



Osservazioni sul numero e sulla computabilità

Luca M Possati

► To cite this version:

| Luca M Possati. Osservazioni sul numero e sulla computabilità. 2017. halshs-01519411v2

HAL Id: halshs-01519411

<https://shs.hal.science/halshs-01519411v2>

Preprint submitted on 21 Jun 2017

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

Osservazioni sul numero e sulla computabilità

Questo saggio formula una semplice domanda: *che cos'è la computazione?* La nozione di algoritmo, macchina di Turing, funzione ricorsiva e Tesi di Church-Turing rappresentano i fondamenti dell'informatica teorica con applicazioni importanti nelle scienze cognitive, nelle grammatiche formali e in altri settori come il *DNA Computing* o il *Quantum Computing*. Lo scopo del saggio non è introdurre alla teoria della computabilità o descrivere le sue principali ramificazioni, bensì sviluppare una *fenomenologia della computabilità*, ossia un nuovo approccio alla computazione che non dipenda dalle nozioni di *effective procedure* o algoritmo. Il cuore della nostra argomentazione sarà una rielaborazione dell'esperimento mentale di Max Black sulle due sfere identiche ma discernibili. E come ogni fenomenologia, anche questa assumerà un andamento *archeologico*, proponendosi di rintracciare le radici non-teoriche e non-linguistiche del gesto logico e della regola come tale.

L'ipotesi proposta è duplice: *a)* l'atto del computare presuppone sempre un'interpretazione dell'identità – e ciò implica che l'identità sia qualcosa di *interpretabile*; *b)* tale interpretazione mette in discussione il principio della identità degli indiscernibili e dà luogo a una struttura logica di natura *paraconsistente* e *dialethetica*, cioè una struttura che ammette contraddizioni vere e che chiameremo *iterazione*. Computabile è anzitutto un tipo di oggetto *inesistente e tecnico*, costruito sull'iterazione. Tale oggetto è il numero. Dunque, tutti i numeri sono per principio computabili, ma non di tutti i numeri possiamo *mostrare* questa fondamentale computabilità. Il numero ci apparirà come un fossile, un insieme di stratificazioni logiche diverse e intrecciate. E in questo sta la sua incredibile plasticità.

Il saggio vuole tracciare dei confini e aprire prospettive di ricerca, lungi dall'aver mire di esaustività. Ci proponiamo di raggiungere due risultati. Il primo è la re-interpretazione della nozione di funzione ricorsiva e la distinzione di due sensi maggiori di computabilità. Il secondo è il chiarimento di uno dei concetti più usati ma meno compresi, l'iterazione. A guidarci è una duplice convinzione: da una parte, che nella computazione e nella ricorsività risieda la forma più elementare, embrionale (sebbene non l'unica) della razionalità (umana e non); dall'altra, che l'iterazione, procedura razionale sottovalutata, disprezzata o ignorata, sia un propulsore essenziale del nostro pensiero e vada considerata a tutti gli effetti come una struttura logica. Di qui una domanda ancor più radicale: può darsi un'iterazione perfetta, un perfetto *self-doubling*, un'identità moltiplicata?

§ 1. *Identità – lineamenti fondamentali*

L'identità è una relazione logica con proprietà specifiche. Tecnicamente, è un predicato a due posti. Il calcolo della deduzione naturale prevede due regole che ne governano l'uso: *l'eliminazione dell'identità* ($E =$) e *l'introduzione dell'identità* ($I =$). In generale, si dice che l'identità è una relazione riflessiva ($x = x$), simmetrica ($x = y \rightarrow y = x$) e transitiva ($x = y \wedge y = u \rightarrow x = u$). Il senso di questa struttura logica si fonda su due principi. Il primo è l'*indiscernibilità degli identici*:

$$\forall x \forall y [(x = y) \rightarrow \forall F (Fx \leftrightarrow Fy)]$$

Se t è identico a s , allora ogni proprietà di t è anche una proprietà di s . Il principio è un'affermazione sulle proprietà a partire dalla constatazione dell'identità pensata come fatto primitivo, elementare e in ultimo luogo inspiegabile¹.

Il secondo principio è l'inverso del primo: l'*identità degli indiscernibili*, e afferma l'identità a partire da una valutazione delle proprietà degli oggetti in questione. Nei termini di una logica di secondo ordine, l'identità degli indiscernibili può essere formulata così:

¹ Bisogna inoltre distinguere l'indiscernibilità degli identici dalla sua versione linguistica, detta sostitutività degli identici *salva veritate*, per cui se due termini denotano la stessa cosa sono sempre intercambiabili negli enunciati in cui compaiono lasciandone inalterato il valore di verità. Ci sono contesti in cui la sostitutività degli identici fallisce, mentre l'indiscernibilità degli identici resta valida. Cfr. R. Cartwright, "Identity and Substitutivity", in M. K. Munitz (dir.), *Identity and Individuation*, New York, New York University Press, 1971.

$$\forall x \forall y [\forall F (Fx \leftrightarrow Fy) \rightarrow (x = y)]$$

È Leibniz a introdurre la questione cruciale: dal fatto che un oggetto x possiede tutte le qualità di un oggetto y , e vice versa, segue necessariamente che x e y sono identici, coincidono, che sono lo stesso oggetto? La risposta di Leibniz è positiva: se due oggetti hanno in comune tutte le qualità e le relazioni, allora essi sono lo stesso oggetto². L'identità qualitativa è la condizione dell'identità numerica. Dio non ha alcuna ragione per creare due oggetti differenti *solo numero*. Ci deve sempre essere almeno una qualità o una relazione che singularizza quel determinato oggetto, anche soltanto la relazione più semplice: il fatto basilare per cui ciascun oggetto ha con se stesso una relazione che altri non possono avere, l'auto-identità.

Black fornisce un contro-esempio che mette fuori gioco l'identità degli indiscernibili³. Attraverso un esperimento mentale, egli mostra che possiamo pensare senza contraddizione la possibilità di un universo composto da due sfere di ferro perfettamente identiche, con le stesse qualità relazionali e non relazionali, *ma distinte*. Le sfere di Black sono identiche, ma due. Non vi è nulla di contraddittorio nella discernibilità degli identici. Il principio *is threatened if there is a situation in which objects of the relevant kind seem to be indiscernible in the relevant respect without being identical*⁴.

Un argomento molto simile a quello di Black è stato proposto da Strawson nel capitolo intitolato *Monads* del suo saggio *Individuals*⁵. Consideriamo una scacchiera. L'universo W è quello limitato dai margini della stessa. Prendiamo due caselle di una scacchiera simmetriche rispetto alla diagonale: entrambe “vedono” la scacchiera nello stesso modo. Le caselle sono al contempo un *luogo* dal quale il resto dell'universo W (nel caso in questione, la scacchiera) è visto ma anche l'oggetto visto. Dato un tale universo W non si possono dare delle descrizioni individuanti che

2 Cfr. G. W. Leibniz, *Monadologie*, § 9 (1714); *Nouveaux essais sur l'entendement humain*, livre II (1703); *Lettres à Clarke*, 2 juin 1716, mi-août 1716, in *Œuvres*, tome I, éd. par L. Prenant, Paris, Aubier-Montaigne, 1972, p. 422-435.

3 Cfr. M. Black, “The Identity of Indiscernibles”, *Mind*, vol. 61, n. 242, 1952, p. 153-164. Impossibile qui dare un quadro esaustivo del dibattito sul tema. Mi limito a rinviare a R. Adams, “Primitive Thisness and Primitive Identity”, *The Journal of Philosophy*, 76, 1979, p. 5-26; K. Hawley, “Identity and Indiscernibility”, *Mind*, 118, 2009, p. 101-109.

4 K. Hawley, “Identity and Indiscernibility”, cit., p. 101.

5 Cfr. P. F. Strawson, *Individuals*, London, Methuen, 1959, p. 117-134.

riescano a differenziare la casella nei termini della scacchiera. I due punti di vista sull'universo *W* sarebbero perciò indiscernibili pur essendo relativi a due individui diversi. È plausibile affermare l'esistenza di particolari distinti ma qualitativamente indiscernibili⁶.

Tra i numerosi critici di Black-Strawson, Della Rocca⁷ pone una questione fondamentale. L'esperimento mentale – dice Della Rocca – non dà alcuna spiegazione della non-identità delle sfere, che viene invece assunta come un fatto primitivo sulla base di una pretesa intuizione evidente. Negare il principio non basta; bisogna fornire anche una giustificazione della non-identità delle sfere in esame. La discernibilità è o non è una qualità? E se non è una qualità, che cos'è? In un universo perfettamente simmetrico non c'è un criterio di individuazione, dunque come facciamo a parlare di due e non di cento sfere? Perché siamo costretti a riconoscere che quelle due sfere sono distinte? Non potremmo dire che sono la stessa sfera in due punti diversi dello spazio? Perché dobbiamo moltiplicare gli oggetti senza una ragione precisa? Black non sa rispondere. Se ne conclude che *if one allows for primitive individuation or a brute fact of non-identity in Black's two-spheres case, then one has no good way to avoid other cases of primitive individuation that are intuitively unacceptable*⁸. Black ammette senza volerlo un principio di individuazione minimale, quello dato dalla locazione: le due sfere sono diverse in quanto occupano posizioni diverse, perciò non sono qualitativamente identiche. Non può mai darsi un caso di perfetta identità tra cose discernibili. Per questo Della Rocca recupera il principio fino a considerarlo una verità necessaria⁹.

Dal canto suo, Hacking ammette l'esistenza di descrizioni di oggetti indiscernibili e anche che vi siano possibili stati di cose in cui queste descrizioni sono soddisfatte da più di un oggetto. Tuttavia, egli afferma che le descrizioni di tali stati di cose (ad esempio le sfere di Black) possono essere re-interpretate come contenenti un solo oggetto in luogo di due indiscernibili. In altre parole,

6 Cfr. su questo punto, M. Carrara, *Impegno ontologico e criteri d'identità. Un'analisi*, Padova, Cluep, 2001.

7 Cfr. M. Della Rocca, "Two Spheres, Twenty Spheres, and the Identity of Indiscernibles", *Pacific Philosophical Quarterly*, 86, 2005, p. 480-492.

8 Ivi, p. 485.

9 Cfr. Z. Garrett, "An Explanation of Complete Colocation of Indiscernibles", *Res Cogitans*, 4, 2013, p. 18-26. L'autore risponde a Della Rocca offrendo una serie di casi in cui l'impossibilità della co-locazione di oggetti materiali è messa a dura prova se non apertamente smentita. Sempre sulla questione della co-locazione, cfr. R. Jeshion, "The Identity of Indiscernibles and the co-location problem", *Pacific Philosophical Quarterly*, 87 (2006), p. 163-176.

c'è differenza nelle descrizioni, ma non nella realtà identificata. Non c'è alcun fatto oggettivo che ci obbliga ad ammettere due entità e non una sola. La semplice distanza non prova nulla: in uno spazio curvo un individuo può essere spazialmente distante da se stesso¹⁰.

Il problema dell'indiscernibilità degli identici è ancor più complesso se consideriamo gli oggetti matematici. Lo strutturalismo *ante rem* di Shapiro¹¹ afferma che i numeri sono posizioni in una struttura, dei *relata* dotati esclusivamente di proprietà relazionali. La struttura, per Shapiro, non è composta da posizioni successivamente occupate da oggetti (strutturalismo *in re*). Al contrario, struttura e posizioni sono oggetti a pieno titoli, anche se oggetti *sui generis*: *mathematical objects are places in structures, and these structures exist independently of any non-mathematical systems that exemplify them*¹². I numeri si identificano soltanto in base alle loro relazioni con gli altri numeri. Il 2 non è 2 perché ha una qualche proprietà intrinseca, bensì per il fatto che viene dopo l'1 e prima del 3. Il numero è quindi un oggetto incompleto, che non ha condizioni di identità fisse, per questo la nostra conoscenza di esso è limitata. Di qui le domande: come identifica i suoi oggetti il matematico? Possiamo trarre conclusioni metafisiche sulla natura dei numeri dall'esame di tali procedure di identificazione (relazioni, funzioni, insiemi ecc.)? In effetti, le diverse strutture che il matematico usa (ad esempio: numeri naturali, numeri razionali, numeri reali, numeri immaginari, ecc.) non sono indipendenti, ma in relazione tra loro (ad esempio: il numero 2 è incastonato in queste diverse strutture, non scompare, e ogni volta viene “letto” da diverse prospettive)¹³. Che relazione c'è tra le strutture?

Burgess¹⁴ e Keränen¹⁵ hanno dimostrato che esistono alcuni oggetti matematici che, pur essendo distinti, godono delle stesse proprietà relazionali e sono indiscernibili. Shapiro si sbaglia: identifica oggetti che in realtà sono distinti. Keränen cita l'esempio dei numeri complessi: i , $-i$.

10 I. Hacking, “The Identity of Indiscernibles”, *Journal of Philosophy*, 72 (9), 1975, p. 249-256.

11 Cfr. S. Shapiro, *Philosophy of Mathematics: Structure and Ontology*, Oxford, Oxford University Press, 1997.

12 S. Shapiro, “Identity, Indiscernibility, and ante rem Structuralism: The Tale of i and $-i$ ”, *Philosophia Mathematica*, 16 (III), 2008, p. 286.

13 F. MacBride, “Structuralism Reconsidered”, in S. Shapiro (ed.), *The Oxford Handbook of Philosophy of Mathematics and Logic*, Oxford, Oxford University Press, 2005, p. 563-589.

14 J. Burgess, Book Review: Stewart Shapiro, *Philosophy of Mathematics: Structure and Ontology*, *Notre Dame Journal of Formal Logic*, 40 (2), 1999, p. 283-291.

15 J. Keränen. “The Identity Problem for Realist Structuralism”, *Philosophia Mathematica*, 9 (3), 2001, p. 308-330; Id., “The Identity Problem for Realist Structuralism II: A Reply to Shapiro”, in F. MacBride (ed.), *Identity and Modality*, Oxford, Oxford University Press, 2006, p. 146-163.

Questi numeri possiedono un automorfismo non triviale, sono strutture non rigide: i , $-i$ godono delle stesse proprietà relazionali, eppure sono numeri diversi. Lo strutturalista *ante rem* non riesce a spiegare perché $i = -i$ non implichi il fatto ch'essi siano lo stesso numero, che siano amalgamati in un solo oggetto. L'identità relativa su cui poggia lo strutturalismo *ante rem* entra in crisi. Shapiro stesso parla di *a metaphysical standoff*¹⁶.

I modi di uscirne mantenendo in piedi lo strutturalismo *ante rem* sono diversi e tutti suscettibili di critiche. In un saggio del 2005 Ladyman¹⁷ afferma che, quando consideriamo oggetti astratti come i numeri non possiamo pretendere di raggiungere una discernibilità assoluta come quella che abbiamo con gli oggetti fisici. Dobbiamo invece accontentarci di ottenere due tipi di discernibilità: *relative discernibility* e *weak discernibility*. Nel primo caso, due oggetti sono detti relativamente discernibili se e solo se esiste una formula con due variabili libere che si applica ad essi soltanto in una direzione, come ad esempio la relazione $>$ applicata ai numeri naturali: $5 > 4$, $3 > 2$, $56 > 42$, ecc. Nel secondo caso, due oggetti saranno debolmente discernibili se e solo se essi soddisfano una relazione non riflessiva a due posti. Ad esempio, i e $-i$ sono debolmente discernibili poiché ad essi è applicabile la relazione “è l'inverso di (e non è uguale a zero)”. I due oggetti possono entrare in una relazione molto semplice, il che implica la possibilità di porre una distanza minima. La stessa cosa avviene per due punti del piano euclideo: essi soddisfano la relazione “la distanza tra di essi è $d > 0$ ”. Tale visione si amplia e si rafforza in un saggio del 2012¹⁸, dove Ladyman, Linnebo e Pettigrew forniscono le definizioni formali di una serie di tipi di discernibilità, dandone un'analisi logica.

MacBride¹⁹ ha contestato questa soluzione: la relazione dell'inverso non basta a distinguere i e $-i$ poiché essi sono lo stesso oggetto ma istanziato due volte e quindi sono perfettamente intercambiabili anche nella relazione dell'inverso. La *weak discernibility* non regge, non può essere

16 S. Shapiro, “Identity, Indiscernibility, and ante rem Structuralism”, cit., p. 289.

17 Cfr. J. Ladyman, “Mathematical Structuralism and the Identity of Indiscernibles”, *Analysis*, 65 (3), 2005, p. 218–221.

18 J. Ladyman, Ø. Linnebo, R. Pettigrew, “Identity and Indiscernibility in Philosophy and Logic”, *The Review of Symbolic Logic*, 5, 2012, p. 162–186.

19 Cfr. F. MacBride, “What Constitutes the Numerical Diversity of Mathematical Objects”, *Analysis*, 66, 2006, p. 63–69.

un criterio di identità per i e $-i$. Come può *distinguere* una relazione tra cose *intercambiabili*?

Ma abbiamo davvero bisogno di uno specifico criterio di identità per distinguere i e $-i$? È la domanda di Shapiro che tuttavia, replicando a Keränen, egli non risolve ma dissolve. Esistono numeri che non conosciamo e non conosceremo mai, ma ai quali possiamo benissimo riferirci. Non potremo mai individuare tutti i punti dello spazio euclideo, eppure possiamo nominarli. Parlare di un tipo di oggetto, distinguendolo da un altro, non implica il fatto che identifichiamo quell'oggetto mediante un determinato criterio di identità, ad esempio postulando una *haecceitas* come fa Keränen. Shapiro rifiuta con forza l'identità degli indiscernibili, almeno per quanto riguarda gli oggetti astratti, e afferma che l'identità è un fatto primitivo, triviale, qualcosa che i matematici usano in continuazione e che perciò presuppongono, non debbono né possono spiegare. In matematica, *identity cannot be defined in full generality, in a non-circular manner*²⁰. Così possiamo dire $i = -i$ anche se non possiamo dare una spiegazione precisa del perché lo diciamo. Ciò non significa sottovalutare il ruolo dell'identificazione in matematica. La pratica matematica presuppone sempre l'identità; gli oggetti matematici sono sempre già identificati. *In presupposing the identity relation, we assume no more than is presupposed in ordinary mathematical practice*²¹.

Poniamoci, ora, due domande:

- 1) Shapiro dissolve o *nasconde* il problema?
- 2) Il rifiuto del principio dell'identità degli indiscernibili è giustificato? È possibile accettare il principio, in quanto “principio robusto e abbastanza generale” che “esprime una condizione necessaria di qualunque relazioni di identità”²², *pur limitandolo*? È possibile, in determinate situazioni, ammettere trasgressioni del principio senza per questo dover abbandonare *in toto* il principio stesso?

Nel prosieguo mi discosterò da un approccio eccessivamente strutturalistico e formalistico, assumendo quale punto di partenza un'ermeneutica del concetto di identità. Seguirò due strade: la prima (§ 2) riprende il contro-esempio di Black modificandolo con la considerazione di oggetti

20 S. Shapiro, “Identity, Indiscernibility, and ante rem Structuralism”, cit., p. 292.

21 Ivi, p. 294.

22 A. Varzi, *Parole, oggetti, eventi e altri argomenti di metafisica*, Roma, Carocci, 2001, p. 65.

inesistenti e tecnici, interamente progettati e creati dalla mente umana; la seconda (§ 3) si concentra invece sul concetto di identità e sul suo funzionamento proponendo un diverso approccio generale alla questione, in particolare al rapporto tra l'identità qualitativa e l'identità numerica. Entrambe le strade, sebbene autonome, s'influenzano. Insieme ci condurranno alla formulazione di un nuovo esperimento mentale sulla linea di Black (§ 4). Chiamerò “iterazione” una specifica trasgressione del principio dell'identità degli indiscernibili e mostrerò che è possibile darne una modellizzazione logica su basi paraconsistenti (§ 5, 6, 7). Da tali premesse svilupperò l'analisi fenomenologica della computazione (§ 8, 9, 10, 11).

Il saggio presenta una struttura ben definita. Le sezioni 2, 3 e 4 possono considerarsi tre parti di una sola introduzione che contiene i lineamenti essenziali di un'ermeneutica dell'identità. Su questa si fondono le altre due parti: la riconsiderazione logica dell'iterazione e l'approccio alla computabilità.

Avremo dunque una precisa scansione delle tesi:

- approccio ermeneutico alla questione dell'identità
- definizione dell'iterazione come trasgressione dell'identità degli indiscernibili
- individuazione, nell'iterazione, della radice logica del numero e della computazione.

§ 2. *Per una critica dell'identità degli indiscernibili*

Riprendiamo daccapo la questione dell'identità degli indiscernibili. Tutto dipende da che cosa intendiamo con “proprietà” o “qualità”. Se intendiamo con “proprietà” ogni possibile predicato attribuibile a x , quindi anche “essere identico a x ” o “essere diverso da y ”, il principio diventa un mero truismo ed è inattaccabile: ogni cosa si differenzia dal resto per il semplice fatto di essere identica a se stessa, ma questo non dice niente dell'oggetto. L'auto-identità non è una proprietà, tantomeno una relazione. Se invece intendiamo con “proprietà” solo le qualità *interessanti*, cioè non triviali, quelle che determinano un certo modo di essere di x e di y , allora la critica di Black riprende fiato.

E tuttavia, anche in questo caso si possono sollevare nuove obiezioni, persino più pesanti, come l'ipotesi delle “qualità nascoste”: potrebbero esistere delle qualità che ancora non conosciamo e che una sfera possiede ma l'altra no. Inoltre: chi stabilisce l'*interesse* di una qualità? A chi il diritto di decidere quali proprietà o relazioni sono ammissibili e quali no? Il termine “qualità”²³ sembra essere troppo vago. D'altronde, questo è l'argomento centrale delle critiche di Quine alla *second order logic*²⁴. L'esperimento di Black è destinato al fallimento.

Se vogliamo dare un colpo netto al principio, dobbiamo cambiare prospettiva, precisare i nostri termini e restringere il campo di azione. D'ora in avanti con “oggetto” o “cosa” intenderemo semplicemente: “portatore di proprietà”. Un “oggetto” *porta* certe proprietà, e in tal modo *soddisfa* certi predicati che designano tali proprietà. Ciò nonostante, anche la proprietà è un oggetto e porta certe altre proprietà. Il rosso ha la proprietà di essere un colore, il colore a sua volta ne ha altre, e così via. L'importante non è il regresso all'infinito in sé, bensì l'insuperabilità dello schema oggetto / proprietà, soggetto / predicato.

Consideriamo primitivo lo schema oggetto / proprietà: [O / P], per cui *se c'è un oggetto, c'è almeno una proprietà che ci parla di quell'oggetto, e se c'è una proprietà, c'è almeno un oggetto che porta quella proprietà, e di cui quella proprietà ci parla*. È il cosiddetto principio di comprensione per oggetti: per qualsiasi descrizione definita *a*, qualche oggetto soddisfa esattamente *a*²⁵. Possiamo parlare degli oggetti soltanto attraverso le loro proprietà.

Ora, fino a quando consideriamo *oggetti materiali*, come fanno Black e i suoi critici, la critica

23 Abbiamo certe proprietà in virtù di quel che siamo (proprietà intrinseche) e altre in virtù dei rapporti che abbiamo col mondo e con gli altri enti che lo popolano (proprietà estrinseche). La distinzione è però altamente problematica. Che cos'è veramente intrinseco? La forma di un oggetto è intrinseca? La costituzione quantica è intrinseca o estrinseca? Chi stabilisce il confine? Il dibattito è amplissimo (cfr. B. Weatherston, D. Marshall, “Intrinsic vs. Extrinsic Properties”, *The Stanford Encyclopedia of Philosophy* (Fall 2014 Edition), Edward N. Zalta (ed.), URL = <http://plato.stanford.edu/archives/fall2014/entries/intrinsic-extrinsic/>). La risposta alla domanda sulla “natura” delle proprietà distingue le varie forme di meinonghianismo contemporanei (cfr. F. Berto, *L'esistenza non è logica*, Roma-Bari, Laterza, 2010, capitoli 6, 7). Nel seguito non mi impegnerò in una discussione su questo aspetto. E questo per due motivi: a) mi limito a parlare di qualità *interessanti* in maniera molto generale, riferendomi a “qualità che ci dicono qualcosa dell'oggetto in questione”, il che mi basta per escludere non solo “essere identico a se stesso” o “diverso da altro” dal novero delle proprietà, ma anche le proprietà modali e temporali, nonché le verità *de dicto*; b) in ogni caso il principio di iterazione – che introdurrò nella sezione 6 – riguarda *tutte* le proprietà degli oggetti considerati, e questo elimina fin dal principio una discussione sulla “natura” delle proprietà.

24 Cfr. W. V. O. Quine, “Whitehead and the Rise of Modern Logic”, in P. A. Schilpp, *The Philosophy of Alfred North Whitehead*, New York, Tudor, 1941, p. 127-163.

25 Cfr. F. Berto, *L'esistenza non è logica*, cit., p. 102 (riprendo la formulazione, ma in una versione più semplificata).

del principio è impossibile, quantomeno si espone a obiezioni sostanziali²⁶. Se proviamo invece a considerare *oggetti immateriali* e molto *semplici*, ovvero oggetti astratti con un *range* di qualità interessanti ristretto e stabilito fin dall'inizio, allora la critica recupera terreno. In altre parole: possiamo sconfiggere il principio producendo appositi oggetti costruiti “in laboratorio”, oggetti *tecnici* e *inesistenti*, che controlliamo totalmente. Con la nostra mente possiamo costruire oggetti che non hanno alcuna collocazione spaziotemporale e ai quali attribuiamo un insieme chiuso di proprietà fisse – facendo così svanire la minaccia della “qualità nascoste”.

Il riferimento a Meinong e alla sua *Gegenstandstheorie*²⁷, nonché a una schiera di autori che ne hanno ripreso e ampliato le tesi in contesti filosofici diversi (Mally, Parsons, Routley, Priest, Zalta, ecc.)²⁸, ha qui un peso specifico essenziale. La tesi generale dei meinonghiani, che qui facciamo nostra, è al contempo semplice e disarmante: non tutto esiste, ci sono entità che non esistono, che vanno a comporre una sfera indifferente all'essere e al non essere. In altre parole, esistere è un predicato primitivo che alcuni oggetti possiedono, altri no. Di contro alla linea neoparmenidea di un Quine²⁹, i meinonghiani affermano che esistere non è una costante logica predicabile universalmente di tutto e non è definibile nella logica elementare standard. Oggettività (“x è un oggetto”), identità (“x è x”) ed esistenza (“x esiste”) non sono nozioni necessariamente connesse. Possiamo allora parlare legittimamente di oggetti non esistenti nel senso di astratti, non fisici, privi di riferimenti spazio-temporali. La matematica è un mondo di oggetti inesistenti che spetta al

26 Cfr. le obiezioni sollevate anche da G. Priest, *One. Being an Investigation into the Unity of Reality and its Parts, including the Singular Object which is Nothingness*, Oxford, Oxford University Press, 2014, p., cit., p. 22-24.

27 Cfr. A. Meinong, *Teoria dell'oggetto*, a cura di E. Coccia, Macerata, Quodlibet, 2003.

28 Cfr. G. Priest, *Towards Non-Being. The Logic and Metaphysics of Intentionality*, Oxford, Oxford University Press, 2005.

29 Cfr. W. V. O. Quine, “On What There Is”, in Id., *From a Logical Point of View*, Cambridge, Harvard University Press, 1961. Quine sostiene che l'affermazione “Non tutto esiste” è contraddittoria e palesemente falsa. Alla domanda “che cosa esiste” bisogna semplicemente rispondere: “tutto”. Per Quine essere è “il valore di una variabile quantificata” (“Designation and Existence”, *Journal of Philosophy*, 1939, p. 708). Quine presuppone che il “riferirsi a” (pensare, rappresentare, nominare, parlare di) implichi l'esistenza e che essere ed esistere coincidano: non è possibile riferirsi a qualcosa che non esiste. Com'è stato dimostrato, tuttavia, questa concezione si autodistrugge. Quine e seguaci (ma anche i precursori: Hume e Kant, Frege e Russell) presuppongono l'identificazione tra esistenza e quantificazione (la proprietà di avere istanze, di essere esemplificati), e quindi una pesante e indebita riduzione logica dell'esistenza. Per questi autori, l'esistenza non è una proprietà autonoma, semmai è una proprietà di secondo livello, o proprietà di proprietà: la proprietà di un concetto di essere istanziato. Cfr. F. Berto, *L'esistenza non è logica*, cit., p. 24-32. Berto mostra tutte le fallacie logiche cui va incontro questo approccio (p. 50-70) e propone una ridefinizione del problema: l'esistenza è un predicato che ha a che fare con l'avere dei poteri causali. Sulle possibili interpretazioni dei quantificatori al di là del loro uso quineano-esistenziale, cfr. M. Plebani, *Introduzione alla filosofia della matematica*, Roma, Carocci, 2011, p. 35-40.

matematico descrivere. E per trattare adeguatamente questi oggetti inesistenti una logica del primo ordine è insufficiente³⁰.

Non esaminiamo nel dettaglio questo vastissimo campo di ricerche. La cosa non è essenziale per il nostro percorso. Ci basti tenere aperta una possibilità: se spezziamo la connessione quineana tra oggetto, identità, esistenza e quantificazione (quindi apriamo anche a un diverso senso del quantificatore), possiamo parlare di oggetti inesistenti, non fisici, e delle loro proprietà. Non avremo bisogno di altro. Nel seguito faremo riferimento a oggetto inesistenti e *tecnic*i, completamente costruiti dalla mente umana. Ciò significa che questi oggetti sono dotati di un insieme di qualità finito, chiuso, che è del tutto sotto il nostro controllo. Gli “oggetti” del nostro esperimento mentale saranno *artefatti*, oggetti chiusi e controllati.

§ 3. Clusterizzazione dell'identità

L'identità è una questione metafisica cruciale³¹. Il nostro obiettivo, qui, non è quello di affrontare la questione dell'identità direttamente e in modo esaustivo, di stabilire se si tratti o meno di una relazione essenziale o contingente, a priori o a posteriori, o quale trattamento logico debba essa ricevere. Vogliamo proporre un *metodo* diverso per trattare la questione, partendo dalla semplice osservazione secondo cui l'identità è *plurivoca*, è un nome sotto il quale si raccolgono strategie concettuali differenti, con regole eterogenee e che possono anche entrare in contrasto tra

30 Cfr. S. Shapiro, *Foundations without Foundationalism. A Case for Second-Order Logic*, Oxford, Oxford University Press, 1991. Il programma fondazionale di Shapiro non esprime affatto un atteggiamento radicalmente anti-fondazionalista, bensì un fondazionalismo che si libera dal predominio della logica di primo ordine e dai suoi eccessi. È un fondazionalismo che parla di *fondazioni*, non di *fondazione*. Secondo Shapiro, *without question, the most widely acclaimed surviving fragment of the foundationalist programme is classical first-order logic*, logica che *is a central component of contemporary views on correct inference, as opposed to subjective certainty*. La connessione proposta porta poi direttamente al dubbio: *given the failure of the foundationalist programmes, one should not be overly confident that first-order logic is the only system worthy of our attention* (p. 35-37). Nel momento in cui dobbiamo descrivere la pratica matematica, la logica del primo ordine si rivela del tutto inadeguata: *first-order languages and semantics are inadequate models of mathematics* (p. 43). La descrizione di importanti aspetti della matematica richiede il passaggio a una logica di secondo ordine. La posizione di Shapiro è significativa poiché egli non attacca affatto i critici della logica di secondo ordine, Quine in primis, ma cerca invece di sfruttarne i suggerimenti, ampliandoli. Emerge così una prospettiva di relativismo logico molto aperta e dinamica, che supera l'*idolo* di una giustificazione ideale a priori, rigorosamente deduttiva – premessa indispensabile di una totale identificazione tra ragionamento e computazione. Criticare il fondazionalismo, in tutte le sue manifestazioni, significherà non solo smettere di porre un confine inamovibile tra logica e matematica, ma anche smentire l'equazione tra calcolo e ragionamento. Per Shapiro, in questo secondo caso, il problema centrale resta il rapporto tra linguaggio formale e linguaggio naturale.

31 Per un orientamento generale: H. Noonan and B. Curtis, “Identity”, *The Stanford Encyclopedia of Philosophy* (2014 Edition), Edward N. Zalta (ed.), URL = <http://plato.stanford.edu/archives/sum2014/entries/identity/>.

loro. Impieghiamo la parola “identico” in modi diversi. Con le etichette “identità qualitativa”, “identità numerica”, “coincidenza spazio-temporale”, “identità generica”, “identità personale”, ecc. ci rifacciamo a strategie concettuali differenti, il cui equilibrio è raggiunto a fatica. Le critiche di Kripke alla quadridimensionalità come criterio di identità³² e gli argomenti tratti dalla meccanica quantistica³³ – indipendentemente dalle interpretazioni che se ne possono dare – ci spingono in questa direzione.

Avanziamo una seconda osservazione: l'identità non si definisce mai da sola. L'identità non è un concetto bensì una *rete* concettuale, un *cluster* che raccoglie altri *cluster*, ciascuno dei quali si definisce non solo attraverso il riferimento ad altri concetti, ma anche attraverso il riferimento ad altri *cluster* identitari. Così l'identità strutturale si definisce mediante il concetto di isomorfismo o somiglianza, l'auto-identità mediante il concetto di riflessione o retro-versione, l'identità analogica tramite la proporzione, l'identità qualitativa tramite il concetto di qualità o proprietà, l'identità numerica tramite la coincidenza sul luogo e un sistema di coordinate fisso. Ogni *cluster* rappresenta una grammatica diversa, uno schematismo diverso, un *metodo* diverso di usare quella parola, quel concetto. Quando parliamo in generale di “identità di x” non facciamo altro che esprimere l'equilibrio raggiunto tra diversi usi di “identità” nel caso di x. L'identità è quindi un termine che non si rifà a una sola grammatica, bensì a una famiglia di grammatiche diverse. La verifica della comprensione di questo termine dipenderà non tanto dalla constatazione empirica, quanto dal confronto incrociato tra i *cluster*. Da questo confronto si definiranno poco alla volta, e daccapo in ogni situazione particolare, i confini dell'identità. Ovvio che poi ci saranno grammatiche più basilari, “rigide”, che sono quelle più ricorrenti e stabili, e altre invece più “fluide”, mobili.

Focalizziamo tre punti molto generali:

- Non esistono criteri d'identità fissi, intendendo con “criterio di identità” un enunciato che ci dice a quali condizioni si può affermare che x e y siano identici (per Quine, ad esempio,

32 Cfr. A. Borghini, C. Hughes, M. Santanbrogio, A. Varzi (a cura di), *Il genio compreso. La filosofia di Saul Kripke*, Roma, Carocci, 2010, p. 159-164.

33 Cfr. J. Ladyman, T. Bigaj, “The Principle of the Identity of Indiscernibles and Quantum Mechanics”, *Philosophy of Science*, 77 (January 2011), pp. 117-136.

l'occupare una stessa regione dello spazio-tempo, mentre per Davidson l'avere le stesse cause³⁴). Ogni *cluster* risponde a un criterio o criteri diversi.

- I *cluster* sono isolati, indipendenti tra loro, ma entrano in contatto, possono interagire e persino fondersi.
- Ogni enunciato identitario è sempre relativo alla relazione tra le varie forme di identità – per questo in alcuni casi uno stesso enunciato funziona, un altro no³⁵.

Consideriamo ora la relazione tra due *cluster*: l'identità qualitativa e l'identità numerica, che sono due interpretazioni di [O/P] inverse ma complementari.

Nell'identità qualitativa si presuppone un gruppo di qualità per poi, a partire da esse, identificare l'oggetto che le soddisfa. Si va dalle qualità all'oggetto. La prospettiva è fenomenologica: l'oggetto è un fascio di qualità senza alcun sostrato puro o collante ontologico sottostante³⁶. Possiamo darne due interpretazioni: monista o pluralista. Identifico Bruto, l'uomo che uccise Cesare, con un fascio di qualità diverso da quello del suo corpo, la massa di carne e ossa. La persona di Bruto ha una storia e delle caratteristiche personali diverse dal corpo di Bruto. Il monista parla di uno stesso oggetto (Bruto) descritto ora in un modo ora in un altro; si danno fasci diversi in base alle possibili diverse descrizioni, ma poi queste descrizioni si fanno corrispondere a un solo individuo. Il pluralista parlerà invece di due oggetti diversi perché a ogni fascio deve corrispondere un individuo. Lo schema di fondo non cambia. L'identità qualitativa parte dalla fissazione di un

34 Cfr. D. Davidson, "The Individuation of Events", in N. Rescher (dir.), *Essays in Honor of Carl G. Hempel*, Dordrecht, Reidel, 1969, p. 216-234.

35 I rappresentanti del relativismo – la tesi per cui l'identità è un concetto relativo – sono molti. Scontato il riferimento al relativismo sortale di Geach (*Reference and Generality*, Ithaca, Cornell University Press, 1962) o a quello concettuale di Putnam (*Reason, Truth and History*, Cambridge, Cambridge University Press, 1981). Non approfondisco queste posizioni. Mi limito a sottolineare che la concezione della relatività dei *cluster* qui proposta si rifà piuttosto al pluralismo radicale di Goodman espresso nel celebre *Ways of Worldmaking*, la cui tesi è che non c'è un mondo, ma tanti mondi attuali, nessuno dei quali omnicomprensivo. Applicando questo discorso all'identità, diremo che non c'è "l'identità", bensì diversi modi di identificare o costruire identità, modi che s'intrecciano e si stratificano. In tale stratificazione emergono strutture ricorrenti, spesso insuperabili, come la coppia soggetto / predicato, nome / aggettivo.

36 Ci si può richiamare qui alla teoria dei fasci, classicamente espressa da Berkeley, Hume e Russell, e rielaborata in modi diversi da H. Hochberg, "Things and Qualities", in W. Capitan, D. Merrill (dir.), *Metaphysics and Explanation*, Pittsburgh, University of Pittsburgh Press, 1964, p. 82-97; H. Castañeda, "Thinking and the Structure of the World", *Philosophia*, 1974, p. 4-40; J. O'Leary-Hawthorne, J. Cover, "A World of Universals", *Philosophical Studies*, 1998, p. 205-219. Non scendo nelle diverse obiezioni che si possono sollevare, in primis sulla questione della "compresenza" delle proprietà – e quindi sull'identità delle proprietà stesse – e sul cambiamento: cfr. A. Varzi, "La natura e l'identità degli oggetti materiali", A. Coliva (dir.), *Filosofia analitica. Temi e problemi*, Roma, Carocci, 2007, p. 8.

range di qualità per arrivare poi all'oggetto. Il punto è che le qualità sono sempre *comuni*, cioè possono appartenere a più oggetti, e perciò non hanno bisogno a loro volta di individuazione. Quando dico “Socrate è umano”, distingo un tipo (l'essere umano) e considero l'individuo che porta il nome “Socrate” come un'istanziatura di quel tipo, un'esemplificazione (più o meno soddisfacente) di quelle proprietà.

L'identità numerica risponde invece al criterio dell'unità o coincidenza. In questa grammatica, “essere identico” vuol dire coincidere con un *locus* (fisico o teorico) determinato da un *frame* di coordinate più o meno fisse. Quella posizione, descritta appunto dalle coordinate, la occupa soltanto quell'oggetto e non un altro, e ciò gli permette di coincidere con se stesso ed essere unico. Bruto è quell'individuo che coincide con una serie di coordinate spazio-temporali, e non può essere altro – è unico. Come sul piano cartesiano, possiamo identificare uno e un solo punto che corrisponde a quel dato. Identità numerica, coincidenza e *frame* di coordinate vanno di pari passo. Si osservi, anche qui, il movimento del *cluster*: un singolo (Bruto) è presupposto, già dato (lo nomino: uso un nome proprio per indicare quell'individuo), e viene poi singolarizzato nel *frame* definito dalle coordinate (la sua storia, il suo percorso spazio-temporale). L'indiscernibilità degli identici riflette esattamente questo approccio: si presuppone l'identità di due oggetti per poi affermare che possiedono tutte le stesse qualità, lo stesso *locus*.

Il *frame* di coordinate può essere inteso in modi diversi. Il teorico delle sostanze, ad esempio, parlerà di “universali”, “specie” o “sorte” quali criteri di identificazione. È il meccanismo dell'identità generica. Se diciamo: “Socrate è un uomo”, usiamo “Socrate” come nome-coordinata singolarizzante, mentre “uomo” è il genere, il *frame* in cui inserisco l'individuo in forza di una sua caratteristica o condizione di appartenenza a quell'insieme, l'insieme degli uomini. Il movimento della proposizione in esame procede secondo un progressivo rafforzamento dell'individuazione di un singolo (Socrate). In una visione più classica, il *locus* dell'oggetto è costituito dall'appartenenza generica e dalla costituzione materiale: la materia è la coordinata singolarizzante. Per Tommaso³⁷, fonte d'individuazione è infatti la *materia signata*, la materia singolare, ostensibile, formata – la

³⁷ Tommaso d'Aquino, *De ente et essentia* (1256), II, 1-4.

forma, o genere, è il limite del cambiamento – e percepibile. Locke³⁸ antepone invece la dipendenza sortale: ogni sostanza *di quel genere* ha un luogo e un tempo, un'esistenza, da cui sono escluse tutte le altre del suo stesso genere. Il quadridimensionalismo di Quine³⁹ non si sposta da questa linea.

§ 4. *Esperimento mentale: la discernibilità degli identici come contraddizione produttiva*

Come suggerisce Hawley, un buon contro-esempio del principio dell'identità degli indiscernibili deve rispettare due condizioni. Deve mostrare che a) un certo *qualitative arrangement* è possibile e b) due cose distinte sono responsabili di tale *arrangement*⁴⁰.

Ipotizziamo di considerare soltanto un oggetto *x* astratto, non fisico, senza alcuna locazione spaziotemporale. A *x* attribuiamo una descrizione minima, due sole proprietà *interessanti*, diciamo: *b* ed *e*. Si noti che non parliamo di *haecceitas* né di proprietà relative, o altro. Parliamo invece di “proprietà” in generale: abbiamo un oggetto, deve esserci una proprietà. E questo indipendentemente dal trattare *b* ed *e* in maniera classificatoria o tipologica. Considero [O/P] nella maniera più formale possibile.

Ipotizziamo di considerare un secondo oggetto, *y*, anch'esso astratto, inesistente, al quale attribuiamo lo stesso *range* di proprietà: *b* ed *e*. I nostri oggetti *x* e *y* sono interamente formali, vuoti: non sono enti fisici né geometrici. Sono schemi di oggetti e insieme compongono uno schema di universo possibile.

Chiediamoci che tipo di rapporto sussiste tra *x* e *y*. *De facto* una distanza c'è, non sono lo stesso oggetto: li chiamiamo con due nomi diversi. È un punto da non sottovalutare: i sostantivi ci consentono di distinguere gli oggetti, realizzano una sorta di discretezza pre-analitica del mondo. Dunque, una distanza c'è. Se ora, però, ci focalizziamo su questa distanza e la consideriamo come tale, ci accorgiamo ch'essa risulta problematica. I due oggetti *x* e *y* non sono oggetti fisici, quindi non hanno una locazione spaziotemporale; non c'è un *frame* di coordinate fisso.

Esaminiamo con attenzione il rapporto tra *x* e *y*. Perché diciamo che sono due oggetti e non

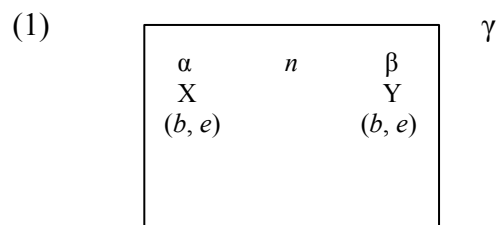
38 Cfr. J. Locke, *An Essay Concerning Human Understanding* (1694), libro II, cap. XXVII, 1-3.

39 Cfr. W. O. Quine, *Word and Object*, Boston, Mit Press, 1960.

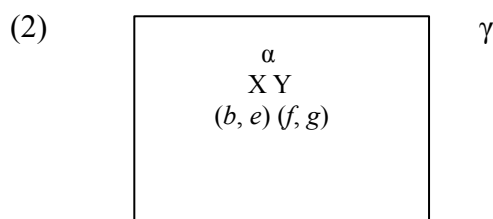
40 K. Hawley, “Identity and Indiscernibility”, cit., p. 102.

uno solo? Perché non diciamo, al contrario, che si tratta di due nomi per un solo oggetto? Che cosa ci impedisce di farlo? Proviamo a ragionare facendo riferimento alle grammatiche dell'identità descritte in precedenza. I due oggetti, x e y , sono due istanziazioni diverse di un solo tipo (b, e) , che è per definizione un *range* di proprietà chiuso, definito, determinato. Un medesimo tipo s'istanzia in due *loci* distinti. Perché introduciamo i *loci*? Perché parlare di posizioni? Abbiamo evocato, alla fine della precedente sezione, la complementarietà di identità qualitativa e identità numerica: non si può considerare l'una senza considerare l'altra, l'una è l'inverso dell'altra. Abbiamo schematizzato la situazione dicendo che i due *cluster* sono come gli estremi di una scala e che nell'uso del linguaggio ordinario tra di essi c'è *un certo* equilibrio. Se tutto questo è vero, allora al nostro esperimento mentale manca qualcosa. Dobbiamo introdurre un terzo oggetto: il *frame*, il quadro di riferimento in cui collocare e singolarizzare x e y . Se possiamo nominare i nostri oggetti, ciò vuol dire che abbiamo uno strumento con cui li distinguiamo. Chiamiamo γ il nostro *frame*. Si compone di tre elementi: due posizioni (α, β) e la distanza tra queste due posizioni (n). Non c'è un rigido sistema di coordinate. Abbiamo semplicemente due posizioni e la distanza tra esse.

Ora, ammesso per definizione che a) x e y condividono tutte le stesse proprietà e b) tali proprietà sono un *set* finito e chiuso, interamente controllato da noi, *allora* siamo costretti a riconoscere che uno stesso tipo s'istanzia in due posizioni diverse del *frame*. L'identità qualitativa resta, mentre quella numerica scompare. I nomi (coordinate singolarizzanti) indicano due posizioni distinte, ma di un tipo identico. I due *cluster* entrano in collisione.



Possiamo anche pensare la situazione inversa:



In questo secondo caso abbiamo una sola posizione nel frame (α), ma due distinti tipi che s'istanzano in questa posizione, tant'è che li chiamiamo con due nomi distinti (x, y). Il modello (2) descrive ad esempio il rapporto tra la statua e il marmo che la compone. Abbiamo due tipi diversi: la statua non è la materia di cui è composta, tanto che non solo risponde a un nome e a una descrizione diversa, ma può anche smettere di coincidere con il marmo. Eppure, la statua e la materia di cui è composta occupano la stessa posizione (almeno per un certo lasso di tempo). Anche in questo caso, i due registri identitari entrano in collisione.

Consideriamo la situazione (1). È un buon contro-esempio? La prima condizione di Hawley è soddisfatta dalla concepibilità stessa della situazione, mentre la seconda dalla possibilità di nominare gli oggetti di cui parliamo con simboli diversi pur considerandoli identici. L'obiettore risponderà ancora che i simboli usati (x, y) sono solo due nomi diversi per uno stesso oggetto. Ma non è affatto così, perché noi continuiamo a parlare di un universo con due oggetti e non uno. Se contiamo i nostri oggetti, la somma darà 2 non 1. Un minimo di identità numerica sussiste ed è il correlato della nostra semplice capacità di distinguere le cose, di riconoscere un confine, una negazione, tra un oggetto e un altro.

Immaginiamo di avere una retta. La dividiamo in due parti uguali e quindi formuliamo la proposizione: “Abbiamo due segmenti uguali”. Che cosa stiamo dicendo? Wittgenstein sostiene che quel che s'intende con tale proposizione è che esiste un sistema mediante il quale si può dimostrare che, dati certi assiomi, questa lunghezza = questa lunghezza. Ciò significa che dagli assiomi di Euclide segue una proposizione secondo la quale una certa cosa è identica a un'altra⁴¹. Ammettiamo

41 Cfr. L. Wittgenstein, *Lezioni sui fondamenti della matematica*, tr. it. di E. Picardi, Torino, Bollati Boringhieri, 1982,

che questo sia vero: gli assiomi di Euclide tracciano i confini di un gioco linguistico e stabiliscono condizioni grazie alle quali definiamo un *range* chiuso di proprietà fisse di rette e segmenti di rette. Posso allora spezzare una retta qualsiasi e considerare due segmenti identici. Li distinguo, ne riconosco i confini, ma hanno le stesse proprietà. Certo, quei segmenti hanno posizioni diverse sul piano, ma nulla impedisce di sovrapporli: a quel punto saranno identici, tranne per il fatto che sono due, non uno. Alla domanda “quanti sono i segmenti?” un osservatore che ha seguito tutto il processo risponderà: “due”. Uno che non ha seguito il processo risponderà: “uno”. Non c'è analogia, ma identità. Eppure esiste una discernibilità che è indipendente dall'identità.

Si danno *artefatti* – oggetti inesistenti, astratti e interamente costruiti dalla nostra mente – per i quali l'identità numerica e l'identità qualitativa non corrispondono. Si spezza l'equilibrio tra queste due forme di identità così com'è sancito dall'identità degli indiscernibili. La complementarità di identità qualitativa e identità numerica e la conseguente scalarità del loro rapporto ci permettono di evitare di dover spiegare daccapo perché possiamo riferirci a- e parlare di- x senza anche riferirci a- e parlare di- y – problema ben presente anche a Shapiro nel caso di $i = -i$.

Presentiamo un altro esempio di collisione tra identità numerica e identità qualitativa. Ipotizziamo – in una prospettiva di realismo modale⁴² – l'esistenza di una pluralità di mondi indiscernibili o quantomeno molto simili tra loro. Esistono in essi individui perfettamente identici, che condividono tutte le proprietà, ma sono distinti, cioè presenti in mondi diversi. Questi individui non sono distinti dalla collocazione spazio-temporale né possono coincidere o sovrapporsi. I mondi

p. 60 (*Wittgenstein's Lectures on the Foundations of Mathematics*, Ithaca, Cornell University Press, 1976).

42 Cfr. D. Lewis, *Counterfactuals*, Oxford, Blackwell, 1973. Quello delle modalità aletiche è uno dei terreni più fertili per la riflessione filosofica. La semantica a mondi possibili è una teoria del significato delle espressioni modali in un linguaggio formale basato appunto sul concetto di mondo possibile. Si valuta la verità o la falsità di ogni enunciato modale attraverso un esame dei mondi possibili – una situazione è possibile vuole dire che c'è almeno un mondo possibile in cui quella situazione è vera; se è necessaria, sarà vera in tutti i mondi possibili. Nello specifico, il realista modale sostiene che la possibilità non è affatto l'esatto contrario dell'attualità. Al contrario, assimila le due cose: possibilità e attualità sono sullo stesso piano, il possibile è solo uno scenario diverso dal nostro mondo. Esistono allora moltissimi mondi, tutti concreti quanto il nostro, composti di persone, cose o altri enti che esistono nello stesso senso in cui diciamo che le persone, le cose o altri enti esistono nel nostro. Con l'assimilazione di possibilità e realtà – se qualcosa è possibile, esiste –, quest'ultima si amplifica a dismisura: tutto è reale perché ci sono infinite situazioni concrete alternative alla nostra. Non ci sono mondi possibili e mondi reali, ma solo *mondi*. La possibilità è quel che esiste in un altro mondo, in un mondo diverso da quello in cui abitiamo. Il realismo modale differisce dall'ersatzismo, secondo cui i mondi possibili sono descrizioni alternative del nostro mondo, e dal combinatorialismo, secondo cui i mondi possibili sono ricombinazioni delle situazioni attuali. Per queste due posizioni, i mondi possibili non sono mondi concreti, esistenti, ma surrogati, derivati, versioni alternative del nostro mondo senza un'esistenza autonoma.

sono isolati e tra di essi sussiste un rapporto che non è spazio-temporale – possono infatti esserci mondi che non hanno spazio-tempo o che non hanno uno spazio-tempo come il nostro. Abbiamo una certa discretezza, ma al contempo identità qualitativa.

C'è una topologia neutrale – pre-logica, di ordine metaforico – che non intacca l'identità qualitativa dei nostri oggetti (x , y) pur separandoli. È allora possibile pensare l'iterazione perfetta, l'identità moltiplicata, e questo non significa mettere in discussione in maniera assoluta il principio dell'identità degli indiscernibili.

Consideriamo ora tre obiezioni.

1) Non sarebbe forse più semplice dire che x e y sono qualitativamente identici ma discernibili in virtù del fatto che hanno due distinti *sostrati*, nel senso che si tratta di due diverse entità portatrici di attributi identici? La nozione di “sostrato puro” è sospetta: se qualcosa come un “sostrato” dev'essere “puro”, scevro di attributi, per essere “portatore” di attributi, allora già in partenza non è poi così “puro” come dovrebbe essere. Questo per non parlare di altri problemi più specifici come la questione del cambiamento e l'identità attraverso il tempo. Attenendoci in partenza a oggetti inesistenti e tecnici, riconducibili a un range di proprietà chiuso e controllabile, abbiamo scelto di non entrare in queste discussioni.

2) La nostra è una soluzione *ad hoc*. Vero, l'esempio è costruito per attaccare il principio senza pietà. Ma c'è dell'altro: l'esempio, a mio giudizio, mostra una situazione logica che chiarifica alcuni fatti matematici basilari.

3) Ammesso che si diano situazioni teoriche in cui l'identità degli indiscernibili è abolita, in tali situazioni emerge una struttura completamente contraddittoria:

- la relazione tra x e y è contraddittoria: l'enunciato “ x è y ” risulta al contempo vero e falso, è *paradossale*. I due, x e y , sono allo stesso tempo identici e diversi. Possiamo nominarli, certo, distinguendoli. Ma in realtà sono identici. Posto che identità e differenza sono l'una la negazione dell'altra, dunque contraddittori, il rapporto tra x e y è contraddittorio. Si dirà: la contraddizione è apparente, dobbiamo distinguere i piani. Dal punto di vista dell'identità

qualitativa sono identici, mentre dal punto di vista dell'identità numerica sono diversi. In realtà, la contraddizione sarebbe apparente se i *cluster* fossero scissi, senza rapporto. Il nostro uso naturale di “stesso” o “identico”, “=”, risulta essere sempre unitario; usiamo “stesso” o “identico” sia per l'identità qualitativa che per l'identità numerica, e queste entrano in contraddizione: l'una è ammessa, l'altra no.

- Anche l'auto-identità di x e di y diventa contraddittoria: l'enunciato “ x è x ” risulta al contempo vero e falso, e lo stesso vale per y . Perciò x è se stesso, x , ma al contempo è diverso da sé perché distinto da y , posto che $x = y$.

L'obiettore ha ragione: x e y costituiscono un universo contraddittorio. E tuttavia, possiamo pensare x e y . La descrizione sopporta la contraddizione. La concepibilità non equivale alla possibilità: possiamo pensare l'impossibile. Ha senso pensare l'identità moltiplicata. Leggeremo allora con occhi diversi la proposizione 5.5302 del *Tractatus*, che dice: *Russell's definition of “=” is inadequate, because according to it we cannot say that two objects have all their properties in common. (Even if this proposition is never correct, it still has sense).*

Chiameremo *iterazione* la discrepanza tra le due grammatiche, identità qualitativa e identità numerica, che abbiamo descritto nel nostro esperimento mentale. Si tratta di una struttura logica precisa, che nega il principio di non contraddizione, ma non è triviale. Ha invece una natura paraconsistente e *dialethetica*.

§ 4. Paraconsistenza e dialetheismo

Da un punto di vista logico, un sistema formale S impiantato sul linguaggio L viene detto consistente o incontraddittorio se non consente mai di dimostrare e refutare allo stesso tempo una formula, ovvero di dimostrare sia una formula che la sua negazione: una contraddizione⁴³. Se invece

⁴³ Definire che cos'è una contraddizione è un problema enorme. La contraddizione ci pone di fronte a una situazione netta, dicotomica: un enunciato e la sua controparte negativa. In termini tecnici, la contraddizione è la congiunzione di due enunciati, uno la negazione dell'altro. Si può anche dire, in una formulazione non collettiva ma distributiva, che la contraddizione è una coppia di enunciati, di cui uno nega l'altro: si elimina il riferimento alla congiunzione. In una prospettiva semantica, la contraddizione sarà la congiunzione (o la coppia) di enunciati che non potranno essere né entrambi veri (subcontrarietà) né entrambi falsi (contrarietà). In ambito metafisico, la contraddizione sarà invece una situazione in cui un oggetto al contempo gode e non gode di una certa proprietà. Di qui l'ipotesi di mondi contraddittori. Il principio di non contraddizione nega la possibilità delle contraddizioni a livello sintattico,

si dà questo caso, allora S è detto inconsistente o contraddittorio. Ed è detto *triviale* se e solo se consente di dimostrare tutte le formule di L. Il *trivialist* crede che tutte le contraddizioni siano vere, e perciò che tutto sia vero. Nel caso in cui, poi, il sistema usato prevede la negazione, esso è detto *banale*, perché dimostra tutto e il contrario di tutto: $a \wedge \neg a \rightarrow b$. Il nesso tra inconsistenza e trivialità è spiegato dalla legge dello pseudo Scoto: *ex contradictione quodlibet*, dalla contraddizione segue qualsiasi cosa. Tecnicamente, nel calcolo della deduzione naturale la legge dello pseudo Scoto si ottiene mediante la regola della *eliminazione della negazione* ($E \neg$) con l'ulteriore passo della *introduzione del condizionale* ($I \rightarrow$). È la versione negativa del paradosso dell'implicazione materiale: il falso, l'assurdo implica qualsiasi cosa. Le logiche paraconsistenti approfondiscono e mettono alla prova questa convinzione.

La legge dello pseudo Scoto è detta “principio di esplosione” per rendere l'idea del potere distruttivo di una contraddizione all'interno di un sistema formale. Se un sistema formale ammette anche una sola contraddizione, le conseguenze sono disastrose e il sistema diventa deduttivamente inutile⁴⁴. Di questo è possibile dare una prova formale nel calcolo della deduzione naturale, come ha dimostrato Popper in un celebre articolo⁴⁵.

Una logica *paraconsistente* evita l'esplosione⁴⁶. Possiamo ammettere contraddizioni, senza per

semantico, ontologico e psicologico. “È impossibile essere e non essere a un tempo” (*Met.* 996b30). Cfr. J. Łukasiewicz, *Del principio di contraddizione in Aristotele*, Macerata, Quodlibet, 2003 (ed. or. 1910).

44 Cfr. F. Berto, *Teorie dell'assurdo. I rivali del principio di non-contraddizione*, Roma, Carocci, 2006, p. 99-100.

45 Cfr. K. Popper, *Conjectures and Refutations*, London, Routledge & Kegan Paul, 1969, p. 540.

46 La paraconsistenza non è un argomento nuovo in filosofia. Già la sillogistica di Aristotele era un perfetto esempio di logica paraconsistente. Anche gli stoici non sembravano riconoscere la necessità dell'esplosività della contraddizione. Tale necessità è invece diventata cruciale nella logica classica. La riscoperta della paraconsistenza è avvenuta nella seconda metà del novecento a partire dalla logica discussiva di Jaśkowski (approccio non aggiuntivo), la strategia a frammentazione di David Lewis, le tesi di Rescher e Brandon (*The Logic of Inconsistency*, Oxford, Blackwell, 1980), i lavori di Newton da Costa, le logiche adattive, le logiche rilevanti e l'approccio di Priest. Da un punto di vista semantico, nella maggior parte delle logiche paraconsistenti, la validità di un argomento è definita nei termini della verità-secondo-una certa interpretazione. Si adottano quindi modelli semantici nei quali è possibile dare un'interpretazione dei termini usati (costanti, variabili, connettivi, quantificatori) in base alla quale a e $\neg a$ possono essere entrambe vere. Una strada è usare la logica a tre valori: “vero”, “falso” e “vero e falso”, la via seguita da Priest. Secondo questa logica, è possibile che a e $\neg a$ siano entrambi veri, posto che a sia “vero e falso”. Il paradosso del mentitore è appunto una proposizione al contempo vera e falsa. Nel suo articolo del 1979 “Logic of Paradox”, *Journal of Philosophical Logic*, 8, 1979, Priest dà una giustificazione di questo punto di vista partendo da un'analisi del primo teorema di incompletezza di Gödel. Se consideriamo S un sistema formale nel quale formalizziamo tutte le nostre procedure dimostrative in matematica, in S l'enunciato di Gödel è indimostrabile. Ciò nonostante, attraverso un semplice ragionamento semantico basato sul T-schema possiamo provare che l'enunciato indimostrabile di Gödel è vero. Il paradosso si scioglie se formalizziamo anche questo semplice ragionamento semantico e lo facciamo rientrare in S. L'enunciato gödeliano diventa dimostrabile. Ma a questo punto, ecco il nocciolo dell'argomento di Priest, S è diventato paraconsistente. L'enunciato è allo stesso tempo vero e falso. Tale modo di procedere, tuttavia, non è immune da critiche, soprattutto dal punto di vista tecnico: cfr. C. S. Chiara, “Priest, the Liar, and Gödel”, *Journal of Philosophical Logic*, 13, 1984, p. 117-124. C'è però una ragione di fondo

questo rendere triviale tutto il nostro pensiero. Già Nietzsche lo riconosceva: l'illogicità è essenziale alla vita e da essa nascono molte cose buone, e solo “gli uomini troppo ingenui” possono credere che tutto vada ricondotto alla coerenza⁴⁷.

Questo tipo di logiche deve però difendersi da un'altra critica radicale, quella di Quine: *change of logic, change of subject*, nel senso che le logiche non-classiche sono sospettate di modificare non solo l'interpretazione, ma anche il significato stesso dei simboli logici. Il logico “deviante” non cambia nulla, o meglio si limita a cambiare argomento. Ad esempio, il simbolo della negazione per il logico tradizionale significa una cosa, mentre per il logico “deviante” un'altra – cambia il referente. Difficile replicare alla critica: siamo al cospetto di uno scontro tra intuizioni e tra vocabolari logici, sul quale appare difficile poter dare soluzioni definitive. Indipendentemente da ciò, le logiche paraconsistenti si identificano più che con un certo uso dei simboli logici, con una filosofia di fondo: esistono motivazioni importanti che ci spingono ad accettare contraddizioni nei nostri sistemi formali, e quindi dobbiamo cambiare qualcosa rispetto alla logica classica.

Quali sono i requisiti metodologici che devono essere soddisfatti affinché una logica possa ammettere contraddizioni e farlo in modo proficuo? In una logica paraconsistente non sono validi il principio di non contraddizione e la legge dello pseudo Scoto. Ci sono contraddizioni, ma non sono esplosive. L'ammissione di contraddizioni non deve trivializzare il sistema in base a una qualche versione della legge dello pseudo Scoto, il che equivale a dire che dobbiamo poterci servire delle

che si può avanzare contro il logico paraconsistente: perché dobbiamo accettare le contraddizioni? Dire che per il fatto che ragioniamo così, che questa è la realtà delle cose, appare una posizione piuttosto conservatrice e forse irrazionale: se non possiamo risolvere il paradosso, lo dobbiamo accettare. Priest risponde a questo genere di obiezioni in “What Is so Bad about Contradictions?”, *Journal of Philosophy*, 8, 1998, p. 410-426, il cui nocciolo è la tesi: se l'obiettivo di qualsiasi processo cognitivo è la verità, non è detto che la contraddizione debba sempre essere un ostacolo a essa, per cui coerenza e razionalità non sempre coincidono.

47 Cfr. F. Nietzsche, *Umano, troppo umano I*, a cura di G. Colli, M. Montinari, Milano, Adelphi, 1965, p. 38. Si noti che anche in Melandri c'è un'intuizione chiarissima di questa possibilità logica tant'è che potremmo considerare la sua logica analogica come una logica della paraconsistenza. Cfr. *La linea e il circolo*, cit., p. 231-232. Il tutto rientra in una precisa interpretazione dei rapporti tra logica e linguaggio che culmina nella valutazione della logica classica come “il risultato dell'assolutizzazione di un equilibrio contingente” (ivi, p. 628), quello tra sintassi linguistica e calcolo. Da sottolineare il fatto che la teoria linguistica melandriana dipende fortemente da quella di Snell e dalla distinzione tra le tre forme originarie: sostantivale (informativa ed estensionale), aggettivale (espressiva e intensionale) e verbale (pragmatica, dinamica, dialettica), che si ritrovano in ogni possibile proposizione. Cfr. B. Snell, *Der Aufbau der Sprache*, Hamburg, 1952. Melandri lega le tre forme di Snell a tre categorie di tropi: la sineddoche, la metonimia e la cataresi, che sono le “metafore fondamentali” distinte da Foucault in *Les mots et les choses* (Paris, Gallimard, 1966). E questo si connette ad altri due punti centrali di *La linea e il circolo*: a) la critica del logicismo inteso in un senso riduzionistico; b) l'idea per la quale la matematica è una razionalità autonoma, indipendente dalla logica, quantomeno dalla logica dell'identità elementare.

contraddizioni per formulare inferenze valide. Ma questo che cosa significa a livello formale? La maggior parte dei logici paraconsistenti cerca di disinnescare lo pseudo Scoto rinunciando a una o più regole di derivazione nel proprio apparato formale. È noto che, nel calcolo della deduzione naturale, la legge dello pseudo Scoto è dimostrata facendo ricorso a quattro regole: “Eliminazione della congiunzione”, “Introduzione della disgiunzione”, “Sillogismo disgiuntivo” e “Introduzione del condizionale”. A giocare un ruolo cruciale nella dimostrazione è la terza, detta anche *modus tollendo ponens* ($\alpha \vee \beta, \neg \alpha / \beta$), il cui comportamento diventa problematico quando si affrontano situazioni inconsistenti. Si è infatti constatato che questa regola, pur valida in sistemi e teorie consistenti, non è in grado di preservare adeguatamente la verità in situazioni contraddittorie. Se supponiamo che α e $\neg \alpha$ sono entrambe vere, allora $\alpha \vee \beta$ sarà comunque vera, anche se β è falsa, in base alle tavole di verità. Di qui, applicando il *modus tollendo ponens*, si dà la possibilità per cui da premesse entrambe vere, $\alpha \vee \beta$, e $\neg \alpha$, possa derivare una conclusione falsa, β . E questo va contro il criterio fondamentale della correttezza logica: non può mai darsi il caso che le premesse siano tutte vere e la conclusione falsa. Anche il logico paraconsistente *deve* rispettare questo criterio.

Il crollo del sillogismo disgiuntivo ha altre gravi conseguenze. Data la definizione classica del condizionale materiale, il *modus tollendo ponens* risulta logicamente equivalente a un'altra regola, ancor più basilare, il *modus ponendo ponens* ($\alpha \rightarrow \beta, \alpha / \beta$; regola di separazione o eliminazione del condizionale). In effetti il condizionale materiale, come connettivo verofunzionale, si definisce nei termini della disgiunzione tra la negazione dell'antecedente e il conseguente ($\neg \alpha \vee \beta$); di qui, l'equivalenza è evidente. Il logico paraconsistente deve rinunciare al *modus tollendo ponens* per disinnescare lo pseudo Scoto, ma non può rinunciare al *modus ponendo ponens*, che esprime una caratteristica inferenziale essenziale del condizionale stesso – un connettivo non conforme a questa regola non è affatto un condizionale. Il logico “deviante” deve allora sviluppare una semantica del condizionale autonoma dalla negazione e dalla disgiunzione in modo tale da evitare l'equivalenza con il *modus tollendo ponens*, che comporterebbe una ricaduta nel principio dello pseudo Scoto, e salvare l'irrinunciabile regola del *modus ponens*⁴⁸. È quel che in letteratura viene chiamato la

48 Cfr. F. Berto e L. Bottai, *Che cos'è una contraddizione*, Roma, Carocci, 2015, p. 54-55; F. Berto, *Teorie dell'assurdo*,

“condizione del *modus ponens*”.

Il logico paraconsistente si muove in bilico, riorientando ogni volta la propria interpretazione e il proprio uso dei connettivi e delle regole inferenziali. L'unico criterio di massima è quello di allontanarsi il meno possibile dalla logica classica. Una logica paraconsistente, essendo comunque un sotto-insieme della logica classica, deve cercare di salvare quel che nella logica standard funziona usandolo al meglio. È questa la “condizione di danno minimo” – anche se tale prescrizione resta piuttosto ambigua⁴⁹ e l'obiezione di Quine è sempre in agguato. Il logico paraconsistente non butta tutta la logica classica, non può farlo. Deve cercare di modificarne attentamente le strutture in modo da *a)* rendere innocue le contraddizioni sintattiche e semantiche (eliminando lo pseudo Scoto), *b)* riuscire a giustificare nel suo sistema i teoremi essenziali dell'approccio classico (è quanto la letteratura chiama *classical recapture*).

Nella logica di da Costa, detta *positive-plus*, si conserva pressoché intatta la parte positiva (senza negazioni) della logica classica, mentre, quando si incontrano contraddizioni, viene alterato il trattamento della negazione. Si fissa allora una base minimale, undici assiomi abbastanza tradizionali, e si introducono due operatori, uno per la consistenza e uno per l'inconsistenza. Se assumiamo *a* come consistente, allora vale il principio di non contraddizione e lo pseudo Scoto, e non potrà mai darsi *a* e non *a* pena l'esplosione. Con l'operatore dell'inconsistenza, diremo invece: se *a* è inconsistente, allora valgono sia *a* che non *a*. Da Costa e i suoi collaboratori hanno sviluppato tutta una serie di livelli logici diversi, sempre più potenti, mostrando come la consistenza e l'inconsistenza si diffondano dai componenti ai composti. Ne emerge una logica molto fluida, che spiega situazioni di fronte alle quali la logica classica non funziona. La negazione è il punto cruciale: da Costa formula alcune clausole semantiche che ne governano il funzionamento e in base alle quali riesce – attenuandone molto il potere – a evitare l'esplosione dello pseudo Scoto e a ottenere la *classical recapture*⁵⁰.

cit., p. 107-108.

49 Cfr. M. Bremer, *An Introduction to Paraconsistent Logic*, Frankfurt a. M., Peter Lang, 2005.

50 Cfr. N. C. A. Da Costa, “Sulla teoria dei sistemi formali contraddittori”, in D. Marconi (a cura di), *La formalizzazione della dialettica. Hegel, Marx e la logica contemporanea*, Torino, Rosenberg & Sellier, 1974.

Il problema del senso della negazione è di centrale importanza anche per il *dialetheism*. Dobbiamo infatti distinguere due gradi di paraconsistenza: debole e forte. A differenziare i due è l'ammissione della realtà di contraddizioni vere: il primo non lo fa, il secondo sì. Esistono teorie inconsistenti – la teoria ingenua degli insiemi, la semantica intuitiva, il calcolo infinitesimale di Leibniz, la teoria atomica di Bohr, i nostri sistemi di credenze quotidiani, ecc. – che sono utili, funzionano. La logica sottostante è una logica paraconsistente. La paraconsistenza debole lo afferma, ma senza ammettere che questa logica faccia riferimento a un modello reale, che insomma esistano stati di cose contraddittori. La paraconsistenza forte, invece, supera questo confine ammettendo la realtà di contraddizioni “vere”: si danno stati di cose che violano il principio di non contraddizione nella sua formulazione logico-semantica. Esistono *alcune* contraddizioni che sono dimostrabili e inevitabili perché reali. Questa forma di paraconsistenza è appunto il *dialetheism*.

Una *dialetheia* è una proposizione di cui tanto l'affermazione quanto la negazione⁵¹ – assunte come operazioni inverse – risultano vere. Il *dialetheism* sostiene che esistono *dialetheia*, ovvero proposizioni paradossali, vere, dalla forma contraddittoria (A e non-A); *a dialetheia is any true statement of the form: α and it is not the case that α [...] our concepts, or some of them anyway, are inconsistent and produce dialetheias*⁵². Esistono stati di cose e oggetti che sono insieme

51 Pur ammettendo alcune contraddizioni vere, il *dialetheist* dev'essere comunque in grado di esprimere l'esclusione di due proprietà incompatibili senza per questo rinunciare a credere che esistono contraddizioni vere. Priest sottolinea la centralità del trattamento della negazione per il *dialetheism* e la delicatezza della questione, data anche la molteplicità di tipi diversi di negazione: la negazione booleana, quella intuizionista, le leggi di De Morgan, etc. ciascuna delle quali risponde a regole e modelli semantici propri. Si tratta allora di capire che cosa intendiamo con “negazione” e con “teoria della negazione” in rapporto alla scelta *dialethetica*. Ovviamente, la negazione logica è ben diversa dalla negazione linguistica: non basta inserire il “non” in una frase qualsiasi per avere, sul piano logico, una negazione. E quindi, *we have a grasp of negation that is independent of the way that 'not' functions, and can use this to determine when 'notting' negates* (G. Priest, *Doubt Truth to Be Liar*, Oxford, Clarendon Press, 2006, p. 77). La negazione di “Qualche uomo è mortale” non è “Qualche uomo non è mortale”, ma “Nessun uomo è mortale”. Il senso logico autentico della negazione è la contraddizione, non la contrarietà né la subcontrarietà. La negazione è *a contradictory-forming operator*. Ci sono allora tre fondamentali leggi della negazione: la legge del terzo escluso, la legge della non contraddizione e la legge della doppia negazione. La negazione del *dialetheist* esprime una contraddizione secca, che divide il mondo in due parti totalmente contrapposte, simmetriche; *a genuine contradictory-forming operator will be one that when applied to a sentence, α , covers all the cases in which α is not true, and it is an operator, $\neg$, such that $\neg \alpha$ is true iff α is not true, i.e. is either false or neither true nor false* (p. 79). Abbiamo $[(\alpha \wedge \neg \alpha) \wedge \neg (\alpha \wedge \neg \alpha)]$, dunque il *dialetheist* non nega il principio di non contraddizione, anzi lo amplifica. Per i particolari tecnici, cfr. soprattutto la discussione della negazione booleana in *Doubt Truth to Be Liar*, cit., p. 88-102. La questione generale della negazione supera i limiti di questo saggio. Per un approccio generale, cfr. il classico L. Horn, *A Natural History of Negation*, Chicago, Chicago University Press, 1989.

52 G. Priest, *In Contradiction. A Study of the Transconsistent*, Oxford, Clarendon Press, 2006 (second edition), p. 4. Cfr. anche G. Priest, F. Berto, “Dialetheism”, *The Stanford Encyclopedia of Philosophy* (Summer 2013 Edition), Edward N. Zalta (ed.), URL = <<https://plato.stanford.edu/archives/sum2013/entries/dialetheism/>>; G. Priest, *An Introduction to Non-Classical Logic*, Cambridge, Cambridge University Press, 2008.

contraddittori e necessari. Il *dialetheism* è una prospettiva metafisica: *the view that some contradictions are true: there are sentences (statements, propositions, or whatever one takes truth-bearers to be) α , such that both α and $\neg \alpha$ are true, that is, such that α is both true and false*⁵³. Come scrive ancora Priest, *dialetheism is a metaphysical view: that some contradictions are true*, mentre *paraconsistency is a property of a relation of logical consequence*⁵⁴. La seconda può sussistere benissimo senza il primo, ma non vice-versa. La prospettiva del *dialetheism* impone una radicale trasformazione del nostro modo di concepire la realtà e la razionalità: *rationality is also intimately connected with dialetheism*⁵⁵.

Per il *dialetheist*, i paradossi sono un banco di prova imprescindibile. In *In Contradiction* Priest parte dalla diagnosi per cui i paradossi non sono semplici errori accidentali o isolati, ma sono generati da una comune condizione teoretica, contraddistinta da due aspetti: l'autoriferimento, la circolarità. Esiste cioè una sola struttura essenziale dei paradossi⁵⁶. Possiamo dividere la totalità degli enunciati in due sotto-insiemi: l'insieme degli enunciati veri e il suo complemento, che non coincide necessariamente *solo* con l'insieme degli enunciati falsi. Il paradosso è un enunciato che si trova in entrambi i sotto-insiemi. Si produce così una situazione bivalente, nel senso che l'enunciato paradossale, nel momento stesso della sua espressione, si apre a due prospettive: l'una in base alla quale è vero, l'altra in base alla quale è falso. Tale duplicità non si risolve semplicemente aumentando i valori di verità e postulando enunciati “veri e falsi”, perché potremo sempre re-interpretare il complemento inserendo in esso enunciati con un quarto valore di verità, e questo

53 G. Priest, *Doubt truth to be a liar*, cit., p. 1.

54 G. Priest, *One*, cit., p. XVIII.

55 G. Priest, *Doubt truth to be a liar*, cit., p. 1.

56 Diremo che un paradosso è un argomento che partendo da premesse intuitivamente vere e attraverso deduzioni intuitivamente accettabili arriva a una contraddizione o a un enunciato altamente controintuitivo. In *Beyond the Limits of Thought* (Cambridge, Cambridge University Press, 1995, p. 2-10), Priest sottolinea che gli enunciati dai quali emergono i paradossi sono enunciati che descrivono alcuni casi limite di quel che può essere concepito o espresso, e mette in rilievo che in tutti i paradossi ricorrono almeno due condizioni basilari: vi è una totalità (quella di tutte le cose esprimibili, descrivibili, ecc.) e un'operazione che genera un oggetto il quale è sia all'interno che all'esterno di questa totalità. Priest chiama queste due condizioni *chiusura* e *trascendenza*, cui aggiunge poi l'autoreferenzialità. Il punto è che un paradosso, pur essendo assurdo, non è affatto inconcepibile: possiamo infatti formularlo, siamo coscienti del suo carattere spiazzante. Importante la considerazione di Melandri: “L'obiezione della tartaruga ad Achille è che nessun sillogismo, se applicato a fatti concreti, può mai essere conclusivo. Infatti, se si vuol formulare esplicitamente l'interpretazione data nella fattispecie della premessa minore (la quale concerne l'accertamento del fatto), non ci salviamo da un regresso all'infinito” (*La linea e il circolo*, cit., p. 364).

renderà possibile replicare il paradosso⁵⁷.

In *In Contradiction* Priest distingue i paradossi in logici, semantici e insiemistici, dedicando poi un capitolo (il terzo della prima parte) ai teoremi di Gödel. L'esempio più immediato è: “io sto mentendo”. Se è vero, sto mentendo, ma chi mente dice il falso, dunque non sto mentendo. Se invece l'affermazione è falsa, non sto mentendo: dunque dico il vero, e cioè mento. Qualsiasi strategia mettiamo in campo per sbrogliare questa situazione, avremo sempre lo stesso risultato: un rafforzamento del paradosso, la contraddizione mi si ripresenta rafforzata. Posso usare i concetti della presunta soluzione per costruire un *revenge liar*, e se cerco di rispondere nuovamente arriverò a un punto tale che la mia risposta si auto-distruggerà.

Il *dialetheist* non cerca di rispondere al paradosso, ma cambia del tutto paradigma: prende come punto di partenza i paradossi, le contraddizioni insuperabili, che accetta come fatti, e costruisce attorno a essi un nuovo punto di vista sulla logica. Egli sostiene che una spiegazione dei paradossi è possibile, fornendo una nuova logica più plastica.

Nella linea *dialethetica* di Priest c'è un altro punto molto rilevante: esistono stati di cose contraddittori che sono non solo veri, ma anche percepibili, *osservabili*. Esistono situazioni contraddittorie che possiamo percepire direttamente come contraddittorie: *we may have perceptual experiences the contents of which are contradictory*⁵⁸. Priest fa l'esempio delle impossibili figure di Escher e della scala di Penrose: l'immaginazione è in grado di costruire situazioni visive non semplicemente paradossali, bensì apertamente in contrasto con la nostra esperienza fisica e contraddittorie in se stesse. *There is an intrigue auditory analogue of the continuously ascending [Penrose's] staircase. It is possible to produce a collection of musical tones which appears to be continuously ascending in pitch whilst, at the same time, never getting any higher – again, a contradictory situation*⁵⁹. Ma Priest fa anche altri due esempi. Il primo è il cosiddetto *waterfall effect*: quando si è sotto l'effetto di droghe o di alcol, capita di vedere l'ambiente circostante muoversi, pur restando fermi – non sapremo mai dire se quella stanza si sta davvero muovendo o

57 Cfr. F. Berto, *Teorie dell'assurdo*, cit., p. 62.

58 G. Priest, “Perceiving Contradictions”, *Australian Journal of Philosophy*, 77, 1999, p. 443.

59 Ivi, p. 441.

se è solo una nostra illusione: possiamo fissare un punto, mantenerlo fermo, anche se tutto il resto continua a “girare”, è la percezione di un movimento stazionario. Il secondo esempio riguarda invece i colori: *Now suppose a subject is shown a field, half of which is red and half of which is green, the two halves being separated by a black line. If the line is then removed, the brain fills in colours in the vacated space, and some observers report seeing that the boundary is now red and green*⁶⁰. E precisa, poche righe dopo: *It might be said that being red and green is not a contradiction. But it is: red and green are complementary colours. It is, hence, a conceptual impossibility for something to be both colours. A feature of complementary colours is that they cant go together. There is no reddy-green, in the same way, for example, that there is a reddy-blue. Something that is red and green, is red and not red*⁶¹.

§ 5. L'iterazione come struttura logica paraconsistente

Poniamo oggetti astratti, inesistenti e perfettamente identici tra loro, ma distinti. La contraddizione non elimina questi oggetti. Chiamiamo iterazione la trasgressione dell'identità degli indiscernibili che in tal modo si compie. La nostra tesi è che l'iterazione è la radice logica, in un senso non standard, di ogni possibile oggetto matematico. Tracciamo così un confine netto tra i numeri, che sono il piano su cui il matematico opera, e la costituzione interna dei numeri. Tutto quel che diremo nelle sezioni 6, 7 e 8 riguarda la costituzione interna del numero *e non il numero*.

AmMESSO l'esperimento mentale della sezione 4, chiameremo “Universo Black”, o U^B , un universo teorico perfettamente iterativo. Distingueremo due principi chiave che lo regolano. Il primo è il *principio di iterazione*. Formuliamolo in un linguaggio di secondo ordine:

$$PI': \forall x \forall y [\forall F (Fx \leftrightarrow Fy), (x \neq y)]$$

Posto che F è un insieme di proprietà *tecnico e chiuso*, dati due oggetti identici con identiche proprietà, essi *non sono* lo stesso oggetto, ma due. Più in generale: dati n oggetti identici, essi *non sono* lo stesso oggetto, ma n . Dato un qualsiasi x – chiameremo *item* questo elemento-unità astratto

⁶⁰ Ivi, p. 442.

⁶¹ Ibid.

– x è se stesso e immediatamente diverso da sé, e quindi molteplice. L'item nasce da un atto creativo della nostra immaginazione: è un oggetto tecnico, interamente dipendente da noi e al quale conferiamo un *range* di qualità fisse e chiuso. Ed è il polo di una relazione, la relazione con la sua copia identica. L'insieme di tutti gli item è U^B . Nella formula abbiamo tre simboli principali: il quantificatore universale $\forall$, l'identità $=$ e la differenza $\neq$ che interpretiamo secondo i risultati del nostro esperimento mentale. C'è poi la virgola, che indica una lettura distributiva, non collettiva, della contraddizione. Facciamo così a meno della congiunzione: $[(x = x), (x \neq x)]$ è una coppia di enunciati di cui uno nega l'altro.

In U^B la contraddizione non è esplosiva, ma funzionale. Due item qualsiasi sono insieme identici e diversi, ma questo non li annulla. Oggetto contraddittorio, l'item ha due facce: *a)* è molteplice ($x \neq x$) perché si sviluppa in una serie, differendo sul piano dell'identità numerica; *b)* è singolo ($x = x$) perché la moltiplicazione non aggiunge nulla alla sua identità qualitativa (c'è sempre uno stesso tipo che s'istanza in maniera identica: un *range* chiuso di qualità astratte, formali). La duplicità dei piani che entrano in collisione (identità numerica / identità qualitativa) rende l'item (x) un oggetto sì contraddittorio, ma fluido, versatile. L'item è *immediatamente* un insieme infinito poiché *immediatamente* si moltiplica all'infinito per PI' . Ma questo infinito è *immediatamente* anche un singolo individuo. Il peso di questo *immediatamente* va calibrato, perché qui sta l'essenza del numero e della computazione. Il nesso paraconsistente del principio impedisce non solo l'esplosione dell'item, ma anche la confusione tra U^B e il singolo item. U^B è un insieme di contraddizioni che si tengono insieme, si moltiplicano, ma non si annullano.

Il secondo principio è il principio di *contraddizione continua* per cui nessuna contraddizione blocca il sistema, bensì lo estende:

$$C^c: \forall x [(x = x), (x \neq x)] \cdot x, x$$

All'apparato simbolico usato in PI' aggiungiamo $\cdot$ che chiamiamo *funtore di combinazione*. È un connettivo molto semplice che procede per eliminazione: il passaggio dalla sinistra alla destra della formula implica la scomparsa di $=$ e $\neq$. Il funtore esprime il distacco tra l'identità qualitativa e

l'identità numerica e l'equilibrio contraddittorio che si crea. Diverso il comportamento delle due virgole: la prima agisce normalmente, mentre la seconda esprime una negazione molto semplice. Questa seconda virgola è il punto più delicato: indica una discernibilità neutrale. Come dicevamo, la contraddizione non distrugge il sistema, bensì lo preserva moltiplicando gli item. Il sistema non è triviale perché da $[(x = x), (x \neq x)]$ non segue qualsiasi cosa, ma sempre e solo la stessa cosa, x , moltiplicata. Si dirà che la nostra minilogica è banale perché monotona. Per la verità, come vedremo, la monotonia non è affatto banale. Anzi, ha un valore costruttivo, euristico.

È un fatto empirico che le particelle elementari subatomiche (bosoni e fermioni) della stessa specie sono identiche ma distinte. Condividono tutte le stesse caratteristiche fisiche, hanno una medesima identità qualitativa, ma sono molteplici e interagiscono tra loro. Identità qualitativa e identità numerica si scindono. Parliamo infatti di tanti elettroni, *e non di un solo elettrone*. Non c'è un *frame* di coordinate con cui singolarizzare queste particelle e considerarle una a una come individui determinati, dato che, in base ai principi della fisica quantistica, esse non possiedono una posizione e una traiettoria fisse. Certo, i fisici hanno elaborato formalismi matematici e metodi statistici per descrivere le particelle e il loro comportamento; hanno cercato di definire un *frame* minimale in cui identificarle, anche solo per poterne parlare, com'è il caso per il teorema spin-statistica. Tuttavia, le particelle rimangono identiche, indistinguibili, e molteplici, tanto che se ne studia la distribuzione e l'interazione. La stessa cosa può dirsi per gli *item subnumerici*: sono identici, ma distinti. Sono le particelle elementari di qualsiasi numero.

§ 6. *Fenomenologia della computazione*

Se consideriamo i postulati di Peano quale base di partenza per ogni buona descrizione fenomenologica del numero⁶², non possiamo non tener conto di tre aspetti essenziali, che la nostra descrizione dovrà giustificare:

(1) Lo statuto dello zero, 0 – perché lo zero è un numero? Come fa l'assenza di quantità a generare un numero?

62 Cfr. E. Mendelson, *Introduzione alla logica matematica*, Torino, Bollati Boringhieri, 1972, p. 129-145.

(2) La profondità del numero, il fatto che in- e a partire da- ogni numero possiamo costruire e inventare infiniti altri numeri. Come si realizza l'operazione del successore? Perché da 0 passo a 1?

(3) Il principio di induzione per cui *tutti* i numeri condividono *tutte* le proprietà dello zero e del successore; esiste una fondamentale continuità qualitativo-strutturale tra i numeri.

Ogni numero si costruisce a partire dallo zero, è una modulazione dello zero. Ma che cos'è lo zero? E perché diciamo che è un numero? E come facciamo a muoverci in U^B ? Consideriamo l'insieme *vuoto*. Ad esso applichiamo due operazioni base che sono PI' : $\forall x[(x = x), (x \neq x)]$; C^c : $\forall x [(x = x), (x \neq x)] \cdot x, x$. La conseguenza è che l'insieme vuoto si moltiplica all'infinito come lo stesso – abbiamo U^B . Dall'applicazione di PI' e C^c nasce perciò una struttura iterativa esponenziale di insiemi vuoti. I numeri sorgono da una limitazione di questa struttura, ossia dall'applicazione di un limite a un'iterazione del tutto formale. Il limite è la scala da cui osservo U^B .

Ma che cosa vuole dire “applicare il limite”? Posto che ammette certe contraddizioni come vere, come fa il *dialetheist* a riconoscere la falsità di altre contraddizioni ed escluderle? Per fermare la catena iterativa bisogna infatti escludere certe contraddizioni e isolarne altre per considerare soltanto queste ultime. L'applicazione del limite è un atto di negazione, ma di una specifica negazione: l'incompatibilità materiale⁶³, per cui il *dialetheist* esclude qualcosa e ferma il concatenarsi delle contraddizioni. Mediante tale operazione, il *dialetheist* può considerare l'insieme delle iterazioni che ha isolato come un unico, singolo oggetto (il 5, il 4, il 3, ecc.) senza per questo cancellarne la natura iterativa interna.

Possiamo allora distinguere tre passi che compongono la procedura logica alla radice di *ogni possibile numero*. Usiamo la sigla “NON” per indicare l'incompatibilità materiale.

63 Cfr. F. Berto e L. Bottai, *Che cos'è una contraddizione*, cit., p. 117-119. La forza dell'incompatibilità materiale proposta da Berto sta nel fatto che si tratta di una forma di negazione che non è definita facendo riferimento alla coppia vero-falso, ma al concetto di esclusione, “concetto implicato dalla nostra esperienza del mondo come agenti, che fronteggiano scelte fra compiere una certa azione o un'altra (qualcosa che riteniamo facciano anche animali non dotati di linguaggi articolati)” (p. 119); è il semplice gesto di esclusione. Questo permette al *dialetheist* di escludere certe contraddizioni come false, e quindi di non ricadere ancora nel trivialismo.

- 1) $\forall x [(x = x), (x \neq x)]$ (iterazione)
- 2) $\forall x [(x = x), (x \neq x)] \cdot x, x$ (contraddizione continua)
- 3) $\text{NON-}\forall x [(x = x), (x \neq x)] \cdot x, x$ (limite)

La paraconsistenza garantisce l'estrema fluidità del meccanismo. Proprio in virtù della sua natura paraconsistente il numero è un oggetto instabile e diventa un mezzo versatile capace di descrivere situazioni anche molto diverse tra loro. L'assumere un punto di vista paraconsistente ci permette di dire che gli item sono al contempo identici e diversi, e quindi *uno e molti*. Il numero ha uno statuto dinamico che “rimbalza” perennemente tra molteplicità (M) e sinteticità (S). Il limite NON ferma il “rimbalzo” e apre il campo per processi logici di altra natura: di primo ordine, di secondo ordine, ecc.

$$M \leftrightarrow S$$

Gli interi negativi (Z) nascono dall'iterazione dei numeri naturali (N), ma in un altro senso. Nel produrli, al principio di iterazione e alla contraddizione continua si legano le due operazioni della funzione di passaggio opposto [$-(n) = -n$; $-(-n) = n$] e il valore assoluto ($|x|$). Non c'è un limite prefissato per un singolo numero, o meglio il limite è costituito da queste due operazioni; semplicemente iteriamo *tutto* l'insieme dei numeri naturali, producendo due serie infinite che vanno in direzioni opposte. Diverso è il passaggio dagli interi ai razionali (Q). In questo caso abbiamo l'iterazione degli interi *negli interi stessi*. E infatti i razionali nascono dal rapporto *interno* tra gli interi. Banalmente: dire 14/11 vuol dire chiedersi “quante volte sta”, quanto si itera l'11 nel 14. Il risultato dell'iterazione dei numeri interi *nei numeri interi* è triplice: *a)* un numero naturale, quando il numeratore è una perfetta iterazione del denominatore; *b)* un numero decimale finito, quando il numeratore è una perfetta iterazione del denominatore perché riusciamo a ridurre lo scarto a zero; *c)* un numero decimale periodico, quando il rapporto tra numeratore e denominatore è sì un'iterazione, ma non del denominatore nel numeratore, bensì di un certo gruppo di numeri: il periodo. I numeri decimali periodici aprono un scenario inedito: l'iterazione degli interi negli interi

stessi produce una nuova iterazione, ancora più profonda, attualmente infinita. Il fatto che le cifre decimali, se sono infinite, debbano comunque iterarsi in modo ciclico è importante: il *loop* è l'unico elemento che distingue i numeri razionali da quelli irrazionali.

Se questo è vero, *che cos'è $\sqrt{2}$* ? che tipo di densità dobbiamo immaginare per giustificare $\sqrt{2}$? Qualsiasi filosofia del numero che voglia dirsi adeguata deve rispondere a questa domanda. In un numero irrazionale la successione delle cifre dopo la virgola è infinita e non periodica, per questo la conoscenza completa di un numero irrazionale è impossibile. Ora, è noto che data una qualsiasi coppia di numeri reali (razionali o irrazionali) esistono infiniti numeri razionali compresi fra essi. Ciò significa che tra due numeri qualsiasi si ripete l'intero infinito insieme di numeri razionali e che ogni numero irrazionale può essere approssimato per eccesso o per difetto. Ma questa resta solo un'approssimazione. Il numero irrazionale è un numero che porta in sé l'intero infinito insieme dei numeri razionali, nel senso che l'intero infinito insieme dei numeri razionali si ripete in esso. *Tutti i numeri in un solo numero*. Il limite c'è, ma è imprendibile. Paradossalmente, i numeri irrazionali sono i numeri più autentici, più originari, perché mostrano la presenza attuale di U^B . Anche i numeri irrazionali sono computabili.

La formazione di *tutti* i numeri gioca sulla “doppia faccia” degli item e la plasticità dell'insieme U^B . Non è altro che un processo di suddivisione e gerarchizzazione di U^B , come nel frattale. È la creazione di una nuova struttura a partire dalla struttura di base.

Atteniamoci per ora soltanto ai numeri naturali. Creiamo collezioni di item a partire da U^B , ma inseriamo un limite che “sospende” l'iterazione. L'1 sarà dunque un insieme formato soltanto dallo 0, l'insieme vuoto, considerato quale suo unico elemento, e da un limite all'iterazione di 0. Il 2 sarà invece un altro livello formato dall'insieme vuoto, dall'insieme 1 (0 + limite) e da un altro limite all'iterazione; così il 3 dallo 0, da 1 e da 2. Abbiamo una successione gerarchica basata sempre sugli stessi elementi (l'insieme vuoto che si moltiplica), che possiamo schematicamente descrivere seguendo la concezione iterativa degli insiemi di Boolos⁶⁴. Le parentesi denotano il limite

⁶⁴ Cfr. G. Boolos, “The Iterative Conception of Set” (1971); “Iteration Again” (1989), in G. Boolos, *Logic, Logic and Logic*, Cambridge, Harvard University Press, 1998, p. 13-29, 88-104. Boolos parte dall'idea che gli elementi di ogni insieme (individui concreti, non insiemi) debbano essere presenti prima che l'insieme si formi. A ogni insieme

dell'iterazione; a destra collochiamo la rappresentazione degli item (in questo caso la funzione limite è indicata da $\langle \rangle$):

$$\begin{array}{ll}
 0 = U^B & \leftarrow o \rightarrow \\
 1 = \{0\} & \langle o \rangle \\
 2 = \{0, \{0\}\} & \langle o o \rangle \\
 3 = \{0, \{0, \{0\}\}\} & \langle o o o \rangle \\
 4 = \{0, \{0, \{0, \{0\}\}\}\} & \langle o o o o \rangle
 \end{array}$$

e così via. Questi insieme-tipo sono i numeri cardinali.

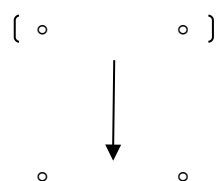
La nostra domanda è *fenomenologico-noematica*, non matematica ma meta-matematica: che rapporto sussiste tra gli 0 contenuti tra le parentesi? Sono perfettamente identici, specificamente identici o soltanto analoghi? E che cosa sono le parentesi? La nostra linea è che questi 0 sono sempre lo stesso oggetto iterato; ogni numero è formato da una serie di copie perfettamente identiche dello 0 che condividono tutte le stesse proprietà. *Ma perché?* Perché invece non diciamo che sono specificamente identici o analoghi? Per il fatto che entrambe queste forme di identità lascerebbero aperta la possibilità di variazioni dello 0, del sussistere di differenze anche minime tra gli 0 e questo distruggerebbe non solo l'unità del numero, ma anche la certezza del procedimento matematico. Il matematico *deve iterare*. La testura interna al numero *deve essere* sempre la stessa: l'iterazione di un'iterazione, in sostanza U^B , cioè una *dialetheia*. E le differenze che sorgono tra un numero e un altro numero sono differenze di gerarchia, di parentesi: ovvero differenze tra numeri, e non differenze nella “materia” di cui sono fatti i numeri. Questo intendevo quando dicevo che il numero è un'immagine di U^B , un punto di vista su una struttura infinita; le differenze, gli scarti esistono tra i punti di vista, non in ciò che vediamo.

corrisponde quindi un livello e ogni insieme si trova a uno stadio successivo rispetto ai suoi elementi. Avremo il livello zero, l'insieme vuoto, dopo di che il livello 1, sul quale si formano tutti gli insiemi che si possono formare a partire dagli insiemi formati al livello precedente – l'insieme vuoto – e un ulteriore individuo, cioè proprio il nuovo insieme formato a livello 1. Con il livello 2 si fa lo stesso. Su ogni livello si recuperano tutti i livelli precedenti iterandoli e considerando l'iterazione come un nuovo insieme che si aggiunge ai precedenti. Dopo tutti questi livelli avremo il livello Omega sul quale si recuperano tutti gli insiemi precedenti e si aggiunge un nuovo individuo, appunto Omega. Boolos dimostra che questa rappresentazione può essere tradotta in una logica del primo ordine con cui è possibile giustificare gli assiomi della teoria classica degli insiemi Zermelo-Fraenkel. Così Boolos mostra che la teoria classica non è soltanto un espediente ad hoc per evitare i paradossi di Russell. Possiamo esporre gli assiomi Zermelo-Fraenkel a partire da una formalizzazione della teoria dei livelli insiemistici in cui l'iterazione gioca un ruolo chiave.

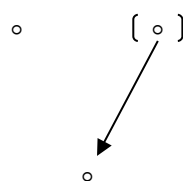
La rappresentazione iterativa dei numeri fondata sulla fluidità di U^B implica un quadruplice “salvataggio”. Essa “salva” dapprima la profondità del numero perché in ogni numero ritroviamo l'item (lo zero), e perciò infiniti altri item con i quali costruire, da- e in- quel numero, infiniti altri numeri con gerarchie diverse. “Salva” la natura ambigua del numero, che insieme è continuo e discreto. “Salva” poi l'induzione matematica perché garantisce la compattezza strutturale dei numeri, che non saranno altro che variazioni su un unico inesauribile tema. Infine, questa rappresentazione “salva” il fenomeno del *nesting* – i numeri sono “annidati” gli uni negli altri – che sta alla radice della ricorsione. Ovvio, ci sono buone ragioni per sostenere l'ipotesi di U^B , ma certo non si tratta di ragioni conclusive. Questa strada ci spiega in modo più chiaro e semplice la struttura interna del numero, quel che la teoria degli insiemi presuppone. Il prezzo dell'ipotesi, tuttavia, è alto: il *dialetheism*. E non è detto che soluzioni migliori non possano essere ottenute a un costo minore.

Ponendoci in qualche modo “al di qua” delle proposizioni primitive di Peano, diremo che *ogni oggetto matematico* sarà composto da tre elementi: *a)* il principio di iterazione, che è insieme principio di scomposizione e di concatenazione infinita; *b)* la funzione parentesi, il limite dell'iterazione; *c)* due orientamenti dell'iterazione: l'addizione e la moltiplicazione, le operazioni fondamentali dell'aritmetica, con tutte le loro proprietà (commutativa, associativa, distributiva). E infatti questi sono i due sviluppi naturali dell'iterazione:

Orientamento moltiplicativo



Orientamento additivo



Dati uno o più individui *qualsiasi*, essi immediatamente si iterano. Posto il limite, bloccato il processo iterativo generale, continuiamo a iterare. Il come dell'iterazione è dato dal primo

orientamento generale, additivo o moltiplicativo, dal quale poi si ottiene il secondo orientamento generale, sottrattivo o divisorio. Tramite l'iterazione possiamo definire tutte le principali operazioni aritmetiche.

Chiariamo con un altro schema. Abbiamo due simboli base: “o” indica l'item (l'insieme vuoto), < > la funzione limite. Abbiamo poi tre operazioni legate alla funzione limite: Z come “nessun limite”, C come “posizione e chiusura del limite”, A come “apertura del limite”. Abbiamo infine due simboli per indicare le due “facce” che l'item, in virtù della sua natura paraconsistente, può assumere: *Id* come “identico a se stesso” ($x = x$); *Dif* come “diverso da se stesso” ($x \neq x$). Il primo livello è sempre U^B .

Per l'orientamento additivo, procediamo in questo modo:

$PI^I: \forall x[(x = x), (x \neq x)]$ $C^C: \forall x [(x = x), (x \neq x)] \bullet x, x$	o	<i>Z, Id, Dif</i>	(o = o), (o ≠ o)
	< o >	<i>C, Id</i>	(o = o)
	o o	<i>A, Dif</i>	(o ≠ o)
	< o o >	<i>C, Id</i>	[(o o) = (o o)]
			

Per l'orientamento moltiplicativo, procediamo in questo modo:

$PI^I: \forall x[(x = x), (x \neq x)]$ $C^C: \forall x [(x = x), (x \neq x)] \bullet x, x$	o	<i>Z, Id, Dif</i>	(o = o), (o ≠ o)
	< o o >	<i>C, Id</i>	[(o o) = (o o)]
	o o o o	<i>A, Dif</i>	[(o o) ≠ (o o)]

$$< o o o o > \quad C, Id \quad [(o o o o) = (o o o o)]$$

.....

Come risulta dallo schema, la funzione limite ha soltanto il ruolo di modulare e articolare le due facce dell'item: non annulla ma regola la paraconsistenza. Il vantaggio è che, da una parte, non cancelliamo affatto la logica classica e l'insiemistica booleana, che sussistono nello spazio definito dal limite – l'incompatibilità materiale “apre una porta” verso il recupero di altre logiche. Dall'altra, possiamo sempre, pur mantenendo un limite chiuso (mettiamo ad esempio: $< o o o o >$), iterare in esso, ovvero considerare il singolo item, “aprirlo”, iterarlo e “chiuderlo”. Da un singolo item possiamo formare infiniti nuovi limiti e limiti all'interno di altri limiti, e così via. Il numero è allora descrivibile come un “concatenamento macchinico” che esprime la complementarità di albero (il limite) e rizoma (U^B), di molare e molecolare, di individuazione e frammentazione dell'individuazione⁶⁵. Attorno a un nucleo inconscio (U^B) si vengono così a stratificare livelli logici diversi, di primo e di secondo ordine, e altro, che talvolta si corrispondono talaltra no.

L'ipotesi di U^B risponde a una domanda cruciale: *come formiamo gli insiemi e in generale gli oggetti matematici?* Per costruire un insieme partiamo da oggetti già dati. Dobbiamo capire come pensiamo questi oggetti e da cosa nascono i loro rapporti. L'insiemistica classica non si occupa di questo punto e presuppone troppo. L'ipotesi di U^B ha soltanto un presupposto: l'insieme vuoto.

Consideriamo dei semplici fatti. Come fa il matematico a sapere che nel momento in cui opera il successore [$s(x) = x + 1$], l'1 aggiunto è sempre lo stesso 1, la stessa quantità, e continuerà sempre a essere la stessa? Come fa a essere sicuro di aggiungere 1 e non, mettiamo, 1,001 o 2 o 3,5? Che cosa intende con quell'1? Se la moltiplicazione è un'addizione ripetuta, che cosa assicura la continuità dell'iterazione? Nel senso: se 5×4 non è altro che $5 + 5 + 5 + 5$, che tipo di relazione sussiste tra questi 5? Si dirà: i quattro 5 sono solo diversi nomi per indicare una stessa cosa, un

65 Cfr. G. Deleuze, F. Guattari, *Mille plateaux. Capitalisme et schizophrénie*, Paris, De Minuit, 1980. Si obietterà che la citazione pare dimenticare tutta la teoria del desiderio e dell'inconscio che sottende l'*Anti-Edipe* nonché *Mille plateaux*. Ma perché siamo così restii al dare al numero una connotazione intensiva e desiderativa? Perché la matematica dev'essere il regno di strutture senza intensità né inconscio? “La scienza sarebbe completamente folle, se la si lasciasse fare, guardate la matematica, non è una scienza, ma un gergo prodigioso e nomadico” (G. Deleuze, F. Guattari, *Mille piani. Capitalismo e schizofrenia*, a cura di M. Carboni, Roma, Castelvecchi, 2010, p. 69).

referente unico, appunto il numero 5. Ma questa idea funziona davvero? I numeri possono essere considerati referenti stabili come le cose materiali per i nomi che usiamo nel linguaggio naturale? Posso moltiplicare un cane per 4? Che cosa denota l'espressione 5^2 ? Se affermiamo che $5 + 5 + 5 + 5$ è soltanto l'iterazione di un nome, vale a dire un insieme di nomi diversi che indicano ogni volta sempre una stessa cosa, il 5, dobbiamo allora essere capaci di rispondere a due domande: a) come facciamo a essere certi che ogni volta consideriamo come referente sempre lo stesso 5? b) Se prendiamo lo *stesso* 5 ogni volta, perché otteniamo 25 e non soltanto sempre quello *stesso* 5? Dove sta l'effetto cumulativo? Nella memoria? Ma il ragionamento matematico non si fonda sullo sforzo mnemonico. Siamo costretti a dire che 5^2 designa tanti diversi 5 perfettamente identici tra loro sintetizzati in un'espressione scritta. Sono sempre una stessa quantità, *ma diverse*. E in che cosa diverse? Se diciamo che tra i 5 di 5^2 c'è un rapporto di isomorfismo, similarità o analogia, per cui essi condividerebbero solo certe proprietà e non altre, allora il matematico non potrà mai – in forza della vaghezza intrinseca alla nozione di analogia – essere certo del suo risultato e non potrà mai mettere con tanta sicurezza il segno = tra 5^2 e 25 o tra 5×4 e 20. In verità, 5^2 è uno schema iterativo: ci dice come e quanto dobbiamo iterare il 5. Quei 5 *devono* condividere tutti le stesse proprietà: sono identici e discernibili. Il $5 + 5 + 5 + 5$ non è altro che un'immagine di U^B , cioè un modo di rappresentare (per limite e orientamento) U^B , dove U^B è una struttura i cui *item* – i nodi che la tengono insieme – non sono affatto numeri; li si chiami operazioni, oggetti o altro, non cambia nulla. Sono semplici rapporti regolati da PI' e C^c . L'item (nel nostro schema lo zero) è un'iterazione pura regolata da principi logici, non matematici; un rapporto logico di cui ci serviamo per pensare il numero. È il non-numero che usiamo per costruire i numeri.

Che senso ha la nostra ipotesi? U^B non è un insieme ordinario, ma è economicamente vantaggioso. Il matematico considera *solo* U^B , ha a che fare con un solo oggetto, una struttura logica semplicissima ma non banale; in ogni suo punto il matematico sa perfettamente dove si trova, sempre nello stesso punto, *eppure il matematico si muove in questa struttura*. Tale muoversi paradossale “sta sotto” l'operare del matematico, è silenzioso e paradossale. Diremo che U^B più che

un insieme è un “ambiente” per insiemi, un “luogo” in cui, grazie al limite e all'orientamento, possiamo costruire sempre nuovi insiemi diversi, perché giochiamo con una “pasta” paraconsistente dinamica, fluida, inesauribile.

§ 7. Fenomenologia della computazione – 2

Ammesso ciò, *che cos'è la computazione?* È la riduzione di un numero alla sua base iterativa, a U^B . Una funzione è ricorsiva se e solo se *esibisce* il rapporto tra l'immagine e quel che l'immagine descrive, tra il numero e l'insieme primitivo U^B . Le funzioni ricorsive, così come gli insiemi, le proprietà e le relazioni ricorsive, non fanno altro che *mostrare* il rapporto tra U^B e l'apparato sintattico-semantico (procedure, formule, teoremi, dimostrazioni) di cui il matematico si serve nel suo lavoro.

La teoria della computabilità⁶⁶ è quel settore della logica matematica in cui vengono indagati concetti quali algoritmo e funzione computabile in modo algoritmico. Una funzione si dice calcolabile se esiste almeno un algoritmo che consente di calcolarne i valori per tutti gli argomenti. Un numero si dice computabile se può essere il valore di una funzione computabile per un certo argomento, ossia se esiste un procedimento che a partire da una base, in un numero finito di passi, dà quel numero come *output*. In altri termini, una funzione (o un numero) è computabile *se e solo se* ad essa corrisponde una *effective procedure*. Questa forma di prova presenta quattro caratteristiche: 1) *executability*, è una procedura che consiste in un numero finito di istruzioni deterministiche (definiscono un solo passo alla volta nel processo), non ambigue e di numero finito; 2) *automaticity*, la procedura non richiede particolari intuizioni né doti creative; 3) *uniformity*, la procedura è sempre la stessa per qualsiasi argomento della funzione; 4) *reliability*, una volta conclusa la procedura genera il valore corretto della funzione per qualsiasi argomento dopo un numero finito di passi. La macchina di Turing, le funzioni ricorsive o il Lambda calcolo di Church sono solo modi diversi di

⁶⁶ Cfr. N. Immerman, "Computability and Complexity", *The Stanford Encyclopedia of Philosophy* (Fall 2011 Edition), Edward N. Zalta (ed.), URL = <<http://plato.stanford.edu/archives/fall2011/entries/computability/>>. Per un inquadramento storico: R. Adams, *An Early History of Recursive Functions and Computability. From Gödel to Turing*, Boston, Docent Press, 1983; B. J. Copeland, C. J. Posy, O. Shagrir (ed.), *Computability. Turing, Gödel, Church, and Beyond*, London-Cambridge, The Mit Press, 2013.

esprimere formalmente il concetto di *effective procedure*⁶⁷.

Gli algoritmi sono *effective procedures*, metodi per la soluzione di problemi, elenchi di istruzioni scritte in un certo linguaggio. Un problema si caratterizza «mediante i dati di cui si dispone all'inizio e i risultati che si vogliono ottenere. Risolvere un problema significa ottenere in uscita i risultati desiderati a partire da un certo insieme di dati presi in ingresso. I dati in ingresso vengono anche detti (valori in) *input* e i risultati in uscita (valori in) *output*»⁶⁸. Il problema può assumere una struttura funzionale posto che la funzione è in generale una relazione, «una corrispondenza tra due insiemi (D e C) che, a ogni elemento di D (detto dominio) preso come argomento, associa come valore uno e un solo elemento di C (detto co-dominio)». In matematica «si è quasi sempre interessati a funzioni in cui dominio e codominio sono insiemi di numeri e la corrispondenza φ è definita mediante operazioni numeriche»⁶⁹. In una funzione, i dati in ingresso corrispondono agli argomenti, mentre quelli in uscita ai valori.

Con le nozioni di algoritmo e di funzione si definiscono i principali concetti della teoria della computabilità: la calcolabilità delle funzioni e la decidibilità di proprietà, relazioni e insiemi. Uno dei risultati più notevoli della teoria della computabilità è che esistono funzioni non computabili, i cui valori non sono calcolabili mediante alcun algoritmo. Si dimostra infatti che l'insieme delle funzioni calcolabili, seppure infinito, è numerabile⁷⁰, mentre l'insieme delle funzioni aritmetiche è un insieme con la cardinalità del continuo. Il primo è “più piccolo” del secondo.

La teoria della ricorsività è il cuore della teoria della computabilità. Il suo obiettivo è dare una determinazione rigorosa del concetto di funzione calcolabile. A rigore, infatti, funzione e algoritmo sono concetti distinti. Possono esserci infiniti algoritmi per una sola funzione. La funzione calcolabile non può essere identificata con l'algoritmo che la calcola: ha bisogno di una caratterizzazione autonoma⁷¹. La teoria della ricorsività mira a caratterizzare la classe delle funzioni

67 Cfr. G. Piccinini, *Physical Computation. A Mechanistic Account*, Oxford, Oxford University Press, 2015, p. 247.

68 M. Frixione, D. Palladino, *Funzioni, macchine, algoritmi. Introduzione alla teoria della computabilità*, Roma, Carocci, 2004, p. 19.

69 Ivi, p. 54-55.

70 Per la dimostrazione, cfr. M. Frixione, D. Palladino, *Funzioni, macchine, algoritmi*, cit., p. 103.

71 Cfr. M. Frixione, D. Palladino, *Funzioni, macchine, algoritmi*, cit., p. 67: “Dal fatto che una stessa funzione può essere calcolata da più algoritmi diversi segue immediatamente che una funzione calcolabile non può essere identificata con un algoritmo che la calcola. Una funzione infatti è definita in maniera del tutto indipendente rispetto

calcolabili aritmetiche, a un argomento, relative al dominio N dei numeri naturali, $\varphi: N \rightarrow N$.

L'idea centrale della teoria della ricorsività è che possiamo dare una caratterizzazione prettamente matematica di funzione calcolabile attraverso quel che si dice *una definizione ricorsiva o induttiva*. Le funzioni calcolabili sono funzioni ricorsive, definibili a partire da alcune funzioni base, molto semplici, mediante operazioni che “conservano” la calcolabilità iniziale. Tale metodologia è stata introdotta negli anni Trenta da matematici e logici quali Gödel, Church, Kleene, Turing, Post e il loro studio si è enormemente sviluppato con l'avvento dei primi calcolatori e dei primi linguaggi di programmazione. Ma le applicazioni di questo concetto – e in realtà di tutta la teoria della computabilità – vanno ben oltre, interessando anche campi come la linguistica, il problema mente-corpo, le scienze cognitive e perfino la genetica.

Che cos'è propriamente una funzione ricorsiva? Una funzione ricorsiva è una funzione che porta “in se stessa” la procedura che la calcola. Possiamo costruirla attraverso una serie di operazioni specifiche che, a partire da funzioni base molto semplici e intuitivamente calcolabili dette *funzioni base*, “conservano” la calcolabilità di tali funzioni. In termini tecnici, si dice che la classe delle *funzioni ricorsive primitive* è la più piccola classe di funzioni aritmetiche contenente le funzioni base e ottenute applicando un numero finito di volte le operazioni di composizione e ricorsione. Tuttavia, la classe delle funzioni ricorsive primitive non esaurisce tutta la classe delle funzioni calcolabili in modo effettivo. Esistono infatti funzioni calcolabili che non sono ricorsive primitive, come ad esempio la funzione Ackermann. È possibile estendere la classe delle funzioni ricorsive primitive a una classe più ampia di funzioni, le funzioni ricorsive generali, che comprende le funzioni ottenute a partire dalle funzioni base applicando un numero finito di volte tre operazioni: composizione, ricorsione e minimalizzazione⁷².

Limitiamoci ora alle funzioni ricorsive primitive. Una funzione si dice *ricorsiva primitiva* se è una funzione base o è ottenuta mediante l'applicazione in un numero finito di volte delle operazioni

ai metodi per computarla”. Se la funzione è calcolabile, ciò non dipende dal metodo che possiamo usare per risolverla. Il che vuol dire che esiste una computabilità intrinseca alla funzione.

72 Cfr. M. Frixione, D. Palladino, *Funzioni, macchine, algoritmi*, cit., p. 180-191.

di composizione e ricorsione alle funzioni base. In generale, la *definizione induttiva*⁷³ di una funzione consta di tre clausole: i) *la base* – si stabilisce il valore della funzione per l'argomento 0; ii) *il passo* – supponendo noto il valore che la funzione assume per un generico argomento n , si definisce il valore della funzione per il successivo di n affidandosi all'induzione; iii) *la conclusione o chiusura*, in cui si raggiunge la definizione ricercata⁷⁴. La teoria della ricorsività ci dice che possiamo usare questo schema per definire l'intera classe delle funzioni ricorsive primitive.

Partiremo allora da alcune *funzioni base*, che sono tre:

(1) La funzione *zero* o z , per cui dato un qualsiasi numero naturale x come argomento, gli assegna sempre e solo lo zero come valore, e dunque avremo $Z(x) = 0$, $Z(0) = 0$, $Z(1) = 0$, ecc. Il decorso dei valori è monotono.

(2) La funzione *successore* o s , per cui dato un numero x come argomento gli assegna come valore il successore, e avremo una situazione di questo tipo: $s(x) = x + 1$; la successione dei valori sarà: $s(0) = 1$, $s(1) = 2$, $s(2) = 3$, $s(3) = 4$, ecc.

(3) Una serie di funzioni dette *di identità*, o *proiezione*. Dato un qualsiasi numero x come argomento, la funzione p^1_1 a un argomento gli assegna come valore sempre quello stesso numero: $p^1_1(x) = x$, e quindi $p^1_1(1) = 1$, $p^1_1(2) = 2$, $p^1_1(3) = 3$, e così via. Lo stesso meccanismo si dà con funzioni a più argomenti. Date due funzioni a due argomenti, p^2_1 e p^2_2 e data una coppia di numeri (x, y) quali argomenti, la prima dà come valore il primo dei due numeri e la seconda il secondo: $p^2_1(x, y) = x$; $p^2_2(x, y) = y$. Diremo: per ogni numero n ci sono n funzioni di identità p^n_i a n argomenti tali che gli assegnano l' i -esimo di tali argomenti come valore. Ad esempio: $p^4_2(2, 6, 7) = 6$. È semplicemente un modo di assegnare a un numero un altro numero di una serie. Ogni volta a quel numero ne associo un altro.

Consideriamo adesso le operazioni che “conservano” la calcolabilità, conducendo da funzioni

73 La definizione è detta anche “derivazione”: possiamo ottenere, per ogni funzione ricorsiva primitiva, una lista di istruzioni che permette di calcolarne i valori. Cfr. M. Frixione, D. Palladino, *Funzioni, macchine, algoritmi*, cit., p. 174-175.

74 Cfr. Cfr. G. Piccinini, *Physical Computation*, cit., p. 279-281; P. Odifreddi and S. Barry Cooper, "Recursive Functions", *The Stanford Encyclopedia of Philosophy* (Fall 2012 Edition), Edward N. Zalta (ed.), URL = <http://plato.stanford.edu/archives/fall2012/entries/recursive-functions/>; P. G. Odifreddi, *Classical Recursion Theory*, vol. I-II, Amsterdam, North Holland, 1989-1999.

calcolabili ad altre funzioni calcolabili. Sono due, per la classe delle funzioni ricorsive primitive: la composizione e la ricorsione.

La composizione⁷⁵ è la generalizzazione a un numero arbitrario di argomenti dell'operazione di composizione di funzioni a un argomento per cui date due funzioni $\varphi: A \rightarrow B$ e $\psi: B \rightarrow C$, tali che il codominio della prima è uguale al dominio della seconda, si può costituire una funzione $\chi: A \rightarrow C$, che sarà detta funzione composta di φ e ψ , posto che per ogni $a \in A$: $\chi(a) = \psi(\varphi(a))$. Dato $a \in A$ si definisce prima $\varphi(a)$ di B , poi si calcola il valore di ψ con argomento $\varphi(a)$, ovvero $\psi(\varphi(a))$, e si trova l'elemento di C che è il valore di χ per l'argomento a . Applicando tale schema a un numero arbitrario di argomenti, procederemo in questo modo: date la funzione $\chi: N^k \rightarrow N$ e le k funzioni $\psi_1, \psi_2, \dots, \psi_k$ di tipo $N^n \rightarrow N$, si dice che la funzione $\varphi: N^n \rightarrow N$ tale che:

$$\varphi(x_1, \dots, x_n) = \chi(\psi_1(x_1, \dots, x_n), \psi_2(x_1, \dots, x_n), \dots, \psi_k(x_1, \dots, x_n))$$

è ottenuta per composizione dalle funzioni precedenti.

La seconda operazione, la ricorsione⁷⁶, è un tipo di induzione matematica: per calcolare il valore di una funzione per un argomento n si utilizza il valore per l'argomento $n - 1$, e così a ritroso. Date le funzioni $\chi: N^n \rightarrow N$ e $\psi: N^{n+2} \rightarrow N$, diciamo che la funzione $\varphi: N^{n+1} \rightarrow N$ tale che:

$$\begin{cases} \varphi(x_1, \dots, x_n, 0) = \chi(x_1, \dots, x_n) \\ \varphi(x_1, \dots, x_n, s(y)) = \psi(x_1, \dots, x_n, y, \varphi(x_1, \dots, x_n, y)) \end{cases}$$

è ottenuta dalle precedenti due funzioni mediante la tecnica ricorsiva.

Si noti che che la definizione induttiva non ha nulla in comune con il cosiddetto metodo assiomatico. Mentre quest'ultimo muove da assiomi (le formule di partenza delle dimostrazioni formali, assunte senza dimostrazione) verso i teoremi (tutte le formule deducibili dagli assiomi mediante regole di inferenza), la definizione induttiva non è affatto un sistema deduttivo, bensì una strategia fondata sulla natura dei numeri e sulle nostre intuizioni fondamentali a proposito dei numeri.

Ora, si dice che le operazioni di composizione e di ricorsione "conservano" la calcolabilità

⁷⁵ M. Frixione, D. Palladino, *Funzioni, macchine, algoritmi*, cit., p. 63, 172-173

⁷⁶ Ivi, p. 169-172.

naturale, immediata delle funzioni iniziali. Ma che cos'è questa "calcolabilità immediata"? E che cosa s'intende con "conservazione della calcolabilità"? Nella definizione induttiva quel che viene conservato è l'insieme delle funzioni base e delle iterazioni elementari che esse veicolano. Una funzione descrive un numero o un rapporto tra numeri. È una forma di descrizione: riflette uno stato del sistema numerico. Ma allora che cosa descrivono le funzioni base? Le funzioni zero, successore e proiezione hanno una forma iterativa. Descrivono alcuni tipi di iterazione molto semplici che ricorrono nell'uso dei numeri: so per principio che qualsiasi argomento associo a zero avrò sempre come risultato lo zero; so per principio che il successore di 1 è 2, e così via. L'induzione matematica presuppone il principio di iterazione. Le funzioni base sono descrizioni molto semplici e tracciano un confine: descrivono U^B e il rapporto di immagine che sussiste tra i numeri e U^B . *La conservazione della calcolabilità è la conservazione della base iterativa, del loop delle funzioni base.* Una funzione è calcolabile se e solo se essa descrive il rapporto tra un certo numero e le funzioni base, ovvero tra un certo numero e U^B . La funzione calcolabile è una descrizione che contiene un'altra descrizione più elementare. Ma non è detto che per ogni numero sia possibile dare una descrizione di questo tipo.

Se questo è vero, avremo due sensi di computabilità. Il primo dice che tutti i numeri sono computabili in quanto schemi iterativi basati sul principio di iterazione e sul principio di contraddizione continua. In questo senso più generale, computabilità è sinonimo di iterabilità. Il numero è un oggetto computabile perché è perfettamente iterabile. Il secondo senso è più ristretto: se tutti i numeri sono computabili perché iterabili, non tutte le funzioni riescono a esibire la natura iterabile del numero, cioè il rapporto tra i numeri e U^B . Le funzioni che esibiscono (descrivono) l'iterabilità dei numeri sono funzioni computabili e la definizione induttiva è il metodo che usiamo per dimostrare che esse sono così. Le funzioni che non esibiscono l'iterabilità non sono computabili. Diciamo *esibire* per indicare il dato fenomenologico, non quello aritmetico. Al matematico non interessa l'iterabilità del numero, ma il risultato del calcolo. Al fenomenologo interessa quel che nell'enunciato matematico *appare*, non tanto quel che viene detto. Il fenomenologo si chiede: che

cosa mostra l'enunciato? Non che cosa dice, ammesso che voglia dire qualcosa, bensì che cosa *mostra* con i suoi simboli. Il fenomenologo assume una posizione *esterna* a qualsiasi simbolo.

In conclusione, la matematica è un insieme di proposizioni sui numeri che sono immagini di U^B . La nostra tesi (2) vuole proporsi come un corollario logico-fenomenologico alla tesi di Church-Turing (1). Formuleremo entrambe in questi termini:

(1) Una funzione è effettivamente calcolabile se e solo se è Turing-calcolabile, ossia ricorsiva generale⁷⁷.

(2) Una funzione è ricorsiva generale se e solo se mostra il rapporto tra un numero e l'ambiente logico U^B . Un numero (n) computabile è un numero cui corrisponde una funzione in grado di mostrare il rapporto tra questo numero e U^B .

Luca M. Possati

⁷⁷ Un importante risultato dimostrato da Turing in diversi contributi è che le funzioni ricorsive generali sono computabili da una macchina di Turing. Church ha poi allargato questo risultato fino alla formulazione della tesi per cui una funzione è effettivamente calcolabile se e solo se è ricorsiva generale. Si identifica quindi il concetto di funzione calcolabile mediante un algoritmo con il concetto rigoroso di funzione ricorsiva generale. Ogni procedimento algoritmico è ricorsivo. Cfr. A. M. Turing, "On Computable Numbers, with an Application to the *Entscheidungsproblem*", *Proceedings of the London Mathematical Society*, 42, 1936/37, p. 230-265; A. Church, "An Unsolvability Problem of Elementary Number Theory", *American Journal of Mathematics*, 58, 1936, p. 345-363; G. Boolos, J. P. Burgess, R. C. Jeffrey, *Computability and Logic*, Cambridge, Cambridge University Press, 2007 (I ed. 1974), p. 71. Resta da sottolineare la problematicità della tesi di Church. Certo, a sostegno di Church va il fatto indiscutibile che tutti i tentativi elaborati per caratterizzare in modo rigoroso le funzioni effettivamente calcolabili si sono rivelati equivalenti. "Tale indipendenza dal sistema formale prescelto è ovviamente un forte elemento a favore della tesi di Church. Si può inoltre dimostrare che le funzioni ricorsive generali sono tutte e sole le funzioni rappresentabili in un sistema formale del primo ordine che comprenda gli assiomi di Peano per l'aritmetica" (M. Frixione, D. Palladino, *Funzioni, macchine, algoritmi*, cit., p. 224). Non è detto però che un giorno non si possa scoprire un procedimento algoritmico, ovvero una funzione calcolabile, che non sia formulabile in maniera ricorsiva. E questo sempre per il fatto "a monte" per cui algoritmo e funzione restano due concetti empirici, intuitivi, di cui si possono elencare soltanto alcune caratteristiche generali. Di qui la necessità di un approfondimento fenomenologico di queste nozioni, che possa dare alcuni parametri in più per capire di che cosa stiamo parlando e valutare meglio la tesi di Church.