tarski, non può essere formalizzata in AF. Dunque, G viene deciso nella metateoria di AF, ovvero in un ambiente in cui abbiamo la possibilità di definire la verità per il linguaggio di AF.

L'argomento di Lucas, pur essendo all'apparenza assai persuasivo, contiene alcuni problemi che sono stati ampiamente dibattuti in letteratura. Nonostante tali difficoltà, l'attrazione di tale argomentazione ha spinto diversi studiosi a tentare completamenti e rielaborazioni del ragionamento lucasiano.⁷⁹ Un argomento simile a quello di Lucas era stato già sostenuto, come si è detto, da Newman, Nagel (1958), ma era stato criticato da Putnam (1961) mediante l'osservazione che il passo L3 è discutibile in quanto il teorema di Gödel per il sistema formale U afferma che G non è decidibile solo se U è coerente. Data la macchina MT che dovrebbe riprodurre tutta l'attività intellettuale dell'uomo, cioè l'insieme degli enunciati di U , è possibile trovare un enunciato indecidibile G tale che si può dimostrare in U che se U è coerente allora G è vero ma non è decidibile da MT. Ma per poter ridurre all'assurdo la tesi meccanicista è necessario dimostrare la coerenza di U . Non è semplice sapere, tuttavia, se la macchina (MT) che dovrebbe simulare le capacità intellettuali dell'uomo produca o meno un insieme coerente di teoremi. Se, tuttavia, per il primo teorema l'enunciato gödeliano è indecidibile e vero solo se il sistema di riferimento è consistente e se per il secondo teorema non è possibile dare una dimostrazione assoluta di coerenza di tale sistema, non resta che dare una dimostrazione relativa di coerenza, ovvero dimostrare che la macchina che ci rappresenta è consistente supposto che lo siamo noi. Lucas, consapevole del problema sollevato da Putnam, costruisce pertanto una serie di argomenti a favore della coerenza dell'intelligenza umana. Ma dare la suddetta dimostrazione

⁷⁸Ricordiamo che è possibile formulare l'Aritmetica Intuitiva (N) come un sistema formalizzato, che chiamiamo Aritmetica Formale (AF) e che nel modello standard di AF $\langle N, +, \times \rangle$ tutti i teoremi di AF sono veri. Torneremo in modo più dettagliato su questo punto nel terzo paragrafo.

⁷⁹Vedi per un'ampia rassegna Fano, Graziani (2013).

relativa di coerenza significherebbe che un uomo è in grado di fare ciò che una macchina o sistema formale non sa fare: ovvero asserire la propria coerenza in modo assoluto. Per tale ragione la costruzione da parte di Lucas di argomentazioni a favore della coerenza della mente umana non va oltre il sostenere un generico principio di indulgenza o carità alla Quine-Davidson nei confronti degli essere umani, ed in tal senso la sua argomentazione gödeliana preserva la debolezza evidenziata. Oltre al problema sollevato da Putnam, studiosi come Douglas Hofstadter (1979), Charles Chihara (1971) e Stewart Shapiro (1998) sottolinearono come anche il punto L4 dell'argomentazione di Lucas non fosse chiaro, poiché non si comprende che cosa voglia dire che l'uomo è in grado di vedere la 'verità' di G e costruirono delle obiezioni agli argomenti lucasiani in merito.⁸⁰

L'argomento di Lucas, nonostante non fosse una novità, produsse un ampio dibattito sulla questione riguardante la possibilità o meno di trovare una confutazione alla tesi meccanicista. Tale disputa coinvolse non solo filosofi, ma anche logici, matematici, informatici ecc. che si confrontarono con l'argomentazione lucasiana sottolineandone aspetti di forza e debolezze. Tra questi contributi alcuni si sono imposti per la loro profondità divenendo un punto di riferimento per la comprensione e lo sviluppo delle argomentazioni pro e contro il meccanicismo. Uno di questi è certamente l'articolo *God, the Devil, and Gödel* del filosofo della matematica Paul Benacerraf (1967) che, mosso dalla convinzione che l'argomento di Lucas non fosse in grado di dimostrare ciò che pretendeva (ovvero che il meccanicismo è una posizione insostenibile), presenta un argomento che, partendo dalle assunzioni che la l'intelligenza umana è al più una Macchina di Turing e che noi conosciamo questa Macchina di Turing, utilizzando i teoremi di Gödel, giunge ad una contraddizione. Benacerraf arriva ad una conclusione diversa da quella di Lucas, ma ugualmente molto interessante.

Benacerraf prova a risolvere i diversi problemi aperti nell'argomento di Lucas, che abbiamo evidenziato nel precedente paragrafo, costruendo una nuova argomentazione. In primo luogo egli osserva che è necessario restringere la nozione di 'capacità intellettuali dell'uomo' mediante l'introduzione di un insieme S definito come 'tutti gli enunciati che io posso derivare e so che sono veri'. In tal modo Benacerraf aggira il problema, posto anche da Chihara, dell'uso del concetto incerto e ambiguo di 'vedere la verità di G'. Mentre, infatti, l'argomento di Lucas si limitava a sostenere che la MT avrebbe dovuto rappresentare le capacità intellettuali umane lavorava in maniera puramente sintattica, per cui poteva produrre un insieme di enunciati incoerente, cioè che contiene una contraddizione. Per contro S è necessariamente coerente, perché i suoi enunciati sono non solo derivabili, ma anche veri. Come è noto, infatti, è possibile dimostrare che se gli enunciati di un qualche sistema formale sono veri in un certo modello, allora quel sistema formale è coerente. In tale modo Benacerraf supera anche l'obiezione di Putnam.

Benacerraf, inoltre, nota un'ulteriore debolezza dell'argomento di Lucas (che verrà ripresa indipendentemente in modo leggermente diverso, ma altrettanto ef-

⁸⁰Vedi Fano, Graziani (2013, pp. 51–53) per approfondimenti.

ficace, da Daniel Dennett): egli si chiede che cosa è che, secondo Lucas, MT non può fare e che invece io posso fare. Io posso trovare una prova semantica di G, cioè, uscendo dal sistema formale U posso individuare un modello di U che renda vera G. Ma siamo sicuri che MT non possa fare la stessa cosa? Come osserverà giustamente Dennett, nel suo lavoro del 1972 (Dennett, 1972), MT è un sistema fisico che quindi, per rappresentare una MT, deve essere adeguatamente interpretato. Ovvero occorre stabilire che cosa sia l'*input*, che cosa l'*output*, come vengano codificati i dati ecc. Per giunta il processo fisico che realizza MT è senz'altro molto complesso, per cui diversamente interpretato potrebbe dare origine a un'altra MT in grado di derivare G. Quindi il passo L3 dell'argomentazione di Lucas è problematico anche sotto questo punto di vista.

Infine, Benacerraf, diversamente da Lucas, distingue chiaramente la parte puramente matematica dei teoremi limitativi, che di per sé non ha alcun significato filosofico, dall'argomentazione filosofica vera e propria che ha bisogno, accanto al teorema di Gödel, anche di quella che egli chiama "una premessa filosofica". Sulla base dell'argomento lucasiano, Benacerraf costruisce un nuovo e più preciso argomento. Cerchiamo di chiarirne i punti in modo schematico.

BI. Si ricordi che S è l'insieme di 'tutti gli enunciati che io posso derivare e che so che sono veri'. S* è la chiusura logica di S nell'ambito di un sistema formale sufficientemente ampio e corretto, ovvero S* contiene tutti gli enunciati derivabili da S in un ragionevole sistema formale. Si noti oltre all'elemento epistemico 'io so' su cui torneremo più avanti, il carattere modale dell'espressione 'posso derivare'. È chiaro che se ci riferissimo solo agli enunciati che io derivo effettivamente, che per quanto sia un insieme grande è di cardinalità finita, non ci sono dubbi che esso sarebbe rappresentabile da una MT. È per questa ragione che occorre introdurre l'espressione 'posso', che tuttavia impedisce una caratterizzazione del tutto precisa dell'insieme S.

BII. Assumiamo che esista un insieme effettivamente enumerabile W_j tale che:

(a) l'enunciato che W_j contiene tutti i teoremi dell'aritmetica (AF) appartiene a S*. In simboli:

$$(28.) \quad 'AF \subseteq W_j' \in S^*$$

Si noti che W_j è effettivamente enumerabile se e solo se esiste una funzione f_j calcolabile e totale che associa a ogni numero naturale un elemento di W_j con eventuali ripetizioni. Inoltre f_j è calcolabile se e solo se esiste un algoritmo che dato in input un possibile argomento di f_j dia in output il rispettivo valore. Tuttavia, se accettiamo la Tesi di Church-Turing,⁸¹ allora l'insieme W_j è effettivamente enumerabile se e solo se esiste una MT di numero j che

⁸¹Ricordiamo qui la sua classica formulazione: la classe delle funzioni effettivamente calcolabili coincide con la classe delle funzioni ricorsive generali (o Turing calcolabili).

calcola la funzione f_j . Intuitivamente (a) afferma che io posso costruire una MT in grado di enumerare tutti i teoremi dell'aritmetica.

(b) L'enunciato secondo cui W_j è contenuto in S^* fa parte di S^* . In simboli:

$$(29.) \quad 'AB \subseteq W_j' \in S^*$$

Intuitivamente (b) afferma che io sono in grado di derivare che l'insieme W_j è un sottoinsieme di S^* .

(c) S^* è un sottoinsieme di W_j . In simboli:

$$(30.) \quad S^* \subseteq W_j$$

Intuitivamente (c) afferma che esiste una MT in grado di produrre tutti gli enunciati veri che io posso derivare. È chiaro che se io posso derivare che W_j è un sottoinsieme di S^* e so che tutti gli enunciati che appartengono a S^* sono veri e inoltre W_j è un sottoinsieme di S^* , allora W_j e S^* coincidono.

BIII. Da (a) a (c), mediante il procedimento di gödelizzazione è possibile derivare una contraddizione in S^* , che invece sappiamo essere consistente. Quindi almeno una delle ipotesi (a) (b) (c) deve essere falsa. Teniamo conto che le ipotesi (a) (b) e (c) sono di natura empirica-filosofica (perché, sebbene si riferiscano ai contenuti di S , cioè a ciò che posso dimostrare ed è vero, S tuttavia è definito in termini modali). È difficile sostenere che (a) sia falsa, cioè che io non sia in grado di mettere a punto un algoritmo che produce tutti i teoremi dell'aritmetica. Dunque i casi sono due: o le mie capacità deduttive non sono rappresentabili da una MT – cioè (c) non è vera – oppure io non sono in grado di derivare la MT che mi rappresenta – cioè (b) non è vera. In conclusione, o le mie capacità deduttive non sono rappresentabili da una macchina di Turing, oppure io non so quale sia questa macchina di Turing.

La dimostrazione di Benacerraf è ovviamente più articolata di quella qui tratteggiata,⁸² qui ci limitiamo a concludere l'esame di questa posizione attraverso una sua citazione:

They [Gödel's Theorems] imply that given any Turing Machine W_j , either I cannot prove that W_j is adequate for arithmetic, or if I am a subset of W_j , then I cannot prove that I can prove everything W_j can. It seems to be consistent with all this that I am indeed a Turing machine, but one with such a complex machine table (program) that I cannot ascertain what it is. In a relevant sense,

⁸²Per alcuni dettagli rimandiamo a Fano, Graziani (2013).

if I am a Turing machine, then perhaps I cannot ascertain which one. In the absence of such knowledge, I can cheerfully go around 'proving' my own consistency, but not in an arithmetic way – not using my own proof predicate. Ignorance in bliss. Of course, I might be an inconsistent Turing Machine. Lucas's protestations to the contrary are not very convincing. (Benacerraf, 1967, p. 29)

L'argomento di Benacerraf, dunque, risolve il problema della coerenza del sistema di enunciati prodotto da MT, ovvero di S^* , limitando il discorso a quegli enunciati che non sono solo derivabili, ma sono anche veri, cioè sono provabili da me in un senso assoluto. Tuttavia, lo stesso Benacerraf ha mostrato che questa nozione porta a una contraddizione senza bisogno di introdurre (a) (b) e (c). Questo punto, spesso lasciato in secondo piano, è invece molto importante per l'argomentazione che Benacerraf vuole costruire e a ragione è stato evidenziato da Charles Chihara (1971) nella sua ricostruzione dell'argomentazione di Benacerraf. Sebbene Chihara sostituisca all'insieme S di Benacerraf l'insieme S' dei numeri di Gödel degli enunciati di AF che io posso provare in senso assoluto, cioè che io posso provare e sono veri (limitando S dunque agli enunciati di AF) e dia all'argomentazione un andamento dimostrativo differente da quello costruito da Benacerraf, sostanzialmente propone un argomento molto simile nell'architettura generale a quello di Benacerraf. Lo stesso Chihara fa notare tale similitudine evidenziando anche il fatto che vi è una difficoltà nel riformulare tale argomento legata al fatto che sia S che S' coinvolgono le 'mie conoscenze', in particolare la nozione di 'prova assoluta' che non si riferisce solamente alla realizzazione di un algoritmo automatico, ma al fatto che 'io sappia' che gli assiomi di AF sono veri e che le sue regole di inferenza sono corrette, cioè mantengono la verità degli assiomi. La presenza di elementi intensionali rende più complicato formalizzare rigorosamente alcune premesse dell'argomentazione. Questo ultimo punto è molto importante per comprendere, come osserveremo, alcuni sviluppi della letteratura in merito, in particolare quelli legati alle ricerche di Reinhardt, di Carlson e di Alexander (cfr. *infra*, Sezione 3.).

Nonostante i molti errori e le oscurità, l'argomentazione di Lucas ebbe un profondo impatto sia sui sostenitori di posizioni meccaniciste che anti-meccaniciste. Un caso emblematico è la riflessione del matematico e fisico inglese Roger Penrose. Questi, come è noto, è convinto che la coscienza che caratterizza l'attività mentale umana non sia rappresentabile nei termini di una Macchina di Turing, anche se non è un'attività che trascenda la fisica, ma per spiegarla occorra una nuova fisica non computabile. Per provare questa tesi Penrose si avvale di argomenti alla Lucas, sebbene per taluni versi molto più raffinati e complessi. In modo particolare, come ormai è stato acquisito in letteratura, Penrose fornisce due argomenti: uno in *The emperor's new mind* del 1989 (Penrose, 1989) e nel secondo capitolo di *Shadows of the Mind*, ed uno nel terzo capitolo di *Shadows of the mind* del 1994 (Penrose, 1994). Sorprendentemente tali argomenti, pur tenendo presente la lezione di Lucas, non fanno riferimento ai sofisticati lavori di Benacerraf e Chihara. In modo particolare è interessante notare come, senza appunto conoscere le riflessioni di Chihara, Penrose costruisca la sua prima argomentazione supponendo che esista una MT capace di simulare tutte le procedure, chiamiamole A , seguite dalla comunità dei matematici per dimostrare teoremi e che (questo è simile a quanto assunto

da Benacerraf) A deve essere corretta, cioè deve produrre solo risultati veri. Dunque l'argomentazione di Penrose migliora in un qualche modo le precedenti in quanto nel definire A, che sostituisce l'insieme S di Benacerraf, non si riferisce a una sola persona, ma all'intera comunità dei matematici, aggirando appunto così alcune delle critiche sollevate da Chihara. Senza entrare in dettaglio,⁸³ Penrose giunge nella sua prima argomentazione a una biforcazione molto simile a quella di Benacerraf:

- Una Ipotesi scettica: i metodi matematici di prova non sono tutti racchiusi in un algoritmo;
- Una Ipotesi agnostica: i metodi matematici di prova sono tutti racchiusi in un algoritmo corretto ma che l'uomo non conoscerà mai con assoluta certezza.

In entrambe le alternative va sottolineato che il problema non è nell'hardware, ma nel software. In pratica, il bivio che le varianti precedenti aprono è equivalente alla premessa (b) del punto B2 di Benacerraf. La presenza di tale alternativa spiega le parti dell'opera di *Shadows of the mind* volte a sciogliere la biforcazione a favore dell'ipotesi scettica e pone le ragioni dello sviluppo di un secondo argomento penrosiano. Sebbene Penrose tornerà anche in altri contributi su tali argomenti (Penrose, 1996), le successive argomentazioni sono da un lato sviluppi e dall'altro raffinamenti di queste prime analisi. Rimandiamo per una interessante analisi dell'argomento di Penrose alle opere di Chalmers (1995) e di Lindström (2000).

Come è stato correttamente sottolineato da Shapiro (1998), un problema fondamentale del dibattito qui considerato è che non è ben chiaro quale debba essere l'esatto contenuto della tesi meccanicista. In realtà tutti gli autori sin qui analizzati si sono confrontati con il problema di definire questo contenuto confutandolo o avvalorandolo. Nonostante ciò, dalle diverse prospettive emerge con chiarezza che qualunque sia tale contenuto, sia il meccanicista che l'anti-meccanicista necessitano di porre alcune idealizzazioni senza le quali non sarebbe possibile costruire alcuna analisi e confronto in merito. Citando Shapiro:

the mechanist claims that there can be a machine whose outputs are the same as those of a human or a group of humans. What sort of machine? What outputs? What aspect of what human? [...] Things get interesting only when we idealize, but things also get murky. (Shapiro, 1998, p. 275)

Le stesse parole, *mutatis mutandis*, potrebbero essere dette per l'anti-meccanicista. Dunque è su queste idealizzazioni che la nostra attenzione va diretta partendo proprio dalle riflessioni dello stesso Gödel sulle implicazioni filosofiche dei suoi teoremi.

2. Il punto di vista di Gödel

Nel 1951 Gödel tenne una delle prestigiose *Gibbs Lectures* per l'American Mathematical Society, il titolo della sua relazione fu *Some basic theorems on the foundations*

⁸³Vedi Bruni (2004) e Fano, Graziani (2013).

of *mathematics and their implications* (Gödel, 1986-2003; Gödel, 1995). I teoremi in questione erano proprio quelli di incompletezza, e le implicazioni filosofiche riguardavano la natura della matematica e le capacità della mente umana. Questa fu una delle poche occasioni ufficiali in cui Gödel espone la sua opinione sulle implicazioni filosofiche dei suoi teoremi. Senza entrare nei dettagli dell'intero testo gödeliano, ciò che qui è per noi interessante è la sua prima parte che è dedicata alla derivazione della tesi della 'incompletabilità della matematica' a partire dai suoi famosi teoremi. Tale tesi era per Gödel in modo particolare sancita dal suo secondo teorema, infatti:

[I]t makes it impossible that someone should set up a certain well defined system of axioms and rules and consistently make the following assertion about it: All of these axioms and rules I perceive (with mathematical certitude) to be correct, and moreover I believe that they contain all of mathematics. If someone makes such a statement he contradicts himself. For if he perceives the axioms under consideration to be correct, he also perceives (with the same certainty) that they are consistent. Hence he has a mathematical insight not derivable from his axioms. (Gödel, 1995, p. 309)

L'idea di Gödel è, dunque, che se qualcuno percepisce con certezza assoluta che un dato sistema formale è corretto,⁸⁴ costui è anche a conoscenza della coerenza del sistema, ovvero conosce la verità dell'enunciato del sistema che stabilisce la coerenza del sistema stesso. Tuttavia per il secondo teorema di Gödel il sistema formale considerato non può dimostrare il proprio enunciato di coerenza, dunque il sistema non cattura tutte le verità aritmetiche e per tale ragione "if someone makes such a statement he contradicts himself". Ma cosa significa ciò? Significa forse che nessun sistema ben definito di assiomi corretti può contenere tutto ciò che è propriamente matematico? Gödel ritiene che tale domanda ha due possibili risposte:

It does, if by mathematics proper is understood the system of all true mathematical propositions; it does not, however if one understands by it the system of all demonstrable mathematical propositions. [...] Evidently no well-defined system of correct axioms can comprise all [of] objective mathematics, since the proposition which states the consistency of the system is true, but not demonstrable in the system. However, as to subjective mathematics it is not precluded that there should exist a finite rule producing all its evident axioms. However, if such a rule exists, we with our human understanding could certainly never know it to be such, that is, we could never know with mathematical certainty that all the propositions it produces are correct; or in other terms, we could perceive to be true only one proposition after the other, for any finite number of them. The assertion, however, that they are all true could at most be known with empirical certainty, on the basis of a sufficient number of instances or by other inductive inferences. If it were so, this would mean that the human mind (in the realm of pure mathematics) is equivalent to a finite machine that, however, is unable to understand

⁸⁴Ovvero 'semanticamente coerente', cioè il sistema dimostra solo cose vere. Si ricorda che la coerenza semantica implica la 'coerenza sintattica' o consistenza.

completely its own functioning. This inability [of man] to understand himself would then wrongly appear to him as its [(the mind's)] boundlessness or inexhaustibility. (Gödel, 1995, pp. 309-310)

La precedente domanda, dunque, pone non solo il problema della inesauribilità o incompletabilità della matematica intesa come totalità delle proposizioni matematiche vere, ma apre anche la questione se la matematica sia in linea di principio inesauribile per la mente umana, ovvero se le capacità dimostrative in matematica della mente umana siano estensionalmente equivalenti ad un certo sistema formale o alla MT a esso associabile (la MT che enumera l'insieme di teoremi del sistema formale corrispondente). La questione, dunque, richiede una riflessione attenta proprio sul rapporto tra quelle che Gödel chiama 'matematica oggettiva' e 'matematica soggettiva'.

Poniamo, innanzitutto, che T sia l'insieme delle verità matematiche esprimibili nell'aritmetica del primo ordine, chiamiamo questa 'l'aritmetica oggettiva', o come Gödel scrive 'matematica oggettiva' ovvero "the body of those mathematical propositions which hold in an absolute sense, without any further hypothesis". Per il Teorema di Tarski T non è definibile nel linguaggio dell'aritmetica, dunque T non è ricorsivamente enumerabile. Definiamo, inoltre, K come l'insieme degli enunciati aritmetici che l'uomo può conoscere o provare in modo assoluto e con certezza matematica, cioè può derivare e sa che sono veri, chiamiamo questa 'l'aritmetica soggettiva' o come Gödel scrive 'matematica soggettiva' che "consists of all those theorems whose truth is demonstrable in some well-defined system of axioms all of whose axioms are recognized to be objective truths and whose rules preserve objective truth".

Qual è, dunque, il rapporto tra K e T ? Citando Feferman (2006) potremmo sintetizzare la risposta gödeliana dicendo che se K fosse uguale a T "then demonstrations in subjective mathematics [were] not confined to any one system of axioms and rules, though each piece of mathematics is justified by some such system. If they do not, then there are objective truths that can never be humanly demonstrated, and those constitute absolutely unsolvable problems". Ovvero, se valesse l'uguaglianza $K=T$, la mente umana non sarebbe equivalente ad alcun sistema formale o MT ad esso associata. Poste, infatti, le caratteristiche di T , per ogni sistema formale vi sarebbe un enunciato dimostrabile dalla mente umana, ma non nel sistema formale. Dunque, il meccanicismo sarebbe senz'altro falso: la non enumerabilità ricorsiva di T implica, infatti, la non esistenza di alcun sistema deduttivo effettivo avente come teoremi tutte e sole le verità dell'aritmetica. Se, diversamente, K non coincidesse con T , e dunque la mente umana fosse equivalente ad un dato sistema formale o alla MT a esso associata, ne seguirebbe l'esistenza di enunciati aritmetici umanamente indecidibili in senso assoluto. Infatti, proprio il secondo Teorema di Incompletezza, come evidenziato da Gödel, permette questa conclusione: la proposizione che esprime la coerenza di K , diciamo Con_K , è vera ma non è dimostrabile all'interno del sistema formale stesso; la negazione di Con_K è falsa e non è provabile in K . Posta l'equivalenza tra mente umana e sistema formale, Con_K non è nemmeno provabile dalla mente umana. Infine, poiché Con_K può essere messa nella forma di un problema diofanteo essa è un problema

assolutamente indecidibile. Tale proposizione è, dunque, una verità inconoscibile. Tali questioni e argomentazioni conducono Gödel all'idea che dai risultati di incompletezza può essere derivata al massimo la seguente disgiunzione:

Either mathematics is incompletable in this sense, that its evident axioms can never be comprised in a finite rule, that is to say, the human mind (even within the realm of pure mathematics) infinitely surpasses the powers of any finite machine, or else there exist absolutely unsolvable diophantine problems of the type specified (where the case that both terms of the disjunction are true is not excluded, so that there are, strictly speaking, three alternatives). (Gödel, 1995, p. 310)

Dunque, seguendo Tieszen (2006) e considerando la traducibilità tra il concetto di sistema formale ben definito e quello di Macchina di Turing, possiamo dire che i teoremi gödeliani mostrano che potrebbe non essere vero che:

- (i) The human mind is a finite machine (a TM) and there are for it no absolutely undecidable Diophantine problems.

The incompleteness theorems show that if we think of the human mind as a TM then there is for each TM some 'absolutely' undecidable Diophantine problem. The denial of the conjunction (i) is, in so many words, Gödel's disjunction. In formulating the negation of (i) Gödel says that the human mind 'infinitely surpasses the powers of any finite machine'. One reason for using such language, I suppose, is that there are denumerably many different Turing machines and for each of them there is some absolutely diophantine problem of the type Gödel mentions. So Gödel's disjunction, understood in this manner, is presumably a mathematically established fact. It is not possible to reject both disjuncts. (Tieszen, 2006, pp. 230–231)

Così la disgiunzione lascia aperte le tre seguenti possibilità:

- (I) la mente umana sorpassa infinitamente i poteri di una macchina finita (MT) e non ci sono problemi diofantei assolutamente irrisolvibili.
- (II) la mente umana sorpassa infinitamente i poteri di una macchina finita (MT) e ci sono problemi diofantei assolutamente irrisolvibili.
- (III) la mente umana non sorpassa infinitamente i poteri di una macchina finita (MT) e ci sono problemi diofantei assolutamente irrisolvibili.

Gödel era convinto (distinguendo tra mente e cervello) che valesse (I), ma era altresì consapevole che i suoi Teoremi di Incompletezza non rendevano impossibile l'esistenza di una procedura meccanica equivalente alla mente umana. Gödel, tuttavia, come abbiamo esposto, riteneva che dai suoi teoremi scendesse che se una simile procedura fosse esistita noi "with our human understanding could certainly never know it to be such, that is, we could never know with mathematical certainty that all the propositions it produces are correct". Ma ciò, posta l'idea gödeliana che "the human mind, in demonstrating mathematical truths, only makes use of evidently true axioms and evidently truth preserving rules of inference

at each stage”, significa appunto che “the human mind (in the realm of pure mathematics) is equivalent to a finite machine that, however, is unable to understand completely its own functioning”.⁸⁵

Come si può notare Gödel non solo fu molto attento nell’affermare che noi non possiamo conoscere con certezza matematica che una certa MT rappresenta K, ma fu anche altrettanto raffinato nel cogliere che nonostante ciò noi potremmo utilizzare altri metodi, ad esempio di natura empirica, per perseguire tale conoscenza:

[W]e could perceive to be true only one proposition after the other, for any finite number of them. The assertion, however, that they are all true could at most be known with empirical certainty, on the basis of a sufficient number of instances or by other inductive inferences.

Nel 1972 Gödel si esprime ulteriormente sulla questione dicendo:⁸⁶

On the other hand, on the basis of what has been proved so far, it remains possible that there may exist (and even be empirically discoverable) a theorem-proving machine which in fact is equivalent to mathematical intuition, but cannot be *proved* to be so, nor even be proved to yield only *correct* theorems of finitary number theory.

Questa formulazione, come si può notare, è significativamente differente da quella del 1951: Gödel in questa nuova dichiarazione sembra ammettere che la mente, almeno nel suo fare matematica, potrebbe essere una macchina e noi potremmo non riconoscere questo fatto o non essere in grado di dimostrarlo.

3. Macchine che conoscono se stesse

La conoscenza del testo di Gödel del 1951 e delle successive riflessioni gödeliane sull’argomento (Wang, 1974) diedero, ovviamente, nuovo impulso al dibattito sul meccanicismo-antimeccanicismo.⁸⁷ Di particolare interesse a tal riguardo sono i lavori di William Reinhardt (Reinhardt, 1985b; Reinhardt, 1985a; Reinhardt, 1986), Stewart Shapiro (1985), Timothy Carlson (2000) e Samuel Alexander (2014) a cui dedicheremo questo paragrafo proprio perché hanno profondamente precisato gli elementi intensionali presenti negli argomenti gödeliani ai quali abbiamo fatto riferimento nelle analisi precedenti. Obiettivo principale di questi lavori è la formalizzazione dei concetti di conoscenza e dimostrabilità (che come abbiamo visto sono alla base delle argomentazioni discusse) in un sistema formale in cui verificare la consistenza di alcune tesi importanti all’interno degli argomenti gödeliani.

Reinhardt prende le mosse, infatti, dal riflettere su un’assiomatizzazione dell’aritmetica contenente una nozione epistemica rappresentata dall’operatore modale

⁸⁵Per un’analisi dell’argomento gödeliano vedi van Atten (2006) e Fano, Graziani (2013); quest’ultimo contiene anche una possibile ricostruzione formale dell’argomento.

⁸⁶Vedi Wang (1974).

⁸⁷Vedi Fano, Graziani (2013).

(informalmente interpretato con) '*it is provable that*'. Utilizzando questa nozione Reinhardt riuscì non solo a provare la consistenza di varianti della Tesi di Church, ma anche a discutere versioni dei Teoremi di Incompletezza e le loro conseguenze epistemologiche.

Per capire tali nuovi argomenti, iniziamo dunque proprio dalla Tesi di Church-Turing la quale, come abbiamo già detto, in una sua versione classica, afferma che:

La classe delle funzioni effettivamente calcolabili coincide con la classe delle funzioni ricorsive generali (o Turing calcolabili).

Le varianti di tale Tesi considerate da Reinhardt sono intese a studiare il legame che esiste tra ciascuna proprietà che si dimostra essere 'intuitivamente debolmente decidibile'⁸⁸ e la MT che ne realizza l'algoritmo di decisione. E le domande centrali sono: si può dimostrare che una tale MT esiste? E se esiste, può essere nota?

Per dare una risposta a tali domande, Reinhardt definisce, appunto, una teoria assiomatica che estende l'Aritmetica di Peano con una nozione epistemica di conoscenza rappresentata da un operatore modale. L'interpretazione dell'operatore è '*it is intuitively provable that*' o, in termini ancora più astratti, '*it can eventually come to be known*'.⁸⁹ Quindi, in un tale sistema, l'insieme di domande plausibili relative al dominio dell'aritmetica include anche questioni sulla conoscenza stessa che il sistema acquisisce. Piuttosto che spiegare precisamente il significato di questa nozione di '*knowability*', l'attenzione si pone sulle proprietà che questa nozione dovrebbe avere, che verranno poi formalizzate tramite assiomi dell'operatore modale. Reinhardt individua le seguenti condizioni di base irrinunciabili per esprimere correttamente un'idea di conoscenza:

- Principio della conseguenza logica: se conosco ϕ e se so che $\phi \rightarrow \psi$ allora ne segue che conosco anche ψ .
- Principio di infallibilità: ciò che è dimostrabile deve essere vero.
- Principio di introspezione: se conosco ϕ allora so di conoscerlo.

L'idea di rappresentare la conoscenza tramite un operatore modale, denominato B da Reinhardt (dal tedesco *Beweissbar*, 'dimostrabile'), oppure K da Carlson (dall'inglese *Knowledge*), deriva innanzi tutto dal fatto che un trattamento sintattico di B in forma di predicato porterebbe a contraddizioni. In particolare, da formule plausibili quali $p \rightarrow \neg B\neg Bp$ e $p \rightarrow \Diamond Bp$ ⁹⁰ si avrebbero comunque conseguenze indesiderate, come ad esempio la dimostrazione di $p \rightarrow Bp$, cioè ogni proprietà vera

⁸⁸Una proprietà P dei numeri naturali è 'intuitivamente debolmente decidibile' se per ogni x , $P(x)$ è dimostrabile quando è vera.

⁸⁹La prima idea di un operatore modale per esprimere il concetto di 'dimostrabilità' è dello stesso Gödel (1933). Vedi Artemov, Beklemishev (2005).

⁹⁰Il diamante $\Diamond$ rappresenta l'operatore modale di 'possibilità' (*possibly*). L'idea che giustifica queste due formule è che se p è valida, allora dovrebbe esistere la possibilità che p possa essere dimostrato, senza negare, ovviamente, che esistono formule valide che non possiamo dimostrare.

deve essere dimostrabile (si leggano in proposito le considerazioni di Reinhardt (1986, p. 433) in merito a Tharp (1973)).

Il sistema formale di Reinhardt, d'ora in avanti denominato Aritmetica Epistemica (AE), estende il linguaggio dell'Aritmetica di Peano $\mathcal{L}$ (i cui simboli non logici sono $0, S, +, \cdot$)⁹¹ con l'operatore modale B . Le formule ben formate del linguaggio ottenuto, chiamato $\mathcal{L}_B$, includono esattamente quelle di $\mathcal{L}$, estese con formule del tipo $B\phi$, laddove ϕ a sua volta è una formula ben formata.

Di seguito, ϕ, ψ denotano formule ben formate in $\mathcal{L}_B$; $x, y \in \mathcal{V}$ rappresentano variabili; una variabile x occorre libera in ϕ se non rientra nel campo d'azione di alcun quantificatore su x . Formule prive di variabili libere sono dette 'enunciati'. Termini e principio di sostitutività sono definiti nel modo classico. In particolare, se t è un termine, $x | t$ rappresenta la sostituzione di x con t , operazione che si estende banalmente anche a $\phi(x | t)$. Un 'assegnamento' è una funzione $s : \mathcal{V} \rightarrow \mathcal{U}$ di dominio l'insieme delle variabili $\mathcal{V}$ e di codominio l'universo di riferimento $\mathcal{U}$. Con $s(x | u)$ si intende la funzione che assegna u alla variabile x e $s(y)$ ad ogni variabile $y \neq x$.

Nella logica di base di AE una struttura $\mathcal{M}$ definita rispetto ad un universo $\mathcal{U}$ consiste di una classica struttura per la parte del Primo Ordine relativa a $\mathcal{L}$, unitamente ad una funzione di interpretazione per B . Questa è una funzione tale che, dato in ingresso un assegnamento s ed una formula modale $B\phi$, restituisce *True* (nel qual caso si scrive $\mathcal{M} \models B\phi[s]$) oppure *False* ($\mathcal{M} \not\models B\phi[s]$). Il soddisfacimento di $\mathcal{M} \models B\phi[s]$ può essere letto, usando le parole di Carlson (2000), nel seguente modo:

ϕ is known when the free variables of ϕ are interpreted according to s .

Ovviamente, l'interpretazione di $B\phi[s]$ non dipende da $s(x)$ se x non occorre libera in ϕ .

L'interpretazione di una formula ϕ secondo la struttura $\mathcal{M}$ segue quindi la classica definizione induttiva della Logica del Primo Ordine. Per brevità, $\mathcal{M} \models \phi$ se $\mathcal{M} \models \phi[s]$ per ogni possibile assegnamento. Diremo inoltre che una formula ϕ è valida se $\mathcal{M} \models \phi$ per ogni $\mathcal{M}$ e che ϕ è conseguenza logica di un insieme di enunciati Σ (denotato $\Sigma \models \phi$) se per ogni struttura $\mathcal{M}$ vale che $\forall \sigma \in \Sigma : \mathcal{M} \models \sigma$ implica $\mathcal{M} \models \phi$.

In AE valgono le condizioni di 'completezza' e 'compattezza'. In particolare, l'insieme di formule valide è ricorsivamente enumerabile e l'insieme $\{\phi : \Sigma \models \phi\}$ è ricorsivamente enumerabile se Σ è ricorsivamente enumerabile.

Veniamo ora alla assiomatizzazione di AE, la cui unica regola di inferenza è il *modus ponens*. In primo luogo, abbiamo gli assiomi di Peano:

1. $\forall x(S(x) \neq 0)$
2. $\forall x \forall y((S(x) = S(y)) \rightarrow (x = y))$

⁹¹Rispettivamente: una costante individuale; una costante funtoriale a un posto; due costanti funtoriali a due posti.

3. $\forall x(x + \mathbf{0} = x)$
4. $\forall x \forall y(x + \mathbf{S}(y) = \mathbf{S}(x + y))$
5. $\forall x(x \cdot \mathbf{0} = \mathbf{0})$
6. $\forall x \forall y(x \cdot \mathbf{S}(y) = x \cdot y + x)$
7. $\forall y_1 \dots \forall y_n((\phi(x \mid \mathbf{0}) \wedge \forall x(\phi \rightarrow \phi(x \mid \mathbf{S}(x)))) \rightarrow \forall x \phi)$

dove 1. e 2. esprimono che $\mathbf{0}$ non è nel codominio di $\mathbf{S}$ e che $\mathbf{S}$ è, come si usa dire, uno-a-uno, ovvero una funzione iniettiva, 3. e 4. (5. e 6.) rappresentano, una volta interpretati nel modello standard, gli schemi per l'addizione (moltiplicazione), mentre 7. è lo schema dell'induzione, tale che ϕ è una formula di $\mathcal{L}_B$ le cui variabili libere appartengono a $\{x, y_1, \dots, y_n\}$.

Passiamo quindi agli assiomi che formalizzano le condizioni imposte sull'operatore modale B presentate prima informalmente. Gli assiomi di base sono dati dalla chiusura universale dei seguenti schemi:

1. $B\forall x\phi \rightarrow \forall xB\phi$
2. $B(\phi \rightarrow \psi) \rightarrow B\phi \rightarrow B\psi$
3. $B\phi \rightarrow \phi$
4. $B\phi \rightarrow BB\phi$

Tali assiomi, che d'ora in avanti chiameremo B1-B4, rispettivamente, rappresentano i principi discussi all'inizio del presente paragrafo (si vedano B2-B4), mentre B1 merita un discorso a parte, legato all'interpretazione della premessa $B\forall x\phi$, la quale asserisce che la validità di ϕ è nota. Tuttavia quando questo è vero, in linea di principio, non significa che siano accessibili (e quindi noti) anche tutti gli oggetti assegnabili alle variabili libere di ϕ . Questo principio di accessibilità viene quindi garantito da B1, il quale asserisce appunto che $B\forall x\phi$ implica la conoscenza di ϕ per ogni elemento dell'universo di riferimento assegnabile a x in ϕ , ovvero $\forall xB\phi$. Ricollegandoci alla funzione semantica di interpretazione di B , si noti come, grazie al principio di accessibilità espresso da B1, valga quanto segue:

$$\mathcal{M} \models B\forall x\phi[s] \text{ implica } \mathcal{M} \models B\phi[s(x \mid u)] \text{ per ogni } u \in \mathcal{U}.$$

In seguito, la 'B-chiusura' di un insieme di enunciati Σ è data da $\Sigma \cup B\Sigma$, dove con la notazione $B\Sigma$ indichiamo l'insieme di formule $B\phi$, con $\phi \in \Sigma$. Quindi, i cosiddetti 'assiomi della conoscenza' sono dati dalla B-chiusura di B1-B4. La teoria assiomatizzata tramite gli assiomi della conoscenza prende il nome di 'Teoria della Conoscenza'. Unendo gli assiomi della conoscenza e la B-chiusura degli assiomi di Peano si ottiene l'assiomatizzazione standard della Teoria della Conoscenza per AE.

Per presentare la formalizzazione di alcune varianti della Tesi di Church-Turing, poniamo che W_e sia un insieme 'semi-decidibile', ovvero ricorsivamente

enumerabile, il cui numero gödeliano è e . A questo punto possiamo utilizzare l'operatore B , che rappresenta la nostra idea di 'intuitivamente dimostrabile', per ragionare sul concetto di 'decidibilità intuitivamente debole' accennata all'inizio (Reinhardt usa l'espressione '*weak B-decidability*' per descrivere la formalizzazione di tale concetto in termini di B). Nel seguito assumiamo che ϕ sia una formula ben formata in $\mathcal{L}_B$ la cui unica variabile libera è x . La prima variante forte della Tesi di Church-Turing in esame è come segue:

$$(31.) \quad \exists e B \forall x (B\phi \leftrightarrow x \in W_e)$$

la cui intuizione è che esiste una MT per la quale è dimostrabile che essa enumera tutti e soli gli assegnamenti di x per i quali è intuitivamente dimostrabile che ϕ è vera. Qui $B\phi$ rappresenta l'idea di *weak B-decidability* richiamata prima, mentre l'appartenenza di x (per ogni x che rende vera ϕ) ad un insieme ricorsivamente enumerabile decreta la calcolabilità, appunto tramite una opportuna MT, del problema della decisione di ϕ . Successivamente, Carlson identifica questo schema nell'intuizione '*I am a Turing machine and I know which one*'. Segue quindi una versione meno forte ottenuta anticipando l'operatore B più esterno a prefissare l'intera espressione:

$$(32.) \quad B \exists e \forall x (B\phi \leftrightarrow x \in W_e)$$

ovvero è dimostrabile che esiste una MT che enumera tutti e soli gli assegnamenti di x che rendono vera ϕ , ma non è dimostrabile quale essa sia. Secondo Carlson ciò equivale '*I know that the set of x for which I know $\phi(x)$ is recursively enumerable*'. Come scrive Reinhardt (1985a, p. 337):

Note that this is like Benacerraf's suggestion (1967) that "it seems to be consistent with all this that I am a Turing machine, but do not know which one."

Infine, rinunciando alla B -chiusura della formula, otteniamo:

$$(33.) \quad \exists e \forall x (B\phi \leftrightarrow x \in W_e)$$

ovvero la MT di cui sopra esiste, cioè ogni proprietà debolmente B -decidibile determina un insieme ricorsivamente enumerabile.

Se il modello di riferimento è quello standard dell'aritmetica, (31.) stabilisce che esiste una MT, e io so quale, della quale si dimostra che essa enumera gli enunciati intuitivamente dimostrabili dell'aritmetica; (32.) implica che è dimostrabile che esiste una tale MT, ma non è noto quale; (33.) stabilisce solo che tale MT esiste. Queste asserzioni e i relativi risultati che andremo ora ad illustrare rappresentano lo sforzo di caratterizzare le conseguenze epistemiche dei teoremi di Gödel secondo il concetto di dimostrabilità intuitiva.

Per presentare i risultati che seguono, assumeremo l'interpretazione standard del modello dell'aritmetica, per il quale B denoterà quindi la conoscenza dell'aritmetica. Formalmente, dato Σ un insieme di enunciati nel linguaggio di AE ,

definiamo la struttura $\mathcal{N}_\Sigma$ sull'universo dei naturali $\mathbb{N}$ tale che, data una formula ben formata ϕ la cui unica variabile libera è x , vale quanto segue:

$$\mathcal{N}_\Sigma \models B\phi[s] \text{ se e solo se } \Sigma \models \phi(x \mid \overline{s(x)})$$

dove $\bar{n}$ denota il 'numerales' del numero naturale n .⁹² Intuitivamente, la struttura $\mathcal{N}_\Sigma$ soddisfa $B\phi$ per l'assegnamento s se e solo se la formula ottenuta sostituendo x in ϕ con il numerale associato a $s(x)$ è conseguenza logica di Σ .

In primo luogo, Reinhardt dimostra che (33.) è consistente in *AE* (Reinhardt, 1985b), ma a questo stesso risultato arriveremo più avanti passando per la validazione di (32.). In seguito Reinhardt (1985a), dimostra anche che in qualunque teoria T in cui (33.) vale, si ha anche:

$$T \models \exists x(\phi \wedge \neg B\phi)$$

che altro non è se non una versione del primo teorema di Gödel nell'ambito della Teoria della Conoscenza. Il risultato più interessante rispetto alla nostra discussione è il seguente (Reinhardt, 1985a):

Teorema 1 (31.) è *inconsistente in AE*.

La dimostrazione si può vedere come un'altra versione del primo teorema di Gödel. Si inizia applicando B1 per ottenere:

$$\exists e \forall x B(B\phi \leftrightarrow x \in W_e)$$

Concentriamoci su $B(B\phi \leftrightarrow x \in W_e)$ e assumiamo $\phi(x) = \neg(x \in W_x)$ ed $x = e$ (quindi in seguito con ϕ intendiamo $\phi(x \mid e)$). Per B3 sappiamo che $B(B\phi \rightarrow \phi)$, mentre per lo schema appena assunto si ha anche $B(B\phi \leftrightarrow \neg\phi)$. Dalla congiunzione di queste due formule, applicando tautologie e distributività, segue $B(\phi \wedge \neg B\phi)$. Ma da questo segue $B\phi \wedge B\neg B\phi$ e applicando di nuovo B3 segue la contraddizione $B\phi \wedge \neg B\phi$.

Possiamo aggiungere che un tale risultato è coerente, ad esempio, con il Teorema di Rice (1953), secondo cui le proprietà non banali⁹³ di insiemi ricorsivamente enumerabili non sono decidibili. Parafrasando nel nostro contesto, non è possibile determinare le MT che enumerano le formule che soddisfano una proprietà non banale.

Il seguente teorema è invece dovuto a Carlson (2000):

Teorema 2 (32.) è *consistente in AE*.

La dimostrazione passa per la definizione di '*knowing machine*', un particolare tipo di macchina che risponde a quesiti su un certo dominio che include anche la conoscenza stessa della macchina. Dimostrare il teorema di cui sopra significa stabilire

⁹²Ad esempio, se $s(x) = 2$ allora $\overline{s(x)} = \mathbf{S}(\mathbf{S}(0))$. Per brevità, scriveremo $\phi^s = \phi(x \mid \overline{s(x)})$.

⁹³Una proprietà è banale se non discrimina gli elementi di un insieme, ovvero vale per tutti o per nessuno di questi.

che una tale macchina esiste, la quale, come vedremo, è definita esattamente da AE integrato con (32.).

In primo luogo, Carlson definisce 'entità' una qualunque teoria T in AE tale che $\mathcal{N}_T$ è un modello della Teoria della Conoscenza. Le entità hanno alcune proprietà interessanti, tra cui:

- Conoscenza: $\sigma \in T$ se e solo se $B\sigma \in T$ per ogni frase σ .
- Estensione: T estende la Teoria della Conoscenza.
- Disgiunzione: se $T \models B\sigma \vee B\psi$, dove σ e ψ sono frasi, allora $T \models \sigma$ oppure $T \models \psi$.
- Esistenza: se x è l'unica variabile libera in ϕ e $T \models \exists x B\phi$ allora esiste un assegnamento s tale che $T \models \phi^s$.

In particolare, la proprietà di conoscenza dipende da B3 e B4. Infatti, per B4 si ha l'implicazione $\mathcal{N}_T \models B\phi \rightarrow BB\phi[s]$ per ogni ϕ ed s ; inoltre, applicando B3 si ottiene la doppia implicazione $\mathcal{N}_T \models B\phi \leftrightarrow BB\phi[s]$ per ogni ϕ ed s , da cui segue $\phi^s \in T$ se e solo se $B\phi^s \in T$. Ma se ϕ è un enunciato, allora $\phi \in T$ se e solo se $B\phi \in T$. La proprietà di conoscenza è importante in quanto rappresenta una condizione necessaria e sufficiente affinché in T valga l'assiomatizzazione standard della Teoria della Conoscenza.

A questo punto, una entità T si definisce '*knowing machine*' se e solo se T è ricorsivamente enumerabile. L'intuizione è che la macchina, per essere tale, deve essere in grado di enumerare gli enunciati che stabilisce essere veri.

Ora, il punto focale della dimostrazione di Carlson è dato da una estensione della definizione di entità che prende in considerazione l'aspetto temporale. In particolare, essendo la nozione di conoscenza atemporale, si introduce il concetto di tempo rispetto all'operatore modale B sostituendolo con un insieme di operatori B_t . Il parametro t appartiene ad un dominio infinito Q di elementi linearmente ordinati e, intuitivamente, rappresenta l'istante di tempo in cui una determinata conoscenza viene acquisita. L'estensione temporale di una entità T prende il nome di 'versione stratificata' di T , caratterizzata dal fatto che ogni qual volta un operatore B_{t_1} si ritrova sintatticamente nel campo d'azione di un operatore B_{t_2} , allora $t_1 \triangleleft t_2$, dove $\triangleleft$ rappresenta la relazione d'ordine per Q .

Esiste quindi un metodo per far collassare la teoria stratificata al fine di recuperare la teoria T . Una serie di risultati mette in evidenza come l'assiomatizzazione per la teoria T dà origine a quella che viene chiamata 'assiomatizzazione stratificata', che poi è alla base della Teoria della Conoscenza Stratificata. A chiudere il cerchio, i risultati relativi al metodo di collassamento permettono di inferire proprietà di T a partire da proprietà della versione stratificata di T .

Poste queste premesse e utilizzando il principio di induzione e ricorsività delle relazioni d'ordine coinvolte, Carlson dimostra che l'assiomatizzazione standard per la versione stratificata di AE è ricorsivamente enumerabile, da cui segue che AE è una *knowing machine*. Un risultato analogo vale per AE integrato con (32.), da cui segue che tale sistema è una *knowing machine*.

Si noti come l'intero apparato si fondi sul fatto che l'entità T sia assiomaticizzabile nel modo standard, risultato indissolubilmente legato alla proprietà di conoscenza, e quindi a $B3$ e $B4$.

Come immediato corollario del teorema precedente, cui per altro è giunto indipendentemente lo stesso Reinhardt (1985b), deriva il seguente risultato.

Teorema 3 (33.) è consistente in AE .

Lo schema (32.) viene anche chiamato '*Strong Mechanistic Thesis*', in quanto rappresenta la più forte posizione meccanicista dimostrabilmente consistente. In tal senso, Reinhardt la ritiene più significativa di (33.) in quanto sebbene (33.) contenga essa stessa una tesi meccanicista, non è comunque prefissata da un operatore intensionale.

Uno studio recente condotto da Alexander (2014) fissa ulteriormente il rapporto tra (31.), (32.) e la assiomatizzazione standard della Teoria della Conoscenza. Come si è visto in precedenza, la formalizzazione di una *knowing machine* consapevole di essere descrivibile da una MT , senza per altro sapere quale, passa per la proprietà di conoscenza, $B3$ e $B4$. In particolare, la B -chiusura di $B3$, $B(B\phi \rightarrow \phi)$, esprime intuitivamente la condizione che la macchina è a conoscenza della propria fattività (qualunque conoscenza la macchina abbia, la macchina sa che tale conoscenza è vera). In generale, una qualunque teoria T si dice *factive* ('fattiva') se $T \vdash B(B\phi \rightarrow \phi)$. Va ricordato che la proprietà di fattività è anche usata da Reinhardt per dimostrare l'inconsistenza di (31.), quindi sembra essere un concetto chiave alla base dei risultati fin qui esposti, e Alexander in effetti lo prova. Offrendo un risultato ortogonale al teorema di Carlson, Alexander definisce una macchina *non factive* che conosce il proprio numero di Gödel, ovvero soddisfa (31.).

Tecnicamente, Alexander elimina dalla assiomatizzazione standard della Teoria della Conoscenza per AE la B -chiusura di $B3$, ovvero la conoscenza della propria infallibilità. Più precisamente, il sistema così ottenuto comprende $B1$ - $B4$ unitamente a $B\phi$ quando ϕ è una istanza di $B1$, $B2$, $B4$, o di un assioma della aritmetica di Peano. Di seguito indichiamo con AE_{mod_fact} la teoria assiomatizzata in tal modo.

Teorema 4 (31.) è consistente in AE_{mod_fact} .

Come nel caso di Carlson, la dimostrazione si costruisce a partire dalla struttura $\mathcal{N}_\Sigma$. In primo luogo, valgono le seguenti proprietà:

- Per ogni Σ , $\mathcal{N}_\Sigma$ soddisfa tutte le istanze di $B2$ e gli assiomi di Peano.
- Se Σ contiene ϕ^s quando ϕ è valida ed s è un qualunque assegnamento (Σ -validity), allora $\mathcal{N}_\Sigma$ soddisfa tutte le istanze di $B1$.
- Se per ogni $\phi \in \Sigma : B\phi \in \Sigma$ (B -closure) e, inoltre, Σ contiene tutte le istanze di $B1$ e $B2$, allora $\mathcal{N}_\Sigma$ soddisfa tutte le istanze di $B4$.

Quindi, per ogni $n \in \mathbb{N}$, si definisce la famiglia di assiomi $\Sigma(n)$ che soddisfa *B-closure* e Σ -*validity* e che contiene come schemi B1, B2, B4, gli assiomi di Peano, ed infine i seguenti, che rappresentano il *kernel* di (31.):

$$(34.) \quad \forall x (B\phi \leftrightarrow \langle x, g(\phi) \rangle \in W_{\bar{n}})$$

dove x è l'unica variabile libera della formula ben formata ϕ di $\mathcal{L}_B$, $g(\phi)$ esprime il numerale associato al numero di Gödel di ϕ , $\langle _, _ \rangle$ denota una biiezione canonica calcolabile da $\mathbb{N}^2$ in $\mathbb{N}$. Essendo *AE* compatta e completa, per la tesi di Church-Turing vale che per ogni $n \in \mathbb{N}$ esiste una funzione calcolabile totale f tale che $W_{f(n)}$ contiene tutti e soli i naturali $\langle m, g(\phi) \rangle$, dove ϕ è come sopra e $\Sigma(n) \models \phi(x \mid \bar{m})$. Inoltre, per il teorema di ricorsione di Kleene,⁹⁴ esiste un naturale n tale che $W_{f(n)} = W_n$. In altri termini, è nota la MT cui si fa riferimento in (34.).

Come conseguenza della costruzione di cui sopra, segue che $\mathcal{N}_{\Sigma(n)}$ soddisfa tutte le istanze delle B-chiusure di B1, B2, B4, degli assiomi di Peano, e delle istanze di (34.). Non resta che dimostrare due proprietà per completare il teorema. La prima è che $\mathcal{N}_{\Sigma(n)}$ soddisfa tutte le istanze di B3 (ma non, si noti, la sua B-chiusura). Per definizione, se $\mathcal{N}_{\Sigma(n)} \models B\phi[s]$ allora $\Sigma(n) \models \phi^s$, ma essendo $\mathcal{N}_{\Sigma(n)} \models \Sigma(n)$ segue anche $\mathcal{N}_{\Sigma(n)} \models \phi^s$. La seconda è che $\mathcal{N}_{\Sigma(n)}$ soddisfa tutte le istanze di (31.). Per costruzione, $\Sigma(n)$ contiene $B\forall x (B\phi \leftrightarrow \langle x, g(\phi) \rangle \in W_{\bar{n}})$ e, come visto precedentemente, un naturale n che la soddisfa esiste, da cui segue il risultato e, di conseguenza, il teorema.

Riassumendo, la condizione di *factivity*, ovvero la consapevolezza della correttezza della propria conoscenza, stabilisce il confine per la validità di due varianti della Tesi di Church. In un caso, abbiamo visto che stante la condizione di *factivity*, un sistema formale sa che esiste una MT che ne rappresenta la conoscenza, ma non è in grado di stabilire quale essa sia. Nell'altro caso, rinunciando alla condizione di *factivity*, si ha che un sistema formale è in grado di dimostrare quale MT ne rappresenta la conoscenza. Si noti che non conoscere la propria fattività non significa che il sistema formale non sia *factive*, ma che semplicemente non è in grado di dimostrarlo.

4. Conclusioni

Le precedenti analisi, modellando gli aspetti intensionali soggiacenti alle argomentazioni gödeliane, consentono dunque di affrontare con maggiore chiarezza la questione delle implicazioni filosofiche dei Teoremi di Gödel.

In primo luogo gli studi di Reinhardt, Carlson e Alexander consentono un chiarimento della tesi meccanicista (come auspicato da Shapiro (1998)): i risultati di consistenza esprimono, infatti, diverse accezioni dell'ipotesi '*mind is mechanical*'. In secondo luogo l'analisi del concetto di dimostrabilità intuitiva (Antonutti,

⁹⁴Per ogni funzione calcolabile totale $f: \mathbb{N} \rightarrow \mathbb{N}$ esiste un input n tale che le due funzioni calcolabili parziali indicizzate da n e $f(n)$, rispettivamente φ_n e $\varphi_{f(n)}$, si comportano nello stesso modo per tutti gli input. In questo senso, n si dice punto fisso di f .

2010); lo studio del rapporto tra computabilità effettiva umana e computabilità algoritmica (Antonutti, 2015; Antonutti, Horsten, in pubblicazione); l'indagine del nesso tra conoscenza/ignoranza da parte di una macchina della propria fattività e ignoranza/conoscenza da parte di una macchina del proprio codice, costituiscono senza dubbio altri importanti chiarimenti che consentono di dare rigore alle idealizzazioni in gioco in tali tipi di argomenti. Da questa prospettiva appare così ora possibile rileggere i vari argomenti gödeliani evitando che "many philosophers dismiss the whole Lucas-Penrose controversy, often by rolling their eyes" (Shapiro, 1998, p. 277) e precisare la conoscenza che gli uomini/macchine possono avere di se stessi.

Bibliografia

- Alexander S. (2014). "A machine that knows its own code". *Studia logica*, 102 (3), pp. 567–576.
- Antonutti M. (2010). "Informal Provability and Mathematical Rigour". *Studia Logica*, 96, pp. 261–272.
- Antonutti M. (2015). "Human effective computability". Manoscritto.
- Antonutti M., L. Horsten (in pubblicazione). "Epistemic Church's thesis and absolute undecidability", in L. Horsten, P. Welch (cur.), *The scope and limits of mathematical knowledge*. Oxford: Oxford University Press.
- Artemov S. N., L. D. Beklemishev (2005). "Provability logic", in D. Gabbay, F. Guenther (cur.), *Handbook of philosophical logic*. Dodrecht: Springer, pp. 189–360.
- Benacerraf P. (1967). "God, the devil and Gödel". *The monist*, 51, pp. 9–32.
- Bruni R. (2004). "Riflessioni sull'incompletezza. I teoremi di Gödel tra logica e filosofia". Tesi di dott. Università degli Studi di Firenze.
- Carlson T. (2000). "Knowledge, machines, and the consistency of Reinhardt's strong mechanistic thesis". *Annals of pure and applied logic*, 105, pp. 51–82.
- Chalmers D. (1995). "Minds, machines, and mathematics". *Psyche*, 2.1.
- Chihara C. (1971). "On alleged refutations of mechanism using Gödel's incompleteness results". *The journal of philosophy*, 69, pp. 507–526.
- Dennett C. (1972). "Review of the freedom of the will". *The journal of philosophy*, 69, pp. 527–531.
- Fano V., P. Graziani (2013). "Mechanical intelligence and Gödelian arguments", in E. Agazzi (cur.), *The legacy of A.M. Turing*. Milano: Franco Angeli, pp. 48–71.
- Feferman S. (2006). "Are there absolutely unsolvable problems? Gödel's dichotomy". *Philosophia mathematica*, 14.2, pp. 134–152.
- Gödel K. (1933). "Eine Interpretation des Intuitionistischen Aussagenkalküls". *Ergebnisse eines Mathematischen Kolloquiums*, 4, pp. 39–40.
- Gödel K. (1986–2003). *Collected works*. Vol. I–V. Oxford: Oxford University Press.
- Gödel K. (1995). "Some basic theorems on the foundations of mathematics and their implications", in K. Gödel (cur.), *Works*. Vol. III. Oxford: Oxford University Press, pp. 304–335.

- Hofstadter D. (1979). *Gödel, Escher, Bach: an eternal golden braid*. New York: Basic books.
- Lindström P. (2000). "Penrose's new argument". *Journal of philosophical logic*, 30, pp. 241–250.
- Lucas J. (1961). "Minds, machine and Gödel". *Philosophy*, 36, pp. 112–127.
- Marciszewski W., R. Murawski (1995). *Mechanization of reasoning in a historical perspective*. Poznań Studies in the Philosophy of the Sciences and the Humanities. Poznań: Rodopi.
- Newman J. R., E. Nagel (1958). *Gödel's proof*. New York: New York University Press.
- Penrose R. (1989). *The emperor's new mind*. Oxford: Oxford University Press.
- Penrose R. (1994). *Shadows of the mind*. Oxford: Oxford University Press.
- Penrose R. (1996). "Beyond the doubting shadow". *Psyche*, 2.23, pp. 89–129.
- Putnam H. (1961). "Minds and machines", in S. Hooks (cur.), *Dimensions of mind*. New York: Collier, pp. 148–179.
- Reinhardt W. (1985a). "Absolute versions of incompleteness theorems". *Noûs*, 19, pp. 317–346.
- Reinhardt W. (1985b). "The consistency of a variant of Church's thesis with an axiomatic theory of an epistemic notation". *Revista colombiana de matemáticas*, 19.1-2, pp. 177–200.
- Reinhardt W. (1986). "Epistemic theories and the interpretation of Gödel's incompleteness theorems". *Journal of philosophical logic*, 15, pp. 427–474.
- Rice H. G. (1953). "Classes of recursively enumerable sets and their decision problems". *Transactions of the American Mathematical Society*, 74.2, pp. 358–366.
- Shapiro S. (1985). *Intensional Mathematics*. Amsterdam: North-Holland.
- Shapiro S. (1998). "Incompleteness, mechanism, and optimism". *The bulletin of symbolic logic*, 4, pp. 273–302.
- Tharp L. (1973). "All truths are absolutely provable". Manoscritto.
- Tieszen R. (2006). "After Gödel: mechanism, reason, and realism in the philosophy of mathematics". *Philosophia mathematica*, 14.2, pp. 229–254.
- Turing A. (1936). "On computable numbers, with an application to the Entscheidungsproblem", *Proceedings of the London Mathematical Society*. Vol. 42, pp. 230–265.
- Turing A. (1950). "Computing machinery and intelligence". *Mind*, 59, pp. 433–460.
- van Atten M. (2006). "Two draft letters from Gödel on self-knowledge of reason". *Philosophia mathematica*, 14, pp. 255–261.
- Wang H. (1974). *From mathematics to philosophy*. New York: Humanities press.
- Webb J. (1980). *Mechanism, mentalism, and metamathematics: an essay on finitism*. Amsterdam: D. Reidel.

Abstracts in English

Claudia Casadio - Lambek Calculus, Linear Logic and Linguistics

This chapter takes linguistic communication into account as a dynamic process, according to non-commutative linear logic. Similar logical grammars for natural language are for example Montague's grammar or categorial grammars (which assign a pivotal role to the relationship between function and argument for classifying linguistic elements). The author aims at obtaining a group of well-built expressions (and clauses) as a result of appropriate demonstrations, by representing syntactic and semantic processes of a language in terms of logical inferences. In this way, the innovative role which can be played by the non-commutative linear logic for studying formal properties of natural languages is underlined. This way was paved by Lambek's formulation of a new kind of calculus for the logic analysis of natural language. Lambek himself defined this calculus 'pregroup calculus' or 'pregroup grammar'. The method of proof nets (introduced to geometrically represent the logical processes implied by a demonstration of linear logic) represents a crucial methodological choice for a proper formal analysis of natural languages.

Cristiano Chesi - Real-time Processing of Complex Sentences

In this work, the author aims at answering an apparently trivial question: why is it so difficult to understand certain sentences? In order to be able to answer, it is previously necessary to understand how to measure the complexity of a sentence objectively and how this complexity correlates with a difficulty due to the processing of these clauses. After a short presentation of the question according to the 'Minimalist Program' framework, the paper tries to analyze why certain sentences result harder to comprehend: there could be unknown words, infrequent words, unfamiliar words, or words or sentences ambiguous at a semantic or structural level. Complexity, however, is measured in this chapter with different and quantifiable structural factors, focusing on 'embedded' relative clauses and cleft sentences. The complexity of processing is measured in terms of a slowdown in reading times, evaluated through eye-tracking techniques. The author accounts for the results obtained by hypothesising an explicit grammatical theory which reflects all the single operations made to build the sentence. This grammatical theory overturns the derivation direction taken for granted in the minimalist approach (top-down rather than bottom-up), so that it is possible to foresee when a sentence will be hard to process and which elements will cause this difficulty.

Alessandro Lenci - Distributional Semantics. A Computational Model of Meaning

This chapter takes into account specific themes such as nature, aims and challenges of the so-called 'distributional semantics'. In distributional semantic, the

meaning of words is studied through a statistical analysis of the linguistic contexts (corpus driven) in which these words occur. Distributional semantics is based on the Distributional Hypothesis. Accordingly, the semantic similarity among two lexemes correlates with the similarity of distribution of these lexemes in linguistic contexts. This aspect allows postulating new definitions for concepts such as 'semantic similarity' of lexemes (that is, the similarity of contextual distributions of two or more lexemes) or such as 'meaning' (that is, a property which emerges from the use of a specific word in connection with the linguistic context). This paper describes how distributional semantics models are generally built and how their performance is generally evaluated. The Distributional Hypothesis offers a definition of semantic similarity in terms of spatial proximity. In this perspective, a critical point for future research is successfully determining which kind of semantic information may be 'captured' on the basis of contextual properties.

Felice Dell'Orletta and Giulia Venturi - ULISSE: a Domain Adaptation Strategy for Syntactic Parsing

This paper deals with Domain Adaptation for automatic syntactic annotation. Until the half of the 1980s, automatic linguistic annotation was based on algorithms built on groups of hand-written rules, defined a priori on the basis of the knowledge of the system to formalise. Subsequently, thanks to the progress of research in the field of Artificial Intelligence and to the development of linguistic resources, algorithms based on machine learning techniques began to be employed. The major difficulties of those algorithms were due to certain aspects of natural language such as ambiguities, diachronic evolutions, or language variations from the original domain of knowledge. More specifically, the issue of Domain Adaptation can be put in the following terms: "can an annotated corpus [which is representative of a specific linguistic variety] be used for the syntactic analysis of a second corpus [which is representative of a different linguistic variety]?" The author answer presenting an algorithm called ULISSE (Unsupervised Linguistically-driven Selection of dEpendency parses), which selects in an optima way the most representative sentences of a new target domain and feed them to the parser in addition to the original training set.

Elisabetta Jezek - Some Linguistic Principles for Creating Lexical Resources

This work analyses some semantic or (more generally) linguistic problems which can be possibly faced when building lexical resources for NLP. It focuses on a specific resource (T-PAS) which aims at creating an inventory of predicate-argument structures for Italian verbs. They are collected through a corpus-driven procedure and manual clustering of distributional information. These structures fall under the shape of verbal patterns, where it is explicit which semantic type is expected for each argument slot. Firstly, semantic types induced by the analysis of lexical sets in corpora must not be confused with ontological categories. A second major problem is that, in specific argument slots of the verb, a number of filler-words do not perfectly match with the expected semantic type for these slots. Moreover,

often members of a lexical set are problematic since they show a weak semantic cohesion. In this paper, the author endeavours to understand whether the semantic information highlighted by all the problems mentioned above must be inserted into T-PAS or not. Future research ought to propose experiments aiming at acquiring predicate-argument structures and lexical sets in a more automated way.

Alessandro Aldini, Vincenzo Fano and Pierluigi Graziani - Some Notes on Gödel's Incompleteness Theorems and the Knowledge of Machines

This work overviews some of the most important points concerning the thesis of the reducibility of human mind to a machine, through the well-known Theorems of Incompleteness of Gödel. The main focus of the authors' discussion is the understanding of which is the relationship between Gödel's Theorems and the knowledge that humans/machines can have of themselves. The famous Turing's 1950 article paved the way for a project the primary aim of which is the construction of machines able to reproduce or at least mechanically simulate the intelligent behaviour of human beings. This project (defined 'mechanistic') has always been criticised. Indeed, Gödel's Theorems have always been seen as an effective instrument to confute the mechanistic thesis both under the 'extensional' perspective (human mind procedures and results can be mechanistic) and under the 'intensional' perspective (the human mind is a peculiar kind of machine). This paper aims at specifying in which terms the knowledge that humans/machines can have of themselves has to be posed, and points out a few clarifications that can be extremely helpful in order to understand with rigour similar idealisations.

Compter parler soigner Between Linguistics and Artificial Intelligence

Proceedings of the Conference “Compter, parler, soigner. Fra linguistica e intelligenza artificiale”
(Pavia, December 15-17, 2014)

Edited by Edoardo Maria Ponti and Marco Budassi

Abstract in English

“Compter parler soigner” gathers the proceedings of a series of conference held at Collegio Ghislieri, in Pavia. It explores the knowledge lying between linguistics and artificial intelligence.

Its aims was finding a balance in representing theories and applications, the level of form as well as that of function, experiments and epistemological foundations. In this way, several disciplines including logic, linguistics, computer science, and neuroscience, were intertwined. This allowed to define a research object that would turn out to be incomplete in case it lacked one of those points of view.

The thread tying all the papers together crosses the development of combinatorial calculus formalisms, the cognitive impact of complex linguistic structures, the mapping of words into vectors in multi-dimensional spaces based on their textual distribution, the adaptation of parsing tools to new domains, the development of linguistic data resources, and finally a clarification of the debate about the question: can machines think?

All of these are attempts to foster a branch of knowledge that is crucial in this era dominated by communication and technology. In this time, a significant part of our lives is virtual, but nonetheless language has not ceased to express all the complexity and the diversity of the human experience by these means.

Edoardo Maria Ponti has a bachelor degree in modern literature at the University of Pavia, where he is currently enrolled to the second year of the master degree in theoretical and applied linguistics. He is fellow of Collegio Ghislieri and a student of the Institute of Advanced Study (IUSS). His interests rely mainly in network analysis, machine learning, and lexical representation.

E-mail: edoardomaria.ponti01@universitadipavia.it

Marco Budassi has a bachelor degree in modern literature at the University of Pavia, where he is currently enrolled to the second year of the master degree in theoretical and applied linguistics. He is fellow of Almo Collegio Borromeo and a student of the Institute of Advanced Study (IUSS). His interests rely mainly in historical linguistics and Celtic languages.

E-mail: marco.budassi01@universitadipavia.it

