

Appunti di Logica Matematica

Silvio Valentini

versione provvisoria e incompleta
3 dicembre 2015

Indice

1	Introduzione	1
2	La nozione di verità	3
2.1	Un esempio matematico	3
2.2	Ancora un esempio matematico	4
2.3	C'era un monaco buddista	4
2.4	Un ultimo esempio matematico	5
2.5	Il gioco del HEX	6
2.6	Un esempio taoista	7
3	L'approccio classico e quello intuizionista	9
3.1	Un po' di proposizioni	10
3.1.1	Proposizioni atomiche	11
3.1.2	I connettivi	12
3.2	Facciamo parlare assieme un classico e un intuizionista	15
3.3	Un linguaggio formale per scrivere di matematica	16
3.4	I quantificatori	17
3.4.1	Il quantificatore universale	17
3.4.2	Il quantificatore esistenziale	18
3.5	La descrizione di una struttura matematica	19
3.5.1	Il problema della descrizione di una struttura	20
3.5.2	Le strutture descritte da un insieme di proposizioni	22
4	Regole per ragionare	25
4.1	Le regole logiche	25
4.1.1	Gli assiomi	25
4.1.2	Congiunzione	26
4.1.3	Disgiunzione	27
4.1.4	Trattare con $\perp$ e $\top$	28
4.1.5	Implicazione	28
4.1.6	Negazione	29
4.1.7	I quantificatori	30
4.2	Regole per la proposizione Id	32
5	La struttura della verità	33
5.1	Reticoli	34
5.2	Algebre di Boole	36
5.2.1	Le tabelline di verità	37
5.3	Algebre di Heyting	38
5.4	Relazioni tra algebre di Boole e algebre di Heyting	40
5.5	Algebre di Boole e di Heyting complete	41

6	Formalizzazione delle regole	43
6.1	Semantica e sintassi	43
6.1.1	Interpretazione delle proposizioni	44
6.2	Le regole	46
6.2.1	I connettivi	47
6.2.2	I quantificatori	52
7	Il teorema di completezza	61
7.1	Introduzione	61
7.2	Il <i>quasi</i> teorema di completezza	61
7.3	Valutazione in uno spazio topologico	64
7.4	Un teorema alla Rasiowa-Sikorski	65
7.5	La completezza del calcolo intuizionista	68
7.6	La completezza del calcolo classico	70
7.6.1	Ridursi solo al vero e al falso	72
8	Il teorema di Compattezza	75
8.1	Soddisfacibilità e compattezza	75
8.2	Alcune conseguenze del teorema di compattezza	75
8.2.1	Inesprimibilità al primo ordine della finitezza	76
8.2.2	Non caratterizzabilità della struttura dei numeri naturali	77
8.2.3	Gli infinitesi	77
8.3	La dimostrazione del teorema di compattezza	78
8.3.1	La definizione dell'ultraprodotto	78
8.3.2	Dimostriamo il teorema di compattezza	81
8.4	Il teorema di completezza forte	82
9	Il teorema di Incompletezza	85
9.1	Alcuni esempi di <i>situazioni</i> di Gödel	86
9.2	Le funzioni ricorsive	88
9.3	Sottoinsiemi ricorsivi e ricorsivamente enumerabili	91
9.4	Teorie e sottoinsiemi dei numeri naturali	92
9.5	La dimostrazione del teorema di incompletezza	93
A	Richiami di teoria degli insiemi	95
A.1	Collezioni e insiemi	96
A.1.1	Il concetto di concetto	97
A.1.2	Insiemi	98
A.2	Relazioni e funzioni	104
A.2.1	Relazioni	104
A.2.2	Funzioni	108
A.3	Insiemi induttivi non numerici	111
A.3.1	Liste	111
A.3.2	Alberi	113
A.4	Cardinalità	113
A.4.1	Insiemi numerabili	114
A.4.2	Insiemi più che numerabili	115
B	Espressioni con arietà	117
B.1	Premessa	117
B.2	Il linguaggio e i suoi elementi	117
B.2.1	Applicazione	117
B.2.2	Astrazione	119
B.2.3	La sostituzione	119
B.2.4	Il processo di calcolo	122
B.3	Teoria dell'uguaglianza	123

B.3.1	Decidibilità dell'uguaglianza tra λ -espressioni	123
B.4	Esercizi risolti	126
C	Dittatori e Ultrafiltri	127
C.1	Il problema	127
C.2	Il Teorema di Arrow	128
C.3	Ultrafiltri e aggregazioni di preferenze	129
C.4	Teorema di Arrow e ultrafiltri	130
C.5	Dimostrazione del teorema di corrispondenza	131
C.6	Dimostrazione del teorema dell'ultrafiltro principale	131
D	Correzione degli esercizi	133
D.1	La verità matematica	133
D.2	L'approccio classico e quello intuizionista	133

Capitolo 1

Introduzione

Lo scopo generale di questo corso è quello di analizzare il concetto di verità matematica e di studiare quali strumenti matematici possono essere utilizzati per trattare in modo conveniente con tale nozione¹.

Inizieremo quindi analizzando alcuni esempi presi sia dall'esperienza matematica che dal linguaggio utilizzato tutti i giorni che mettono in evidenza alcune ambiguità nella nozione di cosa sia una proposizione e, di conseguenza, in quella di verità di una proposizione.

A partire da questa analisi introdurremo le proposizioni che intendiamo utilizzare in questo corso di logica, con una particolare attenzione all'approccio *classico* e a quello *intuizionista*, proponendo per esse le relative nozioni di verità.

Forniremo quindi delle prime regole, il *calcolo dei sequenti*, che ci permetteranno di ragionare senza doverci ogni volta riferire direttamente alla nozione di verità di una proposizione.

Definiremo quindi due strutture matematiche, rispettivamente le *algebre di Boole* e le *algebre di Heyting*, che tentano di formalizzare le nozioni di verità introdotte precedentemente.

Proporremo infine per ciascun approccio un nuovo sistema di calcolo, la *deduzione naturale*, che permette di trattare ad un livello puramente sintattico le nozioni di verità considerate e ne mostreremo la *correttezza*, vale a dire la validità del sistema di regole che devono affermare la verità solamente per proposizioni che sono in effetti vere, e la *sufficienza*, cioè la completezza del sistema di regole che deve essere in grado di affermare la verità di tutte le proposizioni vere che si possono esprimere nel linguaggio che consideriamo.

Studieremo poi i limiti espressivi e dimostrativi dei sistemi che abbiamo proposto. In particolare vedremo che il linguaggio per il quale siamo stati in grado di dimostrare correttezza e sufficienza soffre di forti limiti espressivi e non ci lascia esprimere che una piccola parte delle proposizioni di interesse matematico. Per di più i sistemi dimostrativi a nostra disposizione, cioè calcolo dei sequenti e deduzione naturale, si dimostreranno insufficienti a catturare la nozione di proposizione vera in una struttura che intenda parlare (almeno) dei numeri naturali.

Ringraziamenti

Voglio ringraziare Massimo Corezzola e Susanno Calearo che mi hanno permesso di utilizzare per la preparazione di queste note i loro appunti del corso di Logica Matematica che ho tenuto rispettivamente nell'anno accademico 2005/06 e 2010/11.

¹Come esplicitamente spiegato nel frontespizio, questa versione degli appunti del corso di Logica Matematica si guarda bene dall'essere definitiva; si tratta piuttosto di alcune note ad uso degli studenti che per il momento non sono ancora complete e possono contenere imprecisioni. Qualunque osservazione che possa portare ad un risultato migliore è benvenuta.

Capitolo 2

La nozione di verità

Anche se il nostro scopo è quello di analizzare la nozione di verità per le proposizioni matematiche in questo capitolo non ci limiteremo a queste ma analizzeremo anche alcune frasi che si possono incontrare nel parlato di tutti i giorni.

2.1 Un esempio matematico

Iniziamo comunque con un esempio di carattere matematico:

Dimostrare che esistono due numeri irrazionali tali che elevando il primo al secondo si ottiene un numero razionale

Si può ottenere una dimostrazione molto breve di questa affermazione ragionando nel modo seguente. Si consideri $(\sqrt{2})^{\sqrt{2}}$. Allora abbiamo uno dei due seguenti casi: o $(\sqrt{2})^{\sqrt{2}}$ è razionale, e in questo caso abbiamo finito visto che abbiamo trovato due numeri irrazionali che elevati uno al secondo danno un numero razionale, oppure $(\sqrt{2})^{\sqrt{2}}$ è un numero irrazionale; ma allora elevando $(\sqrt{2})^{\sqrt{2}}$ a $\sqrt{2}$, i.e. un irrazionale elevato a un irrazionale, si ottiene $((\sqrt{2})^{\sqrt{2}})^{\sqrt{2}} = (\sqrt{2})^{\sqrt{2}\sqrt{2}} = (\sqrt{2})^2 = 2$ cioè un numero razionale. Quindi in ogni caso abbiamo ottenuto il risultato desiderato.

Ora è chiaro che (quasi) ogni matematico sarebbe soddisfatto con questa dimostrazione; tuttavia proviamo ad analizzare quel che abbiamo dimostrato per vedere se abbiamo davvero risposto alla domanda che ci era stata fatta, o piuttosto, per provare a capire meglio se la domanda iniziale non racchiudeva una certa dose di ambiguità. Infatti a voler essere precisi con la dimostrazione precedente non abbiamo *esibito* due numeri irrazionali tali che elevando il primo al secondo si ottiene un numero irrazionale; ciò che abbiamo in realtà dimostrato è il fatto che non è possibile che tali numeri non esistano.

Quindi se dimostrare l'esistenza di qualcosa significa esibirne un esempio in realtà non abbiamo davvero risolto l'esercizio proposto mentre se dimostrare l'esistenza significa, più debolmente, far vedere che non può andare sempre male allora la nostra dimostrazione raggiunge il suo scopo (quel che è importante è naturalmente non fare confusione tra le due nozioni di esistenza, come purtroppo accade talvolta).

Si può naturalmente accontentarsi di utilizzare sempre la più debole tra le due nozioni, e rinunciare quindi fin dall'inizio alla speranza di poter utilizzare l'informazione nascosta dentro una affermazione di esistenza di un particolare elemento, ma non si può pretendere poi di poter dire qualcosa sull'elemento che non siamo mai stati in grado di esibire.

Esercizio 2.1.1 *Esibire due numeri irrazionali tali che il primo elevato al secondo sia un numero razionale.*

Esercizio 2.1.2 *(Difficile) Determinare se $(\sqrt{2})^{\sqrt{2}}$ è razionale oppure no rendendo informativa la dimostrazione precedente.*

2.2 Ancora un esempio matematico

Analizziamo ora la seguente proprietà:

*Per ogni successione di numeri naturali $(n_k)_{k \in \omega}$,
esistono due numeri naturali i e j tali che $i < j$ e $n_i \leq n_j$*

La prova più veloce di questa proposizione è forse la seguente. Una successione altro non è che una funzione il cui dominio sono i numeri naturali; nel nostro caso quindi si tratta di una funzione f dai numeri naturali verso i numeri naturali. Si consideri allora l'insieme $\{f(x) \in \omega \mid x \in \omega\}$ dei numeri naturali che costituiscono l'immagine di tale funzione. Per il principio del minimo, tale insieme, essendo non vuoto, ha un elemento minimo $f(i^*)$ in corrispondenza a un qualche numero naturale i^* . La soluzione richiesta è allora semplicemente di porre $i \equiv i^*$ e $j \equiv i^* + 1$.

È chiaro che questa dimostrazione è veloce ed elegante ma lascia un po' con la curiosità di sapere quali siano i valori dove la funzione smette di decrescere. Di nuovo non abbiamo davvero esibito tali valori ma solo mostrato che non può andare sempre male, vale a dire non può succedere che si continui sempre a decrescere, ma non abbiamo dato nessuna indicazione di quali siano tali valori, non li abbiamo cioè saputi esibire.

In realtà in questo caso non è difficile ideare una procedura più informativa per esibire i valori cercati. Una prima idea in questa direzione può essere la seguente. Si consideri il primo valore n_0 nella sequenza, allora o un valore seguente è maggiore di n_0 , e in questo caso abbiamo risolto il nostro problema, o tutti i valori seguenti sono minori di n_0 . Ma in questo secondo caso, per il principio dei cassetti¹, ci devono essere un numero infinito di valori uguali nella sequenza e quindi anche in questo caso avremo una soluzione del nostro problema. Tuttavia anche in questo caso non si tratta ancora di una soluzione sufficientemente informativa (anche se è comunque di qualcosa di più informativo che nel primo caso) e possiamo fare di meglio, dare cioè un metodo che risolve davvero il problema proposto, permettendoci cioè di esibire in un tempo finito, e noto a priori, gli indici richiesti.

Possiamo infatti considerare nuovamente il primo valore n_0 della successione. Abbiamo allora due casi da considerare: o $n_0 \leq n_1$, e allora siamo a posto, o $n_1 < n_0$. In questo secondo caso possiamo proseguire analizzando n_2 e di nuovo siamo nella situazione che o $n_1 \leq n_2$, risolvendo il problema, o $n_2 < n_1 < n_0$. Se continuiamo in questo modo o arriviamo ad un passo in cui $n_k \leq n_{k+1}$ risolvendo il problema o costruiamo una successione finita, che può però essere al massimo lunga n_0 passi, tale che $n_{k+1} < n_k < \dots < n_1 < n_0$ e quindi al massimo in n_0 passi possiamo trovare gli indici richiesti.

Esercizio 2.2.1 *Sia n un qualsiasi numero naturale e $l_0, l_1, l_2 \dots$ sia una successione di n -ple di numeri naturali, i.e. liste finite $(m_1, \dots, m_n)$ con $m_1, \dots, m_n \in \omega$. Dimostrare che esistono due indici i e j tali che $i < j$ e, per ogni $1 \leq k \leq n$, $l_{i_k} \leq l_{j_k}$. (sugg.: abbiamo già visto come fare la dimostrazione quando $n = 1$, possiamo allora continuare con una prova per induzione sulla dimensione n delle n -ple considerate; si noti che si può costruire la prova richiesta utilizzando il principio dei cassetti)*

2.3 C'era un monaco buddista ...

Possiamo introdurre il prossimo esempio con una breve storiella:

C'era un monaco buddista che parte un mattino alle 8:00 dal suo monastero per andare a pregare nel piccolo tempio in cima ad un'alta montagna servito da un solo sentiero praticabile. Vi arriva alle 15:00 e dopo essersi riposato un po' inizia le sue preghiere fino a sera. Dopo un pasto frugale si mette a dormire per ripartire il giorno seguente, sempre alle 8:00 per il viaggio di ritorno al monastero dove giunge alle 13:00, giusto in tempo per il pranzo. Dimostrare che esiste un punto nel sentiero dove il monaco è passato entrambi i giorni alla stessa ora.

Siamo di nuovo di fronte alla richiesta di dimostrare l'esistenza di qualcosa e quindi di nuovo in presenza della solita ambiguità. Tuttavia, a differenza che negli esempi che abbiamo visto prima,

¹Si tratta del principio che afferma che se ho infiniti oggetti e un numero finito di cassetti in qualche cassetto dovrò mettere infiniti oggetti; esso generalizza la comune esperienza che afferma che se ho k cassetti e un numero di oggetti maggiore di k allora in qualche cassetto dovrò mettere almeno due oggetti.

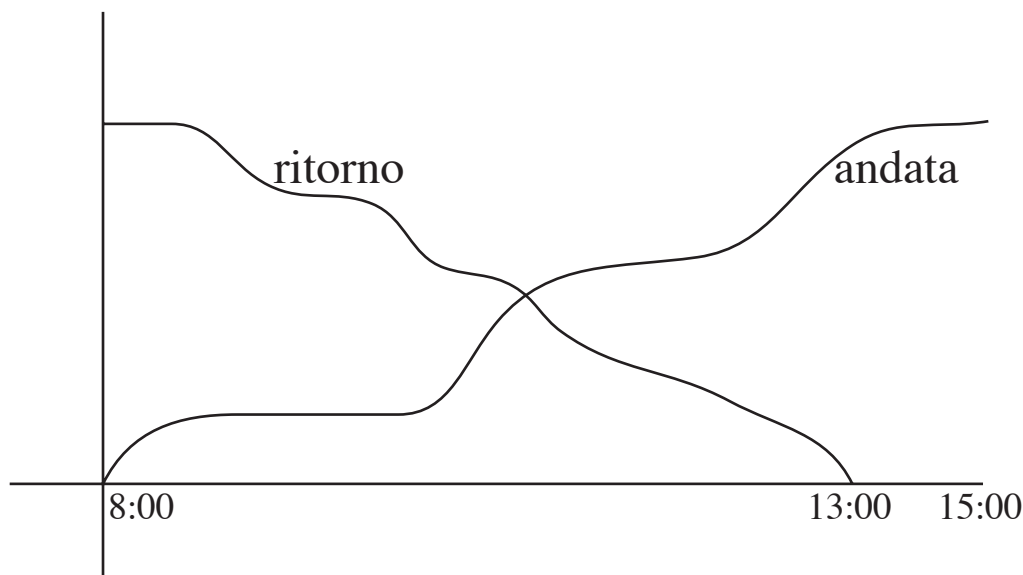


Figura 2.1: i due percorsi sovrapposti

questa volta la situazione è diversa: abbiamo solo la possibilità di fornire la prova che non è possibile che non ci sia il punto richiesto!

La soluzione si *vede* immediatamente facendo semplicemente sovrapporre i due grafici, fatti nei due giorni successivi, che rappresentano la posizione del monaco sul sentiero che sale la montagna e la posizione del monaco che scende dal tempio (si veda la figura 2.1), o se preferite potete immaginare che il secondo giorno il fratello gemello del monaco salga a sua volta al tempio sulla cima della montagna ed è allora chiaro che da qualche parte si incontreranno.

Tuttavia non c'è alcun modo per esibire tale punto visto che non abbiamo nessuna informazione sulla funzione che lega il tempo e la strada percorsa dal monaco.

La situazione è simile a quella del teorema che sostiene che ogni funzione continua f tale che $f(0) = 0$ e $f(1) = 1$ incontra ogni funzione continua g tale che $g(0) = 1$ e $g(1) = 0$, vale a dire che se tracciamo due linee senza staccare la penna dei vertici opposti di un quadrato queste si incontrano. Si tratta di un fatto della cui verità è difficile dubitare, ma assolutamente non costruttivo, cioè non possiamo dire in alcun modo, in mancanza di altre informazioni, dove le due funzioni si incontrano: si tratta di una esistenza che può essere solamente debole.

Esercizio 2.3.1 *Si confronti la situazione sopra descritta con il teorema di Bolzano che afferma che una funzione continua dai numeri reali ai numeri reali che agli estremi di un intervallo chiuso del suo dominio assume valori di segno opposto si annulla all'interno di detto intervallo [Bolz1817].*

2.4 Un ultimo esempio matematico

Un altro caso tipico, e molto importante viste le sue conseguenze in analisi matematica, è quello del teorema di Bolzano-Weierstrass che afferma che ogni successione limitata di numeri reali ha un punto di accumulazione. La dimostrazione di questo teorema si ottiene in genere nel modo seguente. Sia $[a, b]$ l'intervallo che contiene tutti gli infiniti punti della successione. Se lo dividiamo a metà accadrà che almeno una delle due metà contiene infiniti valori della successione (di nuovo una applicazione del principio dei cassetti). Consideriamo quindi la metà di questo nuovo intervallo; di nuovo una delle due metà conterrà un numero infinito di elementi della successione. Proseguendo in questo modo determineremo intervalli sempre più piccoli contenenti infiniti punti della successione e il punto di accumulazione sarà quello determinato dal limite di questa successione di intervalli l'uno incluso nel precedente.

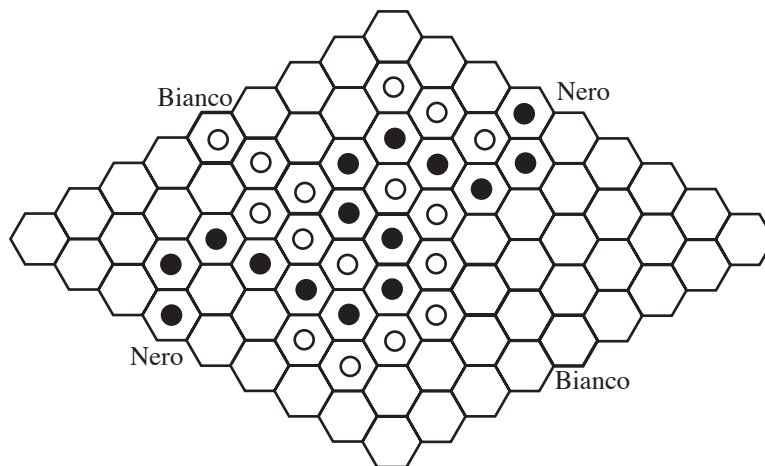


Figura 2.2: i due percorsi sovrapposti

Dovrebbe essere chiaro che si tratta di nuovo di una forma di esistenza che non ci aiuta a costruire il punto di accumulazione visto che in ogni momento del processo che lo determina non abbiamo informazioni di dove tale punto possa essere e che tutte le scelte dettate dal principio dei cassetti sono di fatto non costruttive².

2.5 Il gioco del HEX

Questa volta è un gioco che ci fornisce la possibilità di analizzare una volta di più la nozione di esistenza usata in matematica. Il gioco che ci interessa, noto come Hex [Gar59], si gioca tra due giocatori su una scacchiera costituita da un reticolo 9×9 di esagoni³ (vedi figura 2.2).

I giocatori, il Bianco e il Nero, fanno una mossa a testa posando una pedina del proprio colore su un esagono della scacchiera con l'intento di costruire una linea continua che unisca i due lati etichettati con il loro nome. Vince il giocatore che per primo riesce a unire i due lati con il proprio nome (nella figura 2.2 è rappresentato un percorso vincente per il Nero).

Dovrebbe essere abbastanza chiaro che quando uno dei due giocatori ha completato il suo percorso l'altro non ha più nessuna possibilità di completare il proprio visto che trova ogni cammino sbarrato dal primo giocatore. Basta inoltre pensare appena un po' di più per convincersi che non è possibile che una partita finisca in condizioni di parità, cioè che nessuno dei due giocatori sia in grado di completare il suo percorso, visto che una scacchiera piena di pedine contiene sicuramente un cammino completo per il bianco o per il nero (esercizio: fornire una dimostrazione di questo fatto all'apparenza così banale, ma fare attenzione che si tratta di un teorema equivalente al teorema della curva di Jordan).

Quel che interessa a noi è tuttavia il fatto che esiste una strategia vincente per il primo giocatore ma (fortunatamente!) nessuno è stato finora in grado di esibirne alcuna, e quindi il gioco rimane interessante. Per convincersi dell'esistenza di una strategia vincente per il primo giocatore basta considerare il fatto che, come abbiamo visto, il gioco non può finire in parità, e quindi deve esserci un comportamento vincente per il primo o per il secondo giocatore, e ragionare come segue. Supponiamo che ci sia una strategia vincente per il secondo giocatore; allora il primo giocatore può fare a caso la prima mossa e continuare poi, quando è di nuovo il suo turno di gioco, con la strategia vincente per il secondo giocatore; se ad un certo punto si trovasse nella situazione di dover muovere su una casella che ha già occupato in un momento precedente non dovrà fare altro che muovere nuovamente a caso. Visto che è chiaro che avere una pedina in più sul tavolo di gioco

²Quel che il principio dei cassetti ci assicura è che non è possibile che tutti i cassetti contengano un numero finito di oggetti, ma di nuovo questo non ci aiuta poi molto se quel che vogliamo è esibire un cassetto che contiene infiniti oggetti.

³La dimensione della scacchiera può variare, ma scacchiere troppo piccole rendono il gioco poco interessante mentre scacchiere troppo grandi lo rendono troppo lungo.

non può certo nuocere, ma al più favorire la vittoria, in tal modo il primo giocatore sarebbe sicuro di vincere. Quindi supporre l'esistenza di una strategia vincente per il secondo giocatore porta all'assurda situazione che sarà il primo giocatore a vincere per cui l'ipotesi deve essere sbagliata; ma allora deve esserci una strategia vincente per il primo giocatore.

Siamo quindi nuovamente in presenza di una dimostrazione di esistenza che non esibisce davvero la strategia, cioè quel che si dimostra è il fatto che non è possibile che non ci sia alcuna strategia. Tuttavia prima di cominciare a cercare una strategia vincente è bene essere avvertiti che trovarne una è equivalente a dare un metodo che permetta di esibire un punto fisso di una funzione continua del quadrato unitario in se stesso che sappiamo esistere (di nuovo in modo non costruttivo) per il teorema di punto fisso di Brouwer.

Per un'analisi di carattere matematico, compresa una prova dell'equivalenza tra l'esistenza di una strategia vincente per il primo giocatore e il teorema di punto fisso di Brouwer, si veda [Gal79].

2.6 Un esempio taoista

Per concludere questa prima parte della nostra analisi vediamo che ci possono essere problemi non solo con la comprensione della nozione di verità di una proposizione ma anche con la nozione stessa di cosa possa essere una proposizione.

Supponiamo, ad esempio, che qualcuno pronunci la seguente frase:

Ogni secondo il mondo smette di esistere per un secondo e poi ricomincia da dove aveva smesso

oppure la seguente che potremo interpretare come una forma di 'creazionismo estremo'

Dio ha creato il mondo in questo istante

Possiamo allora chiederci se queste siano frasi dotate di significato e in caso affermativo se siano vere o false. Dovrebbe in ogni caso essere chiaro che si tratta di frasi imbarazzanti. Infatti se la prima di esse venisse considerata da un qualche dio al di fuori del mondo (qualunque cosa questo possa significare) allora essa sarebbe dotata di un valore di verità, corrisponderebbe cioè ad uno stato di realtà il fatto che il mondo smetta e ricominci in modo alternato, e quindi tale dio potrebbe affermare che essa è vera oppure affermare che è falsa. D'altra parte per un qualsiasi essere vivente all'interno del mondo non ci sarebbe alcun modo per verificare quale sia lo stato delle cose, e quindi anche se dotata di valore di verità tale valore di verità sarebbe assolutamente inconoscibile. Ma la situazione sarebbe per altri versi ancora peggiore: visto che noi comunichiamo tramite un linguaggio anche il più formidabile mistico, che potrebbe avere un'intuizione dello stato delle cose, si troverebbe poi totalmente impossibilitato a comunicare in modo credibile i risultati della sua intuizione. Lo stato di verità espresso dalla frase è cioè non solo intrinsecamente inconoscibile ma non esiste neppure la possibilità di accordarsi con gli altri su quale tipo di osservazioni, fisiche o mentali, potrebbero essere sufficienti per accertarne la verità o la falsità⁴.

Per quanto riguarda il creazionismo estremo la situazione è molto simile. Nulla potrebbe infatti impedire ad un dio onnipotente di creare un mondo con tutti i segni di un passato, compresi i nostri ricordi, i libri di storia e tutto quel che può servire per pensare che il mondo esistesse già in passato. Potrebbe quindi essere vero oppure no che un dio ha creato il mondo in questo istante, ma sicuramente non ci sarebbe modo di accertarsene.

In questo senso si potrebbe essere tentati di escludere a priori questo tipo di frasi da uno sviluppo del processo matematico; noi non intendiamo farlo, ma ci sembra doveroso riconoscere che non solo esistono situazioni in cui non si può decidere sul contenuto di verità di una frase ma neppure decidere cosa sarebbe necessario conoscere per poter decidere su tale contenuto di verità.

⁴Notate che non stiamo chiedendoci se tali osservazioni sono verificate oppure no, ma solamente quali osservazioni potremmo mai operare.

Capitolo 3

L'approccio classico e quello intuizionista

Gli esempi della sezione precedente mostrano chiaramente che esistono (almeno) due approcci distinti alla nozione di verità di una proposizione, in particolare per quanto riguarda le proposizioni che hanno a che fare con l'esistenza di un qualche oggetto. Ma la discrepanza tra le possibili nozioni è ancora più profonda se consideriamo in particolare gli ultimi esempi dove è chiaro che nel primo caso siamo di fronte a frasi dotate di un valore di verità mentre nel secondo non succede la stessa cosa.

Volendo semplificare la situazione possiamo stabilire, prima ancora di cercare di capire come si può fare a riconoscere che qualcosa è vero, quali siano le frasi per le quali ha senso chiedersi se sono vere o false. Già a questo livello possiamo riconoscere due diversi atteggiamenti che, tanto per facilitare la loro individuazione, chiameremo l'atteggiamento *classico* o *platonico*, cioè la *matematica della realtà esterna*, e l'atteggiamento *intuizionista*, cioè la *matematica fatta dagli esseri umani*.

La domanda principale a cui vuole rispondere l'approccio classico è come riconoscere la verità delle affermazioni matematiche riguardanti una realtà matematica esterna e immutabile di cui i matematici, al pari dei fisici o dei biologi, devono strappare i segreti ma che ad ogni modo esisterebbe anche indipendentemente dai loro sforzi. Invece l'atteggiamento intuizionista è quello che l'edificio matematico è costruito dai matematici con il loro lavoro e quindi non c'è nulla di esterno da scoprire ma invece è tutto da inventare e quindi da comunicare.

A partire da queste considerazioni, che in realtà estremizzano fin troppo la situazione, diventa possibile proporre la seguente definizione.

Definizione 3.0.1 (Proposizione) *Una proposizione è:*

- (*Classico*) *Una qualsiasi affermazione dotata di un valore di verità,*
- (*Intuizionista*) *Una qualsiasi affermazione per la quale sia possibile accordarsi su cosa ci potrebbe convincere della sua verità.*

Alla luce di questa definizione è chiaro che le frasi che abbiamo utilizzato nel nostro ultimo esempio sono proposizioni per un classico (si tratta naturalmente di proposizioni sulla cui verità non si sa decidere ma che in qualche modo rappresentano uno stato di realtà) mentre non sono neppure proposizioni per un intuizionista e quindi tanto meno ha senso chiedersi se esse siano vere o false visto che non è neppure possibile stabilire con precisione e chiarezza cosa mai ci convincerebbe della loro verità¹.

¹È chiaro che la posizione classica ammette più proposizioni che la posizione intuizionista e che nel primo caso le proposizioni sono in qualche senso indipendenti dagli esseri umani che le devono poi utilizzare, mentre nel secondo caso c'è una qualche forma di arbitrio sull'insieme delle proposizioni e la matematica diventa in realtà un processo sociale dove si dà molto più peso alle difficoltà di comunicazione tra i matematici che potrebbero non riuscire ad accordarsi su cosa sia la maniera di convincersi della verità di una qualche affermazione.

Dopo aver stabilito cosa sia una proposizione possiamo cominciare a chiederci quando una proposizione sia vera e, come c'è da aspettarsi, la risposta sarà diversa per il classico e per l'intuizionista.

Definizione 3.0.2 (Proposizione vera) *Una proposizione è vera se ...*

- (Classico) ... *describe uno stato della realtà,*
- (Intuizionista) ... *sono in possesso di una istanza di quel qualcosa che può convincermi della sua verità su cui ci eravamo accordati nel momento in cui avevamo riconosciuto che una certa affermazione era una proposizione.*

Volendo di nuovo semplificare potremo dire che per un classico una proposizione è vera quando corrisponde allo stato della realtà, mentre l'onere di produrre una prova è solamente un (ingrato?) compito che spetta al matematico per poter spiegare agli altri perchè tale proposizione è vera, mentre per un intuizionista una proposizione è (diventa?) vera solamente nel momento in cui ne possiedo una prova². Visto che la parola prova ha un significato matematico già fissato, d'ora in avanti in questa sezione useremo la parola *ticket* per indicare, nel caso intuizionista, quel che ci convince della verità di una proposizione.

È utile far notare immediatamente che mentre per il classico le proposizioni si dividono automaticamente in vere e false, visto che si tratta di descrivere lo stato di una situazione reale dove le cose stanno o non stanno nel modo descritto dalla proposizione, nel caso dell'intuizionista abbiamo modo di poter esprimere il fatto che una proposizione è vera, dicendo che possediamo il ticket prescritto, ma non c'è modo alcuno per esprimere positivamente il fatto che una proposizione sia falsa: per farlo avremo bisogno degli *anti-ticket* ma non esiste alcun motivo per immaginare di essere in grado di produrre in ogni caso o un ticket o un anti-ticket (vedremo che è possibile esprimere il fatto che la negazione di una proposizione è vera, ma non bisogna confondere questo con il dire che la proposizione originale sia falsa).

Dovrebbe comunque essere chiaro che sia per il classico che per l'intuizionista ci sono affermazioni che, pur essendo state riconosciute come proposizioni, vere non sono; non bisogna cioè far confusione tra il momento in cui si riconosce che una affermazione è una proposizione e il momento in cui la si riconosce come vera. Ad esempio $3 + 2 = 5$ è una proposizione perchè (classico) è una affermazione dotata di un valore di verità o (intuizionista) perchè basta fare il conto per convincersi, ma viene riconosciuta come vera solo quando (classico) ho visto che rappresenta uno stato della realtà o (intuizionista) ho davvero fatto il conto e ho visto che il risultato del fare $3 + 2$ dà come risultato proprio 5. D'altra parte anche $3 + 2 = 6$ è una proposizione perchè (classico) è una affermazione dotata di un valore di verità o (intuizionista) perchè basta fare il conto per convincersi, ma non può venire riconosciuta come vera ne dal classico perchè non rappresenta uno stato della realtà ne dall'intuizionista perchè quando ha davvero fatto il conto ha visto che il risultato del fare $3 + 2$ non dà come risultato 6 (ammesso di essere in grado di distinguere il 5 dal 6!).

3.1 Un po' di proposizioni

Nonostante le differenze di fondo tra l'approccio classico e quello intuizionista sulle nozioni di cosa sia una proposizione e della sua verità, siamo abbastanza fortunati che nella usuale pratica matematica i due approcci spesso si sovrappongono e che si possa sviluppare in entrambi i casi una matematica ricca ed interessante, anche se non sempre coincidente, come abbiamo visto con gli esempi forniti nel capitolo precedente.

Quel che vogliamo fare in questo paragrafo è introdurre un insieme di proposizioni che sia abbastanza espressivo da permetterci di raccontare fatti matematicamente rilevanti e al tempo stesso tali che siano riconosciute come proposizioni sia da un classico che da un intuizionista, pur rimanendo ciascuno dei due all'interno della propria posizione epistemologica.

²Per il momento la nozione di prova sia per il caso classico che per quello intuizionista va intesa in modo informale, ma sarà appunto scopo di questo corso rendere (un po' più) esplicito questo concetto.

3.1.1 Proposizioni atomiche

Cominciamo quindi con le più elementari tra le proposizioni possibili, quelle cioè la cui analisi non possa essere condotta spezzandole in proposizioni più elementari (il nome di proposizioni atomiche viene da qui). Abbiamo in realtà appena incontrato una proposizione atomica, vale a dire $3+2=5$, e non dovrebbe essere quindi difficile capire di cosa si tratta. Ciò di cui stiamo parlando sono quelle proposizioni elementari che corrispondono direttamente ad una relazione matematica. L'esempio che abbiamo visto sopra corrisponde al fatto che la terna $\langle 3, 2, 5 \rangle$ appartiene alla funzione somma (si veda l'appendice A per i richiami sul linguaggio insiemistico). Altri esempi di proposizioni atomiche possono essere la proposizione $\text{Id}(x, y)$ che esprime la relazione identica, cioè la *diagonale* di un insieme, che risulta essere particolarmente interessante perché è una relazione che possiamo definire su un qualsiasi insieme, o la proposizione che esprime la relazione d'ordine tra i numeri naturali.

Naturalmente dal punto di vista classico è immediato verificare che si tratta di proposizioni, visto che si tratta di affermazioni dotate di un valore di verità, ed è pure facile capire quando la proposizione è vera, visto che basta vedere se la relazione descritta corrisponde ad uno stato di realtà. Ovviamente per essere sicuri di essere in grado di operare in questo modo bisogna pensare che la relazione in esame si estenda completamente di fronte a noi nella sua interezza, che potrebbe anche essere infinita, in modo che noi si possa *vedere* se l' n -pla richiesta appartenga oppure no alla relazione; quindi l'approccio classico richiede uno spazio infinito anche se la risposta alle nostre domande può avvenire in un tempo nullo giacché basta guardare il grafico della relazione interessata per saper rispondere (ammesso di riuscire a “vedere” tutto il grafico ...).

D'altra parte quelle considerate sono proposizioni anche da un punto di vista intuizionista perché sappiamo cosa conta come *ticket* della loro verità: basta fare il calcolo che ci porta a vedere se la relazione richiesta è verificata. Naturalmente questo metodo funziona solo per le relazioni per le quali esista la possibilità di fare un calcolo per decidere sulla loro verità (e quindi richiede di sviluppare una qualche teoria delle relazioni calcolabili, si veda il capitolo 9). Inoltre ci vuole del tempo per fare i calcoli; tuttavia non si ha più la necessità dello spazio infinito in cui tenere il grafico delle relazioni. In un certo senso nel passaggio tra l'approccio classico e quello intuizionista si ha uno scambio tra uno spazio infinito e un tempo illimitato; ci troveremo nuovamente di fronte a questa dicotomia quando analizzeremo il rapporto tra semantica e sintassi (vedi capitolo 7).

Conviene trattare a parte, e un po' più in dettaglio, il caso della proposizione di identità Id ; infatti la sua definizione non dipende dal particolare insieme su cui la vogliamo considerare e possiamo quindi proporre per lei dei ticket astratti che non richiedono di fare alcun calcolo. Il primo è il ticket $\text{refl}(t)$ che ci garantisce che un elemento t è uguale a se stesso, cioè che $\text{Id}(t, t)$ è una proposizione vera.

A fianco al ticket refl che ci permette di affermare qualcosa sulla proposizione Id (tramite tale ticket siamo in grado di asserire che la proposizione Id è riflessiva) abbiamo anche bisogno di dire come sia possibile per un intuizionista utilizzare l'informazione che due elementi t e s sono uguali, vale a dire vogliamo essere in grado di dedurre qualcosa dall'informazione che c è un ticket per la proposizione $\text{Id}(t, s)$. A questo scopo possiamo utilizzare il principio della *sostituibilità degli identici*: se t e s sono identici allora possiamo utilizzare s in tutti i casi in cui t andava bene. Per esprimere formalmente questa idea possiamo fare così: sia C una proposizione e supponiamo di aver un ticket d per la proposizione $C[x := t]$ ottenuta dalla proposizione C sostituendovi la variabile x con l'elemento t (per maggiori dettagli sull'operazione di sostituzione si veda l'appendice B). Allora indicheremo con $\text{sost}(c, d)$ il ticket per la proposizione $C[x := s]$ ottenuta sostituendo la variabile x in C con l'elemento s .

C'è un caso speciale di proposizioni atomiche che tornano molto utili nel parlare di matematica. Si tratta della proposizione sempre falsa $\perp$ e della proposizione sempre vera $\top$. Dal punto di vista classico è chiaro che esse non presentano problemi, ma anche dal punto di vista intuizionista possiamo riconoscere queste due affermazioni come proposizioni. La prima come la proposizione per la quale non c'è nulla che possa convincermi della sua verità, e quindi possiamo supporre di poter trasformare un qualsiasi ticket c per $\perp$ in un ticket $R_0(c)$ per una qualsiasi altra proposizione (tanto, se tutto va bene, non ci capiterà mai di avere un ticket per $\perp$). La seconda invece è

la proposizione per la quale c'è sicuramente qualcosa che può convincermi della sua verità (per semplificarci la vita potremo dire che $*$ è il *ticket* canonico per la proposizione $\top$).

Esercizio 3.1.1 *Dimostrare che ld soddisfa la proprietà simmetrica, cioè far vedere che c'è un ticket per $\text{ld}(s, t)$ nell'ipotesi di avere un ticket c per $\text{ld}(t, s)$.*

Esercizio 3.1.2 *Dimostrare che ld soddisfa la proprietà transitiva, cioè far vedere che c'è un ticket per $\text{ld}(t, u)$ nell'ipotesi di avere un ticket c_1 per $\text{ld}(t, s)$ e un ticket c_2 per la proposizione $\text{ld}(s, u)$.*

3.1.2 I connettivi

Introduciamo ora delle particelle linguistiche che servono per connettere, cioè mettere assieme, delle proposizioni al fine di ottenerne di più complesse.

La congiunzione

Il primo modo che può venire in mente per combinare due proposizioni è forse fare la loro congiunzione: se A e B sono le proposizioni che stiamo analizzando allora indicheremo la loro congiunzione con $A \wedge B$. Chiariamo subito che si tratta della proposizione che è vera esattamente quando sono vere sia la proposizione A che la proposizione B , ma come facciamo ad essere sicuri che si tratta di una proposizione?

Come spesso succede, dal punto di vista classico non ci sono problemi visto che dobbiamo solo chiederci se essa ammette valore di verità e la risposta è affermativa visto che abbiamo già detto che $A \wedge B$ è vera esattamente quando lo sono sia A che B ed essendo A e B proposizioni esse ammettono sicuramente valore di verità.

Se supponiamo ora di avere di fronte a noi vari *mondi matematici* per i quali le proposizioni A e B sono significative (si veda il capitolo 6 per una formalizzazione di questa idea) e scriviamo $\mathcal{M} \models C$ per dire che la proposizione C vale nel mondo $\mathcal{M}$ allora la proposizione $A \wedge B$ vale esattamente nell'insieme dei mondi che stanno nell'intersezione tra l'insieme dei mondi dove vale A e quelli dove vale B visto che

$$\mathcal{M} \models A \wedge B \text{ se e solo se } \mathcal{M} \models A \text{ e } \mathcal{M} \models B$$

Se indichiamo quindi con $C^\models$ l'insieme $\{\mathcal{M} \in U \mid \mathcal{M} \models C\}$ dei mondi matematici in cui C è vera abbiamo che $(A \wedge B)^\models \equiv A^\models \cap B^\models$.

Dal punto di vista intuizionista la situazione è un po' diversa e ci richiede di essere un poco più dettagliati visto che dobbiamo anche dire cosa conta come ticket per $A \wedge B$. Ora il ticket per $A \wedge B$ deve garantirci la verità di $A \wedge B$ quando sono vere sia A che B , cioè quando abbiamo un ticket sia per A che per B ; quindi il modo migliore per costruire il ticket desiderato consiste semplicemente nel mettere assieme il ticket per A con il ticket per B . Possiamo dire la cosa in termini leggermente diversi dicendo che se confondiamo una proposizione con l'insieme dei suoi ticket allora $A \wedge B$ è il prodotto cartesiano tra A e B ; quindi se scriviamo $a : A$ per dire che a è un ticket per A e $b : B$ per dire che b è un ticket per B allora avremo $\langle a, b \rangle : A \wedge B$ che dice che la coppia $\langle a, b \rangle$ è un ticket per $A \wedge B$.

Supponiamo di avere ora un generico ticket $c : A \wedge B$; che informazione possiamo trovare in c e come possiamo utilizzarla? Visto che c è un ticket per una congiunzione esso contiene l'informazione necessaria per estrarre un ticket per A ed un ticket per B e potremo quindi utilizzare c tutte le volte che avremo bisogno di un ticket per A e di un ticket per B . Visto che abbiamo già suggerito di pensare ad $A \wedge B$ come al prodotto cartesiano tra A e B viene naturale proporre di utilizzare come metodo per estrarre l'informazione contenuta in c le proiezioni che ogni prodotto cartesiano porta con sé. Se abbiamo quindi un ticket $c : A \wedge B$ indicheremo con $\pi_1(c)$ il ticket per A che c racchiude e con $\pi_2(c)$ quello per B . Ovviamente ci aspettiamo che il contenuto di informazione di $\pi_1(\langle a, b \rangle)$ e quello di a coincidano come pure quello di $\pi_2(\langle a, b \rangle)$ e quello di b .

Esiste tuttavia anche un altro modo per estrarre l'informazione contenuta in c che a fronte di una maggiore complessità offre il vantaggio di una maggiore uniformità di comportamento rispetto ai ticket per gli altri connettivi. Per poterlo definire ci serve spiegare prima come possiamo indicare il ticket per una deduzione che dall'ipotesi della verità di una proposizione A ci permette di ottenere un ticket per una proposizione C . A questo scopo possiamo indicare con una variabile x un generico

ticket per A e ottenere quindi un ticket c per C , che in generale conterrà tale variabile x , analizzando la deduzione di C da A . Ad esempio, nel caso di una deduzione di A da A il ticket sarà $c \equiv x$ mentre nel caso di una deduzione di $A \wedge A$ a partire da A il ticket sarà $c \equiv \langle x, x \rangle$; infine, supponendo che x sia un generico ticket per la proposizione $\perp$ avremo $R_0(x)$ è un ticket per una qualsiasi proposizione C (altri più interessanti esempi potremo vederli più avanti dopo l'introduzione di altri connettivi). In un certo senso c è una legge che ci dice come ottenere un *vero* ticket per C quando ci viene dato un *vero* ticket a per A invece che una variabile x per un ticket: basterà infatti sostituire la variabile x in c con il ticket a (notazione $c[x := a]$, si veda l'appendice B). Possiamo quindi trasformare tale legge in una vera funzione utilizzando la *lambda notazione* che potete trovare nell'appendice B, cioè scrivendo $\lambda x.c$.

Torniamo allora al nostro problema. Supponiamo quindi di possedere un ticket $c : A \wedge B$ e una legge d , in generale contenente le variabili x e y , che dato un generico ticket x per A ed un generico ticket y per B ci fornisce un ticket per la proposizione C . Allora possiamo ottenere un *vero* ticket per C se combiniamo l'informazione c con l'informazione d utilizzando un metodo $\text{split}(c, \lambda x.\lambda y.d)$ in grado di spezzare c nelle sue componenti ed utilizzandole poi al posto di x e y dentro d . In particolare avremo che $\text{split}(\langle a, b \rangle, \lambda x.\lambda y.d)$ dovrà essere uguale a $d[x, y := a, b]$.

Esercizio 3.1.3 Si scriva il ticket per la proposizione $(3 + 2 = 5) \wedge (2 \times 2 = 4)$.

Esercizio 3.1.4 Si definisca il metodo $\pi_1(-)$ utilizzando il metodo $\text{split}(-, -)$.

Esercizio 3.1.5 Si definisca il metodo $\text{split}(-, -)$ utilizzando le due proiezioni.

La disgiunzione

Il secondo metodo che vogliamo considerare per comporre due proposizioni è la disgiunzione: se A e B sono due proposizioni allora intendiamo analizzare la proposizione $A \vee B$ che risulta vera se e solo se almeno una tra le proposizioni A e B è vera.

Che si tratti di una proposizione sia per il classico che per l'intuizionista può essere verificato facilmente. Infatti per quanto riguarda il classico $A \vee B$ ha sicuramente un valore di verità in quanto è vera se e solo se è vera A o è vera B . È facile capire che la proposizione $A \vee B$ vale su tutti i mondi matematici che stanno nell'unione dell'insieme dei mondi in cui vale A con i mondi in cui vale B visto che

$$\mathcal{M} \models A \vee B \text{ se e solo se } \mathcal{M} \models A \text{ o } \mathcal{M} \models B$$

e quindi $(A \vee B)^{\mathcal{F}} \equiv A^{\mathcal{F}} \cup B^{\mathcal{F}}$.

Per quanto riguarda invece l'intuizionista, abbiamo un ticket per $A \vee B$ quando abbiamo un ticket per A oppure quando abbiamo un ticket per B ma, per evitare di perdere dell'informazione che potrebbe tornarci utile in futuro quando cercheremo di scoprire quali sono le conseguenze di $A \vee B$, ci serve anche mantenere l'informazione di quale dei due ci fornisce il ticket per $A \vee B$; quindi l'insieme corrispondente alla proposizione $A \vee B$ è l'unione disgiunta degli insiemi corrispondenti alle proposizioni A e B ; quindi i suoi ticket saranno indicati con $i(a)$ se si tratta di un ticket che viene dal ticket $a : A$ oppure con $j(b)$ se si tratta di un ticket che viene dal ticket $b : B$.

Supponiamo di avere ora un generico ticket c per $A \vee B$; come possiamo utilizzarlo? Essendo c un ticket per $A \vee B$ esso contiene l'informazione del fatto che si tratta di un ticket che viene da un ticket per A , e in questo caso *contiene* un ticket per A , oppure da un ticket per B , e in questo caso esso *contiene* un ticket per B . Possiamo quindi utilizzare c ovunque possiamo ottenere quel che vogliamo sia utilizzando un metodo che trae vantaggio da un ticket per A che utilizzando un metodo che usa un ticket per B . Infatti, in tal caso possiamo prima di tutto analizzare c e, se vediamo che viene da un ticket per A , possiamo poi risolvere il nostro problema utilizzando il primo metodo mentre se vediamo che viene da un ticket per B possiamo utilizzare il secondo metodo. Vediamo quindi come possiamo formalizzare quel che abbiamo detto. Supponiamo di avere un ticket c per $A \vee B$ e due metodi d ed e in grado di fornire rispettivamente un ticket per C sia a partire da un generico ticket x per A che da un generico ticket y per B (in generale accadrà quindi che d conterrà la variabile x ed e conterrà la variabile y). Possiamo allora indicare il ticket per C che desideriamo scrivendo $\text{when}(c, \lambda x.d, \lambda y.e)$. Naturalmente ci aspettiamo che $\text{when}(i(a), \lambda x.d, \lambda y.e)$ sia uguale a $d[x := a]$ e che $\text{when}(j(b), \lambda x.d, \lambda y.e)$ sia uguale a $e[y := b]$.

Esercizio 3.1.6 Si scriva il ticket per la proposizione $(3 + 2 = 5) \vee (2 \times 2 = 5)$.

Esercizio 3.1.7 Si determini un metodo che dato un generico ticket per $A \wedge B$ fornisca un ticket per $A \vee B$.

Una proposizione solamente classica: la negazione

Una proposizione che nasce in modo naturale in ambito classico è la negazione. Se A è una proposizione indicheremo la sua negazione con $\neg A$. Si tratta della proposizione che è vera esattamente quando la proposizione A è falsa e quindi, per il classico, è chiaramente dotata di un valore di verità.

Naturalmente i mondi matematici in cui vale $\neg A$ sono esattamente quelli che stanno nel complemento (rispetto alla totalità considerata U di mondi matematici) di quelli in cui vale A visto che

$$\mathcal{M} \models \neg A \text{ se e solo se } \mathcal{M} \not\models A$$

Perciò $(\neg A)^{\mathcal{F}} \equiv U \setminus A^{\mathcal{F}}$.

Dal punto di vista intuizionista non si tratta invece di una proposizione perché non sappiamo cosa potrebbe essere un ticket che ci permetta di garantirne la verità visto che per l'intuizionista non esistono ticket che certificano la falsità di una proposizione.

Esercizio 3.1.8 Si dimostri che per ogni proposizione A e per ogni mondo matematico $\mathcal{M}$ abbiamo che se $\mathcal{M} \models A$ allora $\mathcal{M} \models \neg\neg A$.

Esercizio 3.1.9 Si dimostri che per ogni proposizione A e per ogni mondo matematico $\mathcal{M}$ abbiamo che $\mathcal{M} \models A \vee \neg A$.

Una proposizione solamente intuizionista: l'implicazione

D'altra parte esiste anche una proposizione che è significativa solamente dal punto di vista intuizionista mentre non ha molto senso per un classico. Si tratta dell'implicazione tra due proposizioni A e B che indicheremo con $A \rightarrow B$. Quel che vogliamo per poter affermare che $A \rightarrow B$ è vera è che quando A è vera sia vera pure B e quindi vogliamo essere sicuri che se abbiamo un ticket per A possiamo avere anche un ticket per B . A questo punto è chiaro che un ticket per $A \rightarrow B$ altro non è che una funzione che trasforma un generico ticket x per A in un ticket b per B ; quindi, dal punto di vista insiemistico, $A \rightarrow B$ altro non è che l'insieme delle funzioni da A verso B ; per denotare un suo elemento scriveremo $\lambda x.b$ se $b : B$ è la legge che fornisce il ticket per B per un generico $x : A$ (si veda la notazione introdotta nell'appendice B).

Visto che il ticket per una implicazione $A \rightarrow B$ è una funzione che trasforma ticket per A in ticket per B non è difficile immaginare cosa possiamo fare quando ci viene fornito un ticket c per $A \rightarrow B$: possiamo semplicemente applicare la funzione c ad un ticket $a : A$ al fine di ottenere un ticket $\text{ap}(c, a)$ per B . Ovviamente $\text{ap}(\lambda x.b, a)$ ha lo stesso contenuto di informazione di $b[x := a]$.

Esercizio 3.1.10 Data una proposizione A , trovare un ticket per $A \rightarrow A$.

Esercizio 3.1.11 Date le proposizioni A e B , trovare un ticket per la proposizione $A \rightarrow (B \rightarrow A)$.

Esercizio 3.1.12 Date le proposizioni A , B e C , trovare un ticket per la proposizione

$$(A \rightarrow B) \rightarrow ((B \rightarrow C) \rightarrow (A \rightarrow C))$$

Esercizio 3.1.13 Date le proposizioni A e B , trovare un ticket per la proposizione

$$(A \wedge (A \rightarrow B)) \rightarrow B$$

Esercizio 3.1.14 Date le proposizioni A , B and C dimostrare che c'è un ticket per la proposizione $(C \wedge A) \rightarrow B$ se e solo se c'è un ticket per la proposizione $C \rightarrow (A \rightarrow B)$.

Esercizio 3.1.15 Date le proposizioni A e B , trovare un ticket per la proposizione

$$(A \wedge B) \rightarrow (A \vee B)$$

3.2 Facciamo parlare assieme un classico e un intuizionista

L'implicazione è quindi per sua natura una operazione dinamica, di trasformazione, e questo spiega perché non esiste una operazione classica davvero equivalente: l'approccio classico predilige lo studio della situazioni per quel che sono e non quando si tratta di qualcosa che evolve nel tempo. Tuttavia, vista l'importanza dell'implicazione in matematica - non esiste in pratica teorema che non abbia la forma di una implicazione - deve esistere una controparte classica dell'implicazione.

Quel che ci serve è quindi di inventare una maniera per un classico e un intuizionista di comunicare anche se essi usano due lingue non completamente compatibili visto che una manca della negazione e l'altra dell'implicazione. Anche in questo caso, come succede in generale quando si traduce una lingua in un'altra, la traduzione non sarà perfetta e quindi potrà dare luogo a qualche incomprensione (ma si vedano gli esercizi 6.2.9 e 6.2.24). Tuttavia l'intuizionista guadagnerà una negazione, e quindi almeno in maniera approssimativa potrà parlare della falsità di una proposizione, e il classico avrà a disposizione un'implicazione, e sarà quindi in grado di dare un senso agli enunciati dei teoremi.

Come al solito, quando in una lingua un termine manca si può tentare di approssimarne il significato, al meglio possibile, utilizzando qualche giro di parole. Cominciamo quindi con l'esprimere intuizionisticamente la negazione. Il problema in questo caso è quello di capire quando la proposizione A è falsa, un concetto che manca nel bagaglio linguistico e concettuale dell'intuizionista. La sua migliore approssimazione è quella di dire che una proposizione è falsa quando assumere che sia vera conduce alla falsità; quindi possiamo porre

$$\neg A \equiv_i A \rightarrow \perp$$

Per quanto riguarda invece la definizione di una proposizione classica che approssimi l'implicazione $A \rightarrow B$ dobbiamo tenere conto della seguente considerazione: quel che si vuole è che $A \rightarrow B$ sia vera se tutte le volte che A è vera allora contemporaneamente anche B è vera e quindi l'unico caso in cui $A \rightarrow B$ è sicuramente falsa accade quando A è vera ma B è falsa. Questo ci porta a proporre la seguente definizione

$$A \rightarrow B \equiv_c \neg(A \wedge \neg B)$$

Questa definizione fa sì che l'implicazione classica non abbia più il significato dinamico che essa aveva per l'intuizionista quanto piuttosto un significato *normativo*, cioè, non è possibile che valga A se non vale anche B , come nella normativa di legge "Se sei ubriaco allora non puoi guidare" che diventa appunto falsa non appena si trova un ubriaco al volante.

Come con ogni forma di traduzione da una lingua ad un'altra non dobbiamo aspettarci che tutto fili liscio³. Infatti il rischio è di trovarsi di fronte a risultati inaspettati se ci si dimentica che la traduzione è solo una approssimazione di quel che si intendeva dire.

Ad esempio il classico potrà restare sorpreso che l'intuizionista non sia in grado di esibire alcun ticket per la proposizione $A \vee \neg A$ se dimentica che non esiste per l'intuizionista un concetto di falso e quindi che l'intuizionista per poter produrre un ticket per tale proposizione dovrebbe essere in grado di produrre un ticket per A o una funzione che permetta di trasformare ogni ticket per A in un ticket per $\perp$ e non c'è alcun motivo perché questo dovrebbe essere possibile per una qualsiasi proposizione A (si consideri ad esempio il caso di una qualsiasi congettura matematica non (ancora?) risolta).

Ma forse ancora più sorpreso potrebbe essere l'intuizionista se dimentica che quando il classico usa l'implicazione non intende dire come si trasforma un ticket, ma solamente esprimere un giudizio su uno stato di realtà. Ad esempio sembra sicuramente strano che un classico possa affermare che queste proposizioni sono vere

- (1) $((A_1 \wedge A_2) \rightarrow B) \rightarrow ((A_1 \rightarrow B) \vee (A_2 \rightarrow B))$
- (2) $(A_1 \rightarrow (A_1 \wedge A_2)) \vee (A_2 \rightarrow (A_1 \wedge A_2))$

se pensiamo all'implicazione come ad un modo di trasformare prove visto che la prima esprimerebbe il miracolo che se la proposizione B è conseguenza di due proposizioni allora in realtà essa sarebbe conseguenza di una solamente tra esse, mentre la seconda, conseguenza della prima, esprime il fatto che la verità di una congiunzione di due proposizioni sarebbe conseguenza di una sola tra esse!

³Tradurre è sempre un po' tradire ...

Esercizio 3.2.1 *Trovare un ticket per la proposizione $(A \wedge \neg A) \rightarrow \perp$*

Esercizio 3.2.2 *Dimostrare che dal punto di vista di un classico $A \rightarrow B$ è vero se e solo se $\neg A \vee B$.*

Esercizio 3.2.3 *Dimostrare che l'implicazione classica soddisfa le condizioni dimostrate sull'implicazione intuizionista negli esercizi 3.1.13 e 3.1.14, naturalmente sostituendo la condizione sull'esistenza di un ticket per una proposizione con la condizione che tale proposizione sia vera.*

Esercizio 3.2.4 *Dimostrare che dal punto di vista classico la negazione e la disgiunzione sono sufficienti per definire tutti i connettivi finora considerati, i.e. $\top$, $\perp$, $\wedge$ e $\rightarrow$.*

Esercizio 3.2.5 *Dimostrare che dal punto di vista classico la proposizione $A \vee \neg A$ è sempre vera. Se ne può costruire un ticket intuizionista?*

Esercizio 3.2.6 *Dimostrare che dal punto di vista classico la proposizione $\neg\neg A \rightarrow A$ è sempre vera. Se ne può costruire un ticket intuizionista?*

Esercizio 3.2.7 *Dimostrare che qualsiasi proposizione intuizionista che abbia un ticket è classicamente vera. Cosa si può dire del viceversa?*

3.3 Un linguaggio formale per scrivere di matematica ...

... e anche di molto altro per quel che conta. Per il momento tuttavia quel che ci interessa è vedere che c'è la possibilità di esprimere in un linguaggio completamente formale, e quindi facile da controllare, le proposizioni che abbiamo introdotto nel paragrafo precedente e, ancora più importante, quelle che introdurremo nei due paragrafi successivi. Questo fatto sarà per noi essenziale quando avremo la necessità di sviluppare qualche metodo che ci garantisca sulla corretta applicazione delle regole di deduzione (si veda il paragrafo 6.2)

Quel che ci serve è semplicemente utilizzare nel caso dei connettivi logici che abbiamo finora introdotto la teoria generale per la scrittura di espressioni che si può trovare nell'appendice B. Secondo tale teoria un linguaggio utile per lavorare con oggetti matematici si può costruire a partire da pochi costrutti (*variabili*, *costanti*, *applicazione* ed *astrazione*) a patto che variabili e costanti abbiano una *arietà* e l'applicazione possa avvenire solamente tra una funzione di arietà $\alpha \rightarrow \beta$ ed un argomento di arietà α per dare un risultato di arietà β .

Nella teoria delle espressioni possiamo a grandi linee distinguere tra espressioni di arietà 0, che non si applicano, e espressioni di arietà $\alpha \rightarrow \beta$ che si possono applicare. È naturale pensare che, così come le funzioni di arietà zero sono gli elementi dell'insieme su cui una funzione è definita, anche nel nostro caso le proposizioni debbano corrispondere alle espressioni di arietà 0 di un linguaggio adatto ad esprimere un linguaggio logico, che contenga cioè i connettivi come costanti.

Per costruire le proposizioni che abbiamo usato finora possiamo quindi utilizzare le seguenti costanti con arietà

$$\begin{array}{ll} \perp : 0 & \top : 0 \\ \wedge : 0 \rightarrow (0 \rightarrow 0) & \vee : 0 \rightarrow (0 \rightarrow 0) \\ \neg : 0 \rightarrow 0 & \rightarrow : 0 \rightarrow (0 \rightarrow 0) \end{array}$$

e le seguenti abbreviazioni (i.e. *zucchero sintattico*):

$$\begin{array}{lll} A \wedge B & \equiv & \wedge(A)(B) \\ \neg A & \equiv & \neg(A) \end{array} \quad \begin{array}{lll} A \vee B & \equiv & \vee(A)(B) \\ A \rightarrow B & \equiv & \rightarrow(A)(B) \end{array}$$

Naturalmente anche le proposizioni atomiche si possono scrivere come particolari espressioni con arietà a patto di introdurre le necessarie costanti con una opportuna arietà. Ad esempio, se vogliamo scrivere la proposizione atomica che afferma che la somma di 3 con 2 è minore di 7 possiamo scrivere, seguendo la usuale pratica matematica, $(3+2) < 7$ ma dobbiamo ricordarci che stiamo utilizzando la seguente definizione

$$(3+2) < 7 \equiv (< (+ (3)(2)) (7))$$

dove 3, 2 e 7 sono costanti di arietà 0, $<$ è una costante di arietà $0 \rightarrow (0 \rightarrow 0)$ e $+$ è una costante di arietà $0 \rightarrow (0 \rightarrow 0)$.

Esercizio 3.3.1 Scrivere come espressioni con arietà le seguenti proposizioni specificando l'arietà delle costanti che esse contengono.

1. $(A \wedge B) \rightarrow (A \vee B)$
2. $(x < 3) \rightarrow (2x < 6)$
3. ...

3.4 I quantificatori

I connettivi che abbiamo considerato finora, sebbene utili, non costituiscono sicuramente un linguaggio matematico sufficiente per esprimere fatti matematici anche semplici: come dire ad esempio che la somma tra due numeri naturali è commutativa?

Quel che ci manca è la possibilità di fare affermazioni esistenziali ed universali, che riguardano cioè l'esistenza di un particolare elemento o tutti gli elementi del dominio di cui vogliamo parlare. In matematica lo strumento che permette di risolvere questo problema di espressività è il ricorso alle *variabili* combinato con l'aggiunta di nuove possibilità di costruire proposizioni.

Quel che le variabili rendono possibile è la scrittura delle *funzioni proposizionali*. In realtà ne abbiamo già (quasi) incontrata una nell'ultimo esercizio, vale a dire $(x < 3) \rightarrow (2x < 6)$. Non si tratta di una funzione proposizionale ma più semplicemente di una proposizione che contiene una variabile. Tuttavia il passo da una espressione ad una funzione è immediato, basta infatti una astrazione. Anche nel nostro caso il passo tra una proposizione che contiene una variabile ed una funzione proposizionale è breve, basta una astrazione per permetterci di ottenere $\lambda x.(x < 3) \rightarrow (2x < 6)$. Ciò che viene denotato da questa nuova scrittura dovrebbe essere chiaro: si tratta di una funzione che ad ogni elemento del dominio associa una proposizione.

Notiamo che nel seguito, con l'intento di rendere più chiara la notazione, useremo varie notazioni per una proposizione, anche contenente variabili, e di conseguenza anche per le funzioni proposizionali. In particolare ricordiamo qui le seguenti

- A per una proposizione che può contenere delle variabili anche se esse non sono esplicitamente indicate
- $A(x)$ per una proposizione in cui vogliamo porre l'accento sulla variabile x ; vale la pena di notare che A è applicata ad x e non può quindi essere una espressione di arietà 0 e perciò non può essere una proposizione; useremo la convenzione che x non appaia libera in A
- $A[x]$ per una proposizione in cui vogliamo mettere l'accento sul fatto che in A può apparire la variabile x ; questa notazione è particolarmente comoda in quanto ci permette di scrivere $A[t]$ per dire $A[x := t]$

3.4.1 Il quantificatore universale

Una volta che abbiamo detto cosa è una funzione proposizionale $\lambda x.A$ possiamo da essa ottenere una nuova proposizione tramite il quantificatore universale. Come si fa sempre in matematica denoteremo tale proposizione con $\forall x.A$ anche se si tratta solo di zucchero sintattico per l'espressione $\forall(\lambda x.A)$, dove $\forall$ è una costante di arietà $(0 \rightarrow 0) \rightarrow 0$, capace quindi di trasformare una funzione proposizionale in una espressione di arietà 0.

Naturalmente, ben più importante del capire come fare a scriverla utilizzando il nostro linguaggio, è verificare che $\forall x.A$ soddisfi le condizioni della definizione 3.0.1.

Per il classico non ci sono problemi particolari. Il significato che si vuole veicolare tramite la scrittura $\forall x.A$ è che $\forall x.A$ è vera se e solo se applicando la funzione proposizionale $\lambda x.A$ ad un qualunque elemento a del dominio si ottiene come risultato una proposizione vera (vale la pena di ricordare che $(\lambda x.A)(a)$ si calcola in $A[x := a]$). Visto che il classico suppone di avere una visione globale del dominio egli non ha, in linea di principio, problemi particolari a *vedere* se $A[x := a]$ è vera su ogni elemento a del dominio, e in questo caso dirà che $\forall x.A$ è vera, altrimenti dirà che essa è falsa. Alla fin fine per il classico la proposizione $\forall x.A$ non è poi molto diversa da una congiunzione fatta tra tutte le proposizioni ottenute applicando la funzione proposizionale $\lambda x.A$ ai

vari elementi del dominio, anche se questi potrebbero essere infiniti. Avremo perciò che $\forall x.A$ sarà vera in ogni mondo matematico che sta nell'intersezione tra tutti i mondi matematici che rendono vera $A[x := a]$ al variare di a nel dominio D visto che

$$\mathcal{M} \models \forall x.A \quad \text{se e solo se} \quad \mathcal{M} \models A[x := a] \quad \text{per ogni } a \text{ del dominio } D$$

cioè $(\forall x.A)^{\mathcal{F}} \equiv \bigcap_{a \in D} A[x := a]^{\mathcal{F}}$.

L'intuizionista sa di non avere una visione globale del dominio, in particolare nel caso questo non sia finito (se il dominio è infinito tale pretesa va sicuramente oltre le capacità umane), e quindi non ha modo di seguire la strada del classico. Ma allora cosa potrebbe essere un ticket per l'espressione $\forall x.A$? Visto che si tratta di vedere se sono tutte vere le proposizioni ottenute applicando una funzione proposizionale a tutti gli elementi del dominio tutto quel che ci serve è una funzione $\lambda x.b$ che vada in *parallelo* alla funzione proposizionale $\lambda x.A$ e che applicata ad un certo elemento a del dominio fornisca un ticket $(\lambda x.b)(a)$, che sappiamo essere uguale a $b[x := a]$, per la proposizione $(\lambda x.A)(a)$, che coincide con $A[x := a]$, che si ottiene applicando la funzione proposizionale allo stesso elemento del dominio.

Naturalmente questo significa anche che un ticket f per $\forall x.A$ si può utilizzare per ottenere un ticket $\text{ap}(f, a)$ per $A[x := a]$ per ogni elemento a del dominio.

Esercizio 3.4.1 *Trovare un ticket per la proposizione $\forall x.(A(x) \rightarrow A(x))$.*

3.4.2 Il quantificatore esistenziale

Data una funzione proposizionale $\lambda x.A$ possiamo ottenere una proposizione anche in un secondo modo, cioè utilizzando il quantificatore esistenziale per ottenere l'espressione $\exists(\lambda x.A)$, dove $\exists$ è una costante di arietà $(0 \rightarrow 0) \rightarrow 0$, che seguendo l'usuale pratica matematica indicheremo con $\exists x.A$.

L'espressione $\exists x.A$ viene riconosciuta come una proposizione dal classico che la usa per dire che esiste un elemento del dominio per cui A vale. Infatti per il classico questa operazione è completamente significativa anche nel caso di un dominio infinito. Si tratta quindi di una operazione non molto diversa da una disgiunzione di tutte le proposizioni ottenute applicando la funzione proposizionale $\lambda x.A$ ai vari, anche infiniti, elementi del dominio D . Avremo perciò che $\exists x.A$ sarà vera in ogni mondo matematico che sta nell'unione tra tutti i mondi matematici che rendono vera $A[x := a]$ al variare di a nel dominio D visto che

$$\mathcal{M} \models \exists x.A \quad \text{se e solo se} \quad \mathcal{M} \models A[x := a] \quad \text{per qualche } a \text{ del dominio } D$$

cioè $(\exists x.A)^{\mathcal{F}} \equiv \bigcup_{a \in D} A[x := a]^{\mathcal{F}}$.

Di nuovo, per l'intuizionista le cose non sono così semplici. La domanda a cui rispondere per sapere cosa è un ticket per $\exists x.A$ è: cosa mi potrebbe convincere della verità di una affermazione esistenziale? Abbiamo visto nel capitolo precedente molti esempi di affermazioni esistenziali e abbiamo fatto notare come ci sia una differenza sostanziale tra la nozione di esistenza che consiste nell'esibire l'oggetto di cui si vuole dimostrare l'esistenza e il limitarsi a far vedere che non è possibile che tale oggetto non ci sia. Naturalmente le due cose coincidono dal punto di vista del classico che per decidere sull'esistenza di un oggetto deve solo guardare la realtà e se succede che non è vero che tutti gli oggetti che ha di fronte (anche una quantità infinita di oggetti!) non soddisfano la proprietà desiderata allora vuol dire che c'è davvero un oggetto che la soddisfa. Nel caso dell'intuizionista questa identità non può essere verificata e quindi per convincersi della verità di una affermazione esistenziale egli non ha altra scelta che richiedere di vedere un oggetto, il *testimone*, ed un ticket che garantisca che la proposizione è vera per tale oggetto. Perciò un ticket per una proposizione esistenziale $\exists x.A$ è una coppia $\langle a, \pi \rangle$ tale che a è un elemento del dominio e π è un ticket per $A[x := a]$.

Visto che un ticket c per $\exists x.A$ è (si riduce a) una coppia possiamo utilizzarlo tramite il costrutto $\text{split}(c, \lambda x.\lambda y.d)$ per ottenere un ticket per C se d è un ticket per C nell'ipotesi che x sia un generico elemento del dominio e y un ticket per A .

È curioso osservare che mentre per il classico il quantificatore universale corrisponde ad una generalizzazione della congiunzione, eventualmente infinita sebbene uniforme visto che si tratta di

istanze della stessa funzione proposizionale, nel caso dell'intuizionista i ticket per una proposizione quantificata universalmente sono delle funzioni, seppure con un codominio che dipende dall'argomento a cui vengono applicate, e quindi questa proposizione corrisponde ad una generalizzazione dell'implicazione.

D'altra parte, mentre per il classico il quantificatore esistenziale corrisponde ad una generalizzazione della disgiunzione per l'intuizionista le cose vanno in maniera molto diversa: i suoi ticket, in quanto coppie costituite da un elemento del dominio e un ticket per la proposizione ottenuta applicando la funzione proposizionale a tale elemento, corrispondono piuttosto a quelli di una strana congiunzione tra un insieme ed una proposizione che dipende dall'elemento scelto in tale insieme.

Esercizio 3.4.2 *Determinare un ticket per la proposizione $A[x := t] \rightarrow \exists x.A$*

Esercizio 3.4.3 *Dimostrare che per un classico $\perp$, $A \rightarrow B$ e $\forall x.A$ sono sufficienti per definire tutti gli altri connettivi e quantificatori.*

Esercizio 3.4.4 *Dimostrare che per un classico la seguente proposizione è sempre vera:*

$$\exists x.(A \rightarrow \forall x.A)$$

Si tratta del famoso principio classico del bevitore: in ogni bar non vuoto c'è un cliente tale che se beve lui tutti bevono; lo stesso principio si può raccontare anche in maniera più attuale: in ogni nazione c'è un elettore tale che se lui vota per il partito X allora tutti votano per il partito X (e quindi è inutile spendere soldi per una campagna elettorale a pioggia e basta concentrare gli sforzi su tale elettore!).

Cosa ne pensa un intuizionista di questo principio? Come fare per dimostrare che non può valere intuizionisticamente?

3.5 La descrizione di una struttura matematica

Vogliamo adesso vedere se i connettivi e i quantificatori finora introdotti sono sufficienti allo scopo di descrivere formalmente oggetti matematicamente significativi. Tanto per fissare le idee supponiamo di voler lavorare formalmente in una teoria che ci permetta di dimostrare (scoprire?) fatti rilevanti sui numeri naturali e le loro proprietà. Prima di tutto dobbiamo decidere di quali aspetti dei numeri naturali vogliamo parlare: in questo primo esempio conviene fare le cose più semplici possibile e quindi ci limiteremo alla relazione di uguaglianza, alle operazioni di somma, prodotto e successore e al numero zero.

Tanto per avere una notazione per rendere chiaro di cosa vogliamo parlare scriveremo

$$\mathcal{N} \equiv \langle \omega, \{=\}, \{+, \cdot, s\}, \{0\} \rangle$$

dove naturalmente ω è l'insieme dei numeri naturali, $=$ è la relazione identica sui numeri naturali, vale a dire l'insieme $\{\langle x, x \rangle \mid x \in \omega\}$, $+$ e $\cdot$ sono rispettivamente le operazioni di somma e prodotto tra numeri naturali, vale a dire le ben note relazioni ternarie, s è l'operazione di successore, vale a dire la relazione binaria definita ponendo $s \equiv \{\langle x, x+1 \rangle \mid x \in \omega\}$ e infine 0 è ovviamente il numero zero.

Quello che abbiamo scritto sopra è solo un particolare esempio della più generale nozione di struttura del primo ordine che nel seguito indicheremo scrivendo

$$\mathcal{U} \equiv \langle U, \mathcal{R}, \mathcal{F}, \mathcal{C} \rangle$$

dove U è il dominio, cioè l'insieme per i cui elementi vogliamo costruire la teoria, $\mathcal{R}$ è un insieme non vuoto (di qualcosa vogliamo pur parlare!) di relazioni elementari sugli elementi di U , $\mathcal{F}$ è un insieme di funzioni sugli elementi di U e $\mathcal{C}$ è un insieme di elementi di U a cui siamo direttamente interessati.

Proseguiamo comunque con il nostro esempio. Il prossimo passo nella costruzione di una descrizione per la struttura $\mathcal{N}$ consiste nella definizione di un linguaggio per parlare degli oggetti matematici rilevanti. Nel nostro caso abbiamo bisogno di una costante linguistica ld di arietà

$0 \rightarrow (0 \rightarrow 0)$ per denotare la relazione identica $=$ che compare nella definizione della struttura, di una costante linguistica **plus** di arietà $0 \rightarrow (0 \rightarrow 0)$ per denotare l'operazione di somma $+$, di una costante linguistica **prod** di arietà $0 \rightarrow (0 \rightarrow 0)$ per denotare l'operazione di prodotto $\cdot$, di una costante linguistica **succ** di arietà $0 \rightarrow 0$ per denotare l'operazione di successore s e infine di una costante linguistica **zero** di arietà 0 per denotare la costante 0 .

Come possiamo ora utilizzare il linguaggio così creato per esprimere fatti rilevanti sulla struttura che tale linguaggio si propone di descrivere?

Tanto per cominciare possiamo osservare che con il linguaggio a nostra disposizione siamo in grado di *nominare* vari elementi e sottoinsiemi dell'insieme dei numeri naturali. Ad esempio siamo in grado di dare un nome ad ogni numero naturale semplicemente scrivendo

$$\underline{n} \equiv \text{succ}^n(\text{zero})$$

dove con $\underline{n}$ intendiamo il nome nel linguaggio del numero n e con $\text{succ}^n(z)$ intendiamo l'applicazione n volte della costante linguistica **succ** alla costante linguistica **zero**. Naturalmente esistono anche altri nomi nel linguaggio per i numeri naturali, ad esempio possiamo considerare $\text{plus}(\text{prod}(\text{succ}(\text{zero})(\text{succ}(\text{zero}))) (\text{zero}))$ come un (complicato!) nome per il numero 1.

Vale la pena di notare che tuttavia il linguaggio da solo non ha alcuna possibilità di sapere che $\text{plus}(\text{prod}(\text{succ}(\text{zero})(\text{succ}(\text{zero}))) (\text{zero}))$ e $\text{succ}(\text{zero})$ sono due nomi diversi per lo stesso numero naturale.

Vediamo ora un esempio di definizione di un sottoinsieme:

$$\text{Pari}(x) \equiv \exists y. \text{Id}(x)(\text{prod}(\text{succ}(\text{succ}(\text{zero}))) (y))$$

e di definizione di una relazione tra numeri naturali:

$$\text{Minore}(x, y) \equiv \exists w. \neg(\text{Id}(w)(\text{zero})) \wedge \text{Id}(\text{plus}(x)(w))(y)$$

Esercizio 3.5.1 *Enunciare con una proposizione il fatto che i numeri che sono multipli di 6 sono esattamente i numeri che sono multipli di 2 e di 3.*

Esercizio 3.5.2 *Trovare una proposizione che definisca il sottoinsieme dei numeri primi.*

Esercizio 3.5.3 *Esprimere la congettura che dice che ogni numero pari è somma di due numeri primi (congettura di Goldbach).*

3.5.1 Il problema della descrizione di una struttura

Dopo avere visto che il linguaggio che abbiamo creato per parlare della struttura $\mathcal{N}$ ci permette in realtà di esprimere fatti significativi su tale struttura, possiamo porci un'altra domanda che risulta molto importante per un matematico che desideri utilizzare il nostro linguaggio per lavorare con i numeri naturali: è possibile utilizzare il nostro linguaggio per caratterizzare la struttura $\mathcal{N}$?

Vediamo di chiarire meglio in che modo si possa utilizzare il linguaggio per caratterizzare una struttura. L'idea è che con il linguaggio possiamo esprimere proposizioni vere nella struttura. Ad esempio, se scriviamo $\forall x. \forall y. \text{Id}(\text{plus}(x)(y))(\text{plus}(y)(x))$ stiamo esprimendo un fatto vero della struttura $\mathcal{N}$ – il fatto che la somma è commutativa – mentre se diciamo che $\exists x. \text{Id}(\text{succ}(x))(\text{zero})$ stiamo esprimendo un fatto che ci aspettiamo non essere vero nella struttura $\mathcal{N}$.

Ora, un linguaggio è sufficiente per caratterizzare la struttura $\mathcal{N}$ se siamo in grado di scrivere, utilizzando tale linguaggio, un insieme di proposizioni che valgano in $\mathcal{N}$ e solo in $\mathcal{N}$. Ad esempio, se scriviamo

$$(*) \quad \forall x. \neg \text{Id}(\text{succ}(x))(\text{zero})$$

abbiamo sicuramente scritto una proposizione vera in $\mathcal{N}$ ma essa non basta sicuramente a caratterizzare $\mathcal{N}$ visto che vale anche nella struttura rappresentata in figura 3.1 dove intendiamo parlare di una struttura su un insieme di due elementi $\{0, 1\}$ dove la funzione successore sia definita ponendo $s(0) = 1$ e $s(1) = 1$ mentre possiamo definire come ci pare somma e prodotto visto che non c'è per il momento nessuna proposizione che parla di **plus** e **prod**.

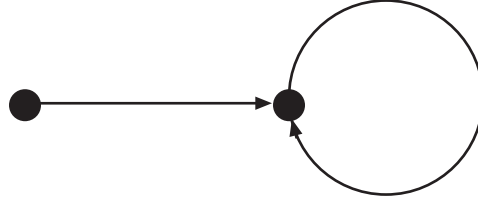


Figura 3.1: struttura su due elementi

Figura 3.2: struttura su due copie di ω

Per evitare questo primo problema potremo allora pensare di aggiungere una proposizione che ci garantisca che la funzione successore è iniettiva:

$$(**) \quad \forall x. \forall y. \text{Id}(\text{succ}(x))(\text{succ}(y)) \rightarrow \text{Id}(x)(y)$$

ma è chiaro che neppure questa proposizione, sebbene valida (nei numeri naturali è vero che la funzione successore è iniettiva), non è sufficiente per discriminare la struttura $\mathcal{N}$ da tutte le altre strutture su cui possiamo interpretare il linguaggio che stiamo considerando. Si consideri ad esempio la struttura considerata nella figura 3.2 costruita con due copie dei numeri naturali da pensarsi una dopo l'altra e in cui le frecce rappresentano la definizione della funzione successore.

Quindi con le due proposizioni che abbiamo trovato finora non siamo riusciti a caratterizzare la struttura $\mathcal{N}$. Potremo naturalmente continuare a provare aggiungendo altre proposizioni, anche proposizioni che coinvolgano **plus** e **prod**, ma conviene provare a ragionare in modo più astratto. Iniziamo ricordando che ω è il minimo insieme che contiene lo 0 ed è chiuso rispetto all'operazione di successore (vedi appendice A.1.2); quindi quel che ci serve è riuscire ad esprimere un'affermazione come “ogni sottoinsieme che contenga 0 e sia chiuso per successore coincide con l'intero insieme” che possiamo (quasi) scrivere come segue

$$(1) \quad \text{per ogni } Y. (0 \in Y) \wedge (\forall x. x \in Y \rightarrow \text{succ}(x) \in Y) \rightarrow \forall x. x \in Y$$

se non fosse per il fatto che la quantificazione più esterna riguarda i sottoinsiemi dell'universo della struttura e non gli elementi di tale universo come accade per il quantificatore universale che abbiamo riconosciuto come capace di costruire proposizioni (un intuizionista potrebbe avere grossi problemi a riconoscere questa quantificazione come dotata di significato visto che dovrebbe essere in grado di considerare la collezione dei sottoinsiemi dei numeri naturali come un insieme, vale a dire dovrebbe essere in grado di generare i suoi elementi).

Quella che stiamo utilizzando non è più una quantificazione del *primo ordine*, cioè sugli elementi dell'universo, ma una quantificazione di secondo ordine, cioè sui sottoinsiemi dell'universo della struttura⁴.

Come abbiamo visto possiamo però utilizzare una proposizione per indicare un sottoinsieme (il nostro esempio riguardava i numeri pari) e tanto più ricco sarà il linguaggio che utilizzeremo tanti più sottoinsiemi di numeri naturali potremo indicare. Allora al posto della proposizione

⁴Per coloro che si chiedono come si vada avanti con gli ordini di quantificazione basta far notare che al primo ordine si quantifica su elementi del dominio U , al secondo ordine su elementi di $\mathcal{P}(U)$, al terzo su elementi di $\mathcal{P}(\mathcal{P}(U))$ e così via, dove $\mathcal{P}(S)$ indica l'insieme delle parti di S (vedi appendice A.4). Tanto per capire che non c'è nulla di strano in queste quantificazioni si tenga conto che nell'usuale definizione di topologia si usa una quantificazione del terzo ordine visto che si richiede che l'unione di ogni famiglia di aperti sia un aperto.

sopra considerata, che utilizza una quantificazione del secondo ordine, possiamo utilizzare infinite proposizioni, una per ciascuna proposizione $A(x)$ che individua un sottoinsieme, scrivendo

$$A(0) \wedge (\forall x. A(x) \rightarrow A(\text{succ}(x)) \rightarrow \forall x. A(x))$$

Naturalmente non si tratta di niente di diverso dalla solita enunciazione della possibilità di fare una prova per induzione.

Tuttavia, neppure questo basta a risolvere davvero il nostro problema di caratterizzazione della struttura $\mathcal{N}$ visto che i sottoinsiemi dei numeri naturali sono una quantità più che numerabile mentre con un qualsiasi linguaggio (ragionevole) non possiamo scrivere più che una quantità numerabile di proposizioni e quindi non siamo in grado di indicare tutti i sottoinsiemi. Per questo motivo le infinite proposizioni del primo ordine non sono sufficienti per sostituire completamente la singola proposizione del secondo ordine (1). Questo fallimento, che per il momento possiamo solo enunciare, è dimostrabile con un vero e proprio teorema relativo alle possibilità di un linguaggio del primo ordine ma non deve scoraggiarci (si veda il capitolo 9). Infatti qualsiasi risultato sui numeri naturali si potrà dimostrare utilizzando solamente una prova finita (l'essere finita è una proprietà richiesta sulle prove!) e quindi solo alcune proposizioni vere saranno necessarie per ottenere quel particolare risultato. Naturalmente sta all'abilità di ogni matematico determinare la teoria giusta che gli permetta di dimostrare i risultati a cui è interessato.

Esercizio 3.5.4 *Si determini una proposizione, scritta nel linguaggio che utilizza solamente Id , succ e zero che sia vera nella struttura dei numeri naturali ma non sia soddisfatta nella struttura di figura 3.2. Si costruisca quindi una struttura che verifica tale proposizione, oltre alle proposizioni $(*)$ e $(**)$ e che non sia isomorfa ai numeri naturali.*

Esercizio 3.5.5 *Un classico non ha difficoltà ad utilizzare quantificatori del secondo ordine perché si tratta solamente di un modo di quantificare su qualcosa che esiste indipendentemente dal fatto che noi ci possiamo quantificare sopra. Un intuizionista ha invece difficoltà ad accettare una proposizione che utilizza un quantificatore universale del secondo ordine perché non è difficile trovarsi di fronte a situazioni circolari: infatti tale proposizione potrebbe essere utilizzata per definire un nuovo sottoinsieme su cui il quantificatore dovrebbe già quantificare! Esibire un esempio di questa situazione spiacevole.*

3.5.2 Le strutture descritte da un insieme di proposizioni

Finora ci siamo concentrati sulla possibilità di utilizzare un linguaggio per descrivere una particolare struttura. Tuttavia, una volta che è comparso sulla scena, il linguaggio può anche giocare un ruolo diverso, e in realtà esso è spesso usato in matematica astratta in tale ruolo.

Scritto infatti un insieme di proposizioni possiamo chiederci se ci sono strutture che rendono vere contemporaneamente tutte le proposizioni in tale insieme (si veda la nozione di soddisfacibilità di un insieme di proposizioni nel capitolo 8). Nella sezione precedente abbiamo visto che questo è possibile quando ci siamo accorti che potevamo interpretare in strutture diverse le stesse proposizioni riguardanti l'operazione di successore. La questione non è più quindi di trovare le proposizioni che caratterizzano una particolare struttura ma piuttosto di trovare i modelli che soddisfano certe proposizioni; si tratta, in qualche senso, del problema duale del precedente.

Tanto per fare un esempio ben noto che viene dall'algebra possiamo considerare un linguaggio che contenga, come nel caso precedente una costante linguistica Id di arietà $0 \rightarrow (0 \rightarrow 0)$ per costruire proposizioni atomiche, una costante linguistica $*$, anch'essa di arietà $0 \rightarrow (0 \rightarrow 0)$, per denotare una operazione binaria e una costante linguistica u di arietà 0 per denotare un particolare elemento. Supponiamo quindi di avere le seguenti tre proposizioni scritte utilizzando tale linguaggio

$$\begin{array}{ll} \text{(associativa)} & \forall x. \forall y. \forall z. \text{Id}(*(*x)(y))(z)(*x)(*y)(z)) \\ \text{(neutro)} & \forall x. \text{Id}(*x)(u)(x) \wedge \text{Id}(*u)(x)(x) \\ \text{(inverso)} & \forall x. \exists y. \text{Id}(*x)(y)(u) \end{array}$$

Non è difficile riconoscere, dietro una notazione un po' pesante, le solite condizioni che definiscono la struttura di gruppo. Le strutture per rendere vere queste proposizioni sono quindi quelle costruite nel modo seguente

$$\langle U, \{=\}, \{\cdot\}, \{1\} \rangle$$

dove $\cdot$ è una operazione binaria sugli elementi di U e 1 è un particolare elemento di U , a patto che $\cdot$ e 1 soddisfino le condizioni richieste dalle proposizioni specificate, e si tratta naturalmente dei gruppi.

Questo nuovo modo di usare il linguaggio porta però con se la necessità di essere molto precisi nella definizione della nozione di interpretazione di una proposizione in una struttura visto che adesso non siamo più in presenza di un linguaggio costruito ad hoc per una struttura. Per il momento quel che possediamo, in virtù della teoria delle espressioni dell'appendice B, è una precisa definizione di come si costruisca un linguaggio, ma per sapere quando l'interpretazione di una proposizione risulta vera in una struttura quel che ci serve è una “matematizzazione” della nozione di verità: è quel che faremo nei prossimi capitoli.

Capitolo 4

Regole per ragionare

Nel capitolo precedente abbiamo introdotto le proposizioni di cui ci vogliamo occupare in questo corso e abbiamo spiegato quando un matematico classico e uno intuizionista considerano come vere tali proposizioni. In linea di principio potremo quindi pensare di poter ragionare su tali proposizioni senza nessun bisogno di aggiungere altro. Tuttavia, specialmente se ricordiamo che possiamo essere in presenza di proposizioni contenenti quantificatori che agiscono su un dominio infinito, lo sviluppo di strumenti per ragionare può essere utile quando non addirittura indispensabile. Faremo le nostre proposte di tali strumenti in questo capitolo e nel capitolo 6 mentre rimandiamo ai capitoli successivi lo studio della *qualità* degli strumenti che proponiamo qui.

4.1 Le regole logiche

Lo strumento principale per ragionare rispetto alla questione della verità delle proposizioni sono le regole logiche. In generale una regola logica presenta un certo numero di premesse ad una conclusione. Una regola risulta corretta se a fronte della verità delle premesse è in grado di garantirci sulla verità della conclusione.

Naturalmente la validità di una regola logica dipende fortemente del significato che diamo alle sue premesse e alla sua conclusione e dal fatto che siamo classici o intuizionisti. In generale le premesse e le conclusioni di una delle regole che noi considereremo avranno la forma seguente

$$\Gamma \vdash B$$

dove Γ è un insieme di proposizioni e B è una proposizione (nel seguito chiameremo *sequente* una scrittura come $\Gamma \vdash B$).

Quando poi vorremo ragionare come un classico *decoreremo* il sequente $\Gamma \vdash B$ con un pedice $\mathcal{M}$, ottenendo $\Gamma \vdash_{\mathcal{M}} B$, e lo *leggeremo* come “se nel mondo matematico $\mathcal{M}$ sono vere tutte le proposizioni in Γ allora anche la proposizione B è vera in $\mathcal{M}$ ”. Diremo quindi che il sequente $\Gamma \vdash B$ vale se esso vale in ogni mondo matematico.

Quando invece vorremmo ragionare come un intuizionista useremo una *decorazione* completamente diversa scrivendo $\bar{x} : \Gamma \vdash b : B$ che *leggeremo* “se $x_1, \dots, x_n, \dots$ sono dei ticket per le proposizioni $A_1, \dots, A_n, \dots$ di Γ rispettivamente allora b è un ticket per la proposizione B ”. In questo caso diremo che il sequente $\Gamma \vdash B$ vale se esiste un ticket b tale che valga la sua decorazione $\bar{x} : \Gamma \vdash b : B$.

Nel seguito per semplificare la notazione scriveremo l'insieme Γ di proposizioni come una lista di proposizioni separate da virgole; in particolare per indicare l'insieme $\Gamma \cup \{A\}$ scriveremo Γ, A . Vediamo allora quali regole corrette possiamo trovare per il classico e l'intuizionista.

4.1.1 Gli assiomi

Prima ancora di cominciare con i connettivi e i quantificatori, le cui regole mostreranno come essi trasportano la verità, bisognerà che ci sia una qualche verità da trasportare, vale a dire che dobbiamo procurarci un qualche punto di partenza. Dobbiamo cioè trovare delle regole che non trasportano la verità ma piuttosto la sanciscono. Si tratta cioè di scoprire degli assiomi. Una possibilità

è la regola seguente (si tratta di una regola con zero premesse, fatta solo con la conclusione)

$$(\text{assioma}) \quad \Gamma, A \vdash A$$

Per vedere che esse vale classicamente in ogni possibile mondo matematico $\mathcal{M}$ bisogna verificare che se tutte le proposizioni in Γ e A sono vere in $\mathcal{M}$ allora A è vera in $\mathcal{M}$ che risulta ovviamente valere.

La cosa interessante è che lo stesso sequente vale anche in una lettura intuizionista ma per vedere questo fatto dobbiamo trovare l'opportuna decorazione intuizionista. In questo caso la cosa non è difficile visto che basta porre

$$(\text{assioma}) \quad \bar{x} : \Gamma, x : A \vdash x : A$$

per ottenere che se x è un generico ticket per la proposizione A allora esso è un ticket per la proposizione A .

4.1.2 Congiunzione

Vogliamo determinare adesso delle regole che ci permettano di dimostrare la verità di una congiunzione. La spiegazione intuitiva che abbiamo dato nel capito precedente ci suggerisce la regola seguente.

$$(\wedge\text{-right}) \quad \frac{\Gamma \vdash A \quad \Gamma \vdash B}{\Gamma \vdash A \wedge B}$$

dove con Γ indichiamo la lista delle proposizioni dalla cui verità la verità di A e di B dipendono (visto il suo ruolo nel seguito ci riferiremo a Γ come al contesto del sequente $\Gamma \vdash C$).

È chiaro che questa regola vale nella lettura classica perchè se $\mathcal{M}$ è un qualsiasi mondo matematico che rende vere tutte le proposizioni in Γ allora la prima premessa ci garantisce che rende vera anche A mentre la seconda ci garantisce che rende vera B . Ma allora, visto il suo significato, anche $A \wedge B$ è vero in $\mathcal{M}$.

La cosa interessante è che la stessa regola vale anche in una prospettiva intuizionista. Se infatti $x_1, \dots, x_n, \dots$ sono dei generici ticket per le proposizioni in Γ allora la prima premessa ci garantisce che esiste un ticket a per la proposizione A mentre la seconda premessa ci garantisce che esiste un ticket b per la proposizione B . Ma allora sappiamo come procurarci un ticket per la proposizione $A \wedge B$: ci basta costruire la coppia $\langle a, b \rangle$.

È sufficiente questa regola per ottenere tutte le proposizioni vere sia pure limitatamente alle proposizioni che utilizzano solo il connettivo $\wedge$? Certamente abbiamo risolto il problema di come fare a dimostrare che una proposizione della forma $A \wedge B$ è vera, non abbiamo però detto nulla su come fare ad usare la conoscenza sul fatto che una certa proposizione della forma $A \wedge B$ è vera, cioè non sappiamo utilizzare $A \wedge B$ quando questa proposizione è una ipotesi. Introduciamo a questo scopo la regola seguente

$$(\wedge\text{-left}) \quad \frac{\Gamma, A, B \vdash C}{\Gamma, A \wedge B \vdash C}$$

che sostanzialmente dice che tutto ciò che è conseguenza di A e di B separatamente è anche conseguenza di $A \wedge B$.

Verificare che si tratti di una regola valida dal punto di vista classico è immediato: se $\mathcal{M}$ rende vere tutte le proposizioni in Γ e anche $A \wedge B$ allora rende vere anche A e B separatamente e quindi la premessa ci assicura che rende vera C .

Anche per quanto riguarda un intuizionista la dimostrazione della validità della regola è diretta. Infatti se $x_1, \dots, x_n, \dots$ sono dei ticket per le proposizioni in Γ e w è un ticket per $A \wedge B$ allora $\text{split}(w, \lambda x. \lambda y. d)$ è il ticket che cerchiamo nell'ipotesi che d sia il ticket per C che abbiamo per ipotesi della validità della premessa della regola.

È utile considerare anche la seguente soluzione anche se è un po' più laboriosa. Supponiamo infatti che $x_1, \dots, x_n, \dots$ siano dei ticket per le proposizioni in Γ e che w sia un ticket per $A \wedge B$. Allora $\pi_1(w)$ è un ticket per A mentre $\pi_2(w)$ è un ticket per B . Se supponiamo ora che il ticket per A nella premessa della regola si chiami x e il ticket per B si chiami y allora il ticket per la proposizione C nella conclusione che stiamo cercando altro non è che $d[x := \pi_1(w), y := \pi_2(w)]$ se d è il ticket per la proposizione C che ci viene fornito dalla premessa.

Esercizio 4.1.1 Dimostrare che la regola $\wedge$ -eliminazione è equivalente alle seguenti due regole

$$(proiezione-1) \quad \frac{\Gamma \vdash A \wedge B}{\Gamma \vdash A} \quad (proiezione-2) \quad \frac{\Gamma \vdash A \wedge B}{\Gamma \vdash B}$$

(sugg.: per fare questo esercizio è utile la seguente regola

$$(weakening) \quad \frac{\Gamma \vdash C}{\Gamma, A \vdash C}$$

da usare dopo averne dimostrato la validità sia classica che intuizionista.)

Esercizio 4.1.2 Dimostrare utilizzando le regole che la congiunzione è commutativa, cioè dimostrare che $A \wedge B \vdash B \wedge A$.

4.1.3 Disgiunzione

Vogliamo ora fare per la disgiunzione l'analogo di quel che abbiamo fatto per la congiunzione. In questo caso è facile trovare le regole per la dimostrazione di una proposizione delle forma $A \vee B$

$$(\vee\text{-right}) \quad \frac{\Gamma \vdash A}{\Gamma \vdash A \vee B} \quad \frac{\Gamma \vdash B}{\Gamma \vdash A \vee B}$$

Che si tratti di regole classicamente valide si vede subito (dimostriamo solo la validità della prima delle due regole essendo la seconda completamente analoga). Se infatti $\mathcal{M}$ è un mondo matematico che rende vere tutte le proposizioni in Γ allora per la premessa dice che anche A è vero in $\mathcal{M}$ e quindi, visto il significato, ne segue che anche $A \vee B$ è vero in $\mathcal{M}$.

Per quanto riguarda la validità intuizionista della stessa regola supponiamo che $x_1, \dots, x_n, \dots$ siano dei ticket per le proposizioni in Γ allora la premessa della regola garantisce l'esistenza di un ticket a per la proposizione A . Ma allora $i(a)$ è un ticket per $A \vee B$.

L'esercizio 4.1.3 è una semplice applicazione delle regole finora introdotte, ma è interessante perchè mostra come dedurre una coppia che sta nella relazione $\vdash$ e che non era considerata direttamente nelle regole fornite: si tratta del primo caso di scoperta di una verità ottenuta attraverso una dimostrazione, cioè un calcolo.

Vediamo ora come trattare una assunzione della forma $A \vee B$. L'idea è che, visto che $A \vee B$ è vero se è vero A o è vero B allora da $A \vee B$ posso derivare tutto quello che posso derivare sia da A che da B . Una regola che esprima questo fatto è la seguente

$$(\vee\text{-left}) \quad \frac{\Gamma, A \vdash C \quad \Gamma, B \vdash C}{\Gamma, A \vee B \vdash C}$$

Per vedere che questa regola è classicamente valida basta osservare che se $\mathcal{M}$ è un mondo matematico che rende vere tutte le premesse in Γ e $A \vee B$ allora o A è vera in $\mathcal{M}$, e in questo caso la prima premessa ci garantisce che in $\mathcal{M}$ vale anche C , o B è vera in $\mathcal{M}$, e in questo caso è la seconda premessa che ci garantisce che in $\mathcal{M}$ vale C .

Vediamo ora che la stessa regola si può giustificare anche dal punto di vista intuizionista. Supponiamo infatti che $x_1, \dots, x_n, \dots$ siano dei ticket per le proposizioni in Γ e che w sia un ticket per $A \vee B$. Ora la premessa sinistra della regola ci garantisce l'esistenza di un ticket d per la proposizione C che dipende, oltre che dai vari ticket $x_1, \dots, x_n, \dots$ per le proposizioni in Γ , dal generico ticket x per la proposizione A mentre la seconda premessa ci garantisce che esiste un ticket e sempre per la proposizione C che dipende dal generico ticket y per la proposizione B . Quindi $\text{when}(w, \lambda x.d, \lambda y.e)$ è un ticket per la proposizione C .

Esercizio 4.1.3 Dimostrare che $A \wedge B \vdash A \vee B$ vale per ogni proposizione A e B .

Esercizio 4.1.4 Dimostrare utilizzando le regole che la disgiunzione è commutativa.

Esercizio 4.1.5 Dimostrare utilizzando le regole la validità della legge di assorbimento, cioè dimostrare che $A \wedge (A \vee B) \vdash A$ e $A \vdash A \wedge (A \vee B)$

Esercizio 4.1.6 Dimostrare con le regole che vale la distributività della disgiunzione sulla congiunzione, cioè $A \vee (B \wedge C) \vdash (A \vee B) \wedge (A \vee C)$.

4.1.4 Trattare con $\perp$ e $\top$

Anche per le proposizioni $\perp$ e $\top$ possiamo trovare regole opportune. Prima di cercarle tuttavia vogliamo far notare che non vogliamo una regola di introduzione a destra per $\perp$ visto che non vogliamo certo dimostrare il falso, né una regola di introduzione a sinistra per $\top$ visto che non ha molto senso assumere il vero. Una volta capito questo dobbiamo ancora decidere che forma dare alle regole $\perp$ -left e $\top$ -right; qui proponiamo le seguenti

$$(\perp\text{-left}) \quad \frac{\Gamma \vdash \perp}{\Gamma \vdash C} \quad (\top\text{-right}) \quad \frac{\Gamma, \top \vdash C}{\Gamma \vdash C}$$

Non bisogna farsi trarre in inganno dal fatto che chiamiamo $\perp$ -left una regola in cui $\perp$ appare a destra e $\top$ -right una regola in cui $\top$ appare a sinistra; infatti a differenza che nelle regole precedenti questa volta la proposizione di cui si parla appare nelle premesse della regola invece che nella conclusione (vedere gli esercizi 4.1.7 e 4.1.8 per capire che la scelta dei nomi è appropriata).

Il significato della prima è che se il falso fosse derivabile allora tutto sarebbe derivabile mentre la seconda dice semplicemente che l'assumere il vero non serve a nulla. Vediamo che si tratta di regole valide.

Dal punto di vista classico la cosa è immediata. Infatti sia $\mathcal{M}$ un mondo matematico che rende vere tutte le proposizioni in Γ . Allora la premessa di $\perp$ -left ci garantisce che $\mathcal{M}$ rende vera anche la proposizione $\perp$, ma questo è impossibile per cui non può esserci mondo matematico che rende vere tutte le proposizioni in Γ e quindi l'implicazione “se $\mathcal{M}$ rende vere tutte le proposizioni in Γ allora rende vera anche C ” vale sicuramente. Veniamo ora alla regola $\top$ -right. Se $\mathcal{M}$ rende vere tutte le proposizioni in Γ allora la premessa della regola ci assicura che rende vero anche C visto che sicuramente $\mathcal{M}$ rende vero $\top$.

Vediamo ora che le regole sono valide anche dal punto di vista intuizionista. Se $x_1, \dots, x_n, \dots$ sono ticket per le proposizioni in Γ allora la premessa della regola $\perp$ -left mi garantisce che esiste un ticket c per $\perp$ e quindi possiamo costruire un ticket $R_0(c)$ per C . D'altra parte, nelle stesse ipotesi, la premessa della regola $\top$ -right mi garantisce che possiamo costruire un ticket per la proposizione C utilizzando i ticket $x_1, \dots, x_n, \dots$ e il ticket $*$ per la proposizione $\top$.

Esercizio 4.1.7 *Dimostrare che la regola $\perp$ -left è equivalente all'assioma $\Gamma, \perp \vdash C$, spiegando in tal modo il suo nome (sugg.: per fare questo esercizio è utile utilizzare la seguente regola*

$$(cut) \quad \frac{\Gamma \vdash A \quad \Gamma, A \vdash C}{\Gamma \vdash C}$$

dopo averne dimostrato la validità sia classica che intuizionista.)

Esercizio 4.1.8 *Dimostrare che la regola $\top$ -right è equivalente all'assioma $\Gamma \vdash \top$ spiegando in tal modo il suo nome.*

Esercizio 4.1.9 *Dimostrare che la regola $\perp$ -left è equivalente all'assioma $\Gamma, \perp \vdash C$ spiegando in tal modo il suo nome.*

4.1.5 Implicazione

Veniamo ora al connettivo intuizionista per eccellenza. Quali regole dobbiamo usare per descriverne il comportamento? Il significato intuizionista della relazione $\vdash$ ci fornisce immediatamente la regola seguente

$$(\rightarrow\text{-right}) \quad \frac{\Gamma, A \vdash B}{\Gamma \vdash A \rightarrow B}$$

Vediamo allora la dimostrazione della validità intuizionista di questa regola. Supponiamo quindi che $x_1, \dots, x_n, \dots$ siano dei ticket per le proposizioni in Γ . Quel che vogliamo ora è un ticket per $A \rightarrow B$. Ma la premessa della regola ci fornisce appunto un metodo b che al variare del ticket x per la proposizione A ci fornisce un ticket per la proposizione B . Il ticket per $A \rightarrow B$ che cerchiamo è quindi la funzione $\lambda x.b$.

Visto che i ticket per una implicazione sono funzioni non è troppo difficile trovare una opportuna regola che permetta di applicare tali funzioni quando per ipotesi ne abbiamo a disposizione una.

$$(\rightarrow\text{-left}) \quad \frac{\Gamma \vdash A \quad \Gamma, B \vdash C}{\Gamma, A \rightarrow B \vdash C}$$

Supponiamo infatti di avere dei ticket $x_1, \dots, x_n, \dots$ per le proposizioni in Γ e che w sia un ticket per $A \rightarrow B$. Allora la premessa sinistra della regola ci assicura che esiste un ticket a per la proposizione A e quindi $\text{ap}(w, a)$ è un ticket per la proposizione B . Ora la premessa destra della regola ci fornisce un metodo c che al variare del ticket y per la proposizione B ci fornisce un ticket per C . Quindi il ticket che cerchiamo per la proposizione C nella conclusione della regola altro non è che $c[y := \text{ap}(w, a)]$.

Esercizio 4.1.10 *Dimostrare che $\vdash A \rightarrow (B \rightarrow A)$.*

Esercizio 4.1.11 *Dimostrare che $\vdash (A \rightarrow (B \rightarrow C)) \rightarrow ((A \rightarrow B) \rightarrow (A \rightarrow C))$.*

4.1.6 Negazione

Siamo infine arrivati al connettivo che ha un significato solo classicamente. Infatti dal punto di vista intuizionista la negazione è un connettivo derivato, definito ponendo $\neg A \equiv A \rightarrow \perp$. Quindi per l'intuizionista le regole che governano la negazione sono quelle derivate dall'implicazione, cioè

$$(\neg\text{-right}) \quad \frac{\Gamma, A \vdash \perp}{\Gamma \vdash \neg A} \quad (\neg\text{-left}) \quad \frac{\Gamma \vdash A}{\Gamma, \neg A \vdash C}$$

Ora, la cosa interessante è che queste regole sono valide anche nell'interpretazione classica anche se questa è completamente diversa da quella intuizionista. Supponiamo infatti che $\mathcal{M}$ sia un mondo matematico che rende vere tutte le proposizioni in Γ ; se a questo punto assumiamo che renda vera anche A allora la premessa della regola $\neg\text{-right}$ ci forza a concludere che dovrebbe rendere vera anche la proposizione $\perp$. Ma questo è impossibile, quindi $\mathcal{M}$ non può rendere vera A e deve quindi rendere vera $\neg A$. Abbiamo quindi dimostrato la correttezza classica della regola $\neg\text{-right}$. Vediamo ora $\neg\text{-left}$. A questo scopo supponiamo che $\mathcal{M}$ sia un mondo matematico che rende vere tutte le proposizioni in Γ e anche $\neg A$. Allora la premessa della regola ci obbliga a concludere che anche A dovrebbe essere vera in $\mathcal{M}$ ma questo è impossibile ($\mathcal{M}$ non può rendere vere contemporaneamente sia A che $\neg A$). Quindi la frase “se $\mathcal{M}$ rende vere tutte le proposizioni in Γ e $\neg A$ allora rende vera C ” vale visto che l'antecedente non può valere.

D'altra parte non possiamo aspettarci che queste regole da sole siano sufficienti per scoprire tutte le proposizioni che sono vere per il classico: possiamo ricordare ad esempio che proposizioni come $A \vee \neg A$ o $\neg\neg A \rightarrow A$ valgono classicamente ma non possiamo aspettarci che abbiano un ticket e quindi non possono essere dimostrabili utilizzando solo regole intuizionisticamente valide. Dobbiamo quindi rafforzarle. Il modo più semplice per ottenere questo risultato è aggiungere direttamente la regola che ci manca

$$(\text{doppia negazione}) \quad \frac{\Gamma \vdash \neg\neg A}{\Gamma \vdash A}$$

La sua validità classica è immediata visto che in ogni mondo matematico $\neg\neg A$ è vera se e solo se è vera A ; il vero problema è capire se la regola della doppia negazione è sufficiente per catturare la verità classica.

Esercizio 4.1.12 *Dimostrare utilizzando le regole valide per il classico che, per ogni proposizione A , $A \vee \neg A$.*

Esercizio 4.1.13 *Dimostrare che le seguenti regole sono equivalenti*

$$\begin{array}{ll}
 (\text{doppia negazione}) & \frac{\Gamma \vdash \neg\neg A}{\Gamma \vdash A} \\
 (\text{terzo escluso}) & \vdash A \vee \neg A \\
 (\text{assurdo}) & \frac{\Gamma, \neg A \vdash \perp}{\Gamma \vdash A} \\
 (\text{consequentia mirabilis}) & \frac{\Gamma, \neg A \vdash A}{\Gamma \vdash A}
 \end{array}$$

4.1.7 I quantificatori

Finita la parte che riguarda i connettivi è venuto ora il momento di occuparci dei quantificatori. Come c'è da aspettarsi i problemi verranno dal trattamento delle variabili.

Quantificatore universale

Come al solito dobbiamo trovare due regole; la prima deve dirci come si fa a dimostrare che una certa proposizione vale per ogni elemento del dominio mentre l'altra deve dirci come usare l'assunzione che una proposizione quantificata universalmente sia vera.

Per quanto riguarda la prima regola l'idea di base è che per dimostrare che qualcosa vale per ogni elemento del dominio senza dover provare tutti i casi, cosa che sarebbe impossibile nel caso di un dominio infinito, l'unica strada è quella di far vedere che la proprietà da dimostrare vale per un generico elemento del dominio per il quale l'unica informazione disponibile è appunto che si tratta di un elemento del dominio e niente altro. Arriviamo quindi alla seguente regola:

$$(\forall\text{-right}^*) \quad \frac{\Gamma \vdash A}{\Gamma \vdash \forall x.A}$$

dove la condizione di non avere alcuna informazione sul particolare elemento x viene resa sintatticamente richiedendo che la variabile x non appaia libera in nessuna proposizione in Γ (abbiamo aggiunto un asterisco nel nome della regola proprio per ricordarci che si tratta di una regola sottoposta ad una condizione¹).

Vediamo ora che si tratta in effetti di una regola valida e che la condizione che x non appaia libera in nessuna proposizione in Γ non è solo una conseguenza della giustificazione della regola che abbiamo fornito sopra ma una vera necessità se vogliamo evitare che la regola ci permetta di fare deduzioni non valide.

Per quanto riguarda la validità classica della regola possiamo ragionare così. Supponiamo che $\mathcal{M}$ sia un mondo matematico che rende vere tutte le proposizioni in Γ . Allora la premessa della regola ci assicura che la proposizione A è vera per il generico elemento x . Quindi in $\mathcal{M}$ vale anche $A[x := a]$ per un generico elemento del dominio a , ma questo vuol proprio dire che in $\mathcal{M}$ è vero $\forall x.A$.

La stessa regola vale anche intuizionisticamente. Infatti, nell'ipotesi di avere un ticket per ogni proposizione in Γ , la premessa della regola ci fornisce un ticket a , che in generale dipende dal generico elemento x del dominio che stiamo considerando, per la proposizione A ; ma allora $\lambda x.a$ è proprio la funzione che cerchiamo come ticket per $\forall x.A$.

Passiamo ora alla regola che determina il modo di usare una proposizione quantificata universalmente. L'idea in questo caso è che dalla verità di una proposizione quantificata universalmente si può dedurre la verità di tutti i suoi esempi particolare, e quindi se qualcosa è deducibile da questi esempi lo è anche dalla proposizione quantificata universalmente. La regola che esprime questa considerazione è la seguente

$$(\forall\text{-left}) \quad \frac{\Gamma, A[x := a] \vdash C}{\Gamma, \forall x.A \vdash C}$$

¹È importante notare che, almeno nel caso in cui Γ è un insieme finito, siamo in grado di verificare davvero se la condizione richiesta è verificata.

dove a è un qualunque elemento del dominio.

Verifichiamone la correttezza classica. Supponiamo che $\mathcal{M}$ sia un mondo matematico che rende vere tutte le proposizioni in Γ e anche $\forall x.A$. Allora in $\mathcal{M}$ è vera $A[x := c]$ per ogni c nel dominio e quindi in particolare è vera $A[x := a]$ che, in virtù della premessa della regola, ci assicura che C è vera.

Si tratta di una regola valida anche intuizionisticamente perchè se abbiamo dei ticket $x_1, \dots, x_n, \dots$ per le proposizioni in Γ e un ticket w per $\forall x.A$, cioè una funzione che applicata ad un qualunque elemento a del dominio ci fornisce un ticket $\mathbf{ap}(w, a)$ per $A[x := a]$, per trovare un ticket per C possiamo usare il metodo c che ci viene fornito dalla premessa della regola. Infatti se z è il ticket per $A[x := a]$ allora il ticket cercato è $c[z := \mathbf{ap}(w, a)]$.

Esercizio 4.1.14 *Dimostrare, utilizzando le regole, che*

$$(\forall x.(A(x) \rightarrow B(x))) \rightarrow ((\forall x.A(x)) \rightarrow (\forall x.B(x)))$$

Trovare invece una struttura che fornisca un contro-esempio per l'implicazione opposta.

Esercizio 4.1.15 *Verificare che la proposizione $\forall v_1.\forall v_2. \mathbf{Id}(v_1, v_2)$ vale in una struttura $\mathcal{U}$ se e solo se il suo universo ha un solo elemento mentre la proposizione $\forall v_1. \mathbf{Id}(v_1, v_1)$ vale in ogni struttura.*

È interessante notare che questo esercizio mostra come si possono caratterizzare le strutture il cui universo ha esattamente un elemento. Come si possono caratterizzare le strutture che hanno esattamente n elementi? (suggerimento: caratterizzare prima le strutture che hanno almeno n elementi e quelle che hanno al più n elementi).

Quantificatore esistenziale

Passiamo ora al quantificatore esistenziale. Cominciamo quindi con la regola che ci permette di provare una proposizione quantificata esistenzialmente. L'idea è naturalmente che una proposizione quantificata esistenzialmente vale quando ho un *testimone* su cui so che essa vale. La regola desiderata è quindi

$$(\exists\text{-right}) \quad \frac{\Gamma \vdash A[x := a]}{\Gamma \vdash \exists x.A}$$

dove a è un qualche elemento del dominio.

Vediamo che questa regola è classicamente valida. Sia $\mathcal{M}$ un qualche mondo matematico tale che tutte le proposizioni in Γ siano vere. Allora la premessa della regola ci garantisce che anche $A[x := a]$ è vera in $\mathcal{M}$ ed esiste quindi un elemento del dominio su cui A è vera, cioè in $\mathcal{M}$ vale $\exists x.A$.

La regola vale anche intuizionisticamente. Infatti se $x_1, \dots, x_n, \dots$ sono ticket per le proposizioni in Γ allora la premessa della regola ci garantisce che possiamo trovare un ticket b per la proposizione $A[x := a]$ e quindi $\langle a, b \rangle$ è un ticket per la proposizione $\exists x.A$.

È adesso il momento di vedere come si utilizza come assunzione una proposizione quantificata esistenzialmente. L'idea questa volta è che la validità di una proposizione quantificata esistenzialmente $\exists x.A$ ci garantisce l'esistenza di un elemento che soddisfa A , ma non ci fornisce (almeno nel caso classico dove si confonde l'esistenza con la mancanza di controesempi, ma anche nel caso intuizionista se non forniamo un sistema di regole che sia in grado di gestire esplicitamente i ticket) nessuna informazione su chi sia tale elemento; quindi possiamo utilizzare $\exists x.A$ come assunzione per derivare una proposizione C solo nel caso in cui siamo in grado di derivare C da una qualunque istanza di A su un generico elemento di cui non conosciamo altro oltre al fatto che è un elemento. La regola cercata è quindi la seguente

$$(\exists\text{-left}^*) \quad \frac{\Gamma, A \vdash C}{\Gamma, \exists x.A \vdash C}$$

dove la presenza dell'asterisco serve a ricordarci che la validità della regola è sottoposta alla condizione che la variabile x non appaia libera né in C né in alcuna proposizione in Γ (il solito modo in cui il calcolo riesce ad esprimere il fatto che non abbiamo ulteriori conoscenze sull'elemento x).

Vediamo di verificarne la validità classica. Sia $\mathcal{M}$ un mondo matematico tale che sono vere tutte le proposizioni in Γ e $\exists x.A$. Esiste allora un elemento a del dominio tale che $A[x := a]$ è

vera in $\mathcal{M}$, ma allora utilizzando la premessa della regola possiamo costruire una nuova prova di $\Gamma, A[x := a] \vdash C$ ripercorrendo tutta la prova data e sostituendo uniformemente a al posto della variabile x (naturalmente, tale sostituzione lungo tutta la prova non modifica Γ e C che per ipotesi non contengono x) e tale nuova prova ci garantisce che anche la proposizione C è vera in $\mathcal{M}$.

In un certo senso la dimostrazione della validità intuizionista della stessa regola è più semplice; questo non deve stupirci troppo visto che l'informazione che ci viene fornita dal sapere che una proposizione esistenziale è vera è molto maggiore nel caso intuizionista che nel caso classico. Supponiamo dunque che $x_1, \dots, x_n, \dots$ e w siano dei ticket per le proposizioni in Γ e per $\exists x.A$; allora $\pi_2(w)$ è un ticket per la proposizione $A[x := \pi_1(w)]$. Ma ora l'ipotesi $\Gamma, A \vdash C$ ci dice che se abbiamo dei ticket $x_1, \dots, x_n, \dots$ per Γ e un ticket y per A possiamo trovare un ticket c per la proposizione C (vale la pena di notare che in c possono comparire sia la variabile x che la variabile y). Quindi sostituendo $\pi_1(w)$ al posto della variabile x e $\pi_2(w)$ al posto di y vediamo che a partire dai ticket $x_1, \dots, x_n, \dots$ per Γ e il ticket $\pi_2(w)$ per $A[x := \pi_1(w)]$ abbiamo un ticket $c[x := \pi_1(w), y := \pi_2(w)]$ per C che è quel che ci richiede la conclusione della regola (vale la pena di notare che quando sostituiamo x con $\pi_1(w)$ nulla cambia in Γ o in C visto che, per ipotesi, questi non contengono x).

Esercizio 4.1.16 *Dimostrare che se $\forall x.(A \rightarrow B)$ allora $(\exists x.A) \rightarrow (\exists x.B)$.*

Esercizio 4.1.17 *Dimostrare utilizzando le regole classiche che se $A \rightarrow \exists x.B$ allora $\exists x.(A \rightarrow B)$.*

Esercizio 4.1.18 *Dimostrare utilizzando le regole classiche che se $(\forall x.A) \rightarrow B$ allora $\exists x.(A \rightarrow B)$ nell'ipotesi che x non appaia libera in B . Dedurre la validità classica del principio del bevitore $\exists x.(A \rightarrow \forall x.A)$.*

4.2 Regole per la proposizione Id

Vista la sua rilevanza in matematica e la sua indipendenza dal particolare dominio che stiamo considerando, non possiamo concludere questo capitolo senza proporre delle regole adatte al trattamento della proposizione per la relazione identica Id .

La prima regola deve garantirci quando possiamo essere sicuri della identità di due elementi. In virtù della riflessività della relazione identica possiamo porre il seguente assioma.

$$(\text{Id-right}) \quad \Gamma \vdash \text{Id}(a, a)$$

È allora ovvio che tale assioma vale classicamente appunto perché la relazione identica è riflessiva, ma esso vale anche intuizionisticamente visto che abbiamo già osservato che abbiamo il ticket $\text{refl}(a)$ che garantisce sulla verità della proposizione $\text{Id}(a, a)$.

Vediamo ora quale può essere una corrispondente regola che ci permetta di utilizzare una ipotesi della forma $\text{Id}(a, b)$.

$$(\text{Id-left}) \quad \frac{\Gamma \vdash C[x := a]}{\Gamma, \text{Id}(a, b) \vdash C[x := b]}$$

Questa regola è infatti valida classicamente visto che se ogni proposizione in Γ è vera in un mondo matematico $\mathcal{M}$ e per di più in tale mondo a e b sono lo stesso elemento, allora l'ipotesi della regola ci garantisce che anche $C[x := a]$ è vero in $\mathcal{M}$ e quindi deve esserlo pure $C[x := b]$ visto che $C[x := b]$ è uguale a $C[x := a]$ perché ottenuto sostituendo in C due cose identiche.

Lo stesso ragionamento ci dovrebbe convincere che la regola è valida anche intuizionisticamente, ma stavolta ci viene richiesto di procurarci anche un ticket per la proposizione $C[x := b]$; ora l'ipotesi ci garantisce che abbiamo un ticket c per la proposizione $C[x := a]$, se assumiamo quindi che w sia un generico ticket per la proposizione $\text{Id}(a, b)$ il ticket cercato altro non è che $\text{sost}(w, c)$.

Esercizio 4.2.1 *Dimostrare utilizzando le regole che la proposizione Id gode della proprietà simmetrica, cioè si può dimostrare che $\forall x.\forall y.\text{Id}(x, y) \rightarrow \text{Id}(y, x)$.*

Esercizio 4.2.2 *Dimostrare utilizzando le regole che la proposizione Id gode della proprietà transitiva, cioè si può dimostrare che $\forall x.\forall y.\forall z.(\text{Id}(x, y) \wedge \text{Id}(y, z)) \rightarrow \text{Id}(x, z)$.*

Esercizio 4.2.3 *Dimostrare utilizzando le regole che se f è un qualsiasi simbolo per funzione diarietà 1 allora $\forall x.\forall y.\text{Id}(x, y) \rightarrow \text{Id}(f(x), f(y))$.*

Capitolo 5

La struttura della verità

Dopo la discussione presentata nel capitolo 3 sulle nozioni di verità matematica, affrontiamo ora il problema di fornire una controparte formale di quelle spiegazioni intuitive. Infatti questo passo sarà necessario nel momento in cui vorremo far vedere che c'è una corrispondenza tra ciò che è vero in una famiglia di strutture e ciò che è dimostrabile formalmente a partire da un certo insieme di proposizioni (vedi capitolo 7). In particolare questa formalizzazione sarà indispensabile per avere qualche strumento che ci permetta di dimostrare che un certo sequente *non* è vero e, nel caso intuizionista, che una certa proposizione non può avere un ticket.

L'idea che vogliamo portare avanti è quella che possiamo considerare un insieme S dei possibili *valori di verità* di una proposizione e metterli tra loro in una qualche relazione d'ordine $x \leq y$ che esprima il fatto che y è almeno tanto vero che x , o equivalentemente che se x è vero allora anche y lo è. Naturalmente ci aspettiamo che nel caso di un matematico classico l'insieme S dovrebbe ridursi a due soli possibili valori di verità, e quindi che tutta la costruzione si riduca alle solite tabelline di verità, mentre nel caso del matematico costruttivo la situazione dovrebbe essere più complessa.

Al fine di capire quali siano le condizioni più naturali da richiedere sulla relazione d'ordine $\leq$ sull'insieme S riprendiamo la nozione di sequente $\Gamma \vdash B$ che abbiamo introdotto nel capitolo precedente e analizziamola nel caso in cui l'insieme Γ sia finito, cioè $\Gamma \equiv A_1, \dots, A_n$. Allora, tanto per cominciare, si vede subito che $A_1, \dots, A_n \vdash B$ vale se e solo se $A_1 \wedge \dots \wedge A_n \vdash B$ e quindi la relazione $\vdash$ che abbiamo definito è di fatto una relazione binaria tra proposizioni.

Se consideriamo ora gli assiomi della forma $A \vdash A$ vediamo che si tratta di una relazione riflessiva mentre la presenza della regola del taglio (cut-rule) ci garantisce che si tratta di una relazione transitiva visto che da $A \vdash B$ e $B \vdash C$ possiamo dedurre $A \vdash C$. Si tratta quindi di un preordine.

Non si tratta tuttavia di una relazione di ordine parziale perchè si vede subito che non è antisimmetrica. Infatti valgono sia il sequente $A \vdash A \wedge A$ che il sequente $A \wedge A \vdash A$ ma A e $A \wedge A$ sono due proposizioni diverse.

Esiste tuttavia una tecnica standard che permette di passare da un preordine ad un ordine parziale; si tratta semplicemente di costruire un quoziente cioè di definire una relazione di equivalenza ponendo

$$A \sim B \text{ se e solo se } A \vdash B \text{ e } B \vdash A$$

e di considerare poi le classi di equivalenza

$$[A]_{\sim} = \{B \mid A \sim B\}$$

per poter infine definire una vera relazione di ordine parziale ponendo

$$[A]_{\sim} \leq [B]_{\sim} \text{ se e solo se } A \vdash B$$

Visto il significato, sia classico che intuizionista della relazione $\vdash$ non è difficile 'leggere' la relazione $[A]_{\sim} \leq [B]_{\sim}$ come *B è almeno tanto vera quanto A oppure se A è vera allora anche B è vera.*

Esercizio 5.0.4 *Dimostrare che la definizione data della relazione $\leq$ tra classi di equivalenza di proposizioni è una buona definizione che non dipende cioè dai rappresentati scelti per definire la classe.*

5.1 Reticoli

La definizione della relazione $\leq$ tra classi di equivalenza di proposizione che abbiamo appena considerato ci suggerisce che se vogliamo fornire una struttura matematica per modellare la verità dobbiamo utilizzare degli ordini parziali. Tuttavia dovrebbe essere immediatamente chiaro che non tutti gli ordini parziali fanno al caso nostro visto che la relazione $\vdash$ tra proposizione permette non solo di definire un ordine parziale tra le classi di equivalenza ma gode anche di proprietà che non sono condivise da tutti gli ordini parziali. Il prossimo passo è quindi quello di cercare di determinare queste proprietà in modo da chiarire completamente quali ordini parziali fanno al caso nostro.

Tanto per cominciare osserviamo che date due proposizioni qualsiasi A e B possiamo dedurre che $A \wedge B \vdash A$ e $A \wedge B \vdash B$ e quindi nella struttura d'ordine parziale che stiamo cercando deve succedere che presi comunque due elementi a e b deve esserci un terzo elemento d minore di entrambi. Ma con le regole che abbiamo studiato nel capitolo precedente è anche possibile dimostrare che, per ogni proposizione C , se $C \vdash A$ e $C \vdash B$ allora $C \vdash A \wedge B$ e quindi, per ogni coppia di elementi a e b non solo deve esistere qualcosa minore di entrambi ma deve esserci anche il più grande tra gli elementi minori di a e b . Una richiesta che dobbiamo quindi fare è che la nostra relazione di ordine parziale $\leq$ ammetta infimo $\inf(a, b)$ per ogni coppia di elementi a e b .

In modo completamente analogo otteniamo che la presenza della disgiunzione tra proposizioni ci obbliga a richiedere che la relazione d'ordine parziale $\leq$ ammetta supremo $\sup(a, b)$ per ogni coppia di elementi a e b . Infatti per ogni coppia di proposizioni A e B abbiamo sia che $A \vdash A \vee B$ e $B \vdash A \vee B$ e che, per ogni proposizione C , se $A \vdash C$ e $B \vdash C$ allora $A \vee B \vdash C$ che implicano che per la relazione d'ordine $\leq$ tra classi di equivalenza di proposizione ci sia il minimo dei maggioranti di due classi qualsiasi.

Una prima approssimazione della struttura d'ordine che cerchiamo ci è quindi fornita dalla struttura di reticolo.

Definizione 5.1.1 (Reticolo: approccio relazionale) *Un reticolo è un insieme parzialmente ordinato $(S, \leq)$ tale che ogni coppia di elementi di S abbia infimo e supremo rispetto a $\leq$.*

A riprova che questa definizione va nella giusta direzione osserviamo che possiamo dare una definizione completamente isomorfa della struttura di reticolo nel modo che segue (si veda l'esercizio 5.1.7).

Definizione 5.1.2 (Reticolo: approccio algebrico) *Un reticolo è un insieme S dotato di due operazioni $\wedge$ e $\vee$ che soddisfino le seguenti condizioni*

$$\begin{array}{lll} \text{(associativa)} & a \wedge (b \wedge c) = (a \wedge b) \wedge c & a \vee (b \vee c) = (a \vee b) \vee c \\ \text{(commutativa)} & a \wedge b = b \wedge a & a \vee b = b \vee a \\ \text{(assorbimento)} & a \wedge (a \vee b) = a & a \vee (a \wedge b) = a \end{array}$$

Questa seconda definizione per la struttura di reticolo mette infatti in evidenza che quel che determina la definizione di reticolo sono alcune proprietà godute dalle classi di equivalenza determinate dalla congiunzione e dalla disgiunzione.

Esistono “in natura” molti esempi di reticoli (si vedano ad esempio quelli ricordati negli esercizi alla fine di questo paragrafo). Tuttavia non ogni reticolo fa la caso nostro visto che oltre a congiunzione e disgiunzione dobbiamo tenere conto anche della presenza della proposizione sempre falsa $\perp$ e della proposizione sempre vera $\top$. Il loro comportamento è regolato dalle regole $\perp$ -left che ci dice che, per ogni proposizione C , $\perp \vdash C$ e $\top$ -right che ci dice che, per ogni proposizione C , $C \vdash \top$. Infatti queste condizioni impongono alla relazione $\leq$ tra classi di equivalenza di proposizioni di avere un minimo elemento, che possiamo indicare con 0 , minore di ogni altro elemento presente nel reticolo ed un massimo elemento, che possiamo indicare con 1 , maggiore di ogni altro elemento presente nel reticolo.

È interessante osservare che anche questa condizione sulla relazione d'ordine si può esprimere come una condizione sulle operazioni $\wedge$ e $\vee$ quando vogliamo considerare la definizione algebrica di reticolo. Infatti si può dimostrare che 1 è il massimo elemento del reticolo se e solo se è l'elemento neutro per l'operazione $\wedge$ e 0 è il minimo elemento del reticolo se e solo se è l'elemento neutro per l'operazione $\vee$ (si vedano gli esercizi 5.1.9 e 5.1.10).

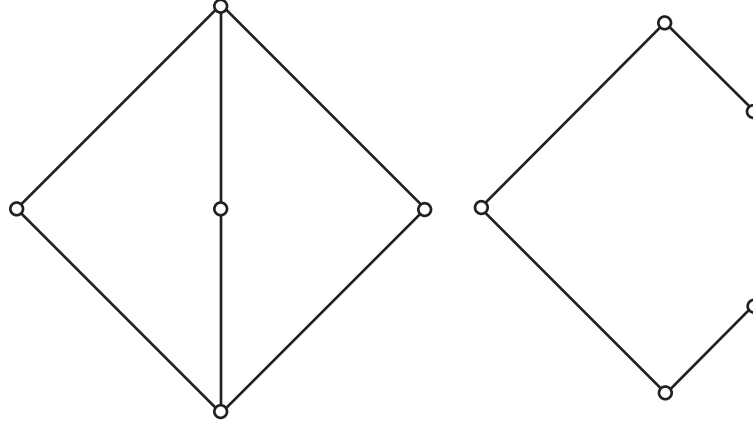


Figura 5.1: reticoli non distributivi

Tuttavia se analizziamo con un po' più di attenzione i connettivi di congiunzione e disgiunzione ci possiamo accorgere che, sia per il classico che per l'intuizionista, vale la condizione che la congiunzione si distribuisce sulla disgiunzione, cioè

$$(\text{distributiva 1}) \quad A \wedge (B \vee C) \vdash (A \wedge B) \vee (A \wedge C) \quad \text{e} \quad (A \wedge B) \vee (A \wedge C) \vdash A \wedge (B \vee C)$$

e la disgiunzione distribuisce sulla congiunzione, cioè

$$(\text{distributiva 2}) \quad A \vee (B \wedge C) \vdash (A \vee B) \wedge (A \vee C) \quad \text{e} \quad (A \vee B) \wedge (A \vee C) \vdash A \vee (B \wedge C)$$

Quindi i reticoli a cui siamo interessati sono solo quelli dove, per ogni a, b e c valgono le seguenti condizioni

$$(\text{distributiva 1}) \quad a \wedge (b \vee c) = (a \wedge b) \vee (a \wedge c)$$

e

$$(\text{distributiva 2}) \quad a \vee (b \wedge c) = (a \vee b) \wedge (a \vee c)$$

mentre esistono anche reticoli che distributivi non sono e questi non sono quindi adatti per fornire una interpretazione completamente corretta.

Per fortuna, in virtù di un famoso teorema di Birkhoff, che non dimostreremo in queste note (vedi [Bir67]), non è difficile riconoscere i reticoli non-distributivi in virtù della possibilità di esprimere in modo relazionale il fatto che un reticolo non è distributivo. Infatti, abbiamo già fatto notare che possiamo rappresentare ogni reticolo con un grafo in cui gli elementi maggiori sono disegnati più in alto (diagrammi di Hasse). Allora un reticolo è non distributivo se e solo se contiene qualche sottoreticolo della forma rappresentata in figura 5.1, dove per sottoreticolo intendiamo un sottoinsieme del reticolo di partenza chiuso per le operazioni di infimo e supremo di tale reticolo.

Esercizio 5.1.3 *Dimostrare che, per qualunque insieme S , l'insieme $\mathcal{P}(S)$ con le operazioni di intersezione e unione è un reticolo distributivo dotato di elemento minimo e massimo.*

Esercizio 5.1.4 *Dimostrare che le operazioni $\min$ e $\max$ definiscono un reticolo su un qualunque insieme linearmente ordinato.*

Esercizio 5.1.5 *Dimostrare che in un reticolo definito come in 5.1.2 vale anche*

$$(\text{idempotenza}) \quad a \wedge a = a \quad a \vee a = a$$

Esercizio 5.1.6 *Dimostrare che in un reticolo definito come in 5.1.2 vale che $x = x \wedge y$ se e solo se $y = x \vee y$.*

Esercizio 5.1.7 *Dimostrare che le due definizioni 5.1.1 e 5.1.2 sono equivalenti (sugg.: per un verso si ponga $a \wedge b \equiv \inf(a, b)$ e $a \vee b \equiv \sup(a, b)$ e nell'altro si definisca $a \leq b \equiv (a = a \wedge b)$)*

Esercizio 5.1.8 Dimostrare che le seguenti sono buone definizioni per le operazioni $\wedge$ e $\vee$ tra classi di equivalenza di proposizioni che non dipendono dai rappresentanti.

$$[A]_{\sim} \wedge [B]_{\sim} \equiv [A \wedge B]_{\sim}$$

$$[A]_{\sim} \vee [B]_{\sim} \equiv [A \vee B]_{\sim}$$

Dimostrare poi che le operazioni così definite soddisfano le condizioni richieste nella definizione 5.1.2 di reticolo.

Esercizio 5.1.9 Dimostrare che l'operazione di infimo è associativa, commutativa e che 1 è il suo elemento neutro.

Esercizio 5.1.10 Dimostrare che l'operazione di supremo è associativa, commutativa e che 0 è il suo elemento neutro.

Esercizio 5.1.11 Dimostrare che, per ogni a e b , $\inf(a, \sup(a, b)) = a$ e $\sup(a, \inf(a, b)) = a$.

Esercizio 5.1.12 Costruire esempi di reticoli dotati di minimo e massimo elemento, aventi solo minimo elemento e aventi solo massimo elemento e senza massimo e minimo elemento.

Esercizio 5.1.13 Costruire esempi di reticoli distributivi e non distributivi.

Esercizio 5.1.14 Esibire un ticket per la proposizione $(A \wedge (B \vee C)) \rightarrow ((A \wedge B) \vee (A \wedge C))$ e un ticket per la proposizione $(A \vee (B \wedge C)) \rightarrow ((A \vee B) \wedge (A \vee C))$.

Esercizio 5.1.15 Dimostrare che se un reticolo soddisfa ad una delle uguaglianze richieste dalla distributività allora esso soddisfa anche all'altra.

Esercizio 5.1.16 Dimostrare che i reticoli rappresentati nella figura 5.1 non sono distributivi.

5.2 Algebre di Boole

Abbiamo quindi capito che le strutture che vanno bene per rappresentare la verità, sia quella classica che quella intuizionista, devono essere dei reticoli distributivi con minimo e massimo elemento. Il problema è adesso che il classico ha anche da rendere conto del connettivo di negazione $\neg(-)$ e per il momento non abbiamo ancora nessuna controparte algebrica per lui. Possiamo naturalmente definire immediatamente una operazione tra le classi di equivalenza delle proposizioni ponendo

$$(\text{complemento}) \quad \nu[A]_{\sim} \equiv [\neg A]_{\sim}$$

e non è difficile dimostrare che si tratta di una buona definizione che non dipende dal rappresentante (si veda l'esercizio 5.2.6).

Le proprietà della negazione che derivano dalle regole che abbiamo presentato nel capitolo precedente sono molte. Ad esempio, abbiamo che $A \vdash \neg\neg A$ e $\neg\neg A \vdash A$ e quindi $\nu\nu[A]_{\sim} = [A]_{\sim}$. Potremo perciò chiedere che nei nostri reticoli si possa definire una operazione $\nu(-)$ tale che, per ogni elemento a , accada che $\nu\nu a = a$.

Ma sappiamo anche che classicamente vale $\vdash A \vee \neg A$ e quindi potremo chiedere che, per ogni a , $a \vee \nu a = 1$ (perchè?). Infine non è difficile dimostrare che $A \wedge \neg A \vdash \perp$ e quindi potremo chiedere che, per ogni a , $a \wedge \nu a = 0$ (perchè?).

Noi definiamo il complemento come segue e vedremo poi che con questa definizione siamo in grado di recuperare anche tutte le informazioni mancanti sulla negazione (vedi gli esercizi alla fine del paragrafo).

Definizione 5.2.1 (Complemento) Sia $(A, \wedge, \vee, 1, 0)$ un reticolo con minimo e massimo elemento e sia $a \in A$. Allora b è un complemento di a se $a \wedge b = 0$ e $a \vee b = 1$.

L'esercizio 5.2.5 mostra che gli elementi di un reticolo possono avere o non avere un complemento e che, nel caso abbiano almeno un complemento, possono averne più di uno. Tuttavia si può dimostrare il seguente teorema.



Figura 5.2: algebra di Boole su 2 elementi

Teorema 5.2.2 (Unicità del complemento) *Sia $(A, \wedge, \vee, 1, 0)$ un reticolo distributivo con minimo e massimo elemento. Allora, se un elemento $a \in A$ ha un complemento ne ha esattamente uno.*

Prova. Supponiamo ci siano due elementi b_1 e b_2 che soddisfano le proprietà richieste per essere complementi di a . Allora essi coincidono in virtù delle seguenti uguaglianze:

$$\begin{aligned}
 b_1 &= b_1 \wedge 1 \\
 &= b_1 \wedge (a \vee b_2) \\
 &= (b_1 \wedge a) \vee (b_1 \wedge b_2) \\
 &= 0 \vee (b_1 \wedge b_2) \\
 &= (a \wedge b_2) \vee (b_1 \wedge b_2) \\
 &= (a \vee b_1) \wedge b_2 \\
 &= 1 \wedge b_2 \\
 &= b_2
 \end{aligned}$$

Siamo quindi giunti alla struttura che meglio descrive la verità classica.

Definizione 5.2.3 (Algebra di Boole) *Una algebra di Boole è un reticolo distributivo con minimo e massimo elemento tale che ogni elemento abbia complemento.*

Visto che in un'algebra di Boole ogni elemento ha complemento e tale complemento è unico, in virtù della distributività, possiamo definire una vera e propria operazione di complementazione che indicheremo con $\nu(-)$.

5.2.1 Le tabelline di verità

Se dovessimo limitarci ad accontentare il classico le algebre di Boole sarebbero la descrizione cercata. Anzi ci potrebbe bastare una sola algebra di Boole quella cioè che ha esattamente due valori di verità: il vero, che possiamo indicare con 1, e il falso, che possiamo indicare con 0. Perciò l'insieme dei valori di verità sarebbe semplicemente $\mathbf{2} \equiv \{0, 1\}$.

Se poi volessimo porre tali valori di verità in un disegno che rappresenti la relazione d'ordine "essere più vero di", ponendo più in alto i valori più veri, ci ritroveremmo con la struttura di verità della figura 5.2

Se adesso volessimo dire come interpretare i connettivi che abbiamo introdotto nel precedente capitolo 2 non dovremo fare altro che ricorrere alle famose tabelline di verità che per completezza ricordiamo qui sotto.

Definizione 5.2.4 (Tabelline di verità) *Le seguenti tabelline definiscono delle operazioni su $\mathbf{2}$ che lo rendono un'algebra di Boole:*

$\wedge$	0	1
0	0	0
1	0	1

$\vee$	0	1
0	0	1
1	1	1

ν	
0	1
1	0

Esercizio 5.2.5 *Trovare esempi di reticoli con minimo e massimo elemento tali che esistano elementi che non hanno complemento, oppure tali che esistano elementi che hanno esattamente un complemento oppure tali che esistano elementi che hanno più di un complemento.*

Esercizio 5.2.6 *Dimostrare che la definizione del complemento sulle classi di equivalenza delle proposizioni è una buona definizione.*

Esercizio 5.2.7 *Dimostrare che, per ogni insieme S , $\mathcal{P}(S)$ è un'algebra di Boole.*

Esercizio 5.2.8 *Costruire un algebra di Boole che non sia (isomorfa a) $\mathcal{P}(S)$ per un qualche insieme S .*

Esercizio 5.2.9 *(Legge della doppia negazione) Dimostrare che in ogni algebra di Boole, per ogni elemento a , accade che $\nu\nu a = a$.*

Esercizio 5.2.10 *Dimostrare che in ogni algebra di Boole, per ogni coppia di elementi a e b , $a \leq b$ se e solo se $\nu b \leq \nu a$.*

Esercizio 5.2.11 *(Legge di De Morgan) Dimostrare che in ogni algebra di Boole, per ogni coppia di elementi a e b , $\nu(a \wedge b) = \nu a \vee \nu b$ (sugg.: visto che il teorema 5.2.2 ci assicura che il complemento è unico, per risolvere questo esercizio basta fare vedere che $\nu a \vee \nu b$ è il complemento di $a \wedge b$).*

Esercizio 5.2.12 *(Legge di De Morgan) Dimostrare che in ogni algebra di Boole, per ogni coppia di elementi a e b , $\nu(a \vee b) = \nu a \wedge \nu b$ (sugg.: simile all'esercizio precedente).*

5.3 Algebre di Heyting

Dobbiamo ora occuparci del caso intuizionista. In questo caso quel che ci serve è una operazione per interpretare l'implicazione. La chiave per capire come definire tale operazione è rendersi conto che l'implicazione è caratterizzata da questa equivalenza intuizionisticamente valida (si vedano gli esercizi 3.1.14 e 5.3.5): $C \wedge A \vdash B$ se e solo se $C \vdash A \rightarrow B$.

Una volta verificato questo fatto, la definizione dell'operazione richiesta è naturale: dobbiamo richiedere che per ogni scelta degli elementi a e b di un reticolo con minimo e massimo elemento si possa trovare un elemento d , che chiameremo una implicazione tra a e b tale che, per ogni c , succeda che $c \wedge a \leq b$ se e solo se $c \leq d$.

Anche in questo caso, come nel caso della negazione, bisogna sincerarsi che quella che stiamo definendo sia una operazione.

Teorema 5.3.1 *Sia $(A, \wedge, \vee, 1, 0)$ un reticolo con minimo e massimo elemento. Allora, dati $a, b \in A$ se esiste una implicazione questa è unica.*

Dimostrazione. Si tratta di una semplice verifica. Supponiamo che d_1 e d_2 siano due possibili implicazioni per a e b . Allora, nel caso che consideriamo la prima implicazione abbiamo che, per ogni c , $c \wedge a \leq b$ se e solo se $c \leq d_1$, mentre se consideriamo la seconda implicazione abbiamo che, per ogni c , $c \wedge a \leq b$ se e solo se $c \leq d_2$. Ma allora se ne deduce che, per ogni c , $c \leq d_1$ se e solo se $c \leq d_2$ che istanziate con d_1 e d_2 al posto di c danno rispettivamente:

1. $d_1 \leq d_1$ se e solo se $d_1 \leq d_2$
2. $d_2 \leq d_1$ se e solo se $d_2 \leq d_2$

Quindi, visto che chiaramente $d_1 \leq d_1$ e $d_2 \leq d_2$ valgono, se ne deduce che

1. $d_1 \leq d_2$
2. $d_2 \leq d_1$

cioè $d_1 = d_2$.

Visto che, se esiste, l'implicazione tra a e b è unica possiamo considerarla una funzione di a e b e quindi usare la seguente notazione $a \rightarrow b$. Arriviamo quindi alla seguente definizione che ci fornisce la struttura adatta ad interpretare i connettivi intuizionisti.

Definizione 5.3.2 (Algebra di Heyting) Una algebra di Heyting è un reticolo distributivo con minimo e massimo elemento tale che l'implicazione sia definita per ogni coppia di elementi.

È utile notare che anche se abbiamo chiesto nella definizione di algebra di Heyting che il reticolo considerato sia distributivo in realtà tale condizione è una conseguenza del fatto che esista sempre l'implicazione.

Teorema 5.3.3 Sia $(A, \wedge, \vee)$ un reticolo che possieda una operazione di implicazione. Allora esso è distributivo.

Dimostrazione. Dobbiamo dimostrare che $a \wedge (b \vee c) = (a \wedge b) \vee (a \wedge c)$ e che $a \vee (b \wedge c) = (a \vee b) \wedge (a \vee c)$ ma in virtù dell'esercizio 5.1.14 è sufficiente dimostrare una sola delle due uguaglianze (si veda comunque l'esercizio 5.3.9). Vediamo quindi di dimostrare la prima tra esse. Un verso è immediato e non richiede neppure di utilizzare le proprietà dell'implicazione. Infatti $a \wedge b \leq a$ e $a \wedge c \leq a$ e quindi $(a \wedge b) \vee (a \wedge c) \leq a$; inoltre $a \wedge b \leq b$ e $a \wedge c \leq c$ e quindi $(a \wedge b) \vee (a \wedge c) \leq b \vee c$ per cui si può concludere che $(a \wedge b) \vee (a \wedge c) \leq a \wedge (b \vee c)$.

D'altra parte $a \wedge b \leq (a \wedge b) \vee (a \wedge c)$ e $a \wedge c \leq (a \wedge b) \vee (a \wedge c)$ e quindi $b \leq a \rightarrow ((a \wedge b) \vee (a \wedge c))$ e $c \leq a \rightarrow ((a \wedge b) \vee (a \wedge c))$. Perciò $b \vee c \leq a \rightarrow ((a \wedge b) \vee (a \wedge c))$ e quindi $a \wedge (b \vee c) \leq (a \wedge b) \vee (a \wedge c)$.

È interessante notare che per i reticoli finiti possiamo dimostrare anche il viceversa (ma si veda la sezione 5.5).

Teorema 5.3.4 In ogni reticolo distributivo finito è possibile definire una operazione di implicazione.

Dimostrazione. Per capire come definire l'operazione di implicazione osserviamo prima di tutto che, in virtù della proprietà che la definisce, l'implicazione tra a e b è il più grande tra gli elementi c tali che $c \wedge a \leq b$. Quindi, se consideriamo l'insieme $\{c \mid c \wedge a \leq b\}$ tutti gli elementi che contiene sono minori o uguali all'implicazione tra a e b che vogliamo definire. Perciò la nostra migliore possibilità è semplicemente definire $a \rightarrow b \equiv \bigvee_{c \wedge a \leq b} c$, dove $\bigvee_{c \wedge a \leq b} c$ è il supremo tra tutti i c tali che $c \wedge a \leq b$ che nel caso di un reticolo finito esiste sicuramente.

Quel che dobbiamo fare ora è verificare che la nostra definizione funzioni, cioè che soddisfi la condizione che definisce l'implicazione. Un verso è facile visto che se $c \wedge a \leq b$ allora ovviamente $c \leq \bigvee_{c \wedge a \leq b} c$. D'altra parte, se $c \leq \bigvee_{c \wedge a \leq b} c$ allora $c \wedge a \leq (\bigvee_{c \wedge a \leq b} c) \wedge a$ e quindi, visto che per ipotesi il reticolo è distributivo, $c \wedge a \leq \bigvee_{c \wedge a \leq b} c \wedge a \leq b$.

Vale la pena di osservare subito che la dimostrazione appena fatta non richiede davvero che il reticolo sia finito ma solamente che il supremo dell'insieme $\{c \mid c \wedge a \leq b\}$ esista e che la distributività dell'infimo tra a e il supremo di tale insieme valga. Si può utilizzare questo fatto per ottenere una vasta classe di algebre di Heyting (si veda l'esercizio 5.3.6).

Esercizio 5.3.5 Dimostrare che $C \wedge A \vdash B$ se e solo se $C \vdash A \rightarrow B$.

Esercizio 5.3.6 Dimostrare che gli aperti di uno spazio topologico sono un'algebra di Heyting.

Esercizio 5.3.7 Dimostrare che in ogni algebra di Heyting, per ogni a e b , vale $(a \rightarrow b) \wedge a \leq b$.

Esercizio 5.3.8 Dimostrare che la condizione che definisce l'operazione di implicazione è equivalente alle seguenti due condizioni:

1. per ogni a , b e c , se $c \wedge a \leq b$ allora $c \leq a \rightarrow b$;
2. per ogni a e b , $(a \rightarrow b) \wedge a \leq b$.

Esercizio 5.3.9 Dimostrare direttamente che in ogni reticolo con implicazione vale la distributività di $\vee$ su $\wedge$, cioè, per ogni a , b e c , vale $a \vee (b \wedge c) = (a \vee b) \wedge (a \vee c)$.



Figura 5.3: algebra di Heyting su 3 elementi

5.4 Relazioni tra algebre di Boole e algebre di Heyting

A questo punto possiamo chiederci quale sia la relazione tra le algebre di Boole e quelle di Heyting. Per rispondere a questa domanda possiamo ricordare che abbiamo già visto che un classico può ricostruire l'implicazione a partire dai connettivi che conosce semplicemente ponendo $A \rightarrow B \equiv \neg(A \wedge \neg B)$. Quindi il primo tentativo per definire una implicazione in un'algebra di Boole è porre

$$a \rightarrow b \equiv \nu(a \wedge \nu b)$$

Vediamo che in effetti questa definizione soddisfa la condizione sulla implicazione in un reticolo. Il primo passo per verificare questo risultato è dimostrare che $\nu(a \wedge \nu b) = \nu a \vee b$ ma, dopo gli esercizi 5.2.9 e 5.2.11, questo è praticamente immediato visto che $\nu(a \wedge \nu b) = \nu a \vee \nu \nu b = \nu a \vee b$.

Quel che dobbiamo verificare diviene quindi che $c \wedge a \leq b$ se e solo se $c \leq \nu a \vee b$. Cominciamo con l'implicazione da sinistra a destra. Supponiamo quindi che $c \wedge a \leq b$ valga. Allora ne segue che

$$c \leq \nu a \vee c = (\nu a \vee c) \wedge 1 = (\nu a \vee c) \wedge (\nu a \vee a) = \nu a \vee (c \wedge a) \leq \nu a \vee b$$

D'altra parte se assumiamo che $c \leq \nu a \vee b$ otteniamo

$$c \wedge a \leq (\nu a \vee b) \wedge a = (\nu a \wedge a) \vee (b \wedge a) = 0 \vee (b \wedge a) = b \wedge a \leq b$$

Quindi abbiamo dimostrato il seguente teorema.

Teorema 5.4.1 *Ogni algebra di Boole è un'algebra di Heyting.*

Non vale invece il viceversa visto che ci sono algebre di Heyting che non sono algebre di Boole. La più semplice è forse quella rappresentata in figura 5.3.

Infatti, non c'è dubbio che si tratta di un'algebra di Heyting visto che è un reticolo distributivo finito linearmente ordinato (dopo l'esercizio 5.1.4 e il teorema 5.3.4 basta verificare la distributività). Tanto per esercitarsi il lettore può verificare che la tabellina dell'implicazione è la seguente

$\rightarrow$	0	a	1
0	1	1	1
a	0	1	1
1	0	a	1

Tuttavia l'elemento a non ha complemento perché tale complemento non può essere 1 perché $1 \wedge a = a \neq 0$, non può essere a perché $a \wedge a = a \neq 0$ e $a \vee a = a \neq 1$ e non può essere 0 perché $0 \vee a = a \neq 1$. Quindi non può trattarsi di una algebra di Boole.

L'esempio precedente, oltre a dimostrare che le algebre di Boole sono un sottoinsieme proprio delle algebre di Heyting, offre anche una interessante lettura logica. Prima di tutto ricordiamo che per un intuizionista non esiste il connettivo di negazione logica; tuttavia ne esiste una approssimazione definita ponendo $\neg A \equiv A \rightarrow \perp$. Questo significa che nel momento in cui interpreteremo in un'algebra di Heyting una proposizione che contiene una negazione dobbiamo utilizzare questa definizione o, equivalentemente, possiamo definire un'operazione di negazione su una algebra di

Heyting ponendo $\nu a \equiv a \rightarrow 0$ (ma non dobbiamo poi aspettarci che si tratti del complemento, che valga cioè che $a \wedge \nu a = 0$ e $a \vee \nu a = 1$). Ora, una proposizione che sia vera per un matematico costruttivo dovrebbe avere valore di verità uguale a 1 in ogni algebra di Heyting indipendentemente dal valore di verità delle sue componenti atomiche (perché? Per una dimostrazione formale di questo fatto si veda il prossimo capitolo). Se ci chiediamo quindi se la proposizione $A \vee \neg A$ vale costruttivamente possiamo rispondere negativamente (come ci aspettiamo!) se siamo in grado di procurarci una opportuna algebra di Heyting e una valutazione di verità per la proposizione A tali che il valore di verità di $A \vee \neg A$ non sia uguale a 1. Se consideriamo allora l'algebra di Heyting di figura 5.3 e pensiamo che il valore di verità di A sia a , corrispondente dal punto di vista intuitivo al fatto che non siamo in grado di decidere se A è vero o falso otteniamo che il valore di verità di $A \vee \neg A$, che per il matematico costruttivo coincide con la proposizione $A \vee (A \rightarrow \perp)$, risulta uguale a $a \vee (a \rightarrow 0) = a \vee 0 = a$. Quindi il valore di verità di $A \vee \neg A$ non è uguale a 1, cioè il valore di verità assunto dalle proposizioni vere, per cui $A \vee \neg A$ non può essere sempre vera dal punto di vista costruttivo. Se le algebre di Heyting costituiscono una interpretazione valida della verità intuizionista (e lo fanno visto che siamo stati sempre attenti di considerare solamente regole giustificate dalla nozione intuizionista di verità) abbiamo trovato una dimostrazione del fatto che una verità classica non è una verità intuizionista. Si tratta naturalmente di un fatto atteso visto che un intuizionista non è certo di grado di trovare un ticket per la proposizione $A \vee \neg A$ visto che questo vorrebbe dire che si è in grado di produrre, per una qualsiasi proposizione A , un ticket per A o un ticket per $\neg A$ che corrisponde un po' troppo ad un desiderio di onnipotenza (o, se si desidera una interpretazione meno psicologica, alla richiesta che ogni proposizione sia decidibile).

Esercizio 5.4.2 *Trovare un'algebra di Heyting ed una interpretazione che non soddisfi la proposizione $\neg A \vee \neg \neg A$.*

Esercizio 5.4.3 *Trovare un'algebra di Heyting ed una interpretazione che non soddisfi la proposizione $\neg \neg A \rightarrow A$.*

Esercizio 5.4.4 *Trovare un'algebra di Heyting ed una interpretazione che non soddisfi la proposizione $((A_1 \wedge A_2) \rightarrow B) \rightarrow ((A_1 \rightarrow B) \vee (A_2 \rightarrow B))$.*

Esercizio 5.4.5 *Trovare un'algebra di Heyting ed una interpretazione che non soddisfi la proposizione $(A_1 \rightarrow (A_1 \wedge A_2)) \vee (A_2 \rightarrow (A_1 \wedge A_2))$.*

5.5 Algebre di Boole e di Heyting complete

Abbiamo quindi risolto il problema di definire una struttura di verità per i connettivi classici e per quelli intuizionisti. È ora il momento di passare ai quantificatori.

Se andiamo a guardare le regole logiche che riguardano i quantificatori che abbiamo introdotto nel precedente capitolo 4 ci accorgiamo che la classe di equivalenza della proposizione $\forall x.A$ è l'infimo dell'insieme delle classi di equivalenza delle proposizioni $A[x := a]$ al variare di a nel dominio che stiamo considerando. Infatti la regola $\forall$ -left ci permette di dimostrare che $\forall x.A \vdash A[x := a]$ vale per ogni a appartenente al dominio che stiamo considerando, e quindi la classe di $\forall x.A$ è minore di ciascuna delle classi di $A[x := a]$; d'altra parte la regola $\forall$ -right ci permette di dimostrare che, per ogni proposizione C tale che $C \vdash A[x := a]$ valga per ogni a nel dominio considerato, abbiamo che $C \vdash \forall x.A$ e quindi la classe di $\forall x.A$ è la più grande tra le classi C minori di tutte le classi di $A[x := a]$ al variare di a nel dominio.

In modo completamente analogo si può vedere che le regole relative al quantificatore esistenziale ci permettono di dimostrare che la classe di $\exists x.A$ è il supremo per l'insieme delle classi di equivalenza delle proposizioni $A[x := a]$ al variare di a nel dominio.

Questo suggerisce che quel di cui abbiamo bisogno è un'operazione che ci permetta di fare, nel caso del quantificatore universale, l'infimo di tutti i valori di verità ottenuti con l'interpretazione delle varie istanze di una certa proposizione e di un'altra operazione che ci permetta di fare, nel caso del quantificatore esistenziale, il supremo degli analoghi valori.

Il problema è però che la definizione di algebra di Boole e di algebra di Heyting non ci garantiscono l'esistenza di altri infimi e supremi a parte quelli binari, e quindi quelli finiti, mentre può ben succedere che una funzione proposizionale sia applicata ad una quantità non finita di oggetti.

La soluzione più banale a questo problema è quella di interpretare solamente in strutture che ci garantiscano la presenza di tutti gli infimi e supremi desiderabili.

Definizione 5.5.1 (Algebre di Boole e Heyting complete) *Un'algebra di Boole o un'algebra di Heyting è completa se ogni suo sottoinsieme U ha infimo e supremo, che indicheremo rispettivamente con $\bigwedge U$ e $\bigvee U$.*

È chiaro che la soluzione di richiedere che infimo e supremo esistano per ogni sottoinsieme è, in un certo senso, esagerata visto che noi abbiamo bisogno solo degli infimi e dei supremi degli insiemi di valori di verità determinati dalle funzioni proposizionali e non di un sottoinsieme arbitrario (si ricordi che le proposizioni, e quindi le funzioni proposizionali, possono essere una quantità numerabile mentre i sottoinsiemi di un'algebra di Boole o di Heyting possono essere una quantità più che numerabile). Pagheremo la semplificazione che abbiamo deciso di fare adesso nel momento in cui dimostreremo il teorema di completezza nel capitolo 7.

Esercizio 5.5.2 *Dimostrare che, per ogni insieme S , $\mathcal{P}(S)$ è un'algebra di Boole completa.*

Esercizio 5.5.3 *Dimostrare che l'insieme degli aperti di uno spazio topologico è un'algebra di Heyting completa (sugg.: attenzione alla scelta corretta dell'infimo di un generico sottoinsieme!).*

Esercizio 5.5.4 *Trovare un esempio di algebra di Boole non completa (sugg.: notare che ogni algebra finita è completa).*

Esercizio 5.5.5 *Sia B una algebra di Boole, non necessariamente completa e si supponga che $(x_i)_{i \in I}$ sia una famiglia di elementi di I dotata di infimo $\bigwedge_{i \in I} x_i$. Dimostrare che allora la famiglia $(\nu x_i)_{i \in I}$ è dotata di supremo e che tale supremo è $\nu \bigwedge_{i \in I} x_i$.*

Esercizio 5.5.6 *Dimostrare che la proposizione $\exists x.(A \rightarrow \forall x.A)$ non vale costruttivamente (sugg.: trovare un dominio e una assegnazione di verità per le proposizioni $A[x := a]$ al variare di a in tale dominio in una opportuna algebra di Heyting che renda il valore di verità di $\exists x.(A \rightarrow \forall x.A)$ diverso da 1).*

Esercizio 5.5.7 *Dimostrare che la proposizione $\forall x.\neg\neg A \rightarrow \neg\neg\forall A$ vale classicamente ma non vale costruttivamente.*

Capitolo 6

Formalizzazione delle regole

Nell'ultimo capitolo abbiamo chiarito quali oggetti matematici useremo per lavorare con il concetto di verità che avevamo informalmente introdotto nel capitolo 2. Quel che vogliamo fare adesso è formalizzare completamente il sistema di regole che avevamo introdotto nel capitolo 4. Anche la versione completamente formalizzata di tale calcolo, che chiameremo ancora *calcolo dei sequenti*, avrà la forma di un insieme di regole che ci permettono di derivare proposizioni vere a partire da proposizioni vere, ma per essere sicuri che questo accada davvero avremo bisogno di essere estremamente precisi sulla nozione di verità dell'interpretazione di una proposizione.

6.1 Semantica e sintassi

Iniziamo precisando la definizione di struttura che abbiamo già introdotto nel capitolo 2, ma che adesso possiamo completare con una algebra che matematizza i possibili valori di verità.

Definizione 6.1.1 (Struttura) Una struttura $\mathcal{U} \equiv (U, \mathcal{R}, \mathcal{F}, \mathcal{C}) + \mathcal{A}$ è costituita dalle seguenti parti:

- Un insieme di valutazione delle proposizioni $\mathcal{A}$ che può essere sia una algebra di Boole che una algebra di Heyting;
- Un insieme non vuoto U detto l'universo della struttura;
- Un insieme non vuoto $\mathcal{R}$ di funzioni totali da elementi di U a valori in $\mathcal{A}$;
- Un insieme, eventualmente vuoto, $\mathcal{F}$ di operazioni, vale a dire funzioni totali, su elementi di U e ad immagine in U ;
- Un sottoinsieme, eventualmente vuoto, $\mathcal{C}$ di U di costanti, vale a dire elementi particolari e designati di U .

Nel capitolo 2 abbiamo già visto qualche esempio di struttura anche se non abbiamo potuto essere abbastanza precisi sull'uso dell'insieme di valutazione delle proposizioni visto che le algebre di Boole e le algebre di Heyting sono state introdotte solamente nell'ultimo capitolo.

Supponendo di avere una struttura di cui vogliamo parlare, abbiamo già suggerito come risolvere il problema di determinare un linguaggio per esprimere le affermazioni che su tale struttura vogliamo fare. Infatti, abbiamo detto che il primo requisito che tale linguaggio deve possedere è che ci siano abbastanza simboli per denotare tutti gli oggetti che appaiono in $\mathcal{R}$, $\mathcal{F}$ e $\mathcal{C}$. Ma questo non è sufficiente per esprimere generiche proprietà di una struttura, proprietà che valgano cioè per ogni elemento della struttura. Quel che ci serve è la possibilità di indicare un generico elemento della struttura, vale a dire che abbiamo bisogno delle *variabili*. Introduciamo quindi nel linguaggio un insieme numerabile (è bene averne sempre di riserva!) di simboli per variabili $v_0, v_1, \dots$ oltre ad un simbolo predicativo P_i^n , di arietà n , per ogni funzione $R_i \in \mathcal{R}$ a n posti, un simbolo funzionale f_j^n , di arietà n , per ogni operazione $F_j^n \in \mathcal{F}$ a n argomenti, ed un simbolo c_k per ogni costante designata appartenente al sottoinsieme $\mathcal{C}$ delle costanti. In questo modo ci sarà nel linguaggio un

simbolo per denotare ogni relazione, ogni operazione ed ogni costante designata della struttura che ci interessa.

Usando i simboli finora introdotti si possono costruire induttivamente delle espressioni, che chiameremo *termini*, nel modo che segue.

Definizione 6.1.2 (Termine) *Le seguenti espressioni di arietà 0 sono termini*

- v , se v è una variabile;
- c , se c è un simbolo per una costante;
- $f_j(t_1, \dots, t_n)$, se $t_1, \dots, t_n$ sono termini e f_j è un simbolo per una funzione di arietà n .

Niente altro è un termine.

Possiamo considerare i termini, almeno quelli che non contengono variabili, come i nomi che si possono usare nel linguaggio per denotare elementi dell'universo della struttura. Infatti

- un simbolo per costante è il nome nel linguaggio della costante che denota;
- un termine complesso, come $f_j(t_1, \dots, t_n)$, è il nome nel linguaggio dell'elemento dell'universo della struttura che si ottiene come risultato dell'applicazione della funzione denotata da f_j agli elementi i cui nomi sono $t_1, \dots, t_n$.

Nel caso in cui un termine contenga delle variabili non possiamo più considerarlo come un nome per un preciso elemento del universo della struttura, ma dobbiamo piuttosto considerarlo come un nome per un generico elemento del universo della struttura che può cambiare, al variare delle possibili interpretazioni delle variabili che il termine stesso contiene. Diventerà un nome per un vero elemento della struttura nel momento in cui fisseremo una qualche mappa σ dall'insieme delle variabili all'universo U della struttura che ci dirà come interpretare ciascuna variabile.

Tuttavia utilizzando solamente i termini non possiamo esprimere affermazioni; ad esempio non possiamo dire quando due termini sono due nomi diversi per lo stesso oggetto dell'universo. Questo è il motivo per cui abbiamo introdotto nel linguaggio anche i simboli predicativi. Utilizzandoli possiamo infatti costruire delle nuove espressioni di arietà 0, vale a dire le *proposizioni*, che abbiamo definito nel capitolo 2.

6.1.1 Interpretazione delle proposizioni

Nell'ultimo capitolo abbiamo definito le algebre di Boole e le algebre di Heyting per *matematizzare* la nozione di verità classica ed intuizionista. Vediamo ora come possiamo utilizzarle per fornire una interpretazione completamente formale delle proposizioni.

Il primo passo consiste naturalmente nel vedere come attribuire un valore di verità alle proposizioni atomiche e, prima ancora, come interpretare i termini. Appena si cerca di formalizzare le spiegazioni intuitive che abbiamo utilizzato finora ci si imbatte in un grave problema: come interpretare termini e proposizioni che contengono variabili libere? Infatti in una struttura ci sono solamente *oggetti matematici* mentre le variabili sono *oggetti linguistici*.

Una volta chiarito dove sta il problema, diamo la definizione formale di cosa intendiamo per interpretazione delle proposizioni di un linguaggio $\mathcal{L}$ in una struttura $\mathcal{U}$.

Definizione 6.1.3 (Interpretazione) *Una valutazione V_I^σ delle proposizioni del linguaggio $\mathcal{L}$ nella struttura $\mathcal{U} \equiv (U, \mathcal{R}, \mathcal{F}, \mathcal{C}) + \mathcal{A}$ per $\mathcal{L}$ consiste di*

- una mappa I che assegna ad ogni simbolo di costante c di $\mathcal{L}$ un elemento di U tra quelli in $\mathcal{C}$ (notazione c_I);
- una mappa I che assegna ad ogni simbolo funzionale f di $\mathcal{L}$ una operazione su U tra quelle di $\mathcal{F}$ (notazione f_I);
- una mappa I che assegna ad ogni simbolo predicativo P di $\mathcal{L}$ una funzione da U verso $\mathcal{A}$ tra quelle di $\mathcal{R}$ (notazione P_I);

- una mappa σ che assegna ad ogni variabile v_i di $\mathcal{L}$ un elemento di U .

A questo punto possiamo rendere la generalità delle proposizioni, espressa nel linguaggio tramite le variabili che nelle strutture non esistono, tramite la generalità delle possibili scelte della mappa σ di interpretazione delle variabili.

Data una interpretazione I del linguaggio $\mathcal{L}$ e una mappa σ per interpretare le variabili in elementi della struttura $\mathcal{U}$ possiamo assegnare un valore di verità a tutte le proposizioni di $\mathcal{L}$.

Da punto di vista pratico la definizione della funzione di interpretazione si svilupperà in due passi: prima vedremo come interpretare i termini in elementi della struttura e poi come interpretare le proposizioni in un valore di verità.

Definizione 6.1.4 (Semantica dei termini) *Sia t un termine del linguaggio $\mathcal{L}$. Allora l'interpretazione di t secondo I , basata sulla interpretazione σ delle variabili (notazione $V_I^\sigma(t)$, o anche più semplicemente t^σ o $\sigma(t)$ quando non vogliamo mettere in evidenza la dipendenza da I) si definisce induttivamente come segue:*

- se t è la variabile x allora $V_I^\sigma(t) \equiv \sigma(x)$
- se t è la costante c allora $V_I^\sigma(t) \equiv c_I$
- se t è il termine $f(t_1, \dots, t_n)$ allora $V_I^\sigma(t) \equiv f_I(V_I^\sigma(t_1), \dots, V_I^\sigma(t_n))$

Poiché i termini sono definiti induttivamente secondo lo stesso schema è chiaro che con questa definizione ogni termine viene interpretato in un elemento dell'universo della struttura.

Possiamo passare allora all'interpretazione delle proposizioni. Abbiamo tuttavia bisogno di introdurre prima una notazione che utilizzeremo nel seguito: indicheremo con $V_I^{\sigma(x/u)}$ l'interpretazione che coincide completamente con l'interpretazione I , basata su σ per quanto riguarda l'interpretazione dei simboli predicativi, dei simboli funzionali, dei simboli di costante e sulla valutazione di tutte le variabili eccetto la variabile x che non viene più interpretata secondo la valutazione σ ma invece nell'elemento $u \in U$.

Siamo ora in grado di introdurre la definizione formale di valutazione di una proposizione. Lo faremo in una sola volta sia per il classico che per l'intuizionista anche se è ovvio che il classico utilizzerà strutture basate su una algebra di Boole completa mentre l'intuizionista utilizzerà strutture basate su una algebra di Heyting completa.

Definizione 6.1.5 (Semantica delle proposizioni) *Sia A una proposizione del linguaggio $\mathcal{L}$. Allora la valutazione di A nella struttura $\langle U, \mathcal{R}, \mathcal{F}, \mathcal{C} \rangle + \mathcal{A}$, secondo l'interpretazione I , basata sull'interpretazione σ delle variabili, (notazione $V_I^\sigma(A)$, o più semplicemente A^σ quando non siamo interessati a mettere in evidenza la dipendenza dall'interpretazione I) è la mappa dalle proposizioni del linguaggio $\mathcal{L}$ all'algebra dei valori di verità definita induttivamente come segue:*

$$\begin{aligned}
\text{se } A \equiv P(t_1, \dots, t_n) & \quad \text{allora} \quad V_I^\sigma(A) \equiv P_I(V_I^\sigma(t_1), \dots, V_I^\sigma(t_n)) \\
\text{se } A \equiv \text{Id}(s, t) & \quad \text{allora} \quad V_I^\sigma(A) \equiv \begin{cases} 1 & \text{se } V_I^\sigma(s) = V_I^\sigma(t) \\ 0 & \text{altrimenti} \end{cases} \\
\text{se } A \equiv \neg B & \quad \text{allora} \quad V_I^\sigma(A) \equiv \nu V_I^\sigma(B) \\
\text{se } A \equiv B \wedge C & \quad \text{allora} \quad V_I^\sigma(A) \equiv V_I^\sigma(B) \wedge V_I^\sigma(C) \\
\text{se } A \equiv B \vee C & \quad \text{allora} \quad V_I^\sigma(A) \equiv V_I^\sigma(B) \vee V_I^\sigma(C) \\
\text{se } A \equiv B \rightarrow C & \quad \text{allora} \quad V_I^\sigma(A) \equiv V_I^\sigma(B) \rightarrow V_I^\sigma(C) \\
\text{se } A \equiv \forall v_i. B & \quad \text{allora} \quad V_I^\sigma(A) \equiv \bigwedge_{u \in U} V_I^{\sigma(x/u)}(B) \\
\text{se } A \equiv \exists v_i. B & \quad \text{allora} \quad V_I^\sigma(A) \equiv \bigvee_{u \in U} V_I^{\sigma(x/u)}(B)
\end{aligned}$$

Anche in questo caso, visto che le proposizioni sono induttivamente definite secondo lo stesso schema è chiaro che con la precedente definizione ad ogni proposizione verrà assegnato un valore di verità.

Vale la pena di osservare il modo in cui viene utilizzata la completezza dell'algebra dei valori di verità allo scopo di interpretare i quantificatori. Nel caso del quantificatore universale l'affermazione che una proposizione $A(x)$ deve valere per ogni elemento della struttura viene resa prendendo l'infimo dei valori che si ottengono variando l'interpretazione della variabile x in tutti i modi possibili e analogamente si opera nel caso del quantificatore esistenziale con la sola differenza che si considera in questo caso il supremo dei valori ottenuti invece che l'infimo.

È interessante notare che la valutazione di una proposizione dipende solo dalla valutazione delle variabili che vi occorrono libere (dimostrarlo!).

È importante sottolineare ancora una volta il fatto che mentre una struttura determina il linguaggio formale da usare per parlarne, una volta fissato tale linguaggio i simboli che in esso compaiono si possono interpretare anche in strutture del tutto differenti e la stessa proposizione può essere interpretata in modi diversi in strutture diverse. In questo senso le proposizioni permettono di discriminare tra le strutture anche se esse sono dello stesso tipo. Questa osservazione è di estremo interesse in quanto suggerisce un possibile uso del linguaggio per descrivere le strutture: un insieme di proposizioni Φ determina le strutture che rendono vere tutte le sue proposizioni per ogni interpretazione delle variabili. A questo punto diventa interessante scoprire, tramite strumenti puramente sintattici, che considerano cioè solamente la manipolazione simbolica, quali sono le proposizioni A che in tali strutture sono vere, la cui verità è cioè conseguenza di quella delle proposizioni in Φ .

Naturalmente, così come esistono proposizioni la cui valutazione dipende dall'interpretazione, ce ne sono alcune, ad esempio $(\forall v_1) \text{Id}(v_1, v_1)$, la cui valutazione è indipendente dalla particolare interpretazione considerata. È chiaro che sono anche queste un interessante oggetto di studio. In realtà le due questioni sono strettamente correlate: infatti le proposizioni vere in ogni interpretazione saranno quelle che sono conseguenza di un insieme vuoto di proposizioni, perché evidentemente tutte le proposizioni nell'insieme vuoto di proposizioni sono soddisfatte in ogni struttura. Vedremo inoltre che, se l'insieme Φ di proposizioni è finito allora il problema della conseguenza logica si può ridurre al problema della verità in ogni struttura.

Esercizio 6.1.6 *Siano B e C due proposizioni. Dimostrare che per ogni interpretazione I in ogni struttura basata su un'algebra di Boole le seguenti coppie di proposizioni hanno la stessa valutazione*

1. $B \wedge C$ e $\neg(B \rightarrow \neg C)$
2. $B \vee C$ e $\neg B \rightarrow C$
3. $\exists v_i. A$ e $\neg \forall v_i. \neg A$

Esercizio 6.1.7 *Dimostrare che la proposizione (vedi esercizio 3.4.4)*

$$\exists x.(A(x) \rightarrow \forall x.A(x))$$

viene valutata in 1 in ogni struttura basata su un'algebra di Boole mentre esiste una algebra di Heyting ed una valutazione in tale algebra di Heyting che è diversa da 1.

6.2 Le regole

D'ora in avanti in questo capitolo lavoreremo per costruirci degli strumenti che ci permettano di scoprire se una proposizione è vera in una struttura senza dover per forza ricorrere alla definizione di interpretazione che abbiamo introdotto nella sezione precedente. Si tratta quindi di passare da una *verifica* ad una *dimostrazione* e, per costruire dimostrazioni dovremo procurarci delle regole che ci permettano di ottenere proposizioni vere quando utilizzate su proposizioni vere, regole cioè che siano in grado di trasportare la verità. Naturalmente, così come ci sono due nozioni di verità, quella classica e quella intuizionista, ci saranno anche regole diverse per i due casi; è forse piuttosto sorprendente il fatto che le regole in pratica coincidano quasi completamente e che la sola differenza

si può ridurre all'aggiungere una regola per trattare il caso della negazione classica rispetto a quella intuizionista.

Di fatto nel capitolo 4 abbiamo già scoperto alcune regole utili per ragionare; il problema è che adesso abbiamo cambiato un po' le regole del gioco, visto che abbiamo completamente separato la semantica dalla sintassi, e dobbiamo quindi vedere come la presenza esplicita delle interpretazioni cambia la nozione di validità delle regole oltre che correggere le regole relative ai quantificatori dove al posto degli elementi del dominio dovranno apparire solo i loro nomi.

Come nel capitolo 4 l'idea di fondo è quella di introdurre la relazione $\vdash$ tra un insieme Γ di proposizioni (adesso che abbiamo parlato di algebre di Boole e di Heyting complete possiamo generalizzare il lavoro fatto nel capitolo 4 al caso di un insieme di proposizioni non necessariamente finito) e una proposizione A il cui significato inteso è che A è una conseguenza di Γ . Visto che questo vale se accade che tutte le volte che tutte le proposizioni in Γ sono valutate in vero anche A viene valutata in vero, la relazione $\vdash$ che stiamo cercando è in realtà quella che dice che, per ogni valutazione $V(-)$ delle proposizioni in una struttura risulta $V(\Gamma) \leq V(A)$, dove con $V(\Gamma)$ intendiamo indicare il valore di verità $\bigwedge_{C \in \Gamma} V(C)$ (si noti che quando l'insieme Γ non è finito ci serve un'algebra completa affinché questo infimo sia definito).

Vale la pena di osservare subito un caso particolarmente interessante della relazione $\vdash$, vale a dire il caso in cui $\Gamma = \emptyset$, cioè $\vdash A$. In questo caso, la definizione di infimo di un insieme mostra che $\bigwedge_{C \in \emptyset} V(C) = 1$ e quindi dire $\vdash A$ significa che, per ogni valutazione $V(-)$, accade che $V(A) = 1$, cioè A va interpretato in *vero* in ogni struttura.

Nella definizione della relazione $\vdash$ siamo sempre ristretti a muoverci tra lo Scilla dell'utilizzare solo regole valide e il Cariddi di non lasciarci sfuggire proposizioni che sono in realtà rese vere da ogni interpretazione. Ad esempio $\{A\} \vdash A$ dovrebbe sicuramente appartenere alla relazione $\vdash$ visto che, per ogni interpretazione $V(-)$, sicuramente vale che $V(A) \leq V(A)$, ma altrettanto sicuramente non possiamo limitarci a definire la relazione $\vdash$ semplicemente dicendo che per ogni proposizione A vale $\{A\} \vdash A$ visto che anche $\{A, B\} \vdash A$ è valida per ogni interpretazione poiché $V(A) \wedge V(B) \leq V(A)$.

Dovrebbe essere chiaro che scrivere solo regole valide è relativamente facile (se non siamo sicuri non prendiamo la regola ...) mentre trovare un sistema di regole completo, sufficiente per scoprire tutte le cose vere in ogni interpretazione è un problema di gran lunga più complesso cui dedicheremo tutto il prossimo capitolo¹.

La strategia che seguiremo per cercare di ottenere la completezza del sistema di regole sarà quella di provare a descrivere, uno alla volta, quel che è necessario per caratterizzare ogni connettivo e quantificatore. Naturalmente nella ricerca delle regole ci faremo aiutare da quel che abbiamo già scoperto nel capitolo 4.

6.2.1 I connettivi

Come d'uso cominceremo con i connettivi visto che essi sono in generale più semplici da trattare dei quantificatori.

Gli assiomi

Gli assiomi che avevamo proposto nel capitolo 4 erano i seguenti

$$(\text{assioma}) \quad \Gamma, A \vdash A$$

Vediamo che essi funzionano anche nel nuovo approccio che stiamo seguendo in questo capitolo. Infatti è immediato verificare che questo assioma è valido in ogni struttura visto che per ogni interpretazione $V(-)$ in una generica struttura si ha che $V(\Gamma) \wedge V(A) \leq V(A)$.

¹In realtà, affinché il sistema di regole sia interessante c'è una richiesta in più: bisogna che sia utilizzabile, cioè bisogna che sia possibile riconoscere in modo effettivo la correttezza dell'applicazione delle regole così che si sia in grado di riconoscere con sicurezza una dimostrazione di $\Gamma \vdash A$. Altrimenti la soluzione del nostro problema sarebbe abbastanza semplice ma anche completamente inutile: basterebbe prendere come assiomi l'insieme $\{\Gamma \vdash A \mid \text{per ogni interpretazione } V(-), V(\Gamma) \leq V(A)\}$ ma in tal caso ci troveremmo nella situazione di non avere nessuna idea di cosa c'è in tale insieme.

Congiunzione

La prima regola sul connettivo $\wedge$ che avevamo considerato nel capitolo 4 era la seguente

$$(\wedge\text{-right}) \quad \frac{\Gamma \vdash A \quad \Gamma \vdash B}{\Gamma \vdash A \wedge B}$$

È facile convincersi che si tratta di una regola valida anche nel nuovo approccio. Infatti, consideriamo una qualsiasi interpretazione $V(-)$. Per ipotesi induttiva, se $\Gamma \vdash A$ allora $V(\Gamma) \leq V(A)$ e se $\Gamma \vdash B$ allora $V(\Gamma) \leq V(B)$; quindi, per le proprietà dell'operazione $\wedge$ in un'algebra di Boole o di Heyting, ne segue che $V(\Gamma) \leq V(A) \wedge V(B) = V(A \wedge B)$.

È utile notare subito che le regole che stiamo definendo, come quelle del capitolo 4, permettono di dimostrare la correttezza di un sequente della forma $\Gamma \vdash A$, i teoremi cioè stabiliscono che la verità di A è una conseguenza della verità di tutte le ipotesi in Γ . Tuttavia in matematica siamo spesso interessati a dimostrare la verità di una proposizione e quindi vogliamo anche costruirci un sistema di regole la cui conclusione sia una proposizione.

Ad esempio nel caso della regola $\wedge$ -right considerata qui sopra quel che sappiamo è che se abbiamo una prova che A è una conseguenza di Γ e una prova che B è una conseguenza di Γ allora deve esserci anche una prova del fatto che $A \wedge B$ è una conseguenza di Γ , ma non abbiamo nessuna informazione su come tale prova sia costruita. Conviene quindi decidere una forma *canonica* per tale prova e questo è il motivo per cui nel seguito scriveremo la seguente forma di composizione di prove

$$\begin{array}{c} \Gamma \quad \Gamma \\ \vdots \quad \vdots \\ \underline{A \quad B} \\ A \wedge B \end{array}$$

ottenuta utilizzando la seguente regola per riuscire a fare l'ultimo passo

$$(\wedge\text{-introduzione}) \quad \frac{A \quad B}{A \wedge B}$$

È interessante notare la forma di questa regola: la conclusione è una proposizione (e sarà così per tutte le regole che considereremo) mentre le premesse sono a loro volta delle proposizioni (e non sarà sempre così).

È sufficiente la regola $\wedge$ -right (o equivalentemente $\wedge$ -introduzione) per ottenere tutte le proposizioni vere sia pure limitatamente alle proposizioni che utilizzano solo il connettivo $\wedge$? Certamente abbiamo risolto il problema di come fare a dimostrare che una proposizione della forma $A \wedge B$ è vera, non abbiamo però detto nulla su come fare ad usare la conoscenza sul fatto che una certa proposizione della forma $A \wedge B$ è vera, cioè non sappiamo utilizzare l'ipotesi $A \wedge B$. Ricordiamo allora che a questo scopo nel capitolo 4 avevamo proposto la seguente regola

$$(\wedge\text{-left}) \quad \frac{\Gamma, A, B \vdash C}{\Gamma, A \wedge B \vdash C}$$

che sostanzialmente dice che tutto ciò che è conseguenza di A e di B separatamente è anche conseguenza di $A \wedge B$.

Verifichiamo che si tratta di una regola valida; anche in questo caso la dimostrazione è banale e richiede solamente l'associatività dell'operazione $\wedge$ in ogni algebra di Boole o di Heyting. Infatti, se $V(\Gamma) \wedge V(A) \wedge V(B) \leq V(C)$ allora $V(\Gamma) \wedge V(A \wedge B) = V(\Gamma) \wedge (V(A) \wedge V(B)) \leq V(C)$.

Anche in questo caso possiamo costruire una dimostrazione *canonica* per la prova che questa regola dichiara esistere; tuttavia rispetto al caso precedente dobbiamo utilizzare una regola con due premesse una delle quali è una proposizione mentre l'altra è una derivazione, cioè la prova che si può ottenere una certa conclusione a partire da certe ipotesi; per indicare la dipendenza della conclusione dalle ipotesi scriveremo tali ipotesi tra parentesi quadre e indicheremo con un indice il fatto che esse costituiscono le ipotesi della deduzione che fa da premessa alla regola che stiamo costruendo:

$$(\wedge\text{-eliminazione}) \quad \frac{A \wedge B \quad \begin{array}{c} \vdots \\ C \end{array} \quad n}{C}$$

Dovrebbe essere facile vedere che questa regola traduce esattamente la regola $\wedge$ -left visto che l'ipotesi a destra è esattamente la premessa della regola $\wedge$ -left e la premessa a sinistra aggiunge la nuova assunzione $A \wedge B$ all'insieme di assunzioni Γ .

Per semplificare la notazione nel seguito non scriveremo nelle regole le assunzioni che non cambiano nel corso della prova; per esempio, nella regola $\wedge$ -eliminazione compaiono esplicitamente quattro assunzioni (anche se in realtà una di esse è un insieme di assunzioni) ma Γ è una assunzione prima di applicare la regola e resta una assunzione anche dopo l'applicazione e quindi, in virtù della convenzione appena fissata, eviteremo di scriverla; d'altra parte $A \wedge B$ è una nuova assunzione e quindi deve sicuramente essere scritta e A e B erano premesse nella derivazione di C da Γ , A e B ma non lo sono più nella nuova derivazione di C e questo fatto deve essere indicato.

Esercizio 6.2.1 *Dimostrare che la regola $\wedge$ -eliminazione è equivalente alle seguenti due regole*

$$(proiezione-1) \quad \frac{A \wedge B}{A} \quad (proiezione-2) \quad \frac{A \wedge B}{B}$$

Esercizio 6.2.2 *Quale è la forma canonica di una prova corrispondente ad un assioma?*

Disgiunzione

Vogliamo ora fare per la disgiunzione l'analogo di quel che abbiamo fatto per la congiunzione. Quanto detto nel capitolo 4 ci suggerisce le regole seguenti.

$$(\vee\text{-right}) \quad \frac{\Gamma \vdash A}{\Gamma \vdash A \vee B} \quad \frac{\Gamma \vdash A}{\Gamma \vdash A \vee B}$$

È facile verificare la correttezza di tali regole. Infatti se $V(\Gamma) \leq V(A)$ allora $V(\Gamma) \leq V(A) \vee V(B) = V(A \vee B)$ e il secondo caso è completamente analogo.

Il passaggio alle prove canoniche si ottiene con le regole seguenti

$$(\vee\text{-introduzione}) \quad \frac{A}{A \vee B} \quad \frac{B}{A \vee B}$$

Vediamo ora come trattare con una assunzione della forma $A \vee B$. Nel capitolo 4 avevamo proposto questa regola

$$(\vee\text{-left}) \quad \frac{\Gamma, A \vdash C \quad \Gamma, B \vdash C}{\Gamma, A \vee B \vdash C}$$

Verifichiamo che si tratta di una regola valida anche nella nuova interpretazione. Per farlo potremo utilizzare il fatto che la disgiunzione si distribuisce sulla congiunzione, ma il modo più facile è utilizzare le proprietà dell'implicazione, che sappiamo essere un connettivo le cui proprietà sono collegate con la validità della proprietà distributiva (si ricordi che l'implicazione c'è sia per le algebre di Heyting che per quelle di Boole). Infatti, le assunzioni delle regole significano che, per ogni valutazione $V(-)$, $V(\Gamma) \wedge V(A) \leq V(C)$ e $V(\Gamma) \wedge V(B) \leq V(C)$ e quindi $V(A) \leq V(\Gamma) \rightarrow V(C)$ e $V(B) \leq V(\Gamma) \rightarrow V(C)$. Ma queste ultime due implicano che $V(A \vee B) = V(A) \vee V(B) \leq V(\Gamma) \rightarrow V(C)$ e quindi ne segue che $V(\Gamma) \wedge V(A \vee B) \leq V(C)$.

Una volta imparata la tecnica non è difficile trovare una opportuna prova canonica anche per questa regola. Basta infatti utilizzare la seguente regola

$$(\vee\text{-eliminazione}) \quad \frac{\begin{array}{c} [A]_1 \quad [B]_1 \\ \vdots \\ A \vee B \end{array} \quad \frac{\begin{array}{c} \dot{C} \\ C \end{array} \quad \begin{array}{c} \dot{C} \\ C \end{array}}{C} 1$$

Esercizio 6.2.3 *Dimostrare utilizzando le nuove regole la validità della legge di assorbimento, cioè dimostrare che $A \wedge (A \vee B) \vdash A$ e $A \vdash A \wedge (A \vee B)$*

Esercizio 6.2.4 *Dimostrare con le nuove regole che vale la distributività di della disgiunzione sulla congiunzione, cioè $A \vee (B \wedge C) \vdash (A \vee B) \wedge (A \vee C)$.*

Trattare con $\perp$ e $\top$

Per quanto riguarda le proposizioni $\perp$ e $\top$ nel capitolo 4 avevamo individuato queste regole.

$$(\perp\text{-left}) \quad \frac{\Gamma \vdash \perp}{\Gamma \vdash C} \quad (\top\text{-right}) \quad \frac{\Gamma, \top \vdash A}{\Gamma \vdash A}$$

Il significato della prima è che se il falso fosse derivabile allora tutto sarebbe derivabile mentre la seconda dice semplicemente che l'assumere il vero non serve a nulla.

Vediamo che si tratta di regole valide. Cominciamo con la regola $\perp$ -left. Sia $V(-)$ una valutazione. Allora, per ipotesi induttiva, $V(\Gamma) \leq V(\perp) = 0$, ma 0 è il minimo elemento dell'algebra e quindi sicuramente vale che $V(\perp) \leq V(C)$ per cui si ottiene $V(\Gamma) \leq V(C)$ per transitività.

D'altra parte se $V(\Gamma) \wedge V(\top) \leq V(A)$ allora sicuramente vale $V(\Gamma) \leq V(A)$ visto che $V(\top) = 1$ e $V(\Gamma) \wedge 1 = V(\Gamma)$.

Trovarne una corrispondente prova canonica non è difficile.

$$(\perp\text{-eliminazione}) \quad \frac{\perp}{C} \quad (\top\text{-introduction}) \quad \frac{}{\top}$$

Implicazione

Veniamo ora al connettivo intuizionista per eccellenza: l'implicazione². Quali regole dobbiamo usare per descriverne il comportamento? Nel capitolo 4 avevamo già individuato una regola $\rightarrow$ -right e una regola $\rightarrow$ -left.

$$(\rightarrow\text{-right}) \quad \frac{\Gamma, A \vdash B}{\Gamma \vdash A \rightarrow B}$$

Visto che la regola discende direttamente dalla condizione che definisce l'implicazione dimostrarne la validità è banale. Infatti, se $V(\Gamma) \wedge V(A) \leq V(B)$ allora $V(\Gamma) \leq V(A) \rightarrow V(B) = V(A \rightarrow B)$.

Costruire ora la prova canonica è immediato

$$(\rightarrow\text{-introduzione}) \quad \frac{\begin{array}{c} [A]_1 \\ \vdots \\ B \end{array}}{A \rightarrow B} 1$$

Veniamo ora alla regola $\rightarrow$ -left. Quella che avevamo individuato è la seguente:

$$(\rightarrow\text{-left}) \quad \frac{\Gamma \vdash A \quad \Gamma, B \vdash C}{\Gamma, A \rightarrow B \vdash C}$$

Per dimostrarne la validità supponiamo che $V(\Gamma) \leq V(A)$. Allora

$$\begin{aligned} V(\Gamma) \wedge V(A \rightarrow B) &\leq V(A \rightarrow B) \wedge V(A) \\ &= (V(A) \rightarrow V(B)) \wedge V(A) \\ &\leq V(B) \end{aligned}$$

visto che in ogni algebra di Heyting H , per ogni coppia di elementi $a, b \in H$, $a \wedge (a \rightarrow b) \leq b$ (dimostrarlo!). Quindi $V(\Gamma) \wedge V(A \rightarrow B) \leq V(B)$. Se consideriamo ora la seconda ipotesi, cioè $V(\Gamma) \wedge V(B) \leq V(C)$, possiamo subito concludere $V(\Gamma) \wedge V(A \rightarrow B) \leq V(C)$ per la transitività della relazione d'ordine.

Per convincerci che abbiamo scelto la regola giusta può essere utile ricordare che in virtù dell'esercizio 5.3.8 la condizione che, per ogni a e b di un'algebra di Heyting, $(a \rightarrow b) \wedge a \leq b$ è equivalente alla implicazione da destra a sinistra della condizione che definisce l'implicazione in un'algebra di Heyting. Facciamo quindi vedere che la regola che abbiamo proposto è sufficiente per derivare la controparte linguistica di tale affermazione. In realtà la cosa è immediata:

$$\frac{A \vdash A \quad A, B \vdash B}{A, A \rightarrow B \vdash B}$$

²Le stesse regole andranno bene in ogni caso anche per la definizione classica del connettivo in termini degli altri connettivi visto che abbiamo già verificato che ogni algebra di Boole è un'algebra di Heyting.

Ora per completare l'analisi dell'implicazione ci manca solo da trovare la prova canonica per questa nuova regola. Una traduzione diretta della regola $\rightarrow$ -left ci porta alla regola seguente

$$(*) \quad \frac{[B]_1 \quad \begin{array}{c} \vdots \\ A \rightarrow B \quad A \quad \dot{C} \end{array}}{C} 1$$

che però è veramente poco maneggevole. Tuttavia una istanza di questa regola è sufficiente per esprimerne tutta la potenza. Consideriamo infatti il caso in cui C sia B ; allora la regola diviene

$$\frac{[B]_1 \quad \begin{array}{c} \vdots \\ A \rightarrow B \quad A \quad \dot{B} \end{array}}{B} 1$$

dove la deduzione di B da B si può ottenere banalmente. Quindi possiamo semplificare la regola in

$$(\rightarrow\text{-eliminazione}) \quad \frac{A \rightarrow B \quad A}{B}$$

che è infatti una regola così famosa che ha anche un nome: *modus ponens*.

Viene lasciata come esercizio la verifica che il modus ponens è sufficiente per derivare la regola (*).

Esercizio 6.2.5 *Dimostrare utilizzando le nuove regole che $\vdash A \rightarrow (B \rightarrow A)$.*

Negazione

Siamo infine arrivati al connettivo che ha un significato solo classicamente. Infatti dal punto di vista intuizionista la negazione è un connettivo derivato, definito ponendo $\neg A \equiv A \rightarrow \perp$, e quindi per l'intuizionista le regole che governano la negazione sono quelle derivate dall'implicazione, cioè

$$\begin{array}{ll} (\neg\text{-right}) \quad \frac{\Gamma, A \vdash \perp}{\Gamma \vdash \neg A} & (\neg\text{-introduzione}) \quad \frac{[A]_1 \quad \begin{array}{c} \vdots \\ \perp \end{array}}{\neg A} 1 \\ (\neg\text{-left}) \quad \frac{\Gamma \vdash A}{\Gamma, \neg A \vdash C} & (\neg\text{-eliminazione}) \quad \frac{A \quad \neg A}{\perp} \end{array}$$

Ma visto che queste regole sono valide in ogni algebra di Heyting e che le algebre di Heyting possono rendere non vere proposizioni che sono vere per il classico (possiamo ricordare ad esempio $A \vee \neg A$ o $\neg\neg A \rightarrow A$) è chiaro che queste regole, benchè valide, non sono sufficienti per esprimere completamente le proprietà della negazione classica. Dobbiamo quindi rafforzarle. Per capire come farlo, vediamo ad esempio la seguente istanza di $\neg$ -introduzione

$$\frac{[\neg A]_1 \quad \begin{array}{c} \vdots \\ \perp \end{array}}{\neg\neg A} 1$$

dove abbiamo utilizzato come assunzione una proposizione negata. La conclusione di questa regola è $\neg\neg A$ e per l'intuizionista non c'è modo di passare da questa proposizione alla proposizione A visto che il fatto di non avere un ticket per $\neg A$ non costituisce certo un ticket per A . Tuttavia per il classico la negazione della falsità di A è sufficiente per garantire la verità di A e quindi se vuole poter dimostrare questo fatto deve rafforzare la regola $\neg$ -introduzione nel modo seguente

$$(\text{assurdo}) \quad \frac{[\neg A]_c \quad \begin{array}{c} \vdots \\ \perp \end{array}}{A} c$$

dove per ricordarci che si tratta di un passaggio valido solo per il classico abbiamo indicato il fatto che la premessa della regola è una deduzione di $\perp$ a partire da $\neg A$ con una lettera c .

La cosa sorprendente è il fatto che questa sola aggiunta al sistema deduttivo dell'intuizionista è sufficiente per dimostrare la verità di tutte le proposizioni classicamente valide. Vediamo a titolo di esempio la prova di $\neg\neg A \rightarrow A$

$$\frac{\frac{[\neg\neg A]_1 \quad [\neg A]_c}{\frac{\perp}{A} c}}{\neg\neg A \rightarrow A} 1$$

Questa volta siamo partiti dalla prova canonica. Vediamo ora quale è la regola corrispondente per la relazione di deducibilità $\vdash$:

$$\frac{\Gamma \vdash \neg\neg A}{\Gamma \vdash A}$$

Riscopriamo cioè la stessa regola che avevamo già incontrato nel capitolo 4.

Esercizio 6.2.6 *Dimostrare che la regola di dimostrazione per assurdo appena introdotta è valida classicamente, i.e. per ogni struttura basata su un'algebra di Boole essa trasporta la verità, ma non intuizionisticamente, i.e. esiste un'algebra di Heyting in cui la deduzione nella premessa della regola è valida ma la conclusione non è interpretata in 1.*

Esercizio 6.2.7 *Dimostrare utilizzando le nuove regole valide per il classico che, per ogni proposizione A , $A \vee \neg A$.*

Esercizio 6.2.8 *Dimostrare che la regola dell'assurdo è equivalente ad aggiungere la seguente regola di doppia negazione*

$$(doppia-negazione) \quad \frac{\neg\neg A}{A}$$

che è a sua volta equivalente ad assumere il seguente assioma su ogni proposizione A

$$(terzo escluso) \quad A \vee \neg A$$

Esercizio 6.2.9 *Dimostrare che la seguente regola*

$$(consequentia mirabilis) \quad \frac{[\neg A]_c \quad \vdots \quad A}{A} c$$

che stabilisce che quando cerchiamo di dimostrare A possiamo sempre usare liberamente l'ipotesi $\neg A$, è equivalente alla regola della dimostrazione per assurdo.

Esercizio 6.2.10 *Sia A una qualunque proposizione scritta utilizzando un linguaggio che consideri solamente i connettivi, vale a dire che A non contiene quantificatori. Dimostrare che A vale classicamente se e solo se $\neg\neg A$ vale intuizionisticamente.*

6.2.2 I quantificatori

Finita la parte che riguarda i connettivi è venuto ora il momento di occuparci dei quantificatori. Il problema in questo caso è che non possiamo semplicemente prendere le regole che abbiamo individuato nel capitolo 4 visto che sia la regola $\forall$ -left che la $\exists$ -right si esprimevano utilizzando direttamente un elemento del dominio e questo non ha più nessun senso in questo nuovo approccio che prevede di collegare la sintassi e la semantica solamente tramite l'interpretazione del linguaggio in una struttura. In generale dovremo riuscire a capire come si possa sostituire un oggetto con il suo nome.

Inoltre non sarà certo una sorpresa scoprire che ci saranno anche problemi che verranno dal trattamento delle variabili che esistono solo linguisticamente ma non hanno nessuna controparte semantica.

Quel che ci servirà nella prova della validità delle regole per i quantificatori che introdurremo nelle prossime sezioni è dimostrare una proprietà che lega l'operazione sintattica di sostituire una variabile con un termine in una proposizione con il modo in cui la proposizione risultante da tale sostituzione verrà interpretata. Iniziamo quindi vedendo come si specializza al caso delle proposizioni la definizione astratta di sostituzione di una variabile con un'espressione in una espressione che abbiamo definito nell'appendice B.

Cominciamo analizzando il seguente esempio. Si consideri la proposizione $((\forall v_1) P(v_1)) \rightarrow P(t)$ dove P è un simbolo predicativo del linguaggio che stiamo considerando e t è un qualsiasi termine; è facile convincersi che si tratta di una proposizione vera in ogni interpretazione visto che il termine t deve comunque essere interpretato in un qualche elemento dell'universo. Ma ciò che ci interessa in questo momento è che per scrivere tale proposizione abbiamo compiuto una operazione sintattica di sostituzione: abbiamo sostituito nella proposizione $P(v_1)$ il termine t al posto della variabile v_1 . Dovrebbe essere immediatamente chiaro che la pura sostituzione testuale, anche se funziona nell'esempio appena visto, non fa, in generale, al caso nostro; si consideri ad esempio la proposizione $(\forall v_1) \text{Eq}(v_1, v_2)$: essa è vera solo in una struttura con un solo elemento. Se a questo punto vogliamo sostituire la variabile v_2 proprio con il termine v_1 ed utilizziamo la semplice sostituzione testuale otteniamo $(\forall v_1) \text{Eq}(v_1, v_1)$ che, come sappiamo, risulta vera in ogni struttura. Tuttavia è chiaro che non era questo quello che si voleva ottenere operando la sostituzione di v_2 con il termine v_1 ; dalla sostituzione si voleva chiaramente che la variabile v_1 svolgesse nella nuova proposizione il ruolo svolto precedentemente dalla variabile v_2 , vale a dire, che quel che si voleva era asserire che comunque si interpreti la variabile v_1 la sua interpretazione coincide con quella di un qualunque elemento dell'universo. Quello che non funziona in questa sostituzione testuale è che la variabile v_1 viene *catturata* dal quantificatore una volta che abbiamo operato la sostituzione. Se vogliamo evitare tale eventualità possiamo usare un piccolo trucco: se c'è pericolo di cattura di una variabile operando una sostituzione, cambiamo il nome della variabile quantificata con un nome *nuovo* che non appaia né nella proposizione dove vogliamo fare la sostituzione né nel termine che vogliamo sostituire; questa operazione è semanticamente giustificata dal fatto che, come abbiamo già osservato, il valore di verità di una proposizione non dipende dall'interpretazione delle variabili quantificate e quindi il loro nome non è rilevante. Diamo quindi la seguente definizione dell'operazione di sostituzione operando per induzione sulla complessità della proposizione in cui vogliamo operare la sostituzione (quel che facciamo altro non è che una istanza dell'algoritmo di sostituzione che appare nell'appendice B).

Definizione 6.2.11 (Sostituzione di una variabile con un termine) *Sia x una variabile e t un termine.*

- *Sia s un termine. Allora la sostituzione della variabile x con il termine t nel termine s è definita induttivamente come segue*
 - *se $s \equiv x$ allora $s[x := t] \equiv t$;*
 - *se $s \equiv y$ allora $s[x := t] \equiv y$, dove stiamo supponendo che y sia una variabile diversa da x ;*
 - *se $s \equiv c$ allora $s[x := t] \equiv c$, dove stiamo supponendo che c sia una costante;*
 - *se $s \equiv f(t_1, \dots, t_n)$ allora $s[x := t] \equiv f(t_1[x := t], \dots, t_n[x := t])$.*
- *Sia A una proposizione. Allora la sostituzione della variabile x con il termine t nella proposizione A è definita induttivamente come segue*
 - *se $A \equiv P(t_1, \dots, t_n)$ allora $A[x := t] \equiv P(t_1[x := t], \dots, t_n[x := t])$;*
 - *se $A \equiv \text{Eq}(t_1, t_2)$ allora $A[x := t] \equiv \text{Eq}(t_1[x := t], t_2[x := t])$;*
 - *se $A \equiv B \wedge C$ allora $A[x := t] \equiv B[x := t] \wedge C[x := t]$;*
 - *se $A \equiv B \vee C$ allora $A[x := t] \equiv B[x := t] \vee C[x := t]$;*
 - *se $A \equiv \neg B$ allora $A[x := t] \equiv \neg(B[x := t])$;*
 - *se $A \equiv B \rightarrow C$ allora $A[x := t] \equiv B[x := t] \rightarrow C[x := t]$;*
 - *se $A \equiv \forall x. B$ allora $A[x := t] \equiv A$;*

- se $A \equiv \forall y. B$ allora $A[x := t] \equiv \forall z. B[y := z][x := t]$, dove stiamo supponendo che y sia una variabile diversa da x e che la variabile z non appaia né in B né in t .
- se $A \equiv \exists x. B$ allora $A[x := t] \equiv A$;
- se $A \equiv \exists y. B$ allora $A[x := t] \equiv \exists z. B[y := z][x := t]$, dove stiamo supponendo che y sia una variabile diversa da x e che la variabile z non appaia né in B né in t .

Per vedere fino in fondo che la definizione che abbiamo dato dell'algoritmo di sostituzione fa proprio quello che vogliamo è conveniente essere un po' più precisi riguardo alle occorrenze delle variabili nelle proposizioni.

Definizione 6.2.12 (Occorrenza libera di una variabile in una proposizione) *Sia A una proposizione e x una variabile che in tale proposizione appare. Allora una occorrenza di x in A è libera se*

- se $A \equiv P(t_1, \dots, t_n)$;
- se $A \equiv B \wedge C$ e tale occorrenza di x è libera in B oppure in C ;
- se $A \equiv B \vee C$ e tale occorrenza di x è libera in B oppure in C ;
- se $A \equiv \neg B$ e tale occorrenza di x è libera in B ;
- se $A \equiv B \rightarrow C$ e tale occorrenza di x è libera in B oppure in C ;
- se $A \equiv \forall v_i. B$ e x è diversa da v_i e l'occorrenza di x considerata è libera in B .
- se $A \equiv \exists v_i. B$ e x è diversa da v_i e l'occorrenza di x considerata è libera in B .

Se una occorrenza di una variabile in una proposizione non è libera si dice che è una occorrenza legata della variabile nella proposizione. Naturalmente una variabile può avere sia occorrenze libere che legate in una proposizione: ad esempio la variabile v_1 ha nella proposizione $(\forall v_1. \text{Eq}(v_1, 0)) \rightarrow \text{Eq}(v_2, v_1)$ una occorrenza legata, la prima, ed una occorrenza libera, la seconda, mentre l'unica occorrenza della variabile v_2 è libera. Diremo ora che una variabile x occorre libera nella proposizione A se in A almeno una delle occorrenze x è libera. Si può allora dimostrare il seguente lemma.

Lemma 6.2.13 *La valutazione della proposizione A dipende solamente dalla valutazione delle variabili che in A occorrono libere, oltre che dalla valutazione delle costanti extra-logiche.*

Lasciamo la sua dimostrazione come esercizio.

Possiamo finalmente dimostrare che la sostituzione che abbiamo introdotto nella definizione qui sopra si comporta bene rispetto alle valutazioni, vale a dire che la valutazione della proposizione $A[x := t]$ dipende dalla valutazione del termine t nello stesso modo in cui la valutazione di A dipendeva dalla valutazione di x . È utile notare che questo lemma è interessante perché per la prima volta abbiamo un collegamento formale tra un'operazione sintattica, i.e. la sostituzione, ed una semantica, i.e. il cambio di valutazione di una variabile.

Teorema 6.2.14 (Correttezza della sostituzione) *Sia A una proposizione, x una variabile, t un termine e σ una valutazione. Allora $V_I^\sigma(A[x := t]) = V_I^{\sigma(x/V_I^\sigma(t))}(A)$.*

La dimostrazione si ottiene per induzione sulla costruzione della proposizione A ma, come al solito, per risolvere il caso in cui A è una proposizione atomica ci serve un risultato analogo sui termini.

Lemma 6.2.15 *Sia s un termine, x una variabile, t un termine e σ una valutazione. Allora $V_I^\sigma(s[x := t]) = V_I^{\sigma(x/V_I^\sigma(t))}(s)$.*

Dimostrazione. La dimostrazione si ottiene per induzione sulla complessità del termine s . Dobbiamo tuttavia considerare due diversi casi quando s è una variabile in corrispondenza ai due casi che appaiono nella definizione di sostituzione (nella dimostrazione per semplificare l'esposizione useremo una notazione semplificata per indicare la valutazione).

- $s \equiv x$. Allora

$$\begin{aligned} s[x := t]^\sigma &\equiv t^\sigma && \text{definizione di sostituzione} \\ &\equiv x^{\sigma(x/\sigma(t))} && \text{definizione di valutazione} \\ &\equiv s^{\sigma(x/\sigma(t))} && \text{perchè } s \equiv x \end{aligned}$$

- $s \equiv y$. Allora

$$\begin{aligned} s[x := t]^\sigma &\equiv y^\sigma && \text{definizione di sostituzione} \\ &\equiv y^{\sigma(x/\sigma(t))} && \text{definizione di valutazione} \\ &\equiv s^{\sigma(x/\sigma(t))} && \text{perchè } s \equiv y \end{aligned}$$

- $s \equiv c$. Allora

$$\begin{aligned} s[x := t]^\sigma &\equiv c^\sigma && \text{definizione di sostituzione} \\ &\equiv c^{\sigma(x/\sigma(t))} && \text{definizione di valutazione} \\ &\equiv s^{\sigma(x/\sigma(t))} && \text{perchè } s \equiv c \end{aligned}$$

- $s \equiv f(t_1, \dots, t_n)$. Allora

$$\begin{aligned} s[x := t]^\sigma &\equiv f(t_1[x := t], \dots, t_n[x := t])^\sigma && \text{definizione di sostituzione} \\ &\equiv f_I(t_1[x := t]^\sigma, \dots, t_n[x := t]^\sigma) && \text{definizione di valutazione} \\ &\equiv f_I(t_1^{\sigma(x/\sigma(t))}, \dots, t_n^{\sigma(x/\sigma(t))}) && \text{ipotesi induttiva} \\ &\equiv f(t_1, \dots, t_n)^{\sigma(x/\sigma(t))} && \text{definizione di valutazione} \\ &\equiv s^{\sigma(x/\sigma(t))} && \text{perchè } s \equiv f(t_1, \dots, t_n) \end{aligned}$$

Con questo abbiamo finito la dimostrazione del lemma e siamo pronti ad affrontare la prova del teorema che avevamo lasciato in sospeso.

- $A \equiv P(t_1, \dots, t_n)$. Allora

$$\begin{aligned} A[x := t]^\sigma &\equiv P(t_1[x := t], \dots, t_n[x := t])^\sigma && \text{definizione di sostituzione} \\ &\equiv P_I(t_1[x := t]^\sigma, \dots, t_n[x := t]^\sigma) && \text{definizione di valutazione} \\ &\equiv P_I(t_1^{\sigma(x/\sigma(t))}, \dots, t_n^{\sigma(x/\sigma(t))}) && \text{lemma precedente} \\ &\equiv P(t_1, \dots, t_n)^{\sigma(x/\sigma(t))} && \text{definizione di valutazione} \\ &\equiv A^{\sigma(x/\sigma(t))} && \text{perchè } A \equiv P(t_1, \dots, t_n) \end{aligned}$$

- $A \equiv B \wedge C$, $A \equiv B \vee C$, $A \equiv \neg B$ o $A \equiv B \rightarrow C$. Questi casi sono abbastanza semplici e vengono lasciati al lettore come esercizio.

- $A \equiv \forall x. B$. Allora

$$\begin{aligned} A[x := t]^\sigma &\equiv (\forall x. B)^\sigma && \text{definizione di sostituzione} \\ &\equiv (\forall x. B)^{\sigma(x/\sigma(t))} && \text{la valutazione di una formula dipende solo dalle variabili che appaiono libere nella proposizione ed } x \text{ non appare libera in } \forall x. B \\ &\equiv A^{\sigma(x/\sigma(t))} && \text{perchè } A \equiv \forall x. B \end{aligned}$$

- $A \equiv \forall y. B$. Allora

$$\begin{aligned}
A[x := t]^\sigma &= (\forall z. B[y := z][x := t])^\sigma && \text{definizione di sostituzione} \\
&= \bigwedge_{u \in U} B[y := z][x := t]^{\sigma(z/u)} && \text{definizione di valutazione} \\
&= \bigwedge_{u \in U} B[y := z]^{\sigma(z/u)(x/\sigma(z/u)(t))} && \text{ipotesi induttiva} \\
&= \bigwedge_{u \in U} B[y := z]^{\sigma(z/u)(x/\sigma(t))} && z \text{ non appare in } t \\
&= \bigwedge_{u \in U} B^{\sigma(z/u)(x/\sigma(t))(y/\sigma(z/u)(x/\sigma(t))(z))} && \text{ipotesi induttiva} \\
&= \bigwedge_{u \in U} B^{\sigma(z/u)(x/\sigma(t))(y/u)} && \sigma[z/u][x/\sigma(t)](z) = u \\
&= \bigwedge_{u \in U} B^{\sigma(x/\sigma(t))(y/u)} && z \text{ non appare in } B \\
&= (\forall y. B)^{\sigma(x/\sigma(t))} && \text{definizione di valutazione} \\
&= A^{\sigma(x/\sigma(t))} && \text{perch  } A \equiv \forall y. B
\end{aligned}$$

- $A \equiv \exists x. B$ e $A \equiv \exists y. B$. Completamente analoghi al caso del quantificatore universale.

Questo lemma   tutto quello che ci serve per poter introdurre le regole sulle proposizioni con quantificatori.

Esercizio 6.2.16 Sia t un termine e I una interpretazione basata sulla valutazione σ delle variabili. Dimostrare per induzione sulla costruzione del termine t che la valutazione di t dipende solo dalla valutazione delle variabili che appaiono in t , vale a dire che due valutazioni che coincidono sulle variabili che appaiono in t valutano t nello stesso modo.

Esercizio 6.2.17 Utilizzando il risultato dell'esercizio precedente dimostrare che la valutazione della proposizione A dipende solamente dalla valutazione delle variabili che in A occorrono libere.

Esercizio 6.2.18 Siano A e B due proposizioni, x una variabile, t un termine e σ una valutazione delle variabili in una opportuna struttura. Dimostrare che $\neg A[x := t]^\sigma = \neg A^{\sigma(x/\sigma(t))}$ e che $(A \rightarrow B)[x := t]^\sigma = (A \rightarrow B)^{\sigma(x/\sigma(t))}$.

Quantificatore universale

Come al solito dobbiamo trovare due regole; la prima deve dirci come si fa a dimostrare che una certa proposizione vale per ogni elemento del dominio mentre l'altra deve dirci come usare l'assunzione che una proposizione quantificata universalmente sia vera.

Per quanto riguarda la prima regola possiamo riprendere la regola proposta nel capitolo 4 visto che in essa non compaiono elementi del dominio.

$$(\forall\text{-right})^* \quad \frac{\Gamma \vdash A}{\Gamma \vdash \forall x. A}$$

dove ricordiamo che la validit  della regola   sottoposta alla condizione che la variabile x non appaia libera in nessuna proposizione in Γ (abbiamo aggiunto un asterisco nel nome della regola proprio per ricordarci che si tratta di una regola sottoposta ad una condizione). Naturalmente la dimostrazione della sua validit  dovr  ora appoggiarsi alla nozione di interpretazione e sar  quindi diversa da quanto abbiamo fatto nel capitolo 4.

Vediamo allora che si tratta in effetti di una regola valida e che la condizione che x non appaia libera in nessuna proposizione in Γ non   solo una conseguenza della giustificazione della regola che abbiamo fornito sopra ma una vera necessit  se vogliamo evitare che la regola ci permetta di fare deduzioni non valide.

Per quanto riguarda la dimostrazione della validit  della regola basta far vedere che, se per ogni valutazione $V(-)$, accade che $V(\Gamma) \leq V(A)$ allora, per ogni valutazione $V^\sigma(-)$, si ha che $V^\sigma(\Gamma) \leq V^\sigma(\forall x. A)$, nell'ipotesi che x non appaia tra le variabili libere di alcuna proposizione in Γ .

Sia allora $V^\sigma(-)$ una qualsiasi interpretazione e consideriamo, per ogni elemento u appartenente al supporto U della struttura per l'interpretazione $V^\sigma(-)$, la nuova interpretazione $V^{\sigma(x/u)}$. Per ipotesi sappiamo che $V^{\sigma(x/u)}(\Gamma) \leq V^{\sigma(x/u)}(A)$ visto che $V(\Gamma) \leq V(A)$ vale per ogni interpretazione $V(-)$ (  interessante notare che usiamo l'ipotesi induttiva su una valutazione che *non*   la stessa per cui dobbiamo provare la conclusione). Tuttavia noi sappiamo che la valutazione di

una proposizione dipende solamente dalla valutazione delle variabili che vi appaiono libere e quindi $V^{\sigma(x/u)}(\Gamma) = V^{\sigma}(\Gamma)$ visto che, per ipotesi, x non appare libera in alcuna proposizione in Γ . Perciò, per ogni elemento $u \in U$, $V^{\sigma}(\Gamma) \leq V^{\sigma(x/u)}(A)$ ma allora ne segue, per definizione di infimo, che $V^{\sigma}(\Gamma) \leq \bigwedge_{u \in U} V^{\sigma(x/u)}(A) = V^{\sigma}(\forall x.A)$.

È ora facile verificare che la condizione sull'occorrenza delle variabili è necessaria. Infatti, in sua mancanza la seguente deduzione sarebbe valida

$$\frac{\text{Eq}(x, 0) \vdash \text{Eq}(x, 0)}{\text{Eq}(x, 0) \vdash \forall x. \text{Eq}(x, 0)}$$

anche se è chiaramente scorretta come si può facilmente verificare utilizzando una struttura su un universo con almeno due elementi e interpretando il segno $\text{Eq}(-, -)$ nella relazione identica.

La prova canonica corrispondente alla precedente regola è ovviamente la seguente

$$(\forall\text{-introduzione})^* \quad \frac{A}{\forall x.A}$$

dove la presenza dell'asterisco segnala il fatto che, come nel caso della regola $\forall\text{-right}$, la regola si può utilizzare solo se la variabile x non appare libera in alcuna ipotesi attiva.

Passiamo ora alla regola che determina il modo di usare una proposizione quantificata universalmente. L'idea in questo caso è che dalla verità di una proposizione quantificata universalmente si può dedurre la verità di tutti i suoi esempi particolare, e quindi se qualcosa è deducibile da questi esempi lo è anche dalla proposizione quantificata universalmente. La regola che esprime questa considerazione è la seguente

$$(\forall\text{-left}) \quad \frac{\Gamma, A[x := t] \vdash C}{\Gamma, \forall x.A \vdash C}$$

dove t è un qualunque termine del linguaggio. Si tratta quindi di una regola molto simile a quella che avevamo proposto nel capitolo 4; la principale differenza sta nel fatto che in A non possiamo sostituire al posto della variabile x (un ente linguistico) un oggetto a del dominio (un ente matematico) e quindi sostituiamo invece il suo nome (di nuovo un ente linguistico). Naturalmente in questo modo corriamo il rischio di perdere dimostrazioni che prima erano possibili visto che non è detto che ogni oggetto matematico abbia un nome nel linguaggio (si consideri ad esempio il caso di una teoria che voglia parlare dei numeri reali: questi sono una quantità più che numerabile mentre i nomi presenti in una teoria *ragionevole* sono una quantità numerabile).

Verifichiamo comunque la correttezza di questa nuova regola. Supponiamo quindi di avere una valutazione $V^{\sigma}(-)$. Allora, per ipotesi induttiva, sappiamo che $V^{\sigma}(\Gamma) \wedge V^{\sigma}(A[x := t]) \leq V^{\sigma}(C)$. Ma il lemma 6.2.14 ci garantisce che $V^{\sigma}(A[x := t]) = V^{\sigma(x/V^{\sigma}(t))}(A)$ e quindi $V^{\sigma}(\forall x.A) \leq V^{\sigma}(A[x := t])$ segue immediatamente visto che $V^{\sigma}(\forall x.A) = \bigwedge_{u \in U} V^{\sigma(x/u)}(A) \leq V^{\sigma(x/V^{\sigma}(t))}(A)$. Quindi $V^{\sigma}(\Gamma) \wedge V^{\sigma}(\forall x.A) \leq V^{\sigma}(C)$ segue facilmente.

Per concludere questa sezione dobbiamo solo decidere la forma canonica della prova la cui esistenza è prevista dalla regola $\forall\text{-left}$. Visto che tutto ciò che dobbiamo dire è che possiamo derivare da $\forall x.A$ tutto ciò che riusciamo a derivare dalle sue istanze ci basta la seguente regola

$$(\forall\text{-eliminazione}) \quad \frac{\forall x.A}{A[x := t]}$$

dove t è un qualunque termine del linguaggio.

Esercizio 6.2.19 *Dimostrare, utilizzando le nuove regole, che*

$$(\forall x.(A(x) \rightarrow B(x))) \rightarrow ((\forall x.A(x)) \rightarrow (\forall x.B(x)))$$

Trovare invece una struttura che fornisca un controesempio per l'implicazione opposta.

Esercizio 6.2.20 *Verificare che la regola $\forall\text{-eliminazione}$ è equivalente alla seguente regola ottenuta direttamente da $\forall\text{-left}$:*

$$\frac{A[x := t]_1 \quad \vdots \quad C}{C} 1$$

dove t è un qualunque termine del linguaggio.

Quantificatore esistenziale

Passiamo ora al quantificatore esistenziale. Il lavoro fatto per il quantificatore universale non sarà del tutto inutile per affrontare il caso di questo nuovo quantificatore visto che ci suggerisce come trattare con le variabili.

Cominciamo quindi con la regola che ci permette di provare una proposizione quantificata esistenzialmente. L'idea è naturalmente di modificare la regola che avevamo proposto nel capitolo 4 sostituendo l'occorrenza di un oggetto con il suo nome. La nuova regola è quindi

$$(\exists\text{-right}) \quad \frac{\Gamma \vdash A[x := t]}{\Gamma \vdash \exists x.A}$$

dove t è un qualche termine del linguaggio. Il problema con questa nuova regola sta naturalmente nel fatto che non è detto che il *testimone* che ci garantisce che un qualche elemento del dominio soddisfa la proposizione A potrebbe non avere un nome.

Vediamo comunque che questa regola è valida. Supponiamo quindi che $V^\sigma(-)$ sia una valutazione in una qualche struttura. Allora, per ipotesi induttiva, sappiamo che $V^\sigma(\Gamma) \leq V^\sigma(A[x := t])$ e quindi possiamo subito concludere che $V^\sigma(\Gamma) \leq V^\sigma(\exists x.A)$ perché, utilizzando il lemma 6.2.14, vediamo subito che $V^\sigma(A[x := t]) = V^{\sigma(x/V^\sigma(t))}(A) \leq \bigvee_{u \in U} V^{\sigma(x/u)}(A) = V^\sigma(\exists x.A)$.

Quale è ora la prova canonica per dimostrare una proposizione quantificata esistenzialmente? Guardando la regola $\exists$ -right abbiamo subito la risposta:

$$(\exists\text{-introduzione}) \quad \frac{A[x := t]}{\exists x.A}$$

dove t è un qualche termine del linguaggio.

È adesso il momento di vedere come si utilizza come assunzione una proposizione quantificata esistenzialmente. A questo scopo possiamo ricordare la regola che abbiamo proposto nel capitolo 4 e vedere se essa è valida nel nuovo approccio.

$$(\exists\text{-left})^* \quad \frac{\Gamma, A \vdash C}{\Gamma, \exists x.A \vdash C}$$

dove la presenza dell'asterisco serve a ricordarci che la validità della regola è sottoposta alla condizione che la variabile x non appaia libera né in C né in alcuna proposizione in Γ (il solito modo in cui il calcolo riesce ad esprimere il fatto che non abbiamo ulteriori conoscenze sull'elemento x).

Supponiamo dunque che $V^\sigma(-)$ sia una qualunque valutazione e consideriamo un qualunque elemento $u \in U$, dove U è il dominio della struttura per la valutazione $V^\sigma(-)$. Allora, per ipotesi induttiva, sappiamo che $V^{\sigma(x/u)}(\Gamma) \wedge V^{\sigma(x/u)}(A) \leq V^{\sigma(x/u)}(C)$ ma l'ipotesi che la variabile x non appaia né in C né in alcuna proposizione in Γ ci permette di dedurre che $V^\sigma(\Gamma) \wedge V^{\sigma(x/u)}(A) \leq V^\sigma(C)$ (vale la pena di notare che ancora una volta, come per la regola $\forall$ -right, abbiamo utilizzato l'ipotesi induttiva su una valutazione che non è quella per la quale stiamo cercando di ottenere la dimostrazione di validità).

Il modo più veloce per concludere è ora utilizzare le proprietà dell'implicazione in un'algebra di Heyting (e quindi in un'algebra di Boole).

Infatti, se $V^\sigma(\Gamma) \wedge V^{\sigma(x/u)}(A) \leq V^\sigma(C)$ allora $V^{\sigma(x/u)}(A) \leq V^\sigma(\Gamma) \rightarrow V^\sigma(C)$ e quindi, per definizione di supremo, $V^\sigma(\exists x.A) = \bigvee_{u \in U} V^{\sigma(x/u)}(A) \leq V^\sigma(\Gamma) \rightarrow V^\sigma(C)$ per cui, sempre per le proprietà dell'implicazione, $V^\sigma(\Gamma) \wedge V^\sigma(\exists x.A) \leq V^\sigma(C)$.

Per finire dobbiamo solo decidere quale sia la prova canonica associata alla regola $\exists$ -left. Vista l'esperienza acquisita con le regole precedenti è immediato verificare che la regola cercata è la seguente:

$$(\exists\text{-eliminazione})^* \quad \frac{\begin{array}{c} [A]_n \\ \vdots \\ \exists x.A \quad C \end{array}}{C} \quad n$$

dove la presenza dell'asterisco serve per ricordarci, come per la regola $\exists$ -left, che la variabile x non deve apparire libera né in C né in alcuna assunzione ancora attiva nella premessa destra della regola.

Esercizio 6.2.21 Dimostrare utilizzando le nuove regole che se $\forall x.(A \rightarrow B)$ allora $(\exists x.A) \rightarrow (\exists x.B)$.

Esercizio 6.2.22 Dimostrare utilizzando le nuove regole classiche che se vale $A \rightarrow \exists x.B$ allora $\exists x.(A \rightarrow B)$.

Esercizio 6.2.23 Dimostrare utilizzando le nuove regole classiche che $(\forall x.A) \rightarrow B$ se e solo se $\exists x.(A \rightarrow B)$ nell'ipotesi che x non appaia libera in B . Dedurne la validità classica del principio del bevitore, i.e. $\exists x.(A \rightarrow \forall x.A)$ e esibirne un contro-esempio intuizionista.

Esercizio 6.2.24 Si consideri la seguente mappa dell'insieme delle proposizioni in sè (traduzione di Gödel).

$$\begin{array}{ll}
 (\perp)^* &= \perp & (P(t_1, \dots, t_n))^* &= \neg\neg P(t_1, \dots, t_n) \\
 (A \wedge B)^* &= A^* \wedge B^* & (A \vee B)^* &= \neg(\neg A^* \wedge \neg B^*) \\
 (A \rightarrow B)^* &= A^* \rightarrow B^* & (\neg A)^* &= \neg A^* \\
 (\forall x.A)^* &= \forall x.A^* & (\exists x.A)^* &= \neg(\forall x.\neg A^*)
 \end{array}$$

Dimostrare che $\Gamma \vdash A$ è dimostrabile classicamente se e solo se $\Gamma^* \vdash A^*$ è dimostrabile intuizionisticamente, dove Γ^* è il risultato dell'applicazione della mappa $(-)^*$ a tutte le proposizioni in Γ .

Capitolo 7

Il teorema di completezza

7.1 Introduzione

Lo scopo di questo capitolo è quello di fornire una prova dei teoremi di adeguatezza dei calcoli che abbiamo presentato nel precedente capitolo, intendiamo cioè far vedere che essi sono sufficienti per scoprire tutte le proposizioni vere in ogni struttura sia nel caso intuizionista che in quello classico.

Inizieremo la prova concentrandoci sul caso intuizionista e mostreremo poi le modifiche necessarie per dimostrare anche la completezza del calcolo classico.

7.2 Il *quasi* teorema di completezza

Iniziamo ricordando le regole del calcolo predicativo intuizionista che abbiamo introdotto nel capitolo precedente (si veda [Takeuti 75]):

(axioms)	$\frac{}{\Gamma, A \vdash A}$		
(indebolimento)	$\frac{\Gamma \vdash C}{\Gamma, A \vdash C}$	(taglio)	$\frac{\Gamma \vdash A \quad \Gamma, A \vdash B}{\Gamma \vdash B}$
($\perp$ -left)	$\frac{}{\Gamma, \perp \vdash A}$	($\top$ -right)	$\frac{}{\Gamma \vdash \top}$
($\wedge$ -left)	$\frac{\Gamma, A, B \vdash C}{\Gamma, A \wedge B \vdash C}$	($\wedge$ -right)	$\frac{\Gamma \vdash A \quad \Gamma \vdash B}{\Gamma \vdash A \wedge B}$
($\vee$ -left)	$\frac{\Gamma, A \vdash C \quad \Gamma, B \vdash C}{\Gamma, A \vee B \vdash C}$	($\vee$ -right)	$\frac{\Gamma \vdash A \quad \Gamma \vdash B}{\Gamma \vdash A \vee B}$
($\rightarrow$ -left)	$\frac{\Gamma \vdash A \quad \Gamma, B \vdash C}{\Gamma, A \rightarrow B \vdash C}$	($\rightarrow$ -right)	$\frac{\Gamma, A \vdash B}{\Gamma \vdash A \rightarrow B}$
($\forall$ -left)	$\frac{\Gamma, A[x := t] \vdash C}{\Gamma, \forall x. A \vdash C}$	($\forall$ -right*)	$\frac{\Gamma \vdash A}{\Gamma \vdash \forall x. A}$
($\exists$ -left*)	$\frac{\Gamma, A \vdash C}{\Gamma, \exists x. A \vdash C}$	($\exists$ -right)	$\frac{\Gamma \vdash A[x := t]}{\Gamma \vdash \exists x. A}$

dove la variabile quantificata non deve apparire libera nella conclusione sia nel caso della regola di $\forall$ -right che in quello della $\exists$ -left.

Nei capitoli precedenti abbiamo visto che possiamo interpretare validamente questo calcolo nel senso che se $\Gamma \vdash A$ è dimostrabile allora $V_I^\sigma(\Gamma) \leq V^\sigma(A)$ vale in ogni struttura basata su una algebra di Heyting completa per ogni interpretazione $V_I^\sigma(-)$. La questione che non abbiamo risolto è però se valga anche l'altra implicazione, se cioè dal fatto che $V_I^\sigma(\Gamma) \leq V^\sigma(A)$ valga per ogni struttura basata su una algebra di Heyting completa e per ogni interpretazione $V_I^\sigma(-)$ sia possibile dedurre che $\Gamma \vdash A$ è dimostrabile con le regole che siamo stati in grado di escogitare finora.

Un primo tentativo che porta *quasi* alla soluzione di questo problema è il seguente.

Tanto per cominciare ricordiamo che ponendo

$$A \dashv\vdash B \equiv A \vdash B \text{ e } B \vdash A$$

otteniamo una relazione di equivalenza tra le proposizioni del calcolo predicativo intuizionista (si veda il capitolo 5). Quindi si possono costruire delle classi di equivalenza ponendo

$$[A]_{\dashv\vdash} \equiv \{B \mid A \dashv\vdash B\}$$

e otteniamo una quantità numerabile di classi di equivalenza sull'insieme delle proposizioni (nell'ipotesi che sia numerabile il numero delle proposizioni). Possiamo ora atteggiare ad algebra di Heyting $\mathcal{P}$ tale insieme di classi di equivalenza ponendo (vedi esercizio 7.2.2):

$$\begin{aligned} O_{\mathcal{P}} &\equiv [\perp]_{\dashv\vdash} \\ 1_{\mathcal{P}} &\equiv [\perp \rightarrow \perp]_{\dashv\vdash} \\ [A]_{\dashv\vdash} \wedge_{\mathcal{P}} [B]_{\dashv\vdash} &\equiv [A \wedge B]_{\dashv\vdash} \\ [A]_{\dashv\vdash} \vee_{\mathcal{P}} [B]_{\dashv\vdash} &\equiv [A \vee B]_{\dashv\vdash} \\ [A]_{\dashv\vdash} \rightarrow_{\mathcal{P}} [B]_{\dashv\vdash} &\equiv [A \rightarrow B]_{\dashv\vdash} \end{aligned}$$

e avremo che

$$[A]_{\dashv\vdash} \leq [B]_{\dashv\vdash} \text{ se e solo se } A \vdash B$$

Inoltre esistono una quantità numerabile di supremi e infimi in corrispondenza dei quantificatori esistenziale e universale (vedi esercizio 7.2.3):

$$\begin{aligned} \bigwedge_{t \in \text{Term}} [A[x := t]]_{\dashv\vdash} &\equiv [\forall x. A]_{\dashv\vdash} \\ \bigvee_{t \in \text{Term}} [A[x := t]]_{\dashv\vdash} &\equiv [\exists x. A]_{\dashv\vdash} \end{aligned}$$

dove Term è l'insieme dei termini del linguaggio che stiamo considerando.

È ora possibile definire una struttura $\mathcal{D}_{\mathcal{P}}$ basata sull'algebra di Heyting $\mathcal{P}$ ponendo

$$\mathcal{D}_{\mathcal{P}} = \langle \text{Term}, \mathcal{R}_{\mathcal{P}}, \mathcal{F}_{\mathcal{P}}, \mathcal{C}_{\mathcal{P}} \rangle + \tau_{\mathcal{P}}$$

dove

- Term è l'insieme dei termini del linguaggio che stiamo considerando,
- $\mathcal{R}_{\mathcal{P}}$ è l'insieme delle funzioni R_P da Term^n verso $\mathcal{P}$, definita in corrispondenza di ogni simbolo predicativo n -rio P del linguaggio ponendo $R_P(t_1, \dots, t_n) = [P(t_1, \dots, t_n)]$,
- $\mathcal{F}_{\mathcal{P}}$ è l'insieme delle funzioni F_f da Term^n verso Term , definite in corrispondenza di ogni segno di funzione n -rio f del linguaggio ponendo $F_f(t_1, \dots, t_n) = f(t_1, \dots, t_n)$
- $\mathcal{C}_{\mathcal{P}}$ è l'insieme delle costanti c del linguaggio.

Possiamo ora definire una valutazione $V^*(-)$ delle proposizioni nella struttura $\mathcal{D}_{\mathcal{P}}$ introducendo l'interpretazione I^* definita ponendo

- $I^*(c) \equiv c$ per ogni costante c
- $I^*(f) \equiv F_f$ per ogni segno di funzione f
- $I^*(P) \equiv R_P$ per ogni simbolo di predicato P .

Se ora supponiamo di porre, per ogni variabile x del linguaggio, $\sigma^*(x) \equiv x$ allora, per ogni termine $t \in \text{Term}$, otteniamo che $V^*(t) = t$ (vedi esercizio 7.2.4). Quindi possiamo definire una valutazione $V^*(-)$ su tutte le proposizioni ponendo

$$V^*(A) \equiv [A]_{\dashv\vdash}$$

Infatti è facile verificare, utilizzando una prova per induzione sulla costruzione delle formule, che con questa definizione otteniamo davvero una valutazione delle formule. Infatti tutti i passaggi necessari sono facili (vedi esercizio 7.2.5), e per questo motivo noi verificheremo qui solo i casi in cui entrano in gioco i termini e le variabili, vale a dire il caso dei predicati atomici e delle proposizioni quantificate (si noti che per ogni formula A e per ogni termine t , $V^*(A[x := t]) = V^{\sigma^*(x/t)}(A)$).

$$\begin{aligned} V^*(P(t_1, \dots, t_n)) &= [P(t_1, \dots, t_n)]_{\dashv\vdash} \\ &= I^*(P)(t_1, \dots, t_n) \\ &= I^*(P)(V^*(t_1), \dots, V^*(t_n)) \end{aligned}$$

$$\begin{aligned} V^*(\exists x.A) &= [\exists x.A]_{\dashv\vdash} \\ &= \bigvee_{t \in \text{Term}} [A[x := t]]_{\dashv\vdash} \\ &= \bigvee_{t \in \text{Term}} V^*(A[x := t]) \\ &= \bigvee_{t \in \text{Term}} V^{\sigma^*(x/t)}(A) \end{aligned}$$

$$\begin{aligned} V^*(\forall x.A) &= [\forall x.A]_{\dashv\vdash} \\ &= \bigwedge_{t \in \text{Term}} [A[x := t]]_{\dashv\vdash} \\ &= \bigwedge_{t \in \text{Term}} V^*(A[x := t]) \\ &= \bigwedge_{t \in \text{Term}} V^{\sigma^*(x/t)}(A) \end{aligned}$$

Possiamo finalmente completare la nostra *quasi* prova del teorema di completezza. Infatti vale la seguente sequenza di implicazioni

$$\begin{aligned} V^*(A_1) \wedge \dots \wedge V^*(A_n) \leq V^*(B) &\quad \text{se e solo se} \quad [A_1]_{\dashv\vdash} \wedge \dots \wedge [A_n]_{\dashv\vdash} \leq [B]_{\dashv\vdash} \\ &\quad \text{se e solo se} \quad [A_1 \wedge \dots \wedge A_n]_{\dashv\vdash} \leq [B]_{\dashv\vdash} \\ &\quad \text{se e solo se} \quad A_1 \wedge \dots \wedge A_n \vdash B \\ &\quad \text{se e solo se} \quad A_1, \dots, A_n \vdash B \end{aligned}$$

Quindi se in ogni struttura e per ogni valutazione $V(-)$ vale che $V(A_1) \wedge \dots \wedge V(A_n) \leq V(B)$ questo varrà in particolare per la struttura $\mathcal{D}_{\mathcal{P}}$ e per la valutazione $V^*(-)$ che abbiamo appena definito, cioè sarà vero che $V^*(A_1) \wedge \dots \wedge V^*(A_n) \leq V^*(B)$ ma, come abbiamo appena visto, questo implica che $A_1, \dots, A_n \vdash B$. Abbiamo quindi *quasi* provato il seguente teorema.

Teorema 7.2.1 (Quasi completezza) *Se, per ogni struttura $\mathcal{D}$ e per ogni valutazione $V(-)$ delle formule del calcolo predicativo intuizionista in $\mathcal{D}$, $V(A_1) \wedge \dots \wedge V(A_n) \leq V(B)$ allora $A_1, \dots, A_n \vdash B$.*

Naturalmente il problema con questa *quasi* dimostrazione è che la struttura $\mathcal{D}_{\mathcal{P}}$ non è basata su una algebra di Heyting completa visto che l'algebra delle classi di equivalenza delle proposizioni ci garantisce l'esistenza dei supremi e degli infimi solo in corrispondenza delle proposizioni quantificate e non di un qualunque insieme di classi di equivalenza e per questo motivo non può essere considerata tra le strutture rilevanti per la dimostrazione del teorema di completezza. Quindi per risolvere questo problema dobbiamo escogitare un metodo che ci permetta di *completare* l'algebra delle classi di equivalenza delle proposizioni.

Esercizio 7.2.2 *Verificare che le operazioni definite sull'insieme delle classi di equivalenza sull'insieme delle proposizioni $\mathcal{P}$ utilizzando la congruenza $\dashv\vdash$ definiscono una algebra di Heyting.*

Esercizio 7.2.3 *Verificare che le operazioni definite sull'insieme delle classi di equivalenza sull'insieme delle proposizioni $\mathcal{P}$ utilizzando la congruenza $\dashv\vdash$ definite ponendo*

$$\begin{aligned} \bigwedge_{t \in \text{Term}} [A[x := t]]_{\dashv\vdash} &\equiv [\forall x.A]_{\dashv\vdash} \\ \bigvee_{t \in \text{Term}} [A[x := t]]_{\dashv\vdash} &\equiv [\exists x.A]_{\dashv\vdash} \end{aligned}$$

definiscono un infimo e un supremo (sugg.: nella dimostrazione che se $C \vdash A[x := t]$ vale per ogni termine t allora vale anche $C \vdash \forall x.A$ bisogna osservare che possiamo sempre trovare una qualche variabile y che non compare in C e quindi sicuramente abbiamo che $C \vdash A[x := y]$ da cui è possibile ottenere $C \vdash \forall y.A[x := y]$, visto che la condizione sulle variabili è rispettata, che ci permette di concludere $C \vdash \forall x.A$ per taglio visto che $\forall y.A[x := y] \vdash \forall x.A$).

Esercizio 7.2.4 Verificare che per ogni termine t del linguaggio $V^*(t) = t$.

Esercizio 7.2.5 Verificare che $V^*(-)$ è una valutazione delle formule costruite con i connettivi proposizionali.

7.3 Valutazione in uno spazio topologico

Il paradigma di algebra di Heyting completa è costituito dagli spazi topologici; quindi possiamo provare a completare una algebra di Heyting immergendola in uno spazio topologico.

Sappiamo già che, preso un qualsiasi spazio topologico τ , possiamo definire una interpretazione valida $V(-)$ che manda le proposizioni in aperti di τ in modo tale che se $\Gamma \vdash B$ allora $V(\Gamma) \subseteq V(B)$ (di fatto possiamo definire una interpretazione valida in ogni algebra di Heyting completa e gli aperti di uno spazio topologico sono un caso speciale di tale situazione generale).

Infatti, se consideriamo una qualunque struttura $\mathcal{D} \equiv \langle D, \mathcal{R}, \mathcal{F}, \mathcal{C} \rangle + \tau$, adatta per il linguaggio $\mathcal{L}$ in cui sono scritte le proposizioni che vogliamo interpretare e dove τ è lo spazio topologico basato su un insieme X di punti sui cui aperti interpreteremo le formule, allora possiamo definire l'interpretazione $V(-)$ nel modo che segue

$$\begin{aligned} V(P(t_1, \dots, t_m)) &= V(P)(V(t_1), \dots, V(t_m)) \\ V(\top) &= X \\ V(\perp) &= \emptyset \\ V(A \wedge B) &= V(A) \cap V(B) \\ V(A \vee B) &= V(A) \cup V(B) \\ V(A \rightarrow B) &= \bigcup \{ \mathcal{O} \in \tau \mid V(A) \cap \mathcal{O} \subseteq V(B) \} \\ V(\forall x.A) &= \text{Int}(\bigcap_{d \in D} V^{[x:=d]}(A)) \\ V(\exists x.A) &= \bigcup_{d \in D} V^{[x:=d]}(A) \end{aligned}$$

dove $V(P)$ è una tra le mappe in $\mathcal{R}$ e con $V^{[x:=d]}$ intendiamo indicare l'interpretazione che coincide con $V(-)$ eccetto per il fatto che la variabile x viene interpretata nell'elemento d del dominio D .

Come sappiamo bene, si può ora dimostrare, per induzione sulla sua complessità, che la valutazione di una qualsiasi proposizione A dipende solo dalla valutazione delle variabili che vi appaiono libere e che, per ogni termine t , $V(A[x := t]) = V^{[x:=V(t)]}(A)$.

Teorema 7.3.1 (Validità) Per ogni valutazione $V(-)$ delle proposizioni di un calcolo intuizionista in una struttura $\mathcal{D}$, se $\Gamma \vdash B$ allora $V(\Gamma) \subseteq V(B)$.

Dimostrazione. La prova si ottiene per induzione sulla lunghezza (finita) della derivazione di $\Gamma \vdash B$. La maggior parte dei casi da considerare sono banali e si ottengono come nel caso generale che abbiamo visto nel precedente capitolo. Mostriamo qui esplicitamente solo il caso delle regole per il quantificatore universale.

Supponiamo quindi che $\Gamma, A[x := t] \vdash C$. Allora, per ipotesi induttiva, otteniamo che, per ogni valutazione $V(-)$, $V(\Gamma) \cap V(A[x := t]) \subseteq V(C)$. Quindi $V(\Gamma) \cap V(\forall x.A) \subseteq V(C)$ segue visto che $V(\forall x.A) = \text{Int}(\bigcap_{d \in D} V^{[x:=d]}(A)) \subseteq \bigcap_{d \in D} V^{[x:=d]}(A) \subseteq V^{[x:=V(t)]}(A) = V(A[x := t])$.

D'altra parte, se $\Gamma \vdash A$ allora, per ipotesi induttiva, otteniamo che $V^{[x:=d]}(\Gamma) \subseteq V^{[x:=d]}(A)$ per ogni $d \in D$. Ma $V^{[x:=d]}(\Gamma) = V(\Gamma)$, visto che, per ipotesi, x non appare in Γ , e quindi $V(\Gamma) \subseteq \bigcap_{d \in D} V^{[x:=d]}(A)$; perciò $V(\Gamma) \subseteq \text{Int}(\bigcap_{d \in D} V^{[x:=d]}(A)) = V(\forall x.A)$ giacché $V(\Gamma)$ è un aperto.

È interessante notare subito che nell'interpretazione del connettivo di implicazione non è necessario considerare tutti gli aperti dello spazio topologico τ ma è sufficiente limitarsi a quelli di una sua base $\mathcal{B}_\tau$ visto che possiamo porre

$$V(A \rightarrow B) = \bigcup \{b \in \mathcal{B}_\tau \mid V(A) \cap b \subseteq V(B)\}$$

giacché $\bigcup \{\mathcal{O} \in \tau \mid V(A) \cap \mathcal{O} \subseteq V(B)\} = \bigcup \{b \in \mathcal{B}_\tau \mid V(A) \cap b \subseteq V(B)\}$. Infatti per dimostrare che $\bigcup \{\mathcal{O} \in \tau \mid V(A) \cap \mathcal{O} \subseteq V(B)\} \subseteq \bigcup \{b \in \mathcal{B}_\tau \mid V(A) \cap b \subseteq V(B)\}$ basta notare che se $x \in \bigcup \{\mathcal{O} \in \tau \mid V(A) \cap \mathcal{O} \subseteq V(B)\}$ allora esiste un aperto $\mathcal{O} \in \tau$ tale che $x \in \mathcal{O}$ e $V(A) \cap \mathcal{O} \subseteq V(B)$; ma $\mathcal{O} = \bigcup \{b \in \mathcal{B}_\tau \mid b \subseteq \mathcal{O}\}$ e quindi esiste un aperto $b \in \mathcal{B}_\tau$ tale che $x \in b$ e $b \subseteq \mathcal{O}$, quindi $V(A) \cap b \subseteq V(A) \cap \mathcal{O} \subseteq V(B)$, da cui si deduce che $x \in \bigcup \{b \in \mathcal{B}_\tau \mid V(A) \cap b \subseteq V(B)\}$; l'altra inclusione è ovvia visto che ogni aperto della base $\mathcal{B}_\tau$ è un aperto di τ .

Analoga situazione si verifica per quanto riguarda l'interpretazione del quantificatore universale. Infatti se U è un sottoinsieme dell'insieme X di punti dello spazio topologico τ e $\mathcal{B}_\tau$ è una base per τ , allora $\text{Int}(U) = \bigcup \{\mathcal{O} \in \tau \mid \mathcal{O} \subseteq U\} = \bigcup \{b \in \mathcal{B}_\tau \mid b \subseteq U\}$ si può dimostrare in modo analogo a quanto fatto qui sopra per l'implicazione. Quindi possiamo porre direttamente

$$V(\forall x.A) = \bigcup \{b \in \mathcal{B}_\tau \mid b \subseteq \bigcap_{d \in D} V^{[x:=d]}(A)\}$$

7.4 Un teorema alla Rasiowa-Sikorski

In questa sezione studieremo i risultati algebrici che ci permetteranno di ottenere il teorema di completezza per il calcolo predicativo intuizionista. Dimostreremo in particolare che ogni algebra di Heyting numerabile si può immergere in una opportuna algebra di Heyting costruita su uno spazio topologico, e quindi completa, in maniera tale che una quantità numerabile di supremi venga rispettata.

Iniziamo quindi ricordando la definizione di algebra di Heyting che avevamo proposto nel capitolo 5: una algebra di Heyting $\mathcal{H} \equiv \langle H, 0, 1, \vee, \wedge, \rightarrow \rangle$ è una struttura tale che $\langle H, 0, 1, \vee, \wedge \rangle$ è un reticolo distributivo con minimo elemento 0 e massimo elemento 1 che possieda una implicazione, cioè, per ogni $y, z \in H$, un elemento $y \rightarrow z$ tale che, per ogni $x \in H$, $x \wedge y \leq z$ se e solo se $x \leq y \rightarrow z$ dove intendiamo che la relazione d'ordine $\leq$ sia quella definita ponendo $x \leq y \equiv (x = x \wedge y)$.

Si noti che la presenza dell'implicazione permette di dimostrare immediatamente che l'operazione $\wedge$ di minimo si distribuisce su tutti i supremi esistenti e non solo su quelli finiti. Infatti, se supponiamo che $T \subseteq H$ abbia supremo in H allora $x \wedge \bigvee \{t : t \in T\}$ è il supremo di $\{x \wedge t : t \in T\}$, cioè anche il sottoinsieme $\{x \wedge t : t \in T\}$ ha supremo e tale supremo coincide con $x \wedge \bigvee \{t : t \in T\}$. Infatti, per ogni $t \in T$, $t \leq \bigvee \{t : t \in T\}$ e quindi $x \wedge t \leq x \wedge \bigvee \{t : t \in T\}$, cioè $x \wedge \bigvee \{t : t \in T\}$ è un maggiorante per tutti gli elementi dell'insieme $\{x \wedge t : t \in T\}$. Inoltre se c'è un maggiorante di tutti gli elementi di $\{x \wedge t : t \in T\}$, cioè se, per ogni $t \in T$, $x \wedge t \leq c$, allora $t \leq x \rightarrow c$ e quindi $\bigvee \{t : t \in T\} \leq x \rightarrow c$ che implica che $x \wedge \bigvee \{t : t \in T\} \leq c$, cioè $x \wedge \bigvee \{t : t \in T\}$ è il minimo dei maggioranti degli elementi in $\{x \wedge t : t \in T\}$.

L'esempio canonico di algebra di Heyting (completa) si ottiene considerando la famiglia degli aperti di un qualsiasi spazio topologico τ su un insieme X , la cui base sia $\mathcal{B}_\tau$, ponendo:

$$\begin{aligned} 0_\tau &\equiv \emptyset \\ 1_\tau &\equiv X \\ \mathcal{O}_1 \wedge_\tau \mathcal{O}_2 &\equiv \mathcal{O}_1 \cap \mathcal{O}_2 \\ \mathcal{O}_1 \vee_\tau \mathcal{O}_2 &\equiv \mathcal{O}_1 \cup \mathcal{O}_2 \\ \mathcal{O}_1 \rightarrow_\tau \mathcal{O}_2 &\equiv \bigcup \{\mathcal{O} \in \tau : \mathcal{O}_1 \cap \mathcal{O} \subseteq \mathcal{O}_2\} = \bigcup \{b \in \mathcal{B}_\tau : \mathcal{O}_1 \cap b \subseteq \mathcal{O}_2\} \end{aligned}$$

Lo scopo del nostro prossimo teorema è quello di far vedere che ogni algebra di Heyting $\mathcal{H}$ numerabile si può immergere in un algebra di Heyting costruita su uno spazio topologico.

Il primo problema da affrontare è ovviamente quello di trovare una qualche collezione di *punti* $\text{Pt}(\mathcal{H})$ su cui definire tale spazio topologico. L'idea per la soluzione di questo problema viene dal fatto di considerare gli elementi di $\mathcal{H}$ come informazioni sui *punti* che vogliamo definire (si veda l'esercizio 7.4.6) e di considerare come *punto* una qualsiasi collezione F di informazioni che abbia le caratteristiche delle collezioni di informazioni che parlano davvero di qualcosa. Allora l'elemento

massimo 1 di $\mathcal{H}$, che corrisponde all'informazione sempre vera, deve appartenere ad F mentre l'elemento minimo 0, che corrisponde all'informazione sempre falsa, non deve appartenere ad F . Inoltre se x e y sono informazioni valide per un certo oggetto F allora anche la loro congiunzione deve essere una informazione valida per lo stesso oggetto, e quindi se $x \in F$ e $y \in F$ allora $x \wedge y \in F$; infine se x è una informazione valida per un qualche oggetto e y è comunque almeno tanto vera che x allora anche y deve essere una informazione valida per lo stesso oggetto, cioè se $x \in F$ e $x \leq y$ allora $y \in F$.

Arriviamo quindi in modo naturale alla definizione di filtro coerente di un'algebra di Heyting.

Definition 7.4.1 (Filtro coerente) *Sia $\mathcal{H}$ una algebra di Heyting. Allora un sottoinsieme F di H è un filtro coerente se:*

$$1 \in F \quad 0 \notin F \quad \frac{x \in F \quad x \leq y}{y \in F} \quad \frac{x \in F \quad y \in F}{x \wedge y \in F}$$

I filtri mancano però di una proprietà importante per caratterizzare i *punti*, vale a dire il fatto che un punto non si può spezzare, cioè se un oggetto soddisfa la disgiunzione tra due proprietà è perché ne soddisfa una. La definizione di filtro coerente va quindi raffinata in quella di filtro primo.

Definition 7.4.2 (Filtro primo) *Sia $\mathcal{H}$ una algebra di Heyting e F un suo filtro. Allora F è un filtro primo se quando $x \vee y \in F$ accade che $x \in F$ o $y \in F$.*

Per ottenere il nostro risultato ci basta concentrarci su una particolare famiglia di filtri primi, vale a dire quelli che *rispettano* una famiglia numerabile di sottoinsiemi di H .

Definition 7.4.3 (Filtro che rispetta un sottoinsieme) *Sia $\mathcal{H}$ una algebra di Heyting, F uno dei suoi filtri e T un sottoinsieme di H dotato di supremo in $\mathcal{H}$. Allora F rispetta T se quando $\bigvee T \in F$ esiste $b \in T$ tale che $b \in F$.*

Il nostro scopo è ora far vedere che in ogni algebra di Heyting esistono filtri coerenti che rispettano una quantità numerabile di sottoinsiemi. In realtà il prossimo teorema mostra non solo che vale tale risultato ma riesce a costruire un filtro che soddisfa anche una ulteriore condizione.

Teorema 7.4.4 *Sia $\mathcal{H}$ una algebra di Heyting e x e y siano due elementi di H tali che $x \not\leq y$. Sia inoltre $T_1, \dots, T_n, \dots$ una famiglia numerabile di sottoinsiemi di H dotati di supremo. Allora esiste un filtro F di $\mathcal{H}$ che rispetta tutti i sottoinsiemi $T_1, \dots, T_n, \dots$ e tale che $x \in F$ e $y \notin F$.*

Dimostrazione. Costruiremo il filtro richiesto con una procedura che richiede una quantità numerabile di passi partendo dal filtro (esercizio: dimostrare che si tratta di un filtro)

$$F_0 \equiv \uparrow\{c_0\} (= \{z \in H \mid c_0 \leq z\})$$

dove $c_0 \equiv x$, che in effetti contiene x e, nelle ipotesi del teorema, non contiene y , estendendolo via via ad un filtro che rispetta tutti i sottoinsiemi $T_1, \dots, T_n, \dots$.

La prima cosa da fare per essere sicuri di ottenere il risultato desiderato è quella di costruire una nuova lista, ancora numerabile, che contenga ciascuno dei sottoinsiemi $T_1, \dots, T_n, \dots$ una quantità numerabile di volte; infatti può ben succedere che il filtro costruito al passo n rispetti il sottoinsieme T_k per il semplice motivo che il suo supremo non gli appartiene mentre esso potrebbe entrare in un passo successivo e quindi dopo ogni passo deve essere possibile verificare per ogni sottoinsieme T_k la condizione che il filtro considerato lo rispetti. Al fine di costruire questa nuova lista di sottoinsiemi possiamo semplicemente seguire il metodo noto come “l'albergo di Hilbert” (si veda ad esempio [LeoTof07]), cioè possiamo considerare la lista $W_1 = T_1$, $W_2 = T_1$, $W_3 = T_2$, $W_4 = T_1$, $W_5 = T_2$, $W_6 = T_3, \dots$.

Supponiamo ora per ipotesi induttiva di aver costruito un elemento c_n tale che $c_n \not\leq y$ e di aver considerato il filtro $F_n \equiv \uparrow\{c_n\}$ degli elementi maggiori di c_n che non contiene quindi y (abbiamo già fatto notare che questo vale per l'elemento c_0 e il filtro F_0) e definiamo l'elemento c_{n+1} nel modo che segue

$$c_{n+1} = \begin{cases} c_n & \text{se } \bigvee W_n \notin F_n \\ c_n \wedge b_n & \text{se } \bigvee W_n \in F_n \end{cases}$$

dove b_n è un elemento di W_n tale che $c_n \wedge b_n \not\leq y$; un tale elemento esiste perché se $\bigvee W_n \in F_n$ allora $c_n \leq \bigvee W_n$; ora se fosse vero che per tutti gli elementi $w \in W_n$ accadesse che $c_n \wedge w \leq y$ allora

$$c_n = c_n \wedge \bigvee W_n = \bigvee_{w \in W_n} c_n \wedge w \leq y$$

che contraddice l'ipotesi induttiva; quindi (prova classica!) deve esistere almeno un elemento $b_n \in W_n$ che soddisfa l'ipotesi richiesta. Possiamo ora porre $F_{n+1} \equiv \uparrow \{c_{n+1}\}$ ed è immediato verificare sia che y non sta nel filtro F_{n+1} e che $F_n \subseteq F_{n+1}$.

Siamo ora in grado di concludere visto che possiamo dimostrare che il filtro F richiesto nell'enunciato del teorema altro non è che $F = \cup_{i \in \omega} F_i$. Infatti, F è un filtro visto che è il risultato dell'unione di una catena di filtri ognuno contenuto nel successivo (esercizio: verificare che questa proprietà vale); inoltre $x \in F$ visto che $x \in F_0 \subseteq F$ mentre $y \notin F = \cup_{i \in \omega} F_i$ perché altrimenti dovrebbe esserci un indice $i \in \omega$ tale che $y \in F_i$ mentre il modo in cui abbiamo definito i filtri della catena esclude che questo possa accadere. Infine F rispetta tutti i sottoinsiemi $T_1, \dots, T_n, \dots$ perché se $\bigvee T_n \in F = \cup_{i \in \omega} F_i$ allora c'è un indice $i \in \omega$ tale che $\bigvee T_n \in F_i$ e quindi, visto che ogni T_n appare una quantità numerabile di volte nella lista $W_1, \dots, W_m, \dots$, per qualche $h \geq i$ deve accadere che $W_h = T_n$ e quindi $\bigvee W_h = \bigvee T_n \in F_i \subseteq F_h$ e perciò esiste $b_h \in T_n$ tale che $b_h \in F_{h+1} \subseteq F$.

Il teorema 7.4.4 si può subito utilizzare per dimostrare l'esistenza di un filtro primo che rispetta una quantità numerabile di supremi se l'algebra di Heyting che stiamo considerando è numerabile.

Corollario 7.4.5 (Teorema fondamentale) *Sia $\mathcal{H}$ una algebra di Heyting numerabile e x e y siano due elementi di H tali che $x \not\leq y$. Sia inoltre $T_1, \dots, T_n, \dots$ una famiglia numerabile di sottoinsiemi di H dotati di supremo. Allora esiste un filtro primo F di $\mathcal{H}$ che rispetta tutti i sottoinsiemi $T_1, \dots, T_n, \dots$ e tale che $x \in F$ e $y \notin F$.*

Dimostrazione. Per ottenere la prova di questo teorema basta osservare che nel caso di un'algebra di Heyting numerabile c'è una quantità numerabile di supremi binari e quindi si può ancora utilizzare il precedente teorema 7.4.4 per costruire un filtro che rispetti la quantità numerabile di supremi binari e la quantità numerabile dei supremi dei sottoinsiemi $T_1, \dots, T_n, \dots$ visto che mettendo assieme due insiemi numerabili si ottiene ancora un insieme numerabile. Ma allora è ovvio che quel che si ottiene è un filtro primo visto che rispetta tutti i supremi binari esistenti.

Può essere anche utile notare che quest'ultimo teorema, quando applicato considerando gli elementi 1 e 0, assicura l'esistenza di un filtro che rispetta una quantità numerabile di supremi visto che la condizione $1 \not\leq 0$ vale sempre.

Il fatto che esistano filtri primi che rispettano una quantità numerabile di supremi è il risultato fondamentale che permette di immergere una algebra di Heyting numerabile nella famiglia degli aperti di uno spazio topologico tramite un morfismo che rispetta tali supremi.

Infatti, se chiamiamo $\text{Pt}(\mathcal{H})$ la collezione di tutti i filtri primi di $\mathcal{H}$ che rispettano i sottoinsiemi $T_1, \dots, T_n, \dots$ allora possiamo costruire la topologia $\tau_{\mathcal{H}}$ su $\text{Pt}(\mathcal{H})$ avente per base $\mathcal{B}_{\tau_{\mathcal{H}}}$ i sottoinsiemi $\text{ext}(x) = \{P \in \text{Pt}(\mathcal{H}) \mid x \in P\}$ per $x \in H$. Infatti è facile verificare che $\mathcal{B}_{\tau_{\mathcal{H}}}$ è una base per uno spazio topologico poiché

$$\begin{aligned} \text{ext}(0) &= \{P \in \text{Pt}(\mathcal{H}) \mid 0 \in P\} = \emptyset \\ \text{ext}(1) &= \{P \in \text{Pt}(\mathcal{H}) \mid 1 \in P\} = \text{Pt}(\mathcal{H}) \\ \text{ext}(x \wedge y) &= \{P \in \text{Pt}(\mathcal{H}) \mid x \wedge y \in P\} \\ &= \{P \in \text{Pt}(\mathcal{H}) \mid x \in P\} \cap \{P \in \text{Pt}(\mathcal{H}) \mid y \in P\} \\ &= \text{ext}(x) \cap \text{ext}(y) \end{aligned}$$

visto che $x \wedge y$ è contenuto nel filtro (primo) P se e solo se $x \in P$ e $y \in P$.

Abbiamo così visto che la mappa ext è una mappa dall'algebra di Heyting $\mathcal{H}$ all'algebra di Heyting (su) $\tau_{\mathcal{H}}$ che rispetta $0_H, 1_H$ e $\wedge_H$. Per vedere che si tratta proprio di un morfismo, che rispetta cioè tutte le operazioni di algebra di Heyting, è conveniente verificare prima di tutto che

$$\text{ext}(x) \subseteq \text{ext}(y) \text{ se e solo se } x \leq y.$$

Possiamo farlo nel modo che segue. Se $x \leq y$ allora, per ogni filtro $P \in \text{ext}(x)$, cioè tale che $x \in P$, otteniamo che $y \in P$, cioè $P \in \text{ext}(y)$; d'altra parte, se $x \not\leq y$ allora per il corollario 7.4.5 esiste un filtro primo P che rispetta tutti i supremi considerati e tale che contiene x e non contiene y , cioè, tale che $P \in \text{ext}(x)$ ma $P \notin \text{ext}(y)$.

Finalmente possiamo dimostrare che $\text{ext}(-)$ è una immersione dell'algebra di Heyting $\mathcal{H}$ nella topologia $\tau_{\mathcal{H}}$. Infatti

$$\begin{aligned} \text{ext}(x \vee y) &= \{P \in \text{Pt}(\mathcal{H}) \mid x \vee y \in P\} \\ &= \{P \in \text{Pt}(\mathcal{H}) \mid x \in P\} \cup \{P \in \text{Pt}(\mathcal{H}) \mid y \in P\} \\ &= \text{ext}(x) \cup \text{ext}(y) \end{aligned}$$

poiché $x \vee y$ è contenuto in un filtro primo P se e solo se $x \in P$ o $y \in P$, e

$$\begin{aligned} \text{ext}(x \rightarrow y) &= \{P \in \text{Pt}(\mathcal{H}) \mid x \rightarrow y \in P\} \\ &= \bigcup \{\text{ext}(z) \mid z \leq x \rightarrow y\} \\ &= \bigcup \{\text{ext}(z) \mid z \wedge x \leq y\} \\ &= \bigcup \{\text{ext}(z) \mid \text{ext}(z \wedge x) \subseteq \text{ext}(y)\} \\ &= \bigcup \{\text{ext}(z) \mid \text{ext}(z) \cap \text{ext}(x) \subseteq \text{ext}(y)\} \\ &= \text{ext}(x) \rightarrow_{\tau_{\mathcal{H}}} \text{ext}(y) \end{aligned}$$

Inoltre, visto il modo in cui abbiamo definito gli elementi di $\text{Pt}(\mathcal{H})$, possiamo dimostrare che $\text{ext}(-)$ rispetta anche tutti i supremi dei sottoinsiemi $T_1, \dots, T_n, \dots$ che stiamo considerando. Infatti

$$\begin{aligned} \text{ext}(\bigvee_{b \in T_i} b) &= \{P \in \text{Pt}(\mathcal{H}) \mid \bigvee_{b \in T_i} b \in P\} \\ &= \bigcup_{b \in T_i} \{P \in \text{Pt}(\mathcal{H}) \mid b \in P\} \\ &= \bigcup_{b \in T_i} \text{ext}(b) \end{aligned}$$

poiché $\bigvee_{b \in T_i} b$ appartiene al filtro P , che per ipotesi rispetta tutti i sottoinsiemi $T_1, \dots, T_n, \dots$, se e solo se esiste un elemento $b \in T_i$ tale che $b \in P$.

Infine è interessante notare che sono rispettati non solo i supremi dei sottoinsiemi $T_1, \dots, T_n, \dots$ ma anche tutti gli infimi esistenti. Infatti

$$\begin{aligned} \text{ext}(\bigwedge_{i \in I} x_i) &= \{P \in \text{Pt}(\mathcal{H}) \mid \bigwedge_{i \in I} x_i \in P\} \\ &= \bigcup \{\text{ext}(z) \mid z \leq \bigwedge_{i \in I} x_i\} \\ &= \bigcup \{\text{ext}(z) \mid z \leq x_i, \text{ per ogni } i \in I\} \\ &= \bigcup \{\text{ext}(z) \mid \text{ext}(z) \subseteq \text{ext}(x_i), \text{ per ogni } i \in I\} \\ &= \bigcup \{\text{ext}(z) \mid \text{ext}(z) \subseteq \bigcap_{i \in I} \text{ext}(x_i)\} \\ &= \text{Int}(\bigcap_{i \in I} \text{ext}(x_i)) \end{aligned}$$

Per concludere la dimostrazione che $\text{ext}(-)$ è una immersione di $\mathcal{H}$ in $\tau_{\mathcal{H}}$ dobbiamo ancora dimostrare che è una mappa iniettiva, vale a dire che se $x \neq y$ allora $\text{ext}(x) \neq \text{ext}(y)$. Ma questo è ovvio visto che se $x \neq y$ allora $x \not\leq y$ o $y \not\leq x$ e quindi $\text{ext}(x) \not\subseteq \text{ext}(y)$ o $\text{ext}(y) \not\subseteq \text{ext}(x)$.

Esercizio 7.4.6 Si consideri l'algebra di Heyting $\mathcal{Q}$ costituita dagli aperti dello spazio topologico la cui base è l'insieme $\{(p, q) \mid p, q \in \mathbb{Q}\}$ degli intervalli aperti di numeri razionali. Chi è la collezione $\text{Pt}(\mathcal{Q})$ dei suoi punti?

7.5 La completezza del calcolo intuizionista

Applichiamo adesso i risultati algebrici che abbiamo ottenuto nella precedente sezione per fare una vera dimostrazione del teorema di completezza per il calcolo predicativo intuizionista.

Visto che $\mathcal{P}$ è una algebra di Heyting numerabile, in virtù del corollario 7.4.5 è possibile definire una sua immersione $\text{ext}(-)$ nell'algebra di Heyting completa degli aperti della topologia $\tau_{\mathcal{P}}$ in modo tale che una quantità numerabile di supremi siano rispettati.

Allora è possibile definire una interpretazione $V^{**}(-)$ delle proposizioni del calcolo predicativo intuizionista nella struttura $\mathcal{D}_{\tau_{\mathcal{P}}}$, uguale alla struttura $\mathcal{D}_{\mathcal{P}}$ eccetto per il fatto che l'algebra di Heyting $\mathcal{P}$ viene sostituita dall'algebra di Heyting completa degli aperti della topologia $\tau_{\mathcal{P}}$, semplicemente ponendo

$$V^{**}(A) \equiv \text{ext}([A]_{\dashv})$$

Infatti è facile vedere che con questa definizione otteniamo davvero una valutazione delle formule. Infatti

$$\begin{aligned} V^{**}(P(t_1, \dots, t_n)) &= \text{ext}([P(t_1, \dots, t_n)]_{\dashv}) \\ &= I^*(P)(t_1, \dots, t_n) \\ &= I^*(P)(V^{**}(t_1), \dots, V^{**}(t_n)) \end{aligned}$$

$$\begin{aligned} V^{**}(A \wedge B) &= \text{ext}([A \wedge B]_{\dashv}) \\ &= \text{ext}([A]_{\dashv} \wedge [B]_{\dashv}) \\ &= \text{ext}([A]_{\dashv}) \cap \text{ext}([B]_{\dashv}) \\ &= V^{**}(A) \cap V^{**}(B) \end{aligned}$$

$$\begin{aligned} V^{**}(A \vee B) &= \text{ext}([A \vee B]_{\dashv}) \\ &= \text{ext}([A]_{\dashv} \vee [B]_{\dashv}) \\ &= \text{ext}([A]_{\dashv}) \cup \text{ext}([B]_{\dashv}) \\ &= V^{**}(A) \cup V^{**}(B) \end{aligned}$$

$$\begin{aligned} V^{**}(A \rightarrow B) &= \text{ext}([A \rightarrow B]_{\dashv}) \\ &= \text{ext}([A]_{\dashv} \rightarrow [B]_{\dashv}) \\ &= \text{ext}([A]_{\dashv}) \rightarrow \text{ext}([B]_{\dashv}) \\ &= V^{**}(A) \rightarrow V^{**}(B) \end{aligned}$$

$$\begin{aligned} V^{**}(\exists x.A) &= \text{ext}([\exists x.A]_{\dashv}) \\ &= \text{ext}(\bigvee_{t \in \text{Term}} [A[x := t]]_{\dashv}) \\ &= \bigcup_{t \in \text{Term}} \text{ext}([A[x := t]]_{\dashv}) \\ &= \bigcup_{t \in \text{Term}} V^{**}(A[x := t]) \\ &= \bigcup_{t \in \text{Term}} V^{**}[x := t](A) \end{aligned}$$

$$\begin{aligned} V^{**}(\forall x.A) &= \text{ext}([\forall x.A]_{\dashv}) \\ &= \text{ext}(\bigwedge_{t \in \text{Term}} [A[x := t]]_{\dashv}) \\ &= \text{Int}(\bigcap_{t \in \text{Term}} \text{ext}([A[x := t]]_{\dashv})) \\ &= \text{Int}(\bigcap_{t \in \text{Term}} V^{**}(A[x := t])) \\ &= \text{Int}(\bigcap_{t \in \text{Term}} V^{**}[x := t](A)) \end{aligned}$$

Possiamo finalmente completare la prova del teorema di completezza. Infatti vale la seguente sequenza di implicazioni

$$\begin{aligned} V^{**}(A_1) \cap \dots \cap V^{**}(A_n) \subseteq V^{**}(B) &\text{ solo se } \text{ext}([A_1]_{\dashv}) \cap \dots \cap \text{ext}([A_n]_{\dashv}) \subseteq \text{ext}([B]_{\dashv}) \\ &\text{ solo se } \text{ext}([A_1]_{\dashv} \wedge \dots \wedge [A_n]_{\dashv}) \subseteq \text{ext}([B]_{\dashv}) \\ &\text{ solo se } \text{ext}([A_1 \wedge \dots \wedge A_n]_{\dashv}) \subseteq \text{ext}([B]_{\dashv}) \\ &\text{ solo se } [A_1 \wedge \dots \wedge A_n]_{\dashv} \leq [B]_{\dashv} \\ &\text{ solo se } A_1 \wedge \dots \wedge A_n \vdash B \\ &\text{ solo se } A_1, \dots, A_n \vdash B \end{aligned}$$

Quindi se in ogni struttura e per ogni valutazione $V(-)$ vale che $V(A_1) \wedge \dots \wedge V(A_n) \leq V(B)$ questo varrà in particolare per la struttura $\mathcal{D}_{\tau_{\mathcal{P}}}$, basata sulla topologia $\tau_{\mathcal{P}}$, e per la valutazione $V^{**}(-)$ che abbiamo appena definito, cioè sarà vero che $V^{**}(A_1) \cap \dots \cap V^{**}(A_n) \subseteq V^{**}(B)$ ma, come abbiamo appena visto, questo implica che $A_1, \dots, A_n \vdash B$. Abbiamo quindi provato il seguente teorema.

Teorema 7.5.1 (Completezza) *Se, per ogni struttura $\mathcal{D}$ e per ogni valutazione $V(-)$ delle formule del calcolo predicativo intuizionista in $\mathcal{D}$, $V(A_1) \wedge \dots \wedge V(A_n) \leq V(B)$ allora $A_1, \dots, A_n \vdash B$.*

Vale la pena di confrontare il risultato di adeguatezza del calcolo dei sequenti che abbiamo ottenuto ora con il risultato di correttezza enunciato nel teorema 7.3.1 visto che essi si differenziano per un piccola, ma sostanziale, dettaglio. La correttezza del calcolo vale infatti in modo indipendente dalla cardinalità dell'insieme di ipotesi Γ che appaiono nel sequente mentre siamo stati in grado di provare l'adeguatezza delle regole solo quando l'insieme di ipotesi Γ è finito. Correggeremo questa mancanza di simmetria nel prossimo capitolo 8.

7.6 La completezza del calcolo classico

Nelle sezioni precedenti abbiamo dimostrato che il sistema di regole per il calcolo predicativo intuizionista è completo. Un risultato completamente analogo si può ottenere per il calcolo predicativo classico rispetto alle strutture basate sulle algebre di Boole complete. Infatti, procedendo in modo analogo al caso intuizionista, possiamo costruire un'algebra di Boole, utilizzando le classi di equivalenza delle proposizioni del linguaggio rispetto alla relazione di equidimostrabilità classica in cui abbiamo naturalmente la possibilità di utilizzare anche le regole per la negazione classica. L'unica novità rispetto al caso precedente sta nel fatto che invece di dover definire l'operazione di implicazione tra classi di equivalenza questa volta dobbiamo definire una operazione di complemento; tuttavia la sua definizione è immediata visto che basta porre

$$\nu[A]_{\dashv} \equiv [\neg A]_{\dashv}$$

Come al solito bisogna ora dimostrare che si tratta di una buona definizione, che non dipende cioè dal rappresentante prescelto nella classe di equivalenza (possiamo ottenere questo risultato utilizzando solo le regole per la negazione intuizionista, vedi esercizio 7.6.3).

Una volta fatto questo è immediato vedere che abbiamo di fatto definito il complemento visto che è facile dimostrare, utilizzando le regole del calcolo classico, che $\nu[A]_{\dashv} \vee [A]_{\dashv} = [\top]_{\dashv}$ e che $\nu[A]_{\dashv} \wedge [A]_{\dashv} = [\perp]_{\dashv}$ (vedi esercizio 7.6.4).

Quindi l'insieme delle classi di equivalenza delle proposizioni determina un'algebra di Boole che possiamo utilizzare per definire una interpretazione $V(-)$ delle proposizioni nel modo ovvio, cioè in modo tale che $V(A) = [A]_{\dashv}$. Infatti per ottenere questo risultato basta estendere la dimostrazione che abbiamo già fatto per il caso intuizionista osservando che

$$V(\neg A) = \nu V(A) = \nu[A]_{\dashv} = [\neg A]_{\dashv}$$

Quel che ci manca da fare è dimostrare che tale algebra di Boole, che in generale non è completa, può essere immersa in un'algebra di Boole completa.

A tal fine, approfittando del fatto che ogni algebra di Boole è anche un'algebra di Heyting, si potrebbe pensare di utilizzare direttamente il metodo che abbiamo già usato nel caso intuizionista. Tuttavia tale approccio non funziona perché, in generale, l'algebra di Heyting associata alla topologia costruita sui *punti* di un algebra di Boole B non è una algebra di Boole (vedi esercizio 7.6.7).

Tuttavia noi sappiamo che l'esempio canonico di algebra di Boole completa è costituito dall'insieme delle parti di un insieme dato (vedi esercizio 5.5.2). Quel che ci serve è quindi individuare un insieme, su cui costruire l'insieme delle parti, che renda facile l'immersione dell'algebra di Boole delle proposizioni. Per fortuna nella ricerca di tale insieme non dobbiamo ricominciare dal nulla ma possiamo sfruttare il lavoro fatto nel caso intuizionista e utilizzare nuovamente l'insieme dei filtri primi che rispettano i supremi esistenti (anche se qualche adattamento sarà necessario e non tutto fila completamente liscio).

Ci sarà utile il seguente lemma.

Lemma 7.6.1 (Ultrafiltro) *Sia B un'algebra di Boole e U uno dei suoi filtri primi. Allora, per ogni $x \in B$, o $x \in U$ o $\nu x \in U$ (nel seguito chiameremo ultrafiltro un filtro che goda di questa proprietà).*

Dimostrazione. Per ipotesi U è un filtro primo e quindi da $x \vee \nu x \in U$, che vale visto che in ogni algebra di Boole $x \vee \nu x = 1$, segue appunto che $x \in U$ o $\nu x \in U$.

In realtà si può dimostrare (vedi esercizio 7.6.5) che in ogni algebra di Boole ultrafiltri e filtri primi coincidono mentre in una algebra di Heyting vale solo l'implicazione che sostiene che ogni ultrafiltro è un filtro primo (vedi esercizio 7.6.6).

Possiamo ora dimostrare che, per ogni algebra di Boole, vale il seguente teorema (le notazioni sono prese dalla sezione precedente)

Teorema 7.6.2 (Completamento di un algebra di Boole numerabile) *Sia B un'algebra di Boole numerabile con una quantità numerabile di sottoinsiemi dotati di supremo che desideriamo*

rispettare e consideriamo la collezione $\text{Pt}(\mathcal{B})$ dei filtri primi di $\mathcal{B}$ che rispettano detti supremi. Allora $\mathcal{P}(\text{Pt}(\mathcal{B}))$ è un'algebra di Boole completa tale che la mappa da $\mathcal{B}$ a $\mathcal{P}(\text{Pt}(\mathcal{B}))$ definita ponendo $\text{ext}(a) = \{P \in \text{Pt}(\mathcal{B}) \mid a \in P\}$ è una immersione che rispetta i supremi considerati.

Dimostrazione. Il risultato che stiamo considerando non è diverso da quello del teorema 7.4.4 e del suo corollario 7.4.5; infatti la sola differenza è che stiamo considerando un'algebra di Boole invece che un'algebra di Heyting, ma noi sappiamo che ogni algebra di Boole è anche un'algebra di Heyting.

Tuttavia, visto che stiamo trattando con un'algebra di Boole, una piccola aggiunta va fatta: dobbiamo verificare che l'immersione rispetta l'operazione di complemento, ma questo risultato è una immediata conseguenza del lemma 7.6.1 visto che, utilizzando tale lemma, possiamo subito ottenere che

$$\text{ext}(\nu a) = \{P \in \text{Pt}(\mathcal{B}) \mid \nu a \in P\} = \mathcal{C}(\{P \in \text{Pt}(\mathcal{B}) \mid a \in P\}) = \mathcal{C}(\text{ext}(a))$$

dove con $\mathcal{C}(A)$ indichiamo il complemento del sottoinsieme A .

Da questo punto in poi la prova della completezza del calcolo classico ricalca in buona parte il lavoro già fatto con la prova della completezza del calcolo intuizionista ma richiede ancora qualche piccolo adattamento. Infatti visto che stiamo ancora utilizzando come *punti* dell'algebra di Boole $\mathcal{B}$ dei filtri primi è immediato vedere che $\text{ext}(1) = \text{Pt}(\mathcal{B})$, $\text{ext}(0) = \emptyset$, $\text{ext}(a \wedge b) = \text{ext}(a) \cap \text{ext}(b)$, $\text{ext}(a \vee b) = \text{ext}(a) \cup \text{ext}(b)$ e abbiamo anche appena visto che $\text{ext}(\nu a) = \mathcal{C}(\text{ext}(a))$.

Per quanto riguarda poi la quantità numerabile di supremi che abbiamo deciso di rispettare si ottiene come prima che $\text{ext}(\bigvee_{b \in T_i} b) = \bigcup_{b \in T_i} \text{ext}(b)$. Tuttavia non riusciamo più a dimostrare che la mappa $\text{ext}(-)$ è una immersione per tutti gli infimi esistenti e quindi non siamo più garantiti nel passo relativo ai quantificatori universali nella dimostrazione del teorema di completezza. Il problema è che non è detto che se un ultrafiltro U di un'algebra di Boole $\mathcal{B}$ contiene tutti gli elementi x_i di una famiglia $(x_i)_{i \in I}$ ne contenga pure l'infimo $\bigwedge_{i \in I} x_i$ (vedi esercizio 7.6.8).

Possiamo tuttavia rimediare a questo problema se ci accontentiamo che la proprietà richiesta valga per una quantità numerabile di infimi. Siano infatti, come prima, $T_1, T_2, \dots$ la quantità numerabile di sottoinsiemi di $\mathcal{B}$ di cui vogliamo rispettare i supremi (che per ipotesi supponiamo esistere) e siano $W_1, W_2, \dots$ una quantità, ancora numerabile, di sottoinsiemi di $\mathcal{B}$ di cui vogliamo rispettare gli infimi (che per ipotesi supponiamo esistere) nel senso che l'ultrafiltro U che vogliamo costruire dovrà essere tale che se, per ogni $w \in W_i$ accade che $w \in U$ allora $\bigwedge_{w \in W_i} w \in U$.

Poniamo allora $W_i^\nu = \{\nu w \in \mathcal{B} \mid w \in W_i\}$ e consideriamo la seguente quantità numerabile di sottoinsiemi di $\mathcal{B}$, $T_1, W_1^\nu, T_2, W_2^\nu, \dots$. Allora in virtù del teorema 7.4.4 e del suo corollario 7.4.5 possiamo costruire un ultrafiltro U che rispetta tutti i supremi degli insiemi $T_1, T_2, \dots$ e tale che se $\bigvee_{z \in W_i^\nu} z \in U$ allora esiste $z \in W_i^\nu$ tale che $z \in U$. Quindi se, per ogni $w \in W_i$, $w \in U$ allora $\bigwedge_{w \in W_i} w \in U$ perché altrimenti, in virtù dell'esercizio 5.5.5, $\nu \bigwedge_{w \in W_i} w = \bigvee_{z \in W_i^\nu} z \in U$ e quindi dovrebbe esistere qualche $z \in W_i^\nu$ tale che $z \in U$, ma allora $\nu w \in U$ per qualche $w \in W_i$ e questo va contro la coerenza di U .

Ma allora abbiamo che

$$\begin{aligned} \text{ext}(\bigwedge_{w \in W_i} w) &= \{U \in \text{Pt}(\mathcal{B}) \mid \bigwedge_{w \in W_i} w \in U\} \\ &= \{U \in \text{Pt}(\mathcal{B}) \mid \text{per ogni } w \in W_i, w \in U\} \\ &= \bigcap_{w \in W_i} \{U \in \text{Pt}(\mathcal{B}) \mid w \in U\} \\ &= \bigcap_{w \in W_i} \text{ext}(w) \end{aligned}$$

A questo punto possiamo proseguire con la prova del teorema di completezza come nel caso precedente tenendo conto del fatto che la quantità di proposizioni quantificate esistenzialmente e quantificate universalmente è numerabile e quindi nell'algebra di Boole delle proposizioni siamo interessati a rispettare solo una quantità numerabile di supremi, in corrispondenza delle proposizioni quantificate esistenzialmente, ed una quantità numerabile di infimi in corrispondenza delle proposizioni quantificate universalmente.

Esercizio 7.6.3 *Dimostrare classicamente che se $A \vdash B$ e $B \vdash A$ allora $\neg A \vdash \neg B$ e $\neg B \vdash \neg A$ (può essere utile osservare che questa proprietà vale anche utilizzando la definizione intuizionista della negazione).*

Esercizio 7.6.4 *Dimostrare classicamente che, per ogni proposizione A , $\nu[A] \vee [A] = [\top]$ e $\nu[A] \wedge [A] = [\perp]$.*

Esercizio 7.6.5 *Dimostrare che ogni ultrafiltro U di un'algebra di Boole è primo.*

Esercizio 7.6.6 *Dimostrare che ogni ultrafiltro U di un'algebra di Heyting è primo ed esibire un controesempio della implicazione opposta.*

Esercizio 7.6.7 *Trovare un esempio di algebra di Boole B tale che l'algebra di Heyting associata alla topologia sui punti di B costruita come nel paragrafo precedente non è una algebra di Boole.*

Esercizio 7.6.8 *Si consideri l'algebra di Boole dei sottoinsiemi finiti e co-finiti dei numeri naturali. Si dimostri che tale algebra ha un ultrafiltro U ed una famiglia di elementi $(x_i)_{i \in I}$ dotata di infimo $\bigwedge_{i \in I} x_i$ tale che, per ogni $i \in I$, $x_i \in U$ ma $\bigwedge_{i \in I} x_i \notin U$.*

7.6.1 Ridursi solo al vero e al falso

Anche se abbiamo fornito una prova di completezza per il calcolo predicativo classico, un lettore classico potrebbe sentirsi non completamente soddisfatto visto che per lui una valutazione dovrebbe dire se una proposizione è vera o falsa, cioè essere una mappa dall'insieme delle proposizioni nell'algebra di Boole $\mathbf{2}$ di due elementi $\{\perp, \top\}$, e non trovarne un valore di verità in una algebra di Boole decisamente complessa come $\mathcal{P}(\text{Pt}(\mathcal{B}))$ dove $\mathcal{B}$ è l'algebra delle classi di equivalenza delle proposizioni.

Per accontentarlo ci toccherà dimostrare alcuni nuovi risultati sulle algebre di Boole. L'idea di fondo, data una generica algebra di Boole $\mathcal{B}$ ed un suo elemento x diverso da 1, è quella di trovare il modo di definire una funzione f_x da $\mathcal{B}$ all'algebra di Boole $\mathbf{2}$ in modo tale che essa rispetti tutte le operazioni e che al tempo stesso mandi in 0 $\in \mathbf{2}$ l'elemento x .

Infatti, supponendo che A sia una qualche proposizione tale che $\text{ext}([A])$ sia diverso da $\text{ext}([\top])$, se siamo in grado di procurarci un omomorfismo f_A dall'algebra di Boole $\mathcal{P}(\text{Pt}(\mathcal{B}))$ all'algebra di Boole di due elementi allora possiamo definire una valutazione delle proposizioni in vero e falso semplicemente ponendo

$$V^*(C) = f_A(\text{ext}([C]))$$

e otterremo che la proposizione non dimostrabile A verrebbe valutata da $V^*(-)$ in 0, cioè l'elemento minimo dell'algebra $\mathbf{2}$.

Vediamo quindi come arrivare a dimostrare tale risultato. Iniziamo dal seguente teorema.

Teorema 7.6.9 (Algebra quoziente) *Sia $\mathcal{B}$ un'algebra di Boole e F un suo filtro. Allora la relazione tra elementi di B definita ponendo*

$$x \sim_F y \equiv (x \rightarrow y) \wedge (y \rightarrow x) \in F$$

è una relazione di equivalenza le cui classi di equivalenza saranno indicate scrivendo, per ogni $x \in B$,

$$[x]_F \equiv \{y \in B \mid x \sim_F y\}$$

Inoltre l'insieme B/F delle classi di equivalenza così ottenute si può atteggiare ad algebra di Boole rispetto alle seguenti operazioni

$$\begin{aligned} 1_{B/F} &\equiv [1]_F & 0_{B/F} &\equiv [0]_F \\ [x]_F \wedge_{B/F} [y]_F &\equiv [x \wedge y]_F & [x]_F \vee_{B/F} [y]_F &\equiv [x \vee y]_F \\ \nu_{B/F}[x]_F &\equiv [\nu x]_F \end{aligned}$$

Dimostrazione. La dimostrazione viene lasciata per esercizio (vedi esercizi 7.6.12, 7.6.13 e 7.6.14).

Può essere utile un breve commento sul significato del teorema che abbiamo appena enunciato. La congruenza $x \sim_F y$ formalizza l'idea che gli elementi x e y sono *uguali* se assumiamo la *teoria*

formalizzata da F . Infatti un filtro è riconoscibile come una teoria visto che contiene le cose vere in ogni caso (condizione: $1 \in F$) e contiene le cose più vere di quelle che già accetta (condizione: se $x \in F$ e $x \leq y$ allora $y \in F$) ed infine è chiuso per congiunzioni (condizione: se $x \in F$ e $y \in F$ allora $x \wedge y \in F$). Naturalmente ci aspettiamo che le cose vere nella teoria siano esattamente gli elementi di F .

Lemma 7.6.10 *Sia $\mathcal{B}$ un'algebra di Boole e F un suo filtro. Allora $[x]_F = [1]_F$ se e solo se $x \in F$.*

Dimostrazione. Basta osservare che $x \in F$ se e solo se $(x \rightarrow 1) \wedge (1 \rightarrow x) \in F$ se e solo se $[x]_F = [1]_F$, visto che $(x \rightarrow 1) \wedge (1 \rightarrow x) = 1 \wedge x = x$.

Il teorema 7.6.9 suggerisce immediatamente, per ogni algebra di Boole $\mathcal{B}$ e ogni suo filtro F , di definire una mappa h_F da $\mathcal{B}$ a B/F ponendo

$$h_F(x) = [x]_F$$

È allora immediato verificare il seguente teorema.

Teorema 7.6.11 (Teorema di omomorfismo) *Sia $\mathcal{B}$ un'algebra di Boole e F un suo filtro. Allora la mappa $h_F : \mathcal{B} \rightarrow B/F$ è una mappa suriettiva che rispetta tutte le operazioni di algebra di Boole. Inoltre $h_F(x) = [1]_F$ per ogni $x \in F$.*

Dimostrazione. Per quanto riguarda la dimostrazione della prima parte dell'enunciato si tratta di una immediata verifica (vedi esercizio 7.6.15).

Inoltre il lemma 7.6.10 ci assicura che $x \in F$ se e solo se $h_F(x) = [x]_F = [1]_F$.

Per trovare la funzione f_x che cerchiamo basta allora trovare un ultrafiltro $U_{\nu x}$ dell'algebra di Boole $\mathcal{B}$ che contenga νx . Infatti, in questo caso l'algebra $B/U_{\nu x}$ avrà esattamente due elementi poiché la mappa $h_{U_{\nu x}}(-)$ manderà un elemento $y \in B$ in $[1]$ se $y \in U_{\nu x}$, come visto nel precedente teorema 7.6.11, oppure in $[0]$ se $y \notin U_{\nu x}$ visto che se $y \notin U_{\nu x}$ allora $\nu y \in U_{\nu x}$ e quindi $\nu h_{U_{\nu x}}(y) = h_{U_{\nu x}}(\nu y) = [1]$ che implica appunto che $h_{U_{\nu x}}(y) = [0]$.

Tuttavia noi sappiamo che un filtro $U_{\nu x}$ contiene, per ogni elemento di un'algebra di Boole, o lui o la sua negazione esattamente quando è un filtro primo (si veda il lemma 7.6.1 e l'esercizio 7.6.5). Quindi per concludere basta far vedere che, per ogni elemento x diverso da 1 esiste un filtro primo $U_{\nu x}$ che contiene la sua negazione. Ma se x è diverso da 1 allora νx è diverso da 0, e perciò in particolare $\nu x \not\leq 0$; quindi possiamo costruire il filtro primo cercato come abbiamo fatto nel corollario 7.4.5, almeno nel caso di algebre di Boole numerabili, ottenendo un filtro primo che contiene νx (e che ovviamente non contiene 0).

Ma il caso delle algebre di Boole numerabili è l'unico che ci interessa visto che noi stiamo lavorando con l'algebra delle classi di equivalenza delle proposizioni.

Esercizio 7.6.12 *Dimostrare che per ogni algebra di Boole $\mathcal{B}$ e per ogni filtro F su tale algebra la relazione $\sim_F$ definita come nel teorema 7.6.9 è una congruenza.*

Esercizio 7.6.13 *Dimostrare che, per ogni algebra di Boole $\mathcal{B}$ e per ogni filtro F su tale algebra di Boole, le operazioni su B/F definite nel teorema 7.6.9 sono ben definite.*

Esercizio 7.6.14 *Dimostrare che, per ogni algebra di Boole $\mathcal{B}$ e per ogni filtro F su tale algebra di Boole, l'algebra B/F con le operazioni definite nel teorema 7.6.9 è un'algebra di Boole.*

Esercizio 7.6.15 *Dimostrare che, per ogni algebra di Boole $\mathcal{B}$ e per ogni filtro F su tale algebra di Boole, la mappa h_F da $\mathcal{B}$ a B/F come definita nel teorema 7.6.11 rispetta tutte le operazioni di algebra di Boole.*

Capitolo 8

Il teorema di Compattezza

L'enunciato del teorema di completezza che abbiamo visto nel capitolo precedente è sicuramente un grande (e forse inaspettato) successo visto che esso stabilisce che possiamo catturare con strumenti sintattici, cioè un sistema di regole induttive che forniscono quindi una nozione di dimostrazione e teorema del tutto controllabile, una nozione completamente sfuggente come quella della verità in ogni struttura per ogni interpretazione.

Bisogna però ammettere che quando si guarda *come* si ottiene la dimostrazione del teorema di completezza non è difficile farsi prendere da una qualche forma di insoddisfazione e ci si sente un po' presi in giro visto che il modello che si considera è sostanzialmente costruito sull'insieme dei *nomi* che il linguaggio permette di formare e la struttura di verità è costruita pesantemente sull'insieme delle proposizioni ordinate tramite la relazione di dimostrabilità.

In questo capitolo vedremo che proprio il fatto che lavoriamo con un linguaggio del primo ordine mentre da un lato ci garantisce che vale il teorema di completezza dall'altro ci porta immediatamente ad alcuni risultati limitativi su quel che si può fare (nel seguito per semplificare la trattazione, lavoreremo sempre nell'ipotesi che tutte le nostre strutture utilizzino l'algebra di Boole a due valori $\{0, 1\}$, ci metteremo cioè in una prospettiva completamente classica).

8.1 Soddisfacibilità e compattezza

Un ruolo centrale per gli sviluppi futuri sarà giocato dalla nozione di soddisfacibilità.

Definizione 8.1.1 (Soddisfacibilità) *Sia Γ un insieme di formule in un linguaggio $\mathcal{L}$. Allora Γ si dice soddisfacibile se esiste una struttura $\mathcal{S}$ per $\mathcal{L}$ ed una interpretazione $V(-)$ tale che, per ogni proposizione $C \in \Gamma$, $V(C) = 1$.*

La soddisfacibilità di un insieme di proposizioni è quindi una nozione del tutto semantica che in linea di principio non ha nulla a che fare con le dimostrazioni. Vedremo tuttavia nella sezione 8.4 che possiamo trovare una nozione sintattica che gli è equivalente.

Il seguente teorema è probabilmente il più importante risultato sulla soddisfacibilità di un insieme di formule; per ora ci limitiamo ad enunciarlo, nei prossimi paragrafi ne illustreremo la rilevanza con una serie di esempi di applicazione mentre rimandiamo la sua dimostrazione al paragrafo 8.3.

Teorema 8.1.2 (Teorema di Compattezza) *Sia Γ un insieme di formule in un linguaggio $\mathcal{L}$ del primo ordine. Allora Γ è soddisfacibile se e solo se ogni suo sottoinsieme finito è soddisfacibile.*

8.2 Alcune conseguenze del teorema di compattezza

Vedremo ora alcune conseguenze del teorema di compattezza che mostrano che se da un lato i linguaggi del primo ordine ci permettono di dimostrare il teorema di completezza dall'altro soffrono di limiti espressivi che non si possono decisamente trascurare.

8.2.1 Inesprimibilità al primo ordine della finitezza

Se ci viene chiesto di definire al primo ordine la struttura di gruppo non abbiamo particolari problemi se non per la fatto che abbiamo a disposizione più soluzioni alternative e potremmo quindi trovarci nell'imbarazzo di decidere quale via seguire. Ad esempio possiamo chiamare gruppo ogni struttura $\langle S, \{=\}, \{\cdot\}, \{1\} \rangle$ che soddisfi i seguenti assiomi (per semplificare la notazione utilizziamo gli stessi segni sia nella semantica che nella sintassi anche se sappiamo che stiamo parlando di cose completamente diverse)

$$\begin{array}{ll} \text{(associatività)} & \forall x. \forall y. \forall z. x \cdot (y \cdot z) = (x \cdot y) \cdot z \\ \text{(neutro)} & \forall x. (x \cdot 1 = x) \wedge (1 \cdot x = x) \\ \text{(inverso)} & \forall x. \exists y. (x \cdot y = 1) \wedge (y \cdot x = 1) \end{array}$$

Il teorema di completezza che abbiamo dimostrato nel capitolo precedente ci garantisce allora che a partire da questi assiomi siamo in grado di dimostrare, utilizzando le regole logiche che abbiamo studiato, ogni proposizione A che si possa esprimere utilizzando l'uguaglianza e l'operazione $\cdot$ — e che valga in ogni gruppo.

Se tuttavia ci venisse chiesto di caratterizzare i gruppi finiti potremmo trovarci in imbarazzo perché potremmo non sapere come fare ad esprimere la proprietà che l'insieme S su cui la struttura si fonda deve essere finito. Non ci vuole tuttavia molto a ricordare ad esempio che S è finito se e solo se ogni funzione iniettiva da S in se è anche suriettiva e quindi una possibile formalizzazione della finitezza del gruppo potrebbe essere la seguente

$$\text{(finitezza)} \quad \forall f. (\forall x \forall y. f(x) = f(y) \rightarrow x = y) \rightarrow (\forall y. \exists x. f(x) = y)$$

Si tratta però di una proposizione (ma sarà davvero una proposizione?) che non quantifica solamente sugli elementi di S ma anche sulle funzioni da S in se e non è quindi una proposizione al primo ordine per cui sappiamo valere il teorema di completezza (e abbiamo quindi il problema di non sapere quali sono le regole *giuste* per trattare con un quantificatore su funzioni).

Sarebbe quindi interessante essere in grado di procurarsi una formalizzazione al primo ordine della finitezza. Proviamo a vedere se questo è possibile.

Sicuramente siamo in grado di caratterizzare le strutture che hanno un certo numero di elementi. Infatti possiamo esprimere il fatto che una struttura ha almeno n elementi scrivendo

$$\begin{aligned} M_n \equiv \exists x_1 \dots \exists x_n. & (x_1 \neq x_2) \wedge (x_1 \neq x_3) \wedge \dots \wedge (x_1 \neq x_n) \\ & (x_2 \neq x_3) \wedge \dots \wedge (x_2 \neq x_n) \\ & \dots \wedge (x_{n-1} \neq x_n) \end{aligned}$$

mentre se vogliamo dire che una struttura ne ha al più n possiamo scrivere

$$\begin{aligned} A_n \equiv \forall x_1 \dots \forall x_n. \forall t. & (x_1 = x_2) \vee (x_1 = x_3) \vee \dots \vee (x_1 = x_n) \vee (x_1 = t) \\ & (x_2 = x_3) \vee \dots \vee (x_2 = x_n) \vee (x_2 = t) \\ & \dots \vee (x_{n-1} = x_n) \vee (x_{n-1} = t) \\ & \dots \vee (x_n = t) \end{aligned}$$

e quindi una struttura ha esattamente n elementi se e solo se rende vera $M_n \wedge A_n$. Si potrebbe quindi sperare di dire che una struttura ha un numero finito di elementi usando una scrittura come $(M_1 \wedge A_1) \vee (M_2 \wedge A_2) \vee (M_3 \wedge A_3) \vee \dots$ ma questa naturalmente non è una proposizione.

Anzi proprio la presenza della proposizione M_n , insieme al teorema di compattezza, ci mette in grado di dimostrare che non ci può essere nessuna formula al primo ordine che esprima la finitezza. Supponiamo infatti che F sia tale proposizione e consideriamo l'insieme di proposizioni

$$\Gamma \equiv \{F, M_1, M_2, \dots, M_n, \dots\}$$

È allora ovvio che Γ non può essere soddisfatto visto che la struttura che lo soddisfa dovrebbe nello stesso momento essere finita ma anche avere un numero illimitato di elementi.

D'altra parte ogni sottoinsieme finito Γ_0 di Γ si può soddisfare in una qualunque struttura su un insieme finito con un numero di elementi maggiore o uguale della proposizione M_n di pedice massimo tra quelle che appaiono in Γ_0 . Ma allora per il teorema di compattezza 8.1.2 dovrebbe

esserci una struttura che soddisfa tutto Γ e questo abbiamo visto che è assurdo. Abbiamo quindi sbagliato a supporre di avere una formula F che esprima la finitezza.

Quindi non possiamo avere una teoria al primo ordine dei gruppi finiti per la quale valga il teorema di completezza: dobbiamo scegliere tra la completezza della teoria e la capacità di esprimere le proprietà che desideriamo!

8.2.2 Non caratterizzabilità della struttura dei numeri naturali

Forse ancora più sorprendente è scoprire che non riusciamo ad avere una teoria in grado di caratterizzare i numeri naturali.

Supponiamo infatti di essere del tutto liberi nella scelta della struttura sui numeri naturali che vogliamo descrivere, possiamo perciò scegliere a piacer nostro quali sono le relazioni, le operazioni e le costanti che desideriamo utilizzare. L'unica ipotesi che facciamo è che tra le relazioni ci sia l'uguaglianza, tra le operazioni ci sia il successore e tra le costanti lo zero (questa richiesta non è certo limitativa sulla teoria che possiamo definire visto che avere un linguaggio più ricco può solo aiutarci a fare una teoria migliore).

Supponiamo ora di essere completamente liberi anche nella scelta degli assiomi che possiamo considerare, anzi possiamo anche pensare di prendere come assiomi *tutte* le proposizioni vere nella struttura dei naturali che siamo in grado di scrivere con i predicati atomici, le operazioni e le costanti che abbiamo deciso di utilizzare. Se indichiamo con $\mathcal{N}$ la nostra struttura possiamo indicare con $\text{Th}(\mathcal{N})$ la teoria dei numeri naturali che ha come assiomi tutte le proposizioni vere in $\mathcal{N}$. In questo senso, $\text{Th}(\mathcal{N})$ è la migliore teoria dei numeri naturali che siamo in grado di esprimere al primo ordine, anche se bisogna ammettere che come teoria è un po' inutile visto che per sapere quali proposizioni ci sono in $\text{Th}(\mathcal{N})$ bisogna già sapere quali proposizioni sono vere in $\mathcal{N}$.

Vediamo ora che addirittura $\text{Th}(\mathcal{N})$ non è in grado di caratterizzare la struttura $\mathcal{N}$. Consideriamo infatti il seguente insieme di proposizioni:

$$\Gamma \equiv \text{Th}(\mathcal{N}) \cup \{x \neq 0, x \neq \text{succ}(0), x \neq \text{succ}^2(0), \dots\}$$

È allora chiaro che se Γ è soddisfacibile la struttura che lo soddisfa non può essere quella dei numeri naturali perché per fornire una interpretazione per la variabile x deve contenere un qualche elemento diverso da 0, diverso da 1, diverso da 2, $\dots$. D'altra parte tale struttura non si può distinguere da $\mathcal{N}$ usando proposizioni del primo ordine visto che deve soddisfare tutte le proposizioni in $\text{Th}(\mathcal{N})$.

Che tale struttura esista è infine una immediata conseguenza del teorema di compattezza; infatti ogni sottoinsieme finito Γ_0 di Γ contiene solo un numero finito di proposizioni del tipo $x \neq \text{succ}^n(0)$, che possono quindi essere tutte soddisfatte nei numeri naturali, e le proposizioni del sottoinsieme $\Gamma_0 \cap \text{Th}(\mathcal{N})$ che si possono pure soddisfare nella struttura dei numeri naturali.

Quindi se vogliamo utilizzare una teoria dei numeri naturali per la quale valga il teorema di completezza, e che sia quindi esprimibile con proposizioni al primo ordine, dobbiamo pagare il prezzo che questa non sarà soddisfatta solo nella struttura dei numeri naturali ma essa sarà la teoria al primo ordine anche di strutture diverse (potremo chiamarle *numeri naturali non-standard*).

8.2.3 Gli infinitesi

Il fatto che esistano strutture non-standard non deve necessariamente essere considerato un fatto negativo. Se consideriamo ad esempio una qualunque teoria al primo ordine per i numeri reali $\text{Th}(\mathcal{R})$ possiamo costruire il seguente insieme di proposizioni

$$\Gamma \equiv \text{Th}(\mathcal{R}) \cup \{0 < x, x < 1, x < \frac{1}{2}, x < \frac{1}{3}, \dots\}$$

Ora ogni sottoinsieme finito Γ_0 di Γ è soddisfacibile perché possiamo rendere vere tutte le proposizione che contiene nella struttura dei numeri reali visto che c'è solo un numero finito di proposizioni della forma $x < \frac{1}{n}$ da rendere vere. Quindi per il teorema di compattezza c'è una struttura che soddisfa tutto Γ , ma questa struttura non sono sicuramente i numeri reali visto che per soddisfare Γ dobbiamo interpretare la variabile x in un elemento maggiore di zero ma minore di $\frac{1}{n}$ per ogni numero naturale n e questo va contro la archimedecità dei numeri reali che sostiene che ogni numero reale può essere superato sommando uno abbastanza volte.

Abbiamo quindi fatto vedere che esiste una struttura che gode di tutte le proprietà al primo ordine dei numeri reali ma contiene anche degli elementi, potremo chiamarli *infinitesimi*, che tra i numeri reali non ci sono (sono maggiori di zero, ma minori di una qualunque quantità prefissata). Nello stesso momento abbiamo anche dimostrato che la proprietà di archimedecità dei reali non può essere esprimibile al primo ordine: se lo fosse essa sarebbe contenuta in $\text{Th}(\mathcal{R})$ e dovrebbe quindi essere vera anche nella struttura con gli infinitesimi la cui esistenza abbiamo dimostrato e che non ne può godere.

8.3 La dimostrazione del teorema di compattezza

Se chiamiamo Γ l'intero insieme di formule che vogliamo soddisfare, l'idea della dimostrazione del teorema di compattezza è di considerare le valutazioni che rendono vere tutte le proposizioni dei vari sottoinsiemi finiti di Γ come degli individui che esprimono le loro opinioni sulle proposizioni che nel linguaggio si possono esprimere e di cercare di costruire una valutazione per l'intero insieme Γ di formule dicendo (in modo molto democratico!) che una proposizione è vera quando è supportata da un numero *grande* di individui; tutto il gioco consisterà quindi nel decidere cosa significa *grande* in questo contesto (si veda comunque l'appendice C per capire che caratterizzare la nozione di sottoinsieme *grande* non è poi così facile e scevro da pericoli e trappole).

8.3.1 La definizione dell'ultraprodotto

Supponiamo di essere nella situazione di avere una famiglia $(\mathcal{D}_i)_{i \in I}$ di strutture dello stesso tipo $\mathcal{D}_i \equiv \langle D_i, \mathcal{R}_i, \mathcal{F}_i, \mathcal{C}_i \rangle$, tutte su un'algebra di Boole di due elementi, e di volere costruire a partire da esse una singola struttura $\mathcal{D} \equiv \langle D, \mathcal{R}, \mathcal{F}, \mathcal{C} \rangle$, che chiameremo l'*ultraprodotto* delle strutture $(\mathcal{D}_i)_{i \in I}$, adatta per lo stesso linguaggio $\mathcal{L}$ e tale che in $\mathcal{D}$ siano soddisfatte tutte le proposizioni (del primo ordine) del linguaggio $\mathcal{L}$ che sono soddisfatte in tutte le strutture $(\mathcal{D}_i)_{i \in I}$.

La prima idea a cui si può pensare per costruire il dominio D della struttura $\mathcal{D}$ è considerare il prodotto diretto di tutte i domini delle varie strutture $(\mathcal{D}_i)_{i \in I}$, cioè di dire che gli elementi d di D sono delle funzioni che associano ad ogni indice $i \in I$ un elemento $d(i) \in D_i$.

Questa idea va bene anche se richiede di procedere con attenzione quando bisogna definire l'interpretazione dei segni predicativi. In questo caso non possiamo infatti limitarci a dare una definizione per componenti visto che in tal modo non saremmo in grado di preservare la soddisfacibilità delle proposizioni del primo ordine (si veda l'esercizio 8.3.4).

Come abbiamo già suggerito, la soluzione di questo problema è quella di interpretare una proposizione A in vero se e solo se essa è vera in un insieme *grande* di strutture $(\mathcal{D}_i)_{i \in I}$.

Per far questo, possiamo lavorare così. Supponiamo di avere, per ogni $i \in I$ una valutazione σ_i delle variabili del linguaggio $\mathcal{L}$ in elementi del dominio D_i della struttura $\mathcal{D}_i$ e una interpretazione V_i dei simboli per costanti, dei simboli per funzione e dei simboli per predicato del linguaggio $\mathcal{L}$ rispettivamente in costanti, funzioni e relazioni della struttura $\mathcal{D}_i$.

Definiamo ora una mappa σ dalle variabili del linguaggio $\mathcal{L}$ in elementi di D in modo tale che

$$\sigma(x)(i) = \sigma_i(x)$$

Nel seguito ci sarà utile la seguente osservazione.

Lemma 8.3.1 *Sia σ la mappa definita come sopra, x una variabile e d sia un qualsiasi elemento di D . Allora, $\sigma(x/d)(z)(i) = \sigma_i(x/d(i))(z)$.*

Dimostrazione. Basta svolgere le definizioni:

- se $z \neq x$, $\sigma(x/d)(z)(i) = \sigma(z)(i) = \sigma_i(z) = \sigma_i(x/d(i))(z)$ mentre
- se $z = x$, $\sigma(x/d)(z)(i) = d(i) = \sigma_i(x/d(i))(z)$.

Quindi l'interpretazione delle variabili $\sigma(x/d)$, che differisce dall'interpretazione σ solo per il fatto che la variabile x viene interpretata nell'elemento $d \in D$, si può identificare con l'interpretazione definita come sopra ma a partire dalle interpretazioni $\sigma_i(x/d(i))$.

Per quanto riguarda poi l'interpretazione V della costante c nella struttura $\mathcal{D}$ ci basta considerare la funzione che associa ad ogni $i \in I$ l'interpretazione di c nella struttura $\mathcal{D}_i$, cioè possiamo porre

$$V(c)(i) \equiv V_i(c)$$

Seguendo la stessa idea possiamo interpretare il simbolo n -rio di funzione f nella funzione $V(f)$ di $\mathcal{F}$ definita per componenti ponendo

$$V(f)(d_1, \dots, d_n)(i) = V_i(f)(d_1(i), \dots, d_n(i))$$

Vale la pena di notare subito che la definizione che abbiamo dato ci porta al seguente lemma.

Lemma 8.3.2 *Sia t un qualsiasi termine del linguaggio $\mathcal{L}$. Allora, per ogni $i \in I$, $V^\sigma(t)(i) = V_i^{\sigma_i}(t)$.*

Dimostrazione. Per induzione sulla costruzione del termine t :

- $t \equiv x$. In questo caso abbiamo che $V^\sigma(x)(i) = \sigma(x)(i) = \sigma_i(x) = V_i^{\sigma_i}(x)$.
- $t \equiv c$. In questo caso abbiamo direttamente che $V^\sigma(c)(i) = V_i^{\sigma_i}(c)$.
- $t \equiv f(t_1, \dots, t_n)$. In questo caso, sfruttando l'ipotesi induttiva, otteniamo che

$$\begin{aligned} V^\sigma(f(t_1, \dots, t_n))(i) &= V^\sigma(f)(V^\sigma(t_1), \dots, V^\sigma(t_n))(i) \\ &= V_i^{\sigma_i}(f)(V^\sigma(t_1)(i), \dots, V^\sigma(t_n)(i)) \\ &= V_i^{\sigma_i}(f)(V_i^{\sigma_i}(t_1), \dots, V_i^{\sigma_i}(t_n)) \\ &= V_i^{\sigma_i}(f(t_1, \dots, t_n)) \end{aligned}$$

Per passare ora all'interpretazione delle proposizioni dobbiamo decidere come formalizzare l'idea intuitiva che una proposizione è vera se è supportata da un insieme grande di *individui* nell'insieme I degli indici. A tal fine, cominciamo osservando che gli unici insiemi di indici a cui siamo interessati sono quelli che caratterizzano le *coalizioni* a supporto di una qualche proposizione. Introduciamo allora la seguente definizione

$$m_A^\sigma \equiv \{i \in I \mid V_i^{\sigma_i}(A) = 1\}$$

per indicare l'insieme di tutti gli indici delle strutture $(\mathcal{D}_i)_{i \in I}$ che, nelle varie interpretazioni $V_i^{\sigma_i}$ soddisfano la proposizione A .

A questo punto siamo pronti per dire cosa intendiamo dicendo che un certo sottoinsieme X di I della forma m_A^σ (questi sono gli unici sottoinsiemi di I che sono interessanti per noi) è *grande*. A tale scopo fissiamo un qualsiasi ultrafiltro U dell'algebra di Boole $\mathcal{P}(I)$ e diciamo che X è *grande rispetto ad U* se $X \in U$, cioè U è l'insieme dei sottoinsiemi grandi di $\mathcal{P}(I)$ (considerando le proprietà richieste su un ultrafiltro ci si può convincere che la cosa, per quanto strana, non è del tutto folle).

L'idea è allora che la valutazione V^σ deve essere tale che per ogni proposizione A di $\mathcal{L}$ vale

$$V^\sigma(A) \equiv (m_A^\sigma \in U)$$

Naturalmente, affinché questo possa accadere, la prima cosa da fare è quella di definire l'interpretazione di un simbolo di predicato P di arietà n in modo che le cose funzionino. Poniamo quindi

$$V^\sigma(P)(d_1, \dots, d_n) = 1 \text{ se e solo se } \{i \in I \mid V_i^{\sigma_i}(P)(d_1(i), \dots, d_n(i)) = 1\} \in U$$

Questo significa ad esempio che considereremo uguali due elementi f e g di D se e solo se $\{i \in I \mid f(i) =_{D_i} g(i)\} \in U$ (esercizio: verificare che si tratta di una relazione di equivalenza).

Siamo allora in grado di dimostrare il seguente fondamentale teorema.

Teorema 8.3.3 (Teorema di Loš) *Sia U un qualsiasi ultrafiltro di $\mathcal{P}(I)$ e sia V^σ definita come sopra sulle proposizioni atomiche. Allora, per ogni proposizione A del linguaggio $\mathcal{L}$, accade che*

$$V^\sigma(A) = 1 \text{ se e solo se } m_A^\sigma \in U$$

Dimostrazione. La dimostrazione è per induzione sulla complessità della proposizione A . Dobbiamo analizzare tutti i casi che troviamo nella definizione 6.1.5.

Cominciamo quindi con il caso di una proposizione atomica $P(t_1, \dots, t_n)$. Dobbiamo quindi dimostrare che $V^\sigma(P(t_1, \dots, t_n)) = 1$ se e solo se $m_{P(t_1, \dots, t_n)}^\sigma \in U$ ma, dopo il lemma 8.3.2, questo è quasi immediato. Infatti

$$\begin{aligned} V^\sigma(P(t_1, \dots, t_n)) = 1 & \quad \text{se e solo se} \quad V^\sigma(P)(V^\sigma(t_1), \dots, V^\sigma(t_n)) = 1 \\ & \quad \text{se e solo se} \quad \{i \in I \mid V_i^{\sigma_i}(P)(V_i^{\sigma_i}(t_1), \dots, V_i^{\sigma_i}(t_n)) = 1\} \in U \\ & \quad \text{se e solo se} \quad \{i \in I \mid V_i^{\sigma_i}(P)(V_i^{\sigma_i}(t_1), \dots, V_i^{\sigma_i}(t_n)) = 1\} \in U \\ & \quad \text{se e solo se} \quad \{i \in I \mid V_i^{\sigma_i}(P(t_1, \dots, t_n)) = 1\} \in U \\ & \quad \text{se e solo se} \quad m_{P(t_1, \dots, t_n)}^\sigma \in U \end{aligned}$$

È inoltre immediato verificare che tutto funziona nel caso la proposizione A sia $\top$ o $\perp$. Infatti $V^\sigma(\top) = 1$ e $m_\top^\sigma \in U$, visto che $m_\top^\sigma = I$, mentre $V^\sigma(\perp) \neq 1$ e $m_\perp^\sigma \notin U$, visto che $m_\perp^\sigma = \emptyset$.

Per quanto riguarda ora i passi induttivi relativi ai connettivi il risultato si ottiene in virtù del fatto che U è un ultrafiltro su $\mathcal{P}(I)$.

Per iniziare, consideriamo il caso di una congiunzione $A \wedge B$. Dobbiamo allora dimostrare che $V^\sigma(A \wedge B) = 1$ se e solo se $m_{A \wedge B}^\sigma \in U$. Ma per ipotesi induttiva sappiamo che $V^\sigma(A) = 1$ se e solo se $m_A^\sigma \in U$ e $V^\sigma(B) = 1$ se e solo se $m_B^\sigma \in U$ e quindi per ottenere il risultato ci basta far vedere che $m_{A \wedge B}^\sigma \in U$ se e solo se $m_A^\sigma \in U$ e $m_B^\sigma \in U$ visto che $V^\sigma(A \wedge B) = 1$ se e solo se $V^\sigma(A) = 1$ e $V^\sigma(B) = 1$. Possiamo farlo così:

$$\begin{aligned} m_{A \wedge B}^\sigma \in U & \quad \text{se e solo se} \quad m_A^\sigma \cap m_B^\sigma \in U \\ & \quad \text{se e solo se} \quad m_A^\sigma \in U \text{ e } m_B^\sigma \in U \end{aligned}$$

dove abbiamo sfruttato il fatto che U è un filtro.

Nel caso di una disgiunzione $A \vee B$ dobbiamo dimostrare invece che $V^\sigma(A \vee B) = 1$ se e solo se $m_{A \vee B}^\sigma \in U$ che in virtù della solita ipotesi induttiva significa dimostrare che $m_{A \vee B}^\sigma \in U$ se e solo se $m_A^\sigma \in U$ o $m_B^\sigma \in U$ visto che $V^\sigma(A \vee B) = 1$ se e solo se $V^\sigma(A) = 1$ o $V^\sigma(B) = 1$.

$$\begin{aligned} m_{A \vee B}^\sigma \in U & \quad \text{se e solo se} \quad m_A^\sigma \cup m_B^\sigma \in U \\ & \quad \text{se e solo se} \quad m_A^\sigma \in U \text{ o } m_B^\sigma \in U \end{aligned}$$

dove abbiamo sfruttato il fatto che U è un filtro primo.

Infine, nel caso di una negazione $\neg A$ dobbiamo dimostrare che $V^\sigma(\neg A) = 1$ se e solo se $m_{\neg A}^\sigma \in U$; visto che $V^\sigma(\neg A) = 1$ se e solo se $V^\sigma(A) \neq 1$ e questo per ipotesi induttiva vale se e solo se $m_A^\sigma \notin U$ per concludere basta far vedere che $m_{\neg A}^\sigma \in U$ se e solo se $m_A^\sigma \notin U$.

$$\begin{aligned} m_{\neg A}^\sigma \in U & \quad \text{se e solo se} \quad \mathcal{C}(m_A^\sigma) \in U \\ & \quad \text{se e solo se} \quad m_A^\sigma \notin U \end{aligned}$$

dove abbiamo di nuovo sfruttato il fatto che U è un ultrafiltro.

Possiamo ora analizzare il caso di una proposizione $\exists x.A$ quantificata esistenzialmente. Dobbiamo allora dimostrare che $V^\sigma(\exists x.A) = 1$ se e solo se $m_{\exists x.A}^\sigma \in U$.

Ma $V^\sigma(\exists x.A) = 1$ se e solo se esiste $d \in D$ tale che $V^\sigma(x/d)(A) = 1$ e, per ipotesi induttiva, questo vale se e solo se $m_A^{\sigma(x/d)} \in U$. Ma in virtù dell'osservazione 8.3.1 sappiamo che $V_i^{\sigma(x/d)_i}(A)$ coincide con $V_i^{\sigma_i(x/d(i))}(A)$ e quindi $m_A^{\sigma(x/d)} = \{i \in I \mid V_i^{\sigma(x/d)_i}(A) = 1\} = \{i \in I \mid V_i^{\sigma_i(x/d(i))}(A) = 1\}$; perciò $V^\sigma(\exists x.A) = 1$ se e solo se esiste $d \in D$ tale che $\{i \in I \mid V_i^{\sigma_i(x/d(i))}(A) = 1\} \in U$.

D'altra parte

$$m_{\exists x.A}^\sigma = \{i \in I \mid V_i^{\sigma_i}(\exists x.A) = 1\} = \{i \in I \mid \text{esiste } d_i \in D_i \text{ tale che } V_i^{\sigma_i(x/d_i)}(A) = 1\}$$

Quindi se supponiamo che $m_{\exists x.A}^\sigma \in U$ e consideriamo un elemento $d \in D$ tale che $d(i) = d_i$ per tutti i $d_i \in D_i$ tali che $V_i^{\sigma_i(x/d_i)}(A) = 1$ otteniamo che

$$m_{\exists x.A}^\sigma = \{i \in I \mid \text{esiste } d_i \in D_i \text{ tale che } V_i^{\sigma_i(x/d_i)}(A) = 1\} \subseteq \{i \in I \mid V_i^{\sigma_i(x/d(i))}(A) = 1\} = m_A^{\sigma(x/d)}$$

Perciò $m_{\exists x.A}^\sigma \in U$ implica che $m_A^{\sigma(x/d)} \in U$ perché U è un filtro. Quindi se $m_{\exists x.A}^\sigma \in U$ allora esiste $d \in D$ tale che $m_A^{\sigma(x/d)} \in U$ e quindi $V^\sigma(\exists x.A) = 1$ vale.

D'altra parte se $V^\sigma(\exists x.A) = 1$ allora esiste $d \in D$ tale che $m_A^{\sigma(x/d)} \in U$ e quindi $m_{\exists x.A}^\sigma \in U$ visto che

$$\begin{aligned} m_A^{\sigma(x/d)} &= \{i \in I \mid V_i^{\sigma(x/d)_i}(A) = 1\} \\ &= \{i \in I \mid V_i^{\sigma_i(x/d(i))}(A) = 1\} \\ &\subseteq \{i \in I \mid \text{esiste } d_i \in D_i \text{ tale che } V_i^{\sigma_i(x/d_i)}(A) = 1\} \\ &= m_{\exists x.A}^\sigma \end{aligned}$$

e U è un filtro.

Non serve infine analizzare il caso di una proposizione quantificata universalmente visto che nell'ambito classico che stiamo considerando possiamo sempre ridurre una proposizione quantificata universalmente ad una negazione di una proposizione quantificata esistenzialmente.

È opportuno notare che la valutazione così costruita soddisfa tutte le proposizioni (del primo ordine) che valgono in tutte le strutture $(\mathcal{D}_i)_{i \in I}$. Infatti, se la proposizione A viene soddisfatta dalla struttura $\mathcal{D}_i$ con l'interpretazione $V_i^{\sigma_i}$ allora i appartiene all'insieme degli indici m_i^σ ; quindi, se una proposizione A viene soddisfatta da ogni struttura $\mathcal{D}_i$ allora $m_i^\sigma = I$ e appartiene quindi ad un qualsiasi ultrafiltro U che noi possiamo considerare; perciò $V^\sigma(A) = 1$.

Esercizio 8.3.4 Si dimostri che in generale il prodotto diretto di due campi non è un campo.

Esercizio 8.3.5 Si costruisca la struttura che rende valide tutte le proposizioni vere sui numeri naturali ma contiene anche un elemento diverso da 0, diverso da 1, ... (si veda la sezione 8.2.2). Determinare chi è tale elemento in questa struttura.

Esercizio 8.3.6 Si costruisca la struttura che rende valide tutte le proposizioni vere sui numeri reali ma contiene anche gli infinitesimi (si veda la sezione 8.2.3). Determinare chi sono tali infinitesimi.

8.3.2 Dimostriamo il teorema di compattezza

Possiamo finalmente concludere la dimostrazione del teorema di compattezza. Ricordiamo che dobbiamo dimostrare che se ogni sottoinsieme finito dell'insieme di proposizioni Γ è soddisfacibile allora l'intero insieme Γ è soddisfacibile.

L'ipotesi stessa del teorema ci suggerisce di utilizzare come famiglia $(\mathcal{D}_i)_{i \in I}$ di strutture quelle che soddisfano, con una opportuna valutazione $V_i^{\sigma_i}$, i vari sottoinsiemi finiti di Γ e quindi come insieme degli indici I proprio i vari sottoinsiemi finiti di Γ , cioè la struttura $\mathcal{D}_i$ con la valutazione $V_i^{\sigma_i}$ è quella che soddisfa tutte le proposizioni che stanno nell'insieme finito i di proposizioni.

Dopo aver visto il teorema 8.3.3 quel che ci manca per costruire la struttura $\mathcal{D}$ che soddisfa tutto Γ è quindi un opportuno ultrafiltro sull'insieme I degli indici.

Consideriamo quindi l'insieme

$$X = \{m_A^\sigma \mid A \in \Gamma\}$$

Vediamo ora che se riusciamo a trovare un ultrafiltro U di $\mathcal{P}(I)$ che contiene X allora l'ultra-prodotto $\mathcal{D}$ e la relativa valutazione V^σ definiti come nel precedente paragrafo 8.3.1 rendono vere tutte le proposizioni in Γ . Infatti supponiamo che A sia una proposizione in Γ , allora $m_A^\sigma \in X$ e quindi, nell'ipotesi che $X \subseteq U$, $m_A^\sigma \in U$ ma questo implica che $V^\sigma(A) = 1$ per il teorema 8.3.3.

Per concludere ci basta quindi far vedere che è possibile costruire un ultrafiltro U che contiene X . Possiamo dimostrare che questo è possibile con una dimostrazione di carattere generale che sfrutta una particolare proprietà dell'insieme X , vale a dire la *proprietà dell'intersezione finita*. Possiamo infatti dimostrare che l'intersezione di un qualunque numero (finito) di elementi di X è non vuota; supponiamo infatti che $m_{A_1}^\sigma, \dots, m_{A_n}^\sigma$ siano un numero finito di elementi di X e quindi che $A_1, \dots, A_n$ siano proposizioni di Γ . Allora l'ipotesi del teorema di compattezza ci garantisce che esiste una struttura $\mathcal{D}_i$ che soddisfa tutte le proposizioni $A_1, \dots, A_n$ e che sta quindi nell'intersezione $m_{A_1}^\sigma \cap \dots \cap m_{A_n}^\sigma$ che non può quindi essere vuota.

Per concludere vediamo quindi che ogniqualevolta abbiamo un'algebra di Boole B e un sottoinsieme X di B che gode della proprietà dell'intersezione finita possiamo costruire un ultrafiltro di B che contiene X .

Teorema 8.3.7 *Sia $\mathcal{B}$ un'algebra di Boole e X un sottoinsieme di B che goda della proprietà dell'intersezione finita. Allora esiste un ultrafiltro U di $\mathcal{B}$ tale che $X \subseteq U$.*

Proof. Utilizzando l'assioma di scelta possiamo ben ordinare l'insieme B e quindi possiamo costruire una lista $x_0, \dots, x_n, \dots, x_\lambda, \dots$ di tutti gli elementi di B in corrispondenza di qualche ordinale α . Allora, osservando che per ogni insieme W di elementi di B

$$\uparrow W \equiv \{z \in B \mid \text{esistono } x_1, \dots, x_n \in W \text{ tali che } x_1 \wedge \dots \wedge x_n \leq z\}$$

è un filtro di $\mathcal{B}$ che non contiene 0_B se e solo W gode della proprietà dell'intersezione finita (esercizio!), possiamo definire induttivamente la seguente catena di filtri consistenti di $\mathcal{B}$:

$$\begin{aligned} U_0 &= \uparrow X \\ U_{\lambda+1} &= \begin{cases} \uparrow(U_\lambda \cup \{x_\lambda\}) & \text{se } \uparrow(U_\lambda \cup \{x_\lambda\}) \text{ non contiene } 0_B \\ U_\lambda & \text{altrimenti} \end{cases} \\ U_\beta &= \bigcup_{\lambda < \beta} U_\lambda \end{aligned}$$

Se poniamo ora $U \equiv \bigcup_{\lambda < \alpha} U_\lambda$ abbiamo trovato l'ultrafiltro desiderato. Infatti, U è un filtro che non contiene 0_B , perché è unione di una catena di filtri che non contengono 0_B , e contiene X perché $X \subseteq U_0 \subseteq U$. Inoltre, per ogni $w \in B$, $w = x_\lambda$ per qualche $\lambda < \alpha$, e quindi se $w \notin U$ significa che $w \notin U_{\lambda+1}$ e questo vuol dire che $\uparrow(U_\lambda \cup \{x_\lambda\})$ contiene 0_B ; quindi esiste qualche elemento $y \in U_\lambda$ tale che $y \wedge w = 0_B$; ma allora $y \leq \nu w$ e quindi $\nu w \in U_\lambda \subseteq U$.

8.4 Il teorema di completezza forte

Come abbiamo già detto la nozione di soddisfacibilità di un insieme di proposizioni è una nozione del tutto semantica; vediamo però che essa ci può essere di aiuto per correggere, insieme al teorema di compattezza che essa porta con sé, quella mancanza di simmetria tra il teorema di validità e il teorema di completezza che avevamo notato alla fine del paragrafo 7.5.

Teorema 8.4.1 (Teorema di completezza forte) *Sia Γ un arbitrario insieme di formule e A una formula. Supponiamo inoltre che per ogni struttura ed ogni interpretazione V^σ accada che $V^\sigma(\Gamma) \leq V^\sigma(A)$. Allora $\Gamma \vdash A$.*

Osserviamo subito che la differenza tra il teorema di completezza che abbiamo già dimostrato e l'enunciato che stiamo considerando ora sta nel fatto che nel teorema di completezza forte non si impone alcun vincolo sulla cardinalità dell'insieme di formule Γ mentre nel teorema di completezza 7.5.1 si richiedeva che l'insieme Γ fosse finito.

Per dimostrare il teorema 8.4.1 faremo uso sia del teorema di completezza 7.5.1 che del teorema di compattezza 8.1.2. Supponiamo infatti che Γ e A siano rispettivamente un insieme di formule e una formula tali che per ogni struttura e ogni interpretazione V^σ accade che $V^\sigma(\Gamma) \leq V^\sigma(A)$. Allora ogni struttura che soddisfa Γ dovrà soddisfare anche A e quindi l'insieme di formule $\Gamma \cup \{\neg A\}$ non potrà essere soddisfacibile. Ne segue quindi, per il teorema di compattezza, che dovrà esistere un sottoinsieme finito Γ_0 di Γ tale che l'insieme $\Gamma_0 \cup \{\neg A\}$ non è soddisfacibile. Ma questo significa che per ogni struttura e per ogni interpretazione V^σ accade che $V^\sigma(\Gamma_0) \leq V^\sigma(A)$ e quindi, per il teorema di completezza, ne possiamo dedurre che $\Gamma_0 \vdash A$. A questo punto per ricavare una dimostrazione di $\Gamma \vdash A$ a partire da quella di $\Gamma_0 \vdash A$ basta indebolire tutti gli assiomi che dipendono dall'insieme di assunzioni Γ_0 ad assiomi che dipendono dall'insieme di assunzioni Γ .

Non dovrebbe ora sorprendere il fatto che il teorema di completezza forte sia un risultato sufficiente per ottenere direttamente il teorema di compattezza e quindi che

$$\text{teorema di completezza} + \text{teorema di compattezza} = \text{teorema di completezza forte}$$

Per dimostrare questo risultato osserviamo prima di tutto che un insieme di proposizioni $\Gamma \cup \{A\}$ è soddisfacibile se e solo se non è vero che tutte le strutture che rendono vero Γ rendono vero anche

$\neg A$, ma quest'ultima affermazione, in virtù del teorema di completezza forte, vale se e solo se non è vero che $\Gamma \vdash \neg A$.

Se consideriamo questa situazione nel caso speciale in cui la proposizione A coincida con $\perp$ viene naturale dare la seguente definizione.

Definizione 8.4.2 (Coerenza) *Sia Γ un insieme di formule in un linguaggio $\mathcal{L}$. Allora Γ si dice coerente se non esiste alcuna dimostrazione di $\Gamma \vdash \perp$.*

Non è allora difficile dimostrare il seguente teorema che stabilisce un legame tra soddisfacibilità e coerenza di un insieme di proposizioni.

Teorema 8.4.3 *Sia Γ un insieme di formule in un linguaggio $\mathcal{L}$. Allora Γ è soddisfacibile se e solo se è coerente.*

Dimostrazione. Prima di tutto osserviamo che Γ è soddisfacibile se e solo se $\Gamma \cup \{\neg \perp\}$ è soddisfacibile. Ma per quanto abbiamo appena osservato $\Gamma \cup \{\neg \perp\}$ è soddisfacibile se e solo se non è vero che $\Gamma \vdash \neg \neg \perp$ o, equivalentemente, se e solo se non è vero che $\Gamma \vdash \perp$, cioè se e solo se Γ è coerente.

Vale forse la pena di osservare che nei risultati che abbiamo considerato non abbiamo posto alcuna ipotesi sul numero di proposizioni presenti nell'insieme Γ . Tuttavia, mentre la nozione di soddisfacibilità di un insieme di proposizioni non dipende chiaramente in alcun modo dalla cardinalità dell'insieme Γ , si potrebbe restare un po' sorpresi nell'osservare che anche la nozione di coerenza non è influenzata dal numero di proposizioni in Γ . In realtà, in tutte le regole logiche che abbiamo considerato il numero di proposizioni che appaiono tra le assunzioni non deve essere necessariamente finito visto che abbiamo considerato assiomi della forma $\Gamma, A \vdash A$ dove Γ è un generico insieme di proposizioni.

Naturalmente una dimostrazione rimane comunque un oggetto di *profondità* finita anche se ogni sequente può essere di *larghezza* arbitraria. Questo naturalmente significa che nel costruire una dimostrazione anche se ipoteticamente possiamo contare su un insieme Γ di ipotesi di cardinalità arbitraria di fatto siamo confinati ad usarne solamente una quantità finita.

Possiamo quindi stabilire immediatamente il seguente teorema (che in realtà è poco più che una semplice osservazione),

Teorema 8.4.4 (Teorema di Compattezza Sintattica) *Il sequente $\Gamma \vdash B$ è dimostrabile se e solo se esiste un sottoinsieme finito Γ_0 di Γ tale che il sequente $\Gamma_0 \vdash B$ è dimostrabile.*

Dimostrazione. È ovvio che una dimostrazione di $\Gamma_0 \vdash B$ si può trasformare in una dimostrazione di $\Gamma \vdash B$ semplicemente trasformando gli assiomi della forma $\Gamma_0 \vdash A$ in assiomi della forma $\Gamma \vdash A$.

D'altra parte una dimostrazione di $\Gamma \vdash B$ utilizza in ogni caso un numero finito di applicazioni delle regole logiche e queste coinvolgono al massimo un sottoinsieme finito Γ_0 di proposizioni di Γ . È allora immediato trasformare la dimostrazione di $\Gamma \vdash B$ in una dimostrazione di $\Gamma_0 \vdash B$ semplicemente cancellando tutte le proposizioni di Γ che non sono utilizzate da alcuna regola.

Questo semplice risultato è la chiave che ci permette di dimostrare il teorema di compattezza a partire dal teorema di completezza forte.

Teorema 8.4.5 (Teorema di Compattezza Semantica) *Sia Γ un insieme di formule in un linguaggio $\mathcal{L}$ del primo ordine. Allora Γ è soddisfacibile se e solo se ogni suo sottoinsieme finito Γ_0 è soddisfacibile.*

Dimostrazione Un verso del teorema è completamente banale: se abbiamo una struttura che rende vere tutte le proposizioni in Γ , essa a maggior ragione rendere vere tutte le proposizioni di un qualunque suo sottoinsieme Γ_0 .

La dimostrazione dell'altra implicazione è invece più complessa, ma abbiamo già detto tutto quel che ci serve per ottenerla. Infatti in virtù del teorema 8.4.3 abbiamo che se Γ non fosse soddisfacibile allora il sequente $\Gamma \vdash \perp$ sarebbe dimostrabile, ma allora per il teorema 8.4.4 solo un sottoinsieme finito Γ_0 di proposizioni di Γ potrebbero essere davvero utilizzate nella dimostrazione

di $\Gamma \vdash \perp$ e quindi potremo costruire anche una dimostrazione di $\Gamma_0 \vdash \perp$, ma questo vorrebbe dire che Γ_0 non sarebbe soddisfacibile.

Questa dimostrazione del teorema di compattezza semantica sembra particolarmente semplice nonostante che si tratti di un teorema di rilevanza eccezionale; ma non ci si deve far trarre in inganno da tale apparente semplicità della dimostrazione che è tale solo perché stiamo utilizzando il teorema di completezza forte che non è semplice per nulla!

Capitolo 9

Il teorema di Incompletezza

Il teorema di incompletezza di Gödel è forse uno dei teoremi della logica matematica che ha avuto maggior successo anche al di fuori della cerchia degli specialisti, naturalmente in una forma che evita le difficoltà insite nella comprensione del suo enunciato e che purtroppo portano spesso ad enunciati che non hanno molto a che fare con il risultato originale (quando non sono proprio del tutto sbagliati!)

Spesso lo si sente enunciare nella forma “ci sono proposizioni vere (dei numeri naturali) che non si possono dimostrare” che suona quanto meno piuttosto strana: come faccio a sapere che si tratta di una proposizione vera se non tramite una dimostrazione? Naturalmente a questo punto del corso siamo abbastanza smaliziati da sapere che le dimostrazioni si fanno all’interno di un ben preciso sistema di regole a partire da un ben preciso insieme di assiomi, e quindi una forma più precisa del risultato precedente potrebbe essere “ci sono proposizioni vere dei numeri naturali che non si possono dimostrare nella teoria T ” che però rischia di diventare o completamente banale (basta prendere come T una teoria dei numeri naturali molto debole) o semplicemente falsa (basta prendere come teoria T quella che ha per assiomi tutte le proposizioni vere sui numeri naturali).

Una forma divulgativa, che non va troppo lontana dal vero risultato di Gödel, potrebbe essere la seguente che abbiamo preso, come quelle usate nei successivi esempi, da [Smullyan81, Smullyan85]. Supponiamo di essere nel paese di Tarskonja in cui gli abitanti si dividono in *cavalieri* che dicono sempre la verità e *furfanti* che mentono sempre. In tale paese ci sono molti club che accolgono i vari membri della società (ad esempio, il club dei pescatori, quello dei giocatori di scacchi, ...). Per quel che riguarda noi la cosa rilevante è che c’è il club “Gödeliani” formato dai *cavalieri confermati* di cui possono far parte solamente i cavalieri dopo che sono stati sottoposti a sofisticati controlli atti ad accertare che siano davvero cavalieri. Ora, un giorno voi incontrate un abitante di Tarskonja che vi dice “Non sono un cavaliere confermato”. Cosa potete dedurne? (si veda l’esercizio 9.1.4 per un esercizio duale).

Non è difficile riconoscere in quel che abbiamo appena descritto la tipica situazione che abbiamo incontrato tante volte durante questo corso. Le frasi pronunciate dagli abitanti di Tarskonja sono le nostre proposizioni; quelle pronunciate dai cavalieri sono quelli che, quando interpretate nella realtà, sono vere mentre quelle pronunciate dai furfanti sono quelle che risultano false quando vengono interpretate. Ora per appartenere al club dei “Gödeliani” bisogna dire solo frasi che siano teoremi di una qualche teoria T sulla realtà. Visto che, se davvero esiste, un abitante di Tarskonja in grado di pronunciare la frase “Non sono un cavaliere confermato” deve essere un cavaliere non confermato (perché?) questo significa che la proposizione deve essere vera ma non dimostrabile nella teoria T .

A questo punto dovrebbe essere chiaro che per dimostrare il teorema di incompletezza ci servirà trovare un modo per raccontare le procedure che portano alla conferma di un cavaliere ed esprimere all’interno della teoria affermazioni che riguardano la teoria stessa.

9.1 Alcuni esempi di *situazioni* di Gödel

Vediamo un secondo esempio che ci aiuterà a capire meglio la situazione. Consideriamo un linguaggio costruito sull'alfabeto costituito dai seguenti segni: S, D, $\neg$, (e). Costruiamo quindi una macchina che stampa *parole* ottenute con tale alfabeto ma la cui uscita filtra solo quelle corrispondenti ad una delle seguenti forme

$$S(X) \quad \neg S(X) \quad SD(X) \quad \neg SD(X)$$

dove X è una qualsiasi parola del linguaggio (queste sono quindi le *formule* del nostro linguaggio).

La semantica delle formule è la seguente:

$S(X)$ è vera	se la macchina è in grado di stampare X
$\neg S(X)$ è vera	se la macchina non è in grado di stampare X
$SD(X)$ è vera	se la macchina è in grado di stampare $X(X)$
$\neg SD(X)$ è vera	se la macchina non è in grado di stampare $X(X)$

La questione che vogliamo risolvere è se sia possibile costruire una macchina corretta e completa. Se pensiamo ad un sistema di assiomi e regole come ad una macchina che produce teoremi allora potremo ad esempio avere

$$\neg S(S) \quad \neg S(\neg) \quad \neg S(D)$$

tra gli assiomi di una macchina corretta visto che S, $\neg$ e D non possono sicuramente comparire tra gli output della macchina visto non sono neppure formule e

$$\frac{X}{S(X)} \quad \frac{SD(X)}{X(X)} \quad \frac{X(X)}{SD(X)}$$

tra le regole.

La questione se sia possibile produrre un sistema di regole che sia corretto e completo rispetto alla semantica che abbiamo dato è però ben più complessa che quella di procurarsi qualche assioma e qualche regola valida, visto che la semantica parla di quel che fa la macchina e si istituisce quindi una sorta di strano circolo tra sintassi e semantica. A riprova di questa difficoltà notiamo ad esempio che la formula

$$\neg SD(\neg SD)$$

è vera se e solo se essa non può apparire in output della macchina. Ma se la macchina è corretta non può verificarsi il caso che essa sia falsa e che la macchina la stampi e quindi se ne deduce che $\neg SD(\neg SD)$ deve essere vera e non stampabile (vedi esercizio 9.1.5). Quindi nelle ipotesi in cui stiamo lavorando non c'è modo di costruire una macchina corretta e completa.

In questo corso abbiamo fatto molta attenzione a tenere ben separate sintassi e semantica mentre per ricreare una situazione simile a quella del presente esempio bisogna che esse siano strettamente collegate, serve cioè una teoria che *parli di se stessa* come le formule fanno rispetto al comportamento della macchina. In realtà possiamo metterci in una situazione analoga se studiamo una teoria (che contiene quella) dei numeri naturali. Infatti, anche se la teoria non parla direttamente del proprio comportamento, i numeri naturali sono abbastanza espressivi da poter codificare molti aspetti della teoria stessa e quindi tramite codifiche e de-codifiche non è impossibile per una teoria cogliere aspetti rilevanti del proprio comportamento (ad esempio poter in qualche modo definire chi sono i propri teoremi).

Vediamo ora un ultimo esempio che sfrutta proprio le capacità espressive dei numeri naturali: la Teoria dei Numeri Straordinari.

Supponiamo quindi di avere un libro, con una quantità numerabile di pagine, tale che ogni pagina contenga un sottoinsieme U_n dei numeri naturali; ad esempio, se pensiamo ad una teoria al primo ordine dei numeri naturali, possiamo costruire un libro come quello richiesto facendo una lista $U_0(x)$, $U_1(x)$, ... di tutte le proposizioni con una variabile libera e di scrivere a pagina n del libro l'insieme dei numeri naturali che rendono vera la proposizione $U_n(x)$.

Diremo quindi che un numero n è *straordinario* se n appartiene all'insieme elencato a pagina n , cioè se $n \in U_n$. Il linguaggio che utilizzeremo per esprimere la nostra teoria dei numeri straordinari conterrà quindi una quantità numerabile di proposizioni ciascuna delle quali sostiene che un

qualche numero è straordinario, ad esempio la proposizione P_h potrebbe sostenere che il numero n è straordinario. Non è sicuramente limitativo supporre che per ogni numero n ci sia (almeno) una proposizione P_h che sostiene che il numero n è straordinario e quindi che esista una mappa che dato n fornisca il numero di pedice, che possiamo indicare con n^* , per tale proposizione atomica.

Possiamo ora ottenere facilmente una teoria corretta e completa dei numeri straordinari prendendo semplicemente come assiomi tutte le proposizioni atomiche $P_h \equiv n \in U_n$ tali che n è davvero un numero straordinario. Il problema con tale teoria è che l'unico modo per capire quali sono i suoi assiomi è quello di guardare direttamente il libro dei sottoinsiemi, ma se siamo capaci di far questo non abbiamo nessun bisogno di una teoria dei numeri straordinari!

Vediamo allora se possiamo imporre qualche condizione sul libro dei sottoinsiemi e sulla teoria che ci interessa per renderla più usabile. Vedremo che in realtà bastano poche condizioni abbastanza ragionevoli per ritrovarci con una teoria che, se corretta, è incompleta.

La prima condizione che chiediamo è la seguente che riguarda solo il libro dei sottoinsiemi

- (C) se A è un sottoinsieme presente nel libro allora anche il suo complemento appare nel libro

Anche la seconda condizione riguarda solo il libro dei sottoinsiemi

- (H) se A è un sottoinsieme presente nel libro allora anche $A^* \equiv \{n \in N \mid n^* \in A\}$ appare nel libro

Nell'ipotesi che il libro dei sottoinsiemi contenga insiemi descrivibili con una proposizione del primo ordine la prima condizione si può soddisfare facilmente perché se il sottoinsieme A è quello descritto dalla proposizione $P_A(x)$ allora il suo complemento può essere descritto dalla proposizione $\neg P_A(x)$.

Un po' più complicato è capire cosa succede con la condizione (H). In questo caso se l'insieme A è descritto dalla proposizione $P_A(x)$ allora l'insieme A^* è descritto dalla proposizione $Q(x) \equiv \exists y. F(x, y) \wedge P_A(y)$, dove $F(x, y)$ è la proposizione che formalizza all'interno della teoria dei numeri naturali che stiamo considerando la funzione f che manda il numero n nel numero di pedice associato n^* della proposizione che sostiene che n è un numero straordinario. Naturalmente questo richiede che la teoria dei numeri naturali che dovremo usare per rendere completamente formale questa dimostrazione del teorema di incompletezza sia in grado di formalizzare davvero la funzione f .

L'ultima condizione che è interessante chiedere è la seguente che, a differenza delle precedenti, riguarda finalmente la teoria che numeri straordinari e non solamente il libro dei sottoinsiemi

- (E) l'insieme dei numeri di pedice delle proposizioni atomiche P_h dimostrabili nella teoria appare nel libro dei sottoinsiemi

Se pensiamo che gli insiemi che appaiono nel libro dei sottoinsiemi siano quelli descritti da qualche proposizione, questa ultima condizione richiede che la teoria dei numeri straordinari sia sufficientemente controllabile visto che una singola proposizione deve essere in grado di descrivere tutti i suoi teoremi.

Naturalmente il problema sarebbe ora quello di verificare se esiste una qualche teoria dei numeri naturali che soddisfi queste tre condizioni (si veda [Smullyan92]), tuttavia per il momento ci basta vedere che esse sono sufficienti per dimostrare che c'è qualche numero straordinario n tale che la teoria non è in grado di dimostrare la sua appartenenza all'insieme U_n .

Teorema 9.1.1 *Se il libro dei sottoinsiemi soddisfa la condizione (H), allora per ogni insieme A che appare in qualche pagina del libro esiste un numero naturale n^* tale che P_{n^*} è vera se e solo se $n^* \in A$.*

Dimostrazione. In virtù della condizione (H) sappiamo che se A è un insieme che appare nel libro dei sottoinsiemi allora anche $B \equiv \{k \in N \mid k^* \in A\}$ compare in qualche pagina del libro. Sia n un numero di pagina tale che $B = U_n$, e consideriamo la proposizione P_{n^*} che asserisce che n è un numero straordinario. Allora otteniamo che P_{n^*} è vera se e solo se $n \in U_n$ se e solo se $n \in B$ se e solo se $n^* \in A$.

Questo risultato è importante perché ci permette di dimostrare il seguente importante risultato dovuto a Tarski.

Teorema 9.1.2 (Teorema di Tarski) *Se un libro dei sottoinsiemi soddisfa alle condizioni (H) e (C) allora l'insieme $V \equiv \{h \in N \mid P_h \text{ è vero}\}$ non compare nel libro dei sottoinsiemi.*

Dimostrazione. Supponiamo che V appaia nel libro. Allora, per la condizione (C), anche il suo complementare V^C dovrebbe apparire nel libro dei sottoinsiemi; quindi, per il teorema precedente, che è conseguenza della sola condizione (H), dovrebbe esistere un numero naturale v^* tale che P_{v^*} è vera se e solo se $v^* \in V^C$, ma questo è assurdo perché $v^* \in V^C$ se e solo se P_{v^*} non è vera.

È interessante notare che questo risultato, che ci dice che la *verità non è esprimibile* con una proposizione aritmetica, si ottiene utilizzando solo proprietà del libro dei sottoinsiemi e non dipende per nulla dalla teoria dei numeri straordinari che stiamo utilizzando. Esso è tuttavia la chiave per una velocissima dimostrazione del teorema di incompletezza.

Teorema 9.1.3 *Se la teoria T dei numeri straordinari che stiamo considerando è corretta e soddisfa le tre condizioni (H), (C) e (E) allora T non è completa, esiste cioè una proposizione vera che non dimostrabile in T .*

Dimostrazione. Dopo il teorema di Tarski, il risultato è immediato: per la condizione (E) l'insieme P_T dei pedici dei teoremi appare in qualche pagina del libro, mentre il teorema di Tarski ci ha fatto vedere che l'insieme V dei pedici delle proposizioni vere non appare. Quindi l'insieme P_T deve essere diverso da V . Ma per la correttezza della teoria deve accadere che ogni proposizione dimostrabile deve essere vera e quindi $P_T \subset V$; perciò deve esistere un pedice h che sta in V ma non in P_T , cioè la proposizione P_h è vera ma non è dimostrabile nella teoria T .

In realtà possiamo anche costruire una dimostrazione del teorema precedente che non sfrutta direttamente il teorema di Tarski. Infatti, dalla condizione (E) segue che l'insieme P_T dei numeri di pedice delle proposizioni dimostrabili nella teoria dei numeri straordinari appare in qualche pagina del libro dei sottoinsiemi; ma allora per la condizione (C) anche il suo complemento P_T^C appare nel libro e quindi per il teorema 9.1.1 esiste un numero c^* tale che P_{c^*} è vera se e solo se $c^* \in P_T^C$, cioè se e solo se c^* è il numero di pedice di una proposizione non dimostrabile nella teoria dei numeri straordinari e quindi se e solo se P_{c^*} non è dimostrabile. Ma allora, se la teoria è corretta, e quindi P_{c^*} dimostrabile implica che P_{c^*} è vera, si deduce subito che P_{c^*} deve essere vera ma non dimostrabile nella teoria T .

Esercizio 9.1.4 *Si supponga che a Tarskonian ci sia anche il club “Berluniani” dei furfanti confermati di cui possono far parte solamente dei furfanti dopo essere stati riconosciuti come tali. Che cosa dovrebbe dire un abitante di Tarskonian per essere riconosciuto come furfante non confermato?*

Esercizio 9.1.5 *Indichiamo che V il fatto che $\neg SD(\neg SD)$ sia vera e con S il fatto che la stessa formula sia stampabile. Dimostrare che se $V \leftrightarrow \neg S$ e $S \rightarrow V$ allora $V \wedge \neg S$.*

Esercizio 9.1.6 *Si supponga di avere un libro dei sottoinsiemi che soddisfa la condizione (E) perché contiene a pagina 8 i pedici delle proposizioni dimostrabili nella teoria dei numeri straordinari, che soddisfa la condizione (C) perché se contiene l'insieme A a pagina n allora il suo complemento si trova a pagina $3 \times n$ e che soddisfa la condizione (H) perché se contiene l'insieme A a pagina n allora contiene l'insieme A^* a pagina $3 \times n + 1$. Dimostrare che la proposizione $75 \in U_{75}$ è vera ma non è dimostrabile.*

9.2 Le funzioni ricorsive

Vediamo ora come formalizzare in una vera dimostrazione le idee che nel paragrafo precedente ci hanno fatto capire quelle che sono le linee principali per una prova del teorema di incompletezza di Gödel.

A questo fine partiremo con la definizione delle funzioni ricorsive. Il nostro scopo è caratterizzare quelle tra le funzioni sui numeri naturali che sono calcolabili; questo ci servirà per riuscire a caratterizzare le teorie le cui dimostrazioni si possono davvero riconoscere come tali, o, se vogliamo restare all'interno della metafora dei *cavaliere confermati*, per riuscire a dare una definizione precisa delle procedure di conferma lecite.

In generale lasceremo alla comprensione personale convincersi che le funzioni che riconosceremo come calcolabili si meritano davvero di essere giudicate tali provando ad esempio a scrivere un programma per calcolatore che le implementa. Seguiremo tuttavia alcune idee guida per capire quali funzioni ammettere all'interno della definizione di funzione calcolabile. Prima di tutto c'è il fatto che stiamo parlando di funzioni sui numeri naturali e che questi sono caratterizzati dalla presenza dello *zero*, della funzione *successore* e dal fatto che, visto che i numeri naturali sono definiti per induzione, le funzioni sui numeri naturali si definiscono per ricorsione; ci servirà poi utilizzare un'altra proprietà tipica dei numeri naturali, il fatto cioè che ogni sottoinsieme non vuoto di numeri naturali ha un elemento minimo. Useremo infine il fatto che le funzioni costituiscono un monoide, sono cioè chiuse rispetto alla operazione di composizione che ammette come elemento neutro la funzione identica.

Presentiamo ora la teoria delle funzioni ricorsive con un approccio tipico nella presentazione di una teoria in matematica, prima diremo cioè quali funzioni riconosciamo direttamente come ricorsive, cioè quali sono gli assiomi della teoria, e poi diremo quali funzioni sono ricorsive nell'ipotesi che altre funzioni lo siano, cioè stabiliremo quali sono le regole della teoria.

I nostri primi assiomi dichiarano quindi che Z , cioè la funzione costante $Z(x) = 0$ che da risultato 0 su ogni input, è una funzione calcolabile, che S , cioè la funzione successore $S(x) = x + 1$, è una funzione calcolabile e che la funzione identica I , cioè $I(x) = x$, è calcolabile.

Visto che intendiamo trattare anche con funzioni di arietà maggiore di uno (le precedenti sono tutte funzioni di arietà uno) converrà generalizzare il caso della funzione identica al caso delle funzioni proiezione I_i^n che ci permettono di recuperare il valore di un particolare argomento tra n , cioè $I_i^n(x_1, \dots, x_n) = x_i$. È evidente che $I = I_1^1$.

Cominciamo ora con le *regole* della nostra teoria delle funzioni ricorsive; la prima regola è direttamente ispirata dall'operazione di composizione tra funzioni e sostiene che se f è una funzione calcolabile ad m argomenti e $g_1, \dots, g_m$ sono m funzioni calcolabili a n argomenti allora anche $h(x_1, \dots, x_n) = f(g_1(x_1, \dots, x_n), \dots, g_m(x_1, \dots, x_n))$ è una funzione calcolabile ad n argomenti. Nel seguito ci sarà utile avere una notazione per la composizione: se F è la notazione per f e $G_1, \dots, G_m$ sono le notazioni per $g_1, \dots, g_m$ allora scriveremo $C[F, G_1, \dots, G_m]$ per indicare la funzione h .

Non è difficile vedere che ci sono molte funzioni calcolabili che non ricadono tra quelle che possiamo ottenere a partire da quelle considerate negli assiomi usando solo la composizione (si veda ad esempio l'esercizio 9.2.1)

Possiamo aggiungere molte altre funzioni se introduciamo una nuova *regola* alla teoria delle funzioni calcolabili: lo schema di ricorsione. Supponiamo quindi che k e g siano due funzioni calcolabili. Allora l'unica funzione f che soddisfa le due seguenti equazioni

$$\begin{cases} f(0, y_1, \dots, y_n) &= k(y_1, \dots, y_n) \\ f(x+1, y_1, \dots, y_n) &= g(x, y_1, \dots, y_n, f(x, y_1, \dots, y_n)) \end{cases}$$

è una funzione calcolabile, che indicheremo con $R[K, G]$ se K e G sono rispettivamente le notazioni per k e g .

È allora chiaro che anche la somma, il prodotto, l'esponenziazione e molte altre funzioni rientrano nel regno delle funzioni calcolabili (si veda l'esercizio 9.2.2), ma siamo davvero arrivati a caratterizzare la classe delle funzioni calcolabili? Una prima osservazione a questo proposito è il fatto che le funzioni ammesse fino a questo momento sono solo funzioni totali, definite cioè per ogni valore del dominio; mancano quindi ad esempio la funzione predecessore, che su zero non è definita, la sottrazione, la divisione e molte altre funzioni che sembrano essere intuitivamente calcolabili. Ci si può quindi chiedere se la totalità della funzione sia una richiesta insita nella caratterizzazione della calcolabilità.

Qualsiasi sia la risposta che diamo a questa domanda dobbiamo però osservare che una caratterizzazione delle funzioni calcolabili del tipo di quella che abbiamo fornito in questo paragrafo, cioè tale che una funzione calcolabile deve avere un programma che la calcola, ci porta inevitabilmente nella situazione che ogni caratterizzazione che preveda che una funzione per essere calcolabile deve essere totale deve essere necessariamente incompleta. Per convincerci di questo possiamo infatti ragionare nel modo seguente. Se, come nel nostro caso, essere calcolabile vuol dire che esiste un programma che descrive la funzione allora è possibile costruire in modo effettivo una lista f_0 ,

$f_1, \dots$ delle funzioni calcolabili (cioè un programma potrà costruire tale lista) e potremo quindi considerare la funzione così definita

$$g(x) \equiv f_x(x) + 1$$

Ora g è sicuramente una funzione totale, visto che tutte le f_x sono per ipotesi funzioni totali, ed è anche una funzione calcolabile, visto che la lista delle funzioni calcolabili che stiamo considerando è costruita in modo effettivo. Ma g non può comparire tra le funzioni calcolabili che abbiamo messo nella lista perché se vi comparisse al posto m avremmo che $f_m(m) = g(m) = f_m(m) + 1$ che è assurdo.

Quindi anche se volessimo solamente caratterizzare tutte le funzioni totali calcolabili saremmo costretti ad ammettere qualche costrutto che ci fornisca anche qualche funzione parziale in modo da aggirare il problema che abbiamo appena visto. Lo facciamo introducendo una nuova regola, la *minimalizzazione*, la cui definizione fa uso del fatto che ogni sottoinsieme non vuoto dei numeri naturali ammette elemento minimo. Questo fatto ci permette infatti di definire una funzione sui numeri naturali ponendo

$$f(x_1, \dots, x_n) = \mu y. g(x_1, \dots, x_n, y) = 0$$

cioè dicendo che il valore di $f(x_1, \dots, x_n)$ è il minimo y tale che $g(x_1, \dots, x_n, y) = 0$ (naturalmente se $g(x_1, \dots, x_n, y)$ non dovesse annullarsi per nessun valore di y la funzione f sarebbe parziale su $x_1, \dots, x_n$). Nel seguito indicheremo la funzione f con $M[G]$ se G è il modo di indicare la funzione calcolabile g .

Se proviamo ora a ripetere la dimostrazione precedente, che ci garantiva che ci sono funzioni totali calcolabili che non appaiono nella lista delle funzioni costruite utilizzando Z, S, I, I_i^n , la composizione e la ricorsione, ci troviamo a dover affrontare un nuovo problema. Infatti la dimostrazione si basava sul fatto di definire una funzione g cambiando in qualche modo il valore della funzione f_x sul numero naturale x , cioè con un argomento *diagonale*. Nel caso precedente questo poteva essere fatto perché, qualsiasi fosse il valore di $f_x(x)$ potevamo sempre incrementarlo di uno. Adesso possiamo ancora definire una nuova funzione g variando in qualche modo il valore di $f_x(x)$ ma per sapere come farlo dobbiamo essere in grado di distinguere il caso in cui $f_x(x)$ è definito dal caso in cui non lo è e se vogliamo poter dire che la funzione g è calcolabile bisogna che sia calcolabile la funzione

$$h(x) \equiv \begin{cases} 0 & \text{se } f_x(x) \text{ è definita} \\ 1 & \text{altrimenti} \end{cases}$$

che ci permette di decidere se f_x su input x è definita oppure no. Infatti in tal caso possiamo definire g ponendo ad esempio

$$g(x) \equiv \begin{cases} f_x(x) + 1 & \text{se } h(x) = 0 \\ 0 & \text{altrimenti} \end{cases}$$

e in tal modo troviamo una funzione che non compare nella lista delle funzioni calcolabili definibili utilizzando Z, S, I, I_i^n , la composizione, la recursione e la minimalizzazione.

Ora le possibilità cui ci troviamo di fronte sono due: o gli assiomi e le regole che abbiamo individuato non sono ancora sufficienti per caratterizzare le funzioni davvero calcolabili o la funzione h non può essere calcolabile. Tuttavia, a differenza del caso precedente, in cui la soluzione che ci permetteva di evitare la situazione spiacevole determinata dall'esistenza di una funzione totale chiaramente calcolabile che non appariva nella lista stava semplicemente nell'aggiungere una qualche regola che ci permettesse di aggiungere anche funzioni parziali calcolabili, questa volta il male è *incurabile* visto che qualsiasi nuova regola o assioma noi si possa escogitare ci sarà sempre una funzione che valuta la parzialità delle funzioni che sono calcolabili rispetto alla nozione di calcolabilità che stiamo esaminando e quindi saremo sempre in grado di produrre una funzione che non appare in tale lista di funzioni calcolabili. Quindi ad un certo punto dovremo accettare il fatto che, qualsiasi sia la nozione di funzione calcolabile che vogliamo utilizzare, la funzione che valuta la parzialità di una funzione calcolabile secondo tale nozione non sarà essa stessa una funzione calcolabile secondo tale nozione di calcolabilità.

Per quanto riguarda poi il momento in cui smettere di cercare nuovi assiomi o nuove regole per costruire la nostra teoria delle funzioni calcolabile possiamo, almeno da un punto di vista strettamente pragmatico, benissimo fermarci qui: nessuno è finora riuscito ad inventarsi una funzione che

fosse calcolabile in un qualche senso intuitivo e che non fosse già presente tra quelli caratterizzabili utilizzando Z , S , I , I_i^n , la composizione, la ricorsione e la minimalizzazione.

Esercizio 9.2.1 *Dimostrare che utilizzando solo Z , S , I , I_i^n e la composizione non si può creare una notazione per la funzione somma.*

Esercizio 9.2.2 *Dimostrare che utilizzando Z , S , I , I_i^n , la composizione e la ricorsione si può creare una notazione per la funzione somma.*

Esercizio 9.2.3 *Dimostrare che utilizzando Z , S , I , I_i^n , la composizione, la ricorsione e la minimalizzazione si può creare una notazione per la funzione parziale sottrazione.*

Esercizio 9.2.4 *Dimostrare che non esiste alcuna funzione calcolabile $h(x, y)$ tale $h(x, y) = 0$ se $f_x(y)$ è definita mentre $h(x, y) = 1$ se $f_x(y)$ non è definita.*

9.3 Sottoinsiemi ricorsivi e ricorsivamente enumerabili

L'aver individuato una teoria delle funzioni calcolabili ci permette di fornire una prima divisione dei sottoinsiemi dei numeri naturali in tre grandi fasce:

- (insiemi ricorsivi) i sottoinsiemi A la cui funzione caratteristica è calcolabile, cioè è calcolabile la funzione c_A tale che, dato un qualsiasi numero naturale n , $c_A(n) = 1$ se $n \in A$ mentre $c_A(n) = 0$ se $n \notin A$.
- (insiemi ricorsivamente enumerabili) i sottoinsiemi A i cui elementi sono quelli che appaiono nell'immagine di una funzione calcolabile, cioè quelli per cui esiste una funzione calcolabile g_A tale che $n \in A$ se e solo se esiste $x \in N$ tale che $n = g_A(x)$.
- (insiemi non ricorsivamente enumerabili) i sottoinsiemi che non coincidono con l'immagine di una funzione calcolabile.

Vediamo subito che i vari gruppi sono tutti non vuoti e che si tratta di gruppi distinti. Riguardo al primo gruppo gli esempi si sprecano: ad esempio il sottoinsieme dei numeri pari e quello dei numeri primi sono esempi di insiemi per i quali siamo in grado di produrre una funzione caratteristica calcolabile (vedi esercizio 9.3.1).

È inoltre chiaro che ogni insieme ricorsivo risulta anche essere ricorsivamente enumerabile. Infatti se A ha una funzione caratteristica c_A calcolabile allora i suoi elementi sono quelli nell'immagine della funzione calcolabile g_A definita ponendo

$$g_A(x) \equiv \begin{cases} \omega & \text{se } c_A(x) = 0 \\ x & \text{altrimenti} \end{cases}$$

dove abbiamo scritto ω per indicare il fatto che la funzione g_A non è definita per quel valore di x .

Esistono comunque, e sono particolarmente interessanti per noi, anche esempi di insiemi che sono ricorsivamente enumerabili ma non ricorsivi. Si tratta di insiemi più difficili da scovare ma possiamo procurarcene immediatamente uno sfruttando il fatto che non esiste una funzione calcolabile in grado di decidere sulla parzialità di una funzione calcolabile. Se definiamo infatti l'insieme K ponendo

$$K \equiv \{n \in N \mid f_n(n) \text{ è definito}\}$$

ci troviamo con un insieme che si può listare, coincide cioè con l'immagine di una funzione calcolabile, ma ha funzione caratteristica non calcolabile. Infatti, per listare K ci basta provare a calcolare $f_0(0)$, $f_1(1)$, $\dots$ inserendo un numero naturale in K man mano che scopriamo che qualche $f_x(x)$ è definito mentre non facciamo nulla per i valori per i quali non siamo stati ancora in grado di capire se $f_x(x)$ è definito oppure no, cioè se consideriamo la funzione calcolabile (perché si tratta di una funzione calcolabile?) definita ponendo

$$g_K(x) \equiv \begin{cases} x & \text{se } f_x(x) \text{ è definita} \\ \omega & \text{altrimenti} \end{cases}$$

otteniamo che K coincide con l'immagine di g_K .

D'altra parte K non ha funzione caratteristica calcolabile visto che poter decidere sull'appartenenza di x a K vuol dire decidere sulla parzialità della funzione f_x sull'argomento x e abbiamo visto che non esiste nessuna funzione calcolabile che possa far questo. Quindi K è un sottoinsieme di N che è ricorsivamente enumerabile ma non ricorsivo.

È immediato inoltre convincersi, utilizzando un argomento di cardinalità, che ci sono anche insiemi di numeri naturali che non sono neppure ricorsivamente enumerabili; infatti i sottoinsiemi dei numeri naturali sono una quantità più che numerabile mentre le funzioni calcolabili, che possiamo usare per listare un insieme, sono un quantità numerabile.

Possiamo tuttavia dare qualche esempio più esplicito di insieme che non è neppure ricorsivamente enumerabile, cioè che non è uguale all'immagine di una funzione calcolabile, se osserviamo che se accade che sia un insieme che il suo complemento sono ricorsivamente enumerabili allora in realtà quell'insieme è ricorsivo. Infatti un metodo effettivo che ci fornisce la funzione caratteristica per tale insieme consiste semplicemente nel listare contemporaneamente sia l'insieme che il suo complemento ed aspettare fino a che il valore desiderato non appare in una delle due liste.

A questo punto dovrebbe essere chiaro come ottenere un insieme che non è neppure listabile: basta considerare il complemento K^C dell'insieme K visto che tale insieme non può essere listabile perché altrimenti K sarebbe ricorsivo mentre abbiamo già osservato che non lo è.

Esercizio 9.3.1 *Dimostrare che la funzione caratteristica dell'insieme dei numeri pari (primi) è calcolabile.*

9.4 Teorie e sottoinsiemi dei numeri naturali

Alcune domande possono sorgere a questo punto: quali sottoinsiemi di numeri naturali sono *naturalmente* collegati ad una teoria? Quali sottoinsiemi dei numeri naturali sono rappresentabili tramite le proposizioni di una teoria?

Per quanto riguarda la prima domanda la chiave per trasformare insiemi di oggetti rilevanti dal punto di vista sintattico, quali ad esempio l'insieme dei termini, quello delle formule, quello delle dimostrazioni o quello dei teoremi, in insiemi di numeri naturali è l'idea di codifica. Non dovrebbe certamente stupire il fatto che possiamo codificare tramite numeri naturali le varie scritture che utilizziamo per esporre una certa teoria del primo ordine (in un certo senso basta scrivere la nostra teoria su un computer per avere automaticamente una codifica visto che tutto ciò che il computer può fare è una manipolazione di numeri in una qualche rappresentazione binaria). Possiamo quindi pensare di avere, associato ad una teoria, l'insieme dei numeri dei suoi termini, l'insieme dei numeri delle sue formule e, con un po' di fatica in più, anche l'insieme delle coppie di numeri costituito dal numero di una prova insieme al numero del teorema dimostrato da tale prova.

Se ci chiediamo ora di che tipo sono questi insiemi non dovrebbe essere difficile accorgersi che, vista la cura con cui abbiamo definito induttivamente cosa dobbiamo intendere per termine e per formula, l'insieme dei numeri dei termini deve essere un insieme ricorsivo e tale deve essere pure l'insieme delle formule. Per capire ad esempio se un certo numero naturale è il numero di una formula possiamo fare così: prima lo decodifichiamo e vediamo quale scrittura quel numero sta codificando e poi guardiamo se tale scrittura è una formula secondo la definizione induttiva di formula che abbiamo dato; è chiaro che la facilità con cui possiamo fare queste operazioni dipende fortemente dal modo in cui abbiamo costruito la codifica delle formule nei numeri ma dovrebbe essere chiaro che in linea di principio l'operazione si può sempre fare (per una codifica effettiva si veda ad esempio [BBJ07]).

Per quanto riguarda invece l'insieme dei numeri delle derivazioni c'è certamente qualche problema per decidere se essi costituiscono un insieme ricorsivo o, equivalentemente, se siamo in grado di riconoscere che un certo oggetto sintattico è una derivazione. Infatti, anche se siamo stati molto attenti a fare in modo da essere in grado di riconoscere effettivamente la corretta applicazione di tutte le regole, non abbiamo alcun controllo sulla effettiva riconoscibilità degli assiomi che ci dicono che $\Gamma \vdash A$ se $A \in \Gamma$ e che richiedono quindi di essere in grado di decidere se una certa proposizione appartiene all'insieme di proposizioni Γ . È allora chiaro che la effettiva riconoscibilità delle derivazioni si riduce alla effettiva riconoscibilità degli assiomi. Se vogliamo avere teorie che siano

davvero utilizzabili, nel senso di essere davvero in grado di riconoscere le derivazioni, dobbiamo quindi richiedere che l'insieme dei numeri degli assiomi sia ricorsivo: in questo caso parleremo di teorie *effettivamente assiomatizzabili*, i cui assiomi si possano cioè riconoscere in modo effettivo.

Se abbiamo una teoria effettivamente assiomatizzabile possiamo chiederci quale sia il genere dell'insieme dei numeri corrispondenti ai suoi teoremi. Si tratta certamente di un sottoinsieme dei numeri corrispondenti alle formule, ma questo non ci dice certamente che si tratta di un insieme ricorsivo! Quel che sappiamo è che un teorema è una formula che ha una dimostrazione nella teoria, quindi l'insieme dei numeri corrispondenti ai teoremi è un insieme ricorsivamente enumerabile visto che per ogni teoria effettivamente assiomatizzabile l'insieme dei numeri delle prove è un insieme ricorsivo (si vedano gli esercizi 9.4.2 e 9.4.3).

Veniamo ora alla seconda questione: la *representabilità* di un sottoinsieme dei numeri naturali all'interno di una teoria. Se abbiamo un qualsiasi linguaggio per una teoria dei numeri naturali possiamo utilizzarlo per *representare* all'interno della teoria particolari insiemi di numeri naturali; ad esempio, ammesso di avere a disposizione una teoria che ci permetta di trattare con il predicato di uguaglianza e con il prodotto, l'insieme dei numeri pari si può esprimere tramite la proposizione $\text{Pari}(x) \equiv \exists y. 2 \times y = x$.

Il seguente teorema, che non dimostreremo (si veda ad esempio [Mat1970]), ci garantisce che se la teoria che stiamo considerando ci permette di trattare con costanti, somme, prodotti e predicati di uguaglianza allora possiamo rappresentare una vasta classe di sottoinsiemi dei numeri naturali.

Teorema 9.4.1 (Davis, Putnam, Robinson, Matiyasevich) *Sia A un sottoinsieme dei numeri naturali. Allora A è ricorsivamente enumerabile se e solo se esiste un polinomio diofanteo $p_A(y_1, \dots, y_n, x)$ a coefficienti interi tale che $k \in A$ se e solo se esistono dei numeri naturali $n_1, \dots, n_n$ tali che $p_A(n_1, \dots, n_n, k) = 0$.*

Quindi l'appartenenza di un numero naturale k ad un insieme ricorsivamente enumerabile A si può equivalentemente esprimere dicendo che una certa equazione polinomiale diofantea $p_A(y_1, \dots, y_n, k) = 0$, dove $p_A(y_1, \dots, y_n, x)$ è un polinomio a coefficienti interi, ha soluzione nei numeri naturali. Ma allora se abbiamo una teoria al primo ordine dei numeri naturali che ci permetta di trattare con le costanti, le somme, i prodotti e il predicato di uguaglianza siamo in grado di trovare, per ogni insieme ricorsivamente enumerabile A , una proposizione $P_A(x)$ che rappresenta A visto che possiamo sempre scrivere l'equazione $p_A(y_1, \dots, y_n, x) = 0$ come una equazione $p_A^1(y_1, \dots, y_n, x) = p_A^2(y_1, \dots, y_n, x)$ tale che i coefficienti dei due polinomi p_A^1 e p_A^2 sono numeri naturali.

Esercizio 9.4.2 *Sia T una teoria dei numeri naturali effettivamente assiomatizzabile. Dimostrare che l'insieme $A \equiv \{n \in N \mid \vdash_T A(n)\}$ dei numeri naturali per cui la proprietà $A(-)$ è dimostrabile è un insieme ricorsivamente enumerabile.*

Esercizio 9.4.3 *Sia $A \times B$ un sottoinsieme ricorsivo di N^2 e sia $C \equiv \{n \in N \mid \exists y. \langle y, n \rangle \in A \times B\}$. Dimostrare che C è un sottoinsieme ricorsivamente enumerabile di N .*

9.5 La dimostrazione del teorema di incompletezza

In questo paragrafo vedremo che una teoria T dei numeri naturali che sia

- corretta,
- effettivamente assiomatizzabile e
- in grado di trattare con le costanti, le somme, i prodotti e il predicato di uguaglianza

deve inevitabilmente essere incompleta, deve cioè esserci una qualche proposizione vera nella struttura dei numeri naturali che la teoria non è in grado di dimostrare (si vedano gli esercizi 9.5.1 e 9.5.2).

Consideriamo infatti nuovamente l'insieme K dei numeri naturali n tali che $f_n(n)$ è definito (in realtà possiamo ottenere la dimostrazione considerando un qualsiasi insieme ricorsivamente

enumerabile che non sia ricorsivo). Abbiamo visto che si tratta di un insieme ricorsivamente enumerabile che non è ricorsivo. Visto che stiamo supponendo di lavorare con una teoria in grado di trattare con le costanti, le somme, i prodotti e il predicato di uguaglianza, in virtù del teorema 9.4.1 abbiamo una proposizione $P_K(n)$ esprimibile nel linguaggio della teoria T che rappresenta l'insieme K . Consideriamo ora i seguenti due insiemi di numeri naturali (per semplicità di notazione non distinguiamo tra i numeri naturali e le loro scritte all'interno della teoria T)

$$K_T^C = \{n \in N \mid \vdash_T \neg P_K(n)\}$$

e

$$K^C = \{n \in N \mid \models_N \neg P_K(n)\}$$

Entrambi intendono individuare i numeri naturali che stanno nel complemento dell'insieme K , ma solo il secondo lo fa davvero perché il primo contiene solo i numeri naturali che stanno nel complemento di K *dal punto di vista* della teoria T .

Visto che stiamo assumendo che la teoria T sia corretta abbiamo immediatamente che $K_T^C \subseteq K^C$ poiché $\vdash_T \neg P_K(n)$ implica $\models_N \neg P_K(n)$.

Per quanto riguarda invece l'altra inclusione, l'ipotesi della effettiva assiomatizzabilità di T implica che essa non può valere; infatti sappiamo che l'insieme K^C , complementare di K , non può essere ricorsivamente enumerabile mentre, in seguito all'esercizio 9.4.2, sappiamo che K_T^C lo è. Quindi deve esistere un numero naturale n^* che sta in K^C ma non sta in K_T^C , cioè deve succedere che $\models_N \neg P_K(n^*)$ ma $\not\vdash_T \neg P_K(n^*)$.

È utile notare che non solo $\not\vdash_T \neg P_K(n^*)$ ma anche $\not\vdash_T P_K(n^*)$ perché altrimenti, per la correttezza di T , dovrebbe valere anche $\models_N P_K(n^*)$ mentre sappiamo che $\models_N \neg P_K(n^*)$ che implica $\not\models_N P_K(n^*)$. Quindi T è incompleta sia nel senso che c'è una proposizione vera nel modello standard dei numeri naturali che T non è in grado di dimostrare ma anche nel senso che c'è una proposizione sulla quale T non è in grado di esprimersi, tale cioè che è *indecidibile* dal punto di vista di T .

Esercizio 9.5.1 *Esibire esempi di teorie che, mancando di una delle caratteristiche richieste per la dimostrazione del teorema di incompletezza, risultano complete.*

Esercizio 9.5.2 *Esibire un esempio di teoria dei numeri naturali che sia corretta, effettivamente assiomatizzabile e abbastanza espressiva da poter trattare con somme, prodotti, costanti e predicati di uguaglianza.*

Appendice A

Richiami di teoria degli insiemi

In questo corso non ci occuperemo di argomenti di matematica che vi sono estranei ma cercheremo piuttosto di ripensare ad alcuni tra gli argomenti di matematica che a livello di scuola superiore sono di solito trattati con poca precisione. In generale saremo più interessati a sollevare problemi piuttosto che a risolverli; infatti risolverli sarà lo scopo dei prossimi corsi di carattere matematico che incontrerete durante i vostri studi. Insomma, quel che vogliamo è che alla fine del corso il vostro modo di affrontare le questioni matematiche sia più maturo e che smettiate di pensare alla matematica come ad una collezione di tecniche da applicare.

Tanto per cominciare, consideriamo ad esempio un insieme numerico che probabilmente supponete di conoscere abbastanza bene, se non altro per averlo usato per anni: l'insieme dei numeri reali.

Sicuramente ricorderete che i numeri reali si possono suddividere in *numeri razionali* e *numeri irrazionali*. I numeri razionali sono di solito associati alle frazioni, anche se non si tratta di una corrispondenza perfetta visto che possiamo trovare un numero razionale in corrispondenza ad ogni frazione mentre lo stesso numero razionale corrisponde in generale a molte frazioni diverse.

Un po' più difficile è forse ricordare che i numeri razionali sono quei numeri la cui espansione decimale è finita o periodica mentre i restanti numeri sono gli irrazionali. È quindi elementare fornire esempi concreti di numeri razionale in quanto basta fare il quoziente di due numeri interi e non troppo difficile neppure fornire qualche esempio concreto di numero irrazionale. Si può ad esempio considerare $0,1011011101111011110\dots$ che non ha uno sviluppo decimale finito o periodico visto che il numero di 1 dopo uno 0 aumenta sempre. Naturalmente, voi conoscete dalla scuola altri numeri irrazionali (sicuramente ricorderete $\sqrt{2}$, $\sqrt{3}$, π , ...) ma come si fa a essere sicuri che il loro sviluppo decimale non è finito o periodico? La domanda è interessante perché pone le prime basi per la necessità dello sviluppo di una accurata formalizzazione e quindi di fare della matematica. Infatti, mentre per fare vedere che un certo numero q è un numero razionale basta trovare una coppia di numeri interi a e b tali che dividendo a per b si ottenga q , per far vedere che un certo numero r *non* è un numero razionale bisogna far vedere che qualsiasi tentativo si faccia di dividere un numero intero per un altro numero intero non si ottiene mai r ; naturalmente questo compito non si può fare per tentativi, visto che c'è il rischio di non arrivare mai alla fine dei tentativi, e bisogna quindi elaborare qualche altra maniera di convincersi che tali numeri interi non esistono. Sorge quindi la necessità di introdurre un'idea completamente nuova, quella di *dimostrazione*. La prova della irrazionalità di alcuni dei numeri precedenti non è troppo difficile ma per altri la dimostrazione non è affatto banale.

Proviamo adesso a porci alcune semplici domande.

Il numero $q = 0,999999\dots$ è senza dubbio un numero razionale visto che è chiaramente periodico. Si tratta di un numero minore, uguale o maggiore di 1?

Mentre si fa presto a decidere che q non è maggiore di 1 forse qualche dubbio rimane rispetto alla decisione se si tratta di un numero minore o di un numero uguale ad 1. Per fortuna la risposta in questo caso non è troppo difficile. Infatti c'è una regoletta che ci dice come associare una frazione ad ogni numero periodico e con tale regoletta non è difficile scoprire che $q = \frac{9}{9}$ e quindi che q e 1 coincidono; si tratta di un risultato forse inaspettato visto che q sembra essere sempre minore di

1 sia pure per quantità che diventano via via minori man mano che si considerano più cifre nella sua espansione decimale.

Questo esempio ci fornisce comunque almeno un paio di considerazioni: un numero razionale (e ancora di più un numero reale) può avere più di una espressione in notazione decimale e inoltre decidere quando due espansioni decimali indicano lo stesso numero non è una operazione che in generale sia banale. Infatti nel caso dei numeri razionali, le cui espansioni decimali sono finite o periodiche, sappiamo cosa fare visto che c'è una regola adatta, ma che fare se ci troviamo di fronte a due numeri irrazionali? Siamo ancora capaci di decidere se due espansioni decimali indicano lo stesso numero reale o due numeri reali diversi?

È evidente che la presenza o meno di una regoletta cambia completamente la situazione: alcune questioni si possono decidere mentre per rispondere ad altre domande non si sa come fare, vale a dire che non c'è un programma che fornisca la risposta.

Consideriamo ora una nuova questione che forse avete dato sempre per scontata: le operazioni.

È chiaro che sappiamo fare le usuali operazioni di somma, sottrazione, prodotto e divisione tra i numeri razionali, ma cosa si può dire delle stesse operazioni tra i numeri reali?

Quel che sappiamo è che un numero reale è irrazionale quando si tratta di una espansione decimale infinita e non periodica, ma allora come si fa ad applicare il metodo che conosciamo per fare la somma al caso di un numero irrazionale con se stesso? Si consideri ad esempio il numero irrazionale $x = 0,10110111011110\dots$ che abbiamo già visto in precedenza. Come si fa a *calcolare* $x + x$ o $x \times x$? Sicuramente non possiamo usare il solito metodo che prevede di cominciare a sommare o moltiplicare dalla cifra più a destra visto che in questo caso tale cifra più a destra non c'è. D'altra parte non possiamo neppure trasformare tale numero in una frazione, ed usare poi il metodo per sommare due frazioni che già conoscete, visto che non c'è nessuna frazione che equivalga ad x .

Sembra quindi che non esista un *metodo* per fare le operazioni tra i numeri reali (nessun programma per un calcolatore permette di fare la somma tra due numeri reali), ma allora cosa vogliono dire queste operazioni?

Naturalmente il problema di non sapere bene cosa vogliamo dire le operazioni tra i numeri reali porta ad altre conseguenze. Ad esempio, $\sqrt{2}$ è *definito* come quel numero reale x tale che $x \times x = 2$, ma abbiamo già detto che se qualcosa come $\sqrt{2}$ esiste esso non è sicuramente un numero razionale e quindi non abbiamo un'idea precisa di cosa voglia dire $x \times x$. Ma allora cosa significa la definizione che abbiamo dato di $\sqrt{2}$?

Naturalmente a partire da quelle che abbiamo proposto nascono altre domande ma sembra chiaro che esiste una forte necessità di ripensare, a partire dai concetti più elementari, la matematica che avete studiato fino ad oggi.

Esercizi

1. Determinare la regola che associa una frazione ad un numero razionale e dimostrare che essa è corretta.
2. Dimostrare che $\sqrt{2}$ non è un numero razionale.
3. Dimostrare che, per ogni numero primo p , $\sqrt{p}$ non è un numero razionale.
4. Dimostrare che esistono due numeri irrazionali x e y tali che x^y è un numero razionale. (*Suggerimento.* Si consideri $\sqrt{2}^{\sqrt{2}}$; se tale numero è razionale abbiamo trovato i due numeri irrazionali cercati; altrimenti si consideri $(\sqrt{2}^{\sqrt{2}})^{\sqrt{2}}$)

A.1 Collezioni e insiemi

L'operazione mentale di pensare insieme, come un'unica entità, più cose è tra le più naturali e si riflette infatti costantemente nel linguaggio naturale: chiamiamo gregge, un unico concetto, una collezione di pecore, chiamiamo classe una collezione di studenti e chiamiamo squadra una collezione di giocatori.

Questa operazione è forse la più semplice tra le operazioni di astrazione che si possono operare sulla realtà. La sua utilità è evidente in quanto rende possibile ridurre la complessità di un mondo in cui il numero di oggetti in gioco cresce facilmente al di là delle nostre possibilità di controllo: ad esempio le classi di una scuola possono essere una ventina mentre il numero degli studenti può essere di cinquecento e, per chi deve fare l'orario delle lezioni, è ben diverso pensare di organizzare venti oggetti piuttosto che cinquecento.

Naturalmente il pensare assieme le cose non è una operazione che viene effettuata in modo completamente casuale ma piuttosto secondo certi criteri che, di volta in volta, possono essere i più diversi: gli studenti di una scuola possono essere raggruppati in classi in relazione alla loro età, ma, quando si tratta di usare la palestra, possono anche essere divisi in due grandi gruppi a seconda del loro sesso. In entrambi i casi i gruppi vengono formati seguendo un certo concetto.

In matematica un simile approccio ha avuto un tale successo che buona parte della matematica si è potuta ridurre alla sola nozione di collezione e alle operazioni che forniscono nuove collezioni a partire da altre collezioni. Nei prossimi paragrafi cominceremo appunto a vedere come tale riduzione possa essere effettuata. Tuttavia prima ancora di cominciare con questo programma di lavoro conviene vedere fino a che punto l'operazione di costruire delle collezioni si possa effettuare in modo sicuro ed esente da paradossi, almeno nel caso in cui si voglia parlare di oggetti matematici. Vogliamo perciò analizzare in modo astratto l'idea del *mettere le cose insieme*: con un facile gioco di parole potremo dire che vogliamo analizzare il *concetto di concetto*.

A.1.1 Il concetto di concetto

Per partire con la nostra analisi, la prima idea che viene in mente, ispirata dalla vita di tutti i giorni, è quella di confondere un concetto con una proprietà: il concetto di gregge è basato sulla proprietà di essere una pecora appartenente ad un certo pastore, quello di classe scolastica sulla proprietà di avere la stessa età, quello di numero naturale pari sulla proprietà di essere divisibile per 2. L'idea sembra quindi buona. Tanto per cominciare, è fuori di dubbio che ogni collezione finita, che possiamo indicare con $\{a_1, \dots, a_n\}$ se essa è formata dagli elementi $a_1, \dots, a_n$, può essere comunque descritta tramite la proprietà di “essere uno dei suoi elementi uguale ad a_1 oppure di essere uguale ad a_2 oppure ... oppure di essere uguale ad a_n ” (notazione $\{x \mid x = a_1 \text{ o } \dots \text{ o } x = a_n\}$ che si legge “la collezione degli elementi x tali che $x = a_1 \text{ o } \dots \text{ o } x = a_n$ ”). Inoltre l'idea funziona bene anche per indicare collezioni che contengono un numero non finito di elementi: in questo caso infatti la notazione precedente, vale a dire $\{a_1, \dots, a_n\}$, non può proprio funzionare e passare attraverso la descrizione tramite una proprietà sembra essere l'unica strada percorribile. In questo modo possiamo ad esempio indicare la collezione dei numeri pari tramite la notazione $\{x \mid x \text{ si divide per } 2\}$ o la collezione dei numeri primi positivi con la notazione $\{x \mid \text{gli unici divisori interi di } x \text{ sono } 1 \text{ e } x \text{ stesso}\}$.

Bisogna tuttavia notare che considerare una collezione di oggetti come un unico ente la rende a sua volta un possibile soggetto di proprietà: ad esempio potrò dire “il gregge è del pastore Silvio” o “la collezione dei numeri pari contiene un numero non finito di elementi”. E quindi possibile che il confondere un concetto con una proprietà possa portare a qualche situazione paradossale. Il primo esempio di tali paradossi si deve a Bertrand Russell. Egli si è posto il problema di analizzare la collezione R i cui elementi sono le collezioni che non sono elementi di se stesse. Osserviamo prima di tutto che se le collezioni possono essere soggetto di una proprietà la frase “le collezioni che non sono elementi di se stesse” è dotata di significato. Ora, è facile riconoscere che, ad esempio, la collezione di tutti i gatti è un elemento di R , poiché essa è un concetto e non è perciò certamente un gatto e non può quindi essere un elemento di se stessa. D'altra parte la collezione di tutte le collezioni con un numero non finito di elementi non è un elemento di R , poiché essa è un elemento di se stessa; infatti, tale collezione ha sicuramente un numero non finito di elementi distinti: ad esempio, essa contiene la collezione dei numeri naturali, quella dei numeri naturali escluso il numero 1, quella dei numeri naturali escluso il numero 2 e così via. Sorge ora, abbastanza naturale, il problema di scoprire se R è oppure no un elemento di se stesso. Ma questa domanda ci conduce direttamente ad una situazione paradossale: se R fosse un elemento di R allora, visto che R è costituito esattamente dalle collezioni che non sono elemento di se stesse, dedurremmo che R non è un elemento di se stesso; d'altra parte se R non è un elemento di se stesso per la definizione stessa di R ci troveremo forzati ad ammettere che R deve essere un elemento di se stesso. Trascrivendo le considerazioni precedenti in simboli utilizzando la notazione $X \in Y$ per indicare che X è un elemento di Y ,

abbiamo

$$R \in R \text{ se e solo se } R \in \{X \mid X \notin X\} \text{ se e solo se } R \notin R$$

Poiché non sono previste altre possibilità, vale a dire che o R è un elemento della collezione R o non lo è, ci troviamo in un paradosso senza via d'uscita.

A.1.2 Insiemi

Il paradosso di Russell che abbiamo visto alla fine della sezione precedente ci obbliga a riflettere sul tipo di collezioni che potranno esserci utilizzate per fare matematica visto che è chiaro che pensare di ammettere tutte le possibili collezioni è una strada troppo pericolosa. Tuttavia analizzando la struttura del paradosso possiamo cercare di capire quali collezioni possiamo ammettere senza correre troppi rischi. Il paradosso sembra nascere da una situazione di autoriferimento in quanto la collezione R è fatta con collezioni tra le quali R stessa può trovarsi. Una via per sfuggire al paradosso è allora quella di abbandonare l'idea di usare tutte le collezioni e di limitarsi ad usare solo quelle formate da elementi che sono già stati riconosciuti come collezioni "sicure". Diremo che se seguiamo questo approccio accetteremo per le collezioni il principio di *buona fondatezza*.

Per evitare di fare confusione, d'ora in avanti chiameremo collezione il risultato di una qualunque operazione mentale consistente nel pensare alle cose assieme (e per le collezioni accetteremo i rischi connessi!) mentre riserveremo il nome di *insieme* per quelle particolari collezioni che possono essere soggetto di proprietà senza dare luogo a paradossi. Naturalmente la nostra speranza è che la collezione degli insiemi si riesca a descrivere per bene e che sia sufficiente per fare tutta la matematica cui siamo interessati. Notate, tanto per cominciare, che il principio di buona fondatezza richiede che la collezione di tutti gli insiemi non sia un insieme.

Esercizi.

1. Dimostrare che la collezione di tutti gli insiemi non è un insieme.
2. Dimostrare che la collezione di tutti gli insiemi infiniti non è un insieme.

Gli assiomi di base

Se vogliamo seguire il principio di buona fondatezza ci ritroviamo all'inizio a non possedere alcun insieme per cui l'unico insieme che possiamo costruire è quello che non ha nessun elemento (si ricordi che le uniche collezioni che ammettiamo come elementi di un insieme sono quelle che sono già state riconosciute come insiemi). Cominciamo perciò con il seguente assioma.

Definizione A.1.1 (Assioma dell'insieme vuoto) *La collezione che non contiene alcun elemento è un insieme che chiameremo insieme vuoto e indicheremo con $\emptyset$.*

Non è difficile rendersi conto che usando solo l'insieme vuoto di matematica se ne fa veramente poca, per cui è bene cercare di costruire nuovi insiemi. Iniziamo dunque introducendo un po' di notazioni che ci saranno utili nel seguito della trattazione.

La più importante relazione tra insiemi è naturalmente la *relazione di appartenenza*, indicata con $X \in Y$, cui abbiamo già accennato nella sezione precedente. La useremo per indicare che X è un elemento dell'insieme Y . La sua negazione verrà denotata con $X \notin Y$ e verrà usata per indicare che l'elemento X non appartiene all'insieme Y . Ad essa segue per importanza la *relazione di sottoinsieme* che può essere definita utilizzando la relazione di appartenenza.

Definizione A.1.2 (Sottoinsieme) *Siano X e Y due insiemi. Scriveremo allora $X \subseteq Y$ come abbreviazione per dire che, per ogni insieme z , se $z \in X$ allora $z \in Y$.*

È interessante notare che se X è un insieme e Y è una collezione i cui elementi appartengono tutti a X allora il principio di buona fondatezza ci spinge a considerare anche Y come un insieme.

Definizione A.1.3 (Assioma del sottoinsieme) *Sia X un insieme e Y una collezione tale che per ogni elemento di Y è anche un elemento di X . Allora Y è un insieme.*

Visto che tutto quello che ci interessa analizzare degli insiemi sono i loro elementi non è difficile accordarsi su fatto che possiamo considerare uguali due insiemi se questi hanno esattamente gli stessi elementi.

Definizione A.1.4 (Estensionalità) *Siano X ed Y due insiemi. Diremo allora che X è uguale a Y se e solo se X e Y hanno gli stessi elementi.*

Usando la definizione di sottoinsieme è immediato verificare che l'insieme X è uguale all'insieme Y se e solo se $X \subseteq Y$ e $Y \subseteq X$.

Possiamo ora tornare alla situazione cui eravamo arrivati prima delle disavventure dovute al paradosso di Russell per vedere come si può, almeno in parte, recuperare l'idea di definire un insieme tramite una proprietà pur rispettando il principio di buona fondatezza. Come abbiamo detto l'idea che ci deve guidare è che possono essere elementi di un insieme solo quelle collezioni che abbiamo già riconosciuto come insiemi. Questo ci porta immediatamente al seguente assioma.

Definizione A.1.5 (Assioma di separazione) *La collezione di tutti gli elementi di X che soddisfano la proprietà $P(-)$ è un insieme che indicheremo con $\{x \in X \mid P(x)\}$.*

Infatti, la collezione $\{x \in X \mid P(x)\}$, essendo composta solamente di elementi di un insieme, e quindi già riconosciuti come insiemi, è una collezione che il principio di buona fondatezza ci assicura essere “non pericolosa”.

Possiamo immediatamente utilizzare l'assioma di *separazione* per definire una prima operazione tra insiemi.

Definizione A.1.6 (Intersezione binaria di insiemi) *Siano X e Y due insiemi. Diremo intersezione tra l'insieme X e l'insieme Y l'insieme così definito*

$$X \cap Y \equiv \{x \in X \mid x \in Y\}$$

È ovvio che se X e Y sono due insiemi allora, come conseguenza dell'assioma di *separazione*, anche $X \cap Y$ è un insieme.

Tuttavia, l'assioma di separazione ci permette di ottenere solamente insiemi che sono sottoinsiemi di altri insiemi e avendo noi, per ora, a disposizione solo l'insieme vuoto tutto ciò che possiamo ottenere è solo l'insieme vuoto. Introduciamo quindi il più semplice dei modi per costruire nuovi insiemi.

Definizione A.1.7 (Assioma della coppia) *Siano X e Y due insiemi. Allora la collezione i cui elementi sono gli insiemi X e Y è un insieme che indicheremo con $\{X, Y\}$.*

Il principio di buona fondatezza ci assicura che questo assioma non è pericoloso e possiamo quindi usarlo per produrre nuovi insiemi. Tramite l'assioma della coppia possiamo produrre insiemi con uno (perché?) o due elementi quale ad esempio $\{\emptyset, \{\emptyset\}\}$. Tuttavia, se il nostro scopo è quello di rifare la matematica usando solo gli insiemi, conviene sforzarci un po' per ottenere almeno insiemi con un numero qualsiasi di elementi (vorremmo pur ottenere i numeri naturali!). Introduciamo quindi il seguente assioma.

Definizione A.1.8 (Assioma dell'unione binaria) *Siano X e Y due insiemi. Allora la collezione i cui elementi sono elementi di X o elementi di Y è un insieme che indicheremo con $X \cup Y$.*

Per giustificare il fatto che l'assioma dell'unione binaria ci permette di definire un insieme “sicuro” dobbiamo di nuovo rifarci al principio di buona fondatezza (perché?).

Usando l'assioma dell'unione è possibile finalmente costruire insiemi con un numero qualsiasi, purché finito, di elementi.

Esercizi

1. Dimostrare che due insiemi X e Y sono uguali se e solo se $X \subseteq Y$ e $Y \subseteq X$.

2. Dimostrare che esiste un unico insieme vuoto.
3. Dimostrare che l'unico sottoinsieme dell'insieme vuoto è l'insieme vuoto.
4. Definire, utilizzando gli assiomi fin qui introdotti, un insieme con un unico elemento.
5. Dimostrare che $X \cap Y = X$ se e solo se $X \subseteq Y$
6. Dimostrare che $W \subseteq X$ e $W \subseteq Y$ se e solo se $W \subseteq X \cap Y$.
7. Dimostrare che $X \cup Y = Y$ se e solo se $X \subseteq Y$
8. Dimostrare che $X \subseteq W$ e $Y \subseteq W$ se e solo se $X \cup Y \subseteq W$.
9. Dimostrare che con gli assiomi fin qui considerati possiamo costruire insiemi finiti con un numero arbitrario di elementi.
10. Dimostrare che se X è un insieme con n elementi e Y è un insieme con m elementi allora l'insieme $X \cap Y$ è vuoto se e solo se l'insieme $X \cup Y$ ha $n + m$ elementi.

I numeri naturali

Naturalmente gli insiemi finiti non sono sufficienti per fare la matematica. Il prossimo passo è quindi quello di definire una collezione i cui elementi siano una ragionevole ricostruzione dei numeri naturali: l'idea di fondo è che dobbiamo costruire, per ogni numero naturale n , un insieme che abbia esattamente n elementi. Sappiamo quindi subito come partire, visto che di insiemi con 0 elementi ce n'è uno solo:

$$0 \equiv \emptyset$$

Il prossimo passo è quello di capire come trovare un insieme che abbia un unico elemento. Anche questo è facile e il problema è piuttosto quello di scegliere un insieme particolare tra tutti gli infiniti insiemi con un unico elemento. Una scelta standard, che risulterà vantaggiosa, è la seguente:

$$1 \equiv \{\emptyset\}$$

cioè $1 \equiv \{0\}$.

A questo punto abbiamo già costruito due insiemi diversi e quindi ottenere un insieme che rappresenti il numero 2 è facile. Infatti basta porre:

$$2 \equiv \{0, 1\}$$

che ci dà anche l'idea generale su come procedere:

$$n \equiv \{0, 1, \dots, n-1\}$$

Possiamo esprimere questa definizione di infiniti oggetti in modo più compatto tramite una definizione *induttiva* dei numeri naturali:

$$\begin{array}{ll} \text{(Base)} & 0 \equiv \emptyset \\ \text{(Passo)} & n+1 \equiv n \cup \{n\} \end{array}$$

che ci dice (base) come partire e (passo) come ottenere un nuovo elemento, il successivo numero naturale, a partire da un elemento già costruito.

Definizione A.1.9 (Assioma dei numeri naturali) *La collezione i cui elementi sono gli insiemi costruiti tramite la precedente definizione induttiva è un insieme che indicheremo con ω .*

Anche in questo caso il principio di buona fondatezza ci assicura che questa collezione non è pericolosa perché costituita solamente con elementi che sono insiemi. Tuttavia, l'averla riconosciuta come insieme è un grande passo avanti: possediamo infatti il nostro primo insieme infinito! Vale la pena di sottolineare che questa definizione dei numeri naturali si fonda pesantemente sull'intuizione che già si possiede dei numeri naturali; non sarebbe infatti possibile, in mancanza di

questa intuizione, afferrare l'idea fondamentale che sostiene che la definizione induttiva dei numeri naturali produce ad ogni passo un nuovo insieme ma soprattutto che il processo di produzione di nuovi insiemi è discreto (cioè ha senso parlare di “step”) e non ha termine.

Prima di concludere questo capitolo, visto che adesso possediamo insiemi con infiniti elementi, è utile rafforzare l'assioma dell'unione binaria a questa nuova situazione.

Definizione A.1.10 (Assioma dell'unione arbitraria) *Sia I un insieme di insiemi. Allora la collezione i cui elementi sono gli elementi di un qualche $X \in I$ è un insieme che indicheremo con $\bigcup_{X \in I} X$.*

È facile rendersi conto che l'assioma dell'unione binaria altro non è che un'istanza di questo assioma nel caso in cui l'insieme I abbia due elementi. Possiamo giustificare questo nuovo assioma facendo appello al principio di buona fondatezza visto che anche in questo caso gli elementi della collezione $\bigcup_{X \in I} X$ sono già stati riconosciuti come elementi di un qualche insieme.

Esercizi

1. Dimostrare che esistono infiniti insiemi con un unico elemento.
2. Dimostrare che i numeri naturali sono uno diverso dall'altro.
3. Trovare una diversa rappresentazione insiemistica per i numeri interi.
4. Dimostrare che $X \subseteq W$ per ogni $X \in I$ se e solo se $\bigcup_{X \in I} X \subseteq W$.

Induzione

La definizione induttiva dei numeri naturali ci mette immediatamente a disposizione una potente tecnica per dimostrare le proprietà dei numeri naturali. Basta infatti notare che ogni numero naturale è *raggiungibile*, a partire dallo 0, applicando un numero sufficiente di volte l'operazione di “passare al successivo” per rendersi conto che ogni proprietà che vale per lo 0 e che vale per $n + 1$ qualora valga per n deve inevitabilmente valere per tutti i numeri naturali.

Definizione A.1.11 (Principio di induzione) *Sia $P(-)$ una proprietà sui numeri naturali. Allora, se vale $P(0)$ e, per ogni numero naturale n , $P(n)$ implica $P(n + 1)$ allora P vale per ogni numero naturale.*

Vedremo più avanti alcune applicazioni del principio di induzione. Illustriamone comunque immediatamente una applicazione che permette di ottenerne una (apparente) generalizzazione che risulterà molto utile.

Definizione A.1.12 (Principio di induzione completa) *Sia $P(-)$ una proprietà sui numeri naturali. Allora, se, per ogni numero naturale n , $P(n)$ è conseguenza di $P(0), \dots, P(n - 1)$, allora $P(-)$ vale per ogni numero naturale.*

Dimostrazione. Introduciamo la seguente nuova proposizione

$$Q(n) \equiv \text{per ogni numero } m \text{ tale che } 0 \leq m \leq n, P(m) \text{ vale}$$

È allora ovvio che se vale $Q(n)$ allora vale anche $P(n)$ visto che se $P(m)$ vale per ogni numero tale che $0 \leq m \leq n$ allora $P(-)$ vale in particolare per n .

Vediamo allora di dimostrare che $Q(-)$ vale per ogni n utilizzando il principio di induzione.

- (Base) Dobbiamo dimostrare che vale $Q(0)$, cioè che per ogni numero m tale che $0 \leq m \leq 0$ vale $P(m)$, cioè che vale $P(0)$. Ora, per ipotesi, noi sappiamo che, per ogni numero naturale n , $P(n)$ è conseguenza di $P(0), \dots, P(n - 1)$ che in particolare nel caso $n = 0$ dice che $P(0)$ è vero in quanto conseguenza di un insieme vuoto di ipotesi.

- (Passo) Dobbiamo dimostrare che se assumiamo che $Q(n)$ valga allora anche $Q(n+1)$ vale. Ora, dire che $Q(n)$ vale significa che, per ogni numero m tale che $0 \leq m \leq n$, vale $P(m)$ ma, per ipotesi, questo implica che $P(n+1)$ vale e quindi, per ogni $0 \leq m \leq n+1$, vale $P(m)$, cioè vale $Q(n+1)$.

Esercizi

1. Dimostrare per induzione che, per ogni numero naturale n , $2^n \geq n+1$.
2. Dimostrare per induzione che per ogni numero naturale n

$$1 + 2 + 3 + \dots + n = \frac{n \times (n+1)}{2}$$

3. Trovare una dimostrazione della precedente uguaglianza che *non* utilizzi il principio di induzione.
4. Dimostrare per induzione che n rette dividono il piano in un numero di regioni minore o uguale a 2^n .
5. Trovare l'errore nella seguente prova per induzione che “dimostra” che tutti i bambini hanno gli occhi dello stesso colore. Sia $P(n)$ la proprietà che afferma che tutti i bambini in un gruppo di n bambini hanno gli occhi dello stesso colore. Allora

- (Base) Ovvio, infatti bisogna andare a verificare che tutti i bambini in un gruppo di 0 bambini hanno gli occhi dello stesso colore.
- (Passo) Supponiamo che tutti i bambini in un gruppo di n bambini abbiano gli occhi dello stesso colore e dimostriamo che lo stesso vale per tutti i bambini di un gruppo di $n+1$ bambini. Consideriamo quindi un gruppo di $n+1$ bambini e togliamone uno; otteniamo in tal modo un gruppo di n bambini e quindi, per ipotesi induttiva, tutti hanno gli occhi dello stesso colore. Togliamone ora un altro ottenendo di nuovo un gruppo di n bambini che, per ipotesi induttiva, hanno nuovamente gli occhi dello stesso colore. Ma allora anche i bambini che avevamo tolto hanno gli occhi dello stesso colore.

6. (algoritmo veloce di moltiplicazione) Dimostrare, utilizzando il principio di induzione completa, che il seguente algoritmo per moltiplicare due numeri naturali è corretto. Dati due numeri n e m se $n = 0$ allora rispondiamo 0. Altrimenti, se n è pari calcoliamo il prodotto di $\frac{n}{2}$ con $2 * m$ mentre se n è dispari calcoliamo il prodotto di $\frac{n-1}{2}$ con $2 * m$ e aggiungiamo m al risultato. Possiamo scrivere il precedente algoritmo nel modo seguente

$$\left\{ \begin{array}{ll} 0 \times m & = 0 \\ (n+1) \times m & = \text{if } \text{Pari}(n+1) \\ & \text{then } \frac{n+1}{2} \times (2 * m) \\ & \text{else } (\frac{n}{2} \times (2 * m)) + m \end{array} \right.$$

L'interesse in questo algoritmo sta nel fatto che in un calcolatore la divisione e la moltiplicazione per 2 sono operazioni primitive estremamente veloci.

7. (Torre di Hanoi) Il gioco della torre di Hanoi consiste nello spostare n dischi di diverso diametro, infilati ordinatamente in un piolo l'uno sull'altro, verso un secondo piolo utilizzando un terzo piolo e con il vincolo che un disco non deve mai essere posto sopra un disco di diametro maggiore. La leggenda narra che ad Hanoi ci sono dei monaci che stanno spostando dall'inizio dell'universo 64 dischi e che l'universo finirà quando anche l'ultimo disco sarà spostato.

Trovare una funzione f che fornisca il minimo numero di mosse necessarie per spostare n dischi e dimostrare per induzione che, per ogni $n \in \omega$, $f(n) \geq 2^{n-1}$ (e possiamo quindi stare tranquilli sulla durata dell'universo!).

8. Sia $\leq$ l'usuale relazione d'ordine tra i numeri naturali ω . Un sottoinsieme S di ω possiede un elemento minimo se esiste un elemento $m \in S$ tale che, per ogni $x \in S$, $m \leq x$. Dimostrare, utilizzando il principio di induzione completa, che ogni sottoinsieme non vuoto S di ω ha elemento minimo. (Sugg.: porre $P(n) \equiv n \notin S$ e dimostrare che se S non ha elemento minimo allora $S = \emptyset$)
9. Dimostrare che se ogni sottoinsieme non vuoto di ω ha minimo allora vale il principio di induzione. (Sugg.: supponiamo che $P(-)$ sia la proprietà in esame e consideriamo il sottoinsieme $S \equiv \{x \in \omega \mid \text{non vale } P(x)\}$; ora, se $P(-)$ non valesse per ogni n allora S sarebbe non vuoto e quindi dovrebbe avere minimo, ma allora ...)

Somma e prodotto

Naturalmente i numeri naturali non ci interessano solo in quanto essi sono il nostro primo esempio di insieme infinito ma anche perché sono il nostro primo esempio di insieme numerico, un insieme cioè su cui possiamo definire le *solite* operazioni. Una operazione infatti altro non è che una funzione totale (per una definizione formale di cosa è una funzione si veda il prossimo capitolo, per ora ci accontentiamo di quel che avete imparato alle scuole superiori). Nel nostro caso le operazioni che vogliamo definire sui numeri naturali sono la somma ed il prodotto (attenzione: sottrazione e divisione non sono operazioni sui numeri naturali visto che non si tratta di funzioni totali).

Per definire tali operazioni useremo in modo essenziale il fatto che i numeri naturali sono un insieme induttivo: visto che ogni numero naturale si può raggiungere applicando un numero sufficiente di volte l'operazione di successore per definire una operazione o una proprietà sui numeri naturali è sufficiente definirla per lo zero e per il successore di un qualsiasi numero naturale.

Cominciamo quindi ricordando l'operazione di successore di un numero naturale ponendo

$$S(n) \equiv n + 1 (\equiv n \cup \{n\})$$

Possiamo allora definire la somma tra due numeri naturali come segue

$$\begin{cases} m + 0 &= m \\ m + S(n) &= S(m + n) \end{cases}$$

e, una volta che sappiamo come fare la somma, possiamo definire il prodotto ponendo

$$\begin{cases} m \times 0 &= 0 \\ m \times S(n) &= m + (m \times n) \end{cases}$$

È interessante notare che queste definizioni si possono facilmente programmare in un qualsiasi linguaggio di programmazione che permetta l'uso della ricorsione.

Naturalmente quelle che abbiamo definito altro non sono che le solite operazioni di somma e prodotto. Si potrebbe ora dimostrare che esse godono delle usuali proprietà cui siete abituati che ricordiamo qui sotto.

Definizione A.1.13 (Proprietà della somma) *La somma tra numeri naturali è una operazione che gode delle seguenti proprietà:*

$$\begin{array}{ll} (\text{commutativa}) & m + n = n + m \\ (\text{associativa}) & m + (n + k) = (m + n) + k \\ (\text{elemento neutro}) & m + 0 = m = 0 + m \end{array}$$

Definizione A.1.14 (Proprietà del prodotto) *Il prodotto tra numeri naturali è una operazione che gode delle seguenti proprietà:*

$$\begin{array}{ll} (\text{commutativa}) & m \times n = n \times m \\ (\text{associativa}) & m \times (n \times k) = (m \times n) \times k \\ (\text{elemento neutro}) & m \times 1 = m = 1 \times m \\ (\text{distributiva}) & m \times (n + k) = (m \times n) + (m \times k) \end{array}$$

Esercizi

1. Sulla falsariga della somma, del prodotto e delle seguenti definizioni, definire per ricorsione altre operazioni sui numeri naturali:

- (predecessore)

$$\begin{cases} \text{pred}(0) &= 0 \\ \text{pred}(n+1) &= n \end{cases}$$

- (esponente)

$$\begin{cases} x^0 &= 1 \\ x^{n+1} &= x \times x^n \end{cases}$$

- (fattoriale)

$$\begin{cases} 0! &= 1 \\ n+1! &= (n+1) \times n! \end{cases}$$

A.2 Relazioni e funzioni

Adesso che abbiamo ricostruito i numeri naturali all'interno di un ambiente insiemistico il prossimo passo è recuperare il concetto di funzione.

Naturalmente nel linguaggio comune, e anche nella matematica che avete già visto, avrete già sentito parlare di funzioni (esempio: la temperatura è *funzione* dell'ora del giorno, prendere l'ombrello è *funzione* delle previsioni del tempo, $\sin(x)$ è una *funzione* trigonometrica, ...).

Quel che vogliamo fare qui è analizzare in modo astratto il concetto di funzione.

A.2.1 Relazioni

Il primo passo per ottenere tale definizione è quello di formalizzare il concetto di relazione all'interno della teoria degli insiemi.

Definizione A.2.1 (Coppia ordinata) Siano x e y due insiemi. Diremo coppia ordinata di x e y l'insieme $\langle x, y \rangle \equiv \{x, \{x, y\}\}$.

Si verifica immediatamente che, per ogni coppia di insiemi x e y , la collezione $\langle x, y \rangle$ è un insieme.

Tramite la definizione di coppia ordinata abbiamo voluto introdurre un meccanismo che ci permetta di distinguere il primo e il secondo elemento all'interno di una coppia di elementi; infatti l'*estensionalità* ci costringe a riconoscere come uguali i due insiemi $\{x, y\}$ e $\{y, x\}$ e quindi se vogliamo distinguere un primo elemento nella coppia dobbiamo, in qualche modo, dichiararlo tale. Questo è esattamente ciò che si ottiene tramite la definizione di coppia ordinata; infatti, si specifica appunto che, nella coppia ordinata $\langle x, y \rangle$, x è il primo elemento di una coppia di elementi costituita da $\{x, y\}$.

Lo scopo della definizione è quello di poter dimostrare che $\langle x_1, y_1 \rangle = \langle x_2, y_2 \rangle$ se e solo se $x_1 = x_2$ e $y_1 = y_2$ (vedi esercizi).

Ora, una relazione tra gli elementi dei due insiemi X e Y può essere specificata dicendo quali elementi dell'insieme X sono in relazione con quali elementi dell'insieme Y . Possiamo esprimere questa situazione all'interno della teoria degli insiemi dicendo che una relazione si identifica con la collezione delle coppie ordinate $\langle x, y \rangle$ il cui primo elemento $x \in X$ è in relazione con il secondo elemento $y \in Y$. Come primo passo, questo ci spinge verso la seguente definizione.

Definizione A.2.2 (Assioma del prodotto cartesiano) Siano X e Y due insiemi. Allora la collezione i cui elementi sono coppie ordinate il cui primo elemento sta in X e il cui secondo elemento sta in Y è un insieme che chiameremo prodotto cartesiano di X e Y e indicheremo con $X \times Y$.

Come al solito, possiamo fare ricorso al principio di buona fondatezza per assicurarci della non pericolosità di questa definizione.

È ora evidente che ogni relazione R tra due insiemi X e Y è un sottoinsieme di $X \times Y$ e quindi è a sua volta un insieme.

Esercizi

1. Siano x e y due insiemi. Verificare, utilizzando gli assiomi introdotti nei paragrafi precedenti, che la coppia ordinata $\langle x, y \rangle$ è un insieme.
2. Siano x_1, x_2, y_1 e y_2 degli insiemi. Dimostrare che $\langle x_1, y_1 \rangle = \langle x_2, y_2 \rangle$ se e solo se $x_1 = x_2$ e $y_1 = y_2$.
3. Supponiamo che X sia un insieme con n elementi e Y sia un insieme con m elementi. Quanti sono gli elementi di $X \times Y$?

Proprietà delle relazioni

Esiste per le relazioni tutto un catalogo di nomi in funzione delle particolari proprietà cui una relazione può soddisfare. Ad esempio, dato un insieme X , una relazione $R \subseteq X \times X$ si dice

- (identità) se $R \equiv \{\langle x, x \rangle \in X \times X \mid x \in X\}$;
- (totale) se $R \equiv X \times X$;
- (riflessiva) se, per ogni $x \in X$, $\langle x, x \rangle \in R$;
- (transitiva) se, per ogni $x, y, z \in X$, se $\langle x, y \rangle \in R$ e $\langle y, z \rangle \in R$ allora $\langle x, z \rangle \in R$;
- (simmetrica) se, per ogni $x, y \in X$, se $\langle x, y \rangle \in R$ allora $\langle y, x \rangle \in R$;
- (antisimmetrica) se, per ogni $x, y \in X$, se $\langle x, y \rangle \in R$ e $\langle y, x \rangle \in R$ allora $x = y$;
- (tricotomica) se, per ogni $x, y \in X$, $\langle x, y \rangle \in R$ o $\langle y, x \rangle \in R$ o $x = y$;
- (preordine) se R è *riflessiva* e *transitiva*;
- (di equivalenza) se R è *riflessiva*, *transitiva* e *simmetrica*;
- (ordine parziale) se R è *riflessiva*, *transitiva* e *antisimmetrica*;
- (ordine totale) se R è un *ordine parziale tricotomico*

È utile ricordare anche la seguente operazione sulle relazioni.

Definizione A.2.3 (Relazione inversa) Sia R una relazione tra l'insieme X e l'insieme Y . Allora chiameremo *relazione inversa* di R la relazione definita ponendo

$$R^- \equiv \{\langle y, x \rangle \in Y \times X \mid \langle x, y \rangle \in R\}$$

Esercizi

1. Sia U un insieme e sia R la relazione tra sottoinsiemi di U definita ponendo $\langle X, Y \rangle \in R$ se e solo se $X \subseteq Y$. Dimostrare che R è una relazione d'ordine parziale.
2. Sia $\leq$ la relazione tra elementi di ω definita ponendo $\langle x, y \rangle \in \leq$ se e solo se $x = y$ oppure $x \in y$. Dimostrare che $\leq$ è una relazione d'ordine totale.
3. Sia X un insieme. Dimostrare che la relazione totale è una relazione di equivalenza.
4. Sia X un insieme e R una relazione su X . Dimostrare che se R è simmetrica allora $R = R^-$.

Rappresentazione cartesiana di una relazione

Data una relazione R tra gli insiemi X e Y , risulta molto conveniente la possibilità di darne una rappresentazione grafica che rende a volte esplicite proprietà che non sono evidenti quando la relazione è rappresentata come insieme di coppie.

Possiamo infatti supporre di rappresentare ciascuno degli insiemi X e Y su una retta e di porre le due rette in posizione l'una ortogonale all'altra. Allora l'insieme dei punti del piano individuato dalle rette X e Y rappresenta il prodotto cartesiano $X \times Y$ tra X e Y e una relazione R tra X e Y si può quindi rappresentare come un particolare sottoinsieme di punti di tale piano.

Esercizi

1. Rappresentare sul piano cartesiano una relazione identica. Di quale punti del piano si tratta?
2. Rappresentare nel piano cartesiano una relazione simmetrica. Cosa si nota?
3. Rappresentare sul piano cartesiano $\omega \times \omega$ le relazioni $\leq$ e $\leq^-$.

Relazioni di equivalenza e partizioni

Le relazioni di equivalenza saranno particolarmente rilevanti negli sviluppi futuri. Una relazione di equivalenza gode infatti delle tre principali proprietà che uno si aspetta dalla relazione Id_X di identità tra elementi dell'insieme X , vale a dire la relazione tra gli elementi di un insieme X definita dall'insieme di coppie i cui elementi sono identici:

$$\text{Id}_X \equiv \{\langle x, x \rangle \in X \times X \mid x \in X\}$$

È infatti immediato verificare che la relazione Id_X è in realtà una relazione di equivalenza, vale a dire è una relazione *riflessiva*, *simmetrica* e *transitiva*, ma la cosa più interessante è il fatto che queste proprietà sono caratteristiche di una relazione di identità nel senso che data una qualsiasi relazione di equivalenza noi possiamo sempre costruire un opportuno insieme su cui tale relazione di equivalenza permette di definire una relazione di identità.

Per mostrare questo importante risultato introduciamo prima di tutto la nozione di *partizione* di un insieme.

Definizione A.2.4 (Partizione) *Sia X un insieme. Allora una partizione su X è un insieme I di sottoinsiemi di X tali che*

- per ogni $Y \in I$, $Y \neq \emptyset$
- $X = \bigcup_{Y \in I} Y$
- se Y e Z sono due elementi di I tali che $Y \neq Z$ allora $Y \cap Z = \emptyset$

Intuitivamente, una partizione di un insieme X è quindi una maniera di dividere X in pezzi separati non vuoti. Infatti la terza condizione sostiene che o due pezzi sono uguali o sono disgiunti. Per di più, dato un qualsiasi elemento $x \in X$, in virtù della seconda e della terza condizione, possiamo sempre individuare quell'unico elemento $Y_x \in I$ tale che $x \in Y_x$.

Ora è immediato verificare che una partizione I su un insieme X determina sempre una relazione di equivalenza R_I tra gli elementi di X definita ponendo, per ogni $x, y \in X$:

$$x R_I y \equiv y \in Y_x$$

vale a dire che $x R_I y$ vale se e solo se x e y appartengono allo stesso elemento della partizione I (verificare per esercizio che R è una relazione di equivalenza).

La cosa più interessante è però che anche una relazione di equivalenza determina una partizione, vale a dire che partizioni e relazioni di equivalenza sono due modi diversi di parlare della stessa cosa.

Definizione A.2.5 (Classe di equivalenza) Sia R una relazione di equivalenza su un insieme X e sia x un elemento di X . Allora la classe di equivalenza di x è il sottoinsieme di X definito da

$$[x]_R \equiv \{y \in X \mid x R y\}$$

Possiamo allora dimostrare il seguente teorema.

Teorema A.2.6 Sia R una relazione di equivalenza sull'insieme X . Allora l'insieme

$$X/R \equiv \{[x]_R \mid x \in X\}$$

delle classi di equivalenza di R è una partizione su X .

Dimostrazione. La prima cosa da osservare è che $\{[x]_R \mid x \in X\}$ è un insieme in quanto, per ogni $x \in X$, $[x]_R$ è un insieme, in quanto sottoinsieme di X , e quindi il principio di buona fondazione ci assicura che non corriamo nessun rischio (si veda comunque il prossimo paragrafo dove vedremo che la collezione $\mathcal{P}(X)$ di tutti i sottoinsiemi di un dato insieme X è un insieme; in virtù dell'assioma del sottoinsieme, questo fatto chiaramente significa che $\{[x]_R \mid x \in X\}$ è un insieme in quanto sottoinsieme di $\mathcal{P}(X)$).

Osserviamo ora che, in virtù della *riflessività* di R , nessuna classe di equivalenza è vuota in quanto, per ogni $x \in X$, $x \in [x]_R$.

Come conseguenza di quanto appena detto si vede subito che ogni elemento di X appartiene ad una qualche classe di equivalenza, vale a dire la *sua* classe di equivalenza, e quindi l'unione di tutte le classi di equivalenza è sicuramente uguale ad X .

Consideriamo ora due classi di equivalenza $[x]_R$ e $[y]_R$. Dobbiamo allora fare vedere che o $[x]_R = [y]_R$ o $[x]_R \cap [y]_R = \emptyset$. Supponiamo allora che $[x]_R \cap [y]_R \neq \emptyset$. Allora esiste un elemento $z \in X$ tale che $z \in [x]_R$ e $z \in [y]_R$, vale a dire che $x R z$ e $y R z$. Ma allora $z R y$ segue per *simmetria* e quindi si ottiene $x R y$ per *transitività*.

A conferma che partizioni e relazioni di equivalenza sono due modi alternativi di presentare lo stesso concetto, non è difficile vedere che le due costruzioni che abbiamo proposto sono una l'inversa dell'altra. Infatti se R è una relazione di equivalenza, I_R è la partizione associata ad R e R_{I_R} la relazione di equivalenza associata alla partizione I_R , allora, per ogni $x, y \in X$,

$$\begin{array}{ll} x R_{I_R} y & \text{sse } y \in [x]_R \\ & \text{sse } x R y \end{array}$$

D'altra parte, per ogni partizione I , la partizione I_{I_R} associata alla relazione di equivalenza R_I associata ad I coincide con la partizione I stessa. Infatti, per ogni $Y \subseteq X$,

$$\begin{array}{ll} Y \in I_{R_I} & \text{sse } (\exists x \in X) Y = [x]_{R_I} \\ & \text{sse } (\exists x \in X) Y = \{z \in X \mid x R_I z\} \\ & \text{sse } (\exists x \in X) Y = \{z \in X \mid z \in Y_x\} \\ & \text{sse } (\exists x \in X) Y = Y_x \\ & \text{sse } Y \in I \end{array}$$

È ora possibile vedere che data una relazione di equivalenza R su un insieme X essa induce una relazione identica sull'insieme X/R . Infatti possiamo porre, per ogni coppia di elementi $[x]_R, [y]_R \in X/R$,

$$[x]_R \text{ Id}_{X/R} [y]_R \equiv x R y$$

Dobbiamo ora verificare che questa è una *buona definizione*, cioè una definizione corretta che non dipende dai particolari rappresentanti usati nelle classi di equivalenze. Supponiamo quindi che $[x_1]_R = [x_2]_R$ e $[y_1]_R = [y_2]_R$. Allora ricaviamo subito che $x_1 R x_2$ e $y_1 R y_2$. Se ora supponiamo che $[x_1]_R \text{ Id}_{X/R} [y_1]_R$ otteniamo subito che $x_1 R y_1$ e quindi, utilizzando le proprietà *simmetrica* e *transitiva* della relazione R , si vede immediatamente che $x_2 R y_2$ e naturalmente questo implica che $[x_2]_R \text{ Id}_{X/R} [y_2]_R$. Adesso è immediato verificare che $\text{Id}_{X/R}$ è una relazione di identità, vale a dire che se $[x]_R \text{ Id}_{X/R} [y]_R$ allora $[x]_R = [y]_R$. Infatti $[x]_R \text{ Id}_{X/R} [y]_R$ significa che $x R y$ che vuol dire appunto che $[x]_R = [y]_R$.

Esercizi

1. Sia X un insieme e I sia una sua partizione. Verificare che la relazione R_I definita ponendo

$$x R_I y \equiv (Y_x = Y_y)$$

è una relazione di equivalenza.

2. Considerate l'insieme dei numeri interi. Possiamo farne una partizione considerando il sottoinsieme di tutti i numeri positivi, il sottoinsieme costituito dal solo numero 0 e il sottoinsieme costituito da tutti i numeri negativi. Trovare la relazione di equivalenza che corrisponde a tale partizione.
3. Considerate la relazione di equivalenza che considera uguali due ore x e y quando esse differiscono per un multiplo di 12 ore, vale a dire quando esiste un numero intero k tale che $x - y = 12 * k$. Quali e quante sono le classi di equivalenza associate a tale relazione?
4. (Costruzione dei numeri interi) Si supponga di voler costruire un insieme che possa essere identificato come l'insieme dei numeri interi pur avendo a disposizione, come nel nostro caso, solamente i numeri naturali. Possiamo allora indicare l'intero i tramite una coppia (m, n) di numeri naturali tali che, *se fossimo in grado di fare sempre la sottrazione tra due numeri naturali*, $i = m - n$. Naturalmente questo approccio soffre di due problemi: senza gli interi *non* possiamo fare sempre la sottrazione tra due numeri naturali e, inoltre, lo stesso intero viene ad essere rappresentato da infinite coppie di numeri naturali.

Per risolvere entrambi questi problemi possiamo considerare sull'insieme $\omega \times \omega$ delle coppie ordinate di numeri naturali consideriamo la seguente relazione

$$\langle m, n \rangle \cong \langle p, q \rangle \equiv (m + q) = (p + n)$$

Dimostrare che $\cong$ è una relazione di equivalenza sull'insieme $\omega \times \omega$.

Possiamo ora considerare l'insieme delle classi di equivalenza $\omega \times \omega / \cong$ come l'insieme Z degli interi.

Infatti all'interno di Z possiamo facilmente ritrovare una copia dei numeri naturali se al numero naturale n associamo la classe di equivalenza determinata dalla coppia $\langle n, 0 \rangle$, ma possiamo anche trovare i *numeri negativi* visto che possiamo rappresentare il *vecchio* numero intero $-n$ tramite la classe di equivalenza il cui rappresentante è $\langle 0, n \rangle$.

Inoltre possiamo definire su Z la usuale operazione di somma ponendo

$$[\langle m, n \rangle]_{\cong} + [\langle p, q \rangle]_{\cong} \equiv [\langle m + p, n + q \rangle]_{\cong}$$

Verificare che si tratta di una *buona definizione* che non dipende dai rappresentanti usati.

La novità rispetto al caso dei numeri naturali è che su Z possiamo finalmente definire anche la sottrazione ponendo

$$[\langle m, n \rangle]_{\cong} - [\langle p, q \rangle]_{\cong} \equiv [\langle m + q, n + p \rangle]_{\cong}$$

Si verifichi anche in questo caso che si tratta di una buona definizione e che

$$[\langle m, 0 \rangle]_{\cong} - [\langle n, 0 \rangle]_{\cong} \equiv [\langle m, n \rangle]_{\cong}$$

5. Sulla falsariga dell'esercizio precedente fornire una costruzione per l'insieme Q dei numeri razionali. (sugg.: non scordatevi delle frazioni!)

A.2.2 Funzioni

Lo scopo principale di questa sezione è la definizione del concetto di funzione. Dopo aver introdotto nella sezione precedente le relazioni possiamo utilizzarle per scegliere al loro interno le particolari relazioni che sono funzioni.

Definizione A.2.7 Siano X e Y due insiemi. Allora una relazione $R \subseteq X \times Y$ è una funzione se e solo se, per ogni $x \in X$, se $\langle x, y \rangle \in R$ e $\langle x, z \rangle \in R$ allora $y = z$.

Un'attenta lettura della precedente definizione rivela che le funzioni sono quelle relazioni da X ad Y tali che un elemento x dell'insieme X di partenza (nel seguito X sarà chiamato anche *dominio della funzione*) è in relazione con al più con un solo elemento y dell'insieme Y di arrivo (nel seguito Y sarà chiamato anche *codominio della funzione*). Questo fatto si può anche esprimere sostenendo che la scelta di x *determina* la scelta di y , vale a dire che il valore di y è *funzione* di x . Il fatto che ci sia al più un solo elemento $y \in Y$ che è in relazione con $x \in X$ ci permette di semplificare un po' la notazione per le funzioni rispetto a quella per le relazioni. Infatti, potremo scrivere $R(x) = y$ al posto di $\langle x, y \rangle \in R$ e dire che y è l'*immagine* di x tramite la funzione R .

Proprietà delle funzioni

Così come nel caso delle relazioni abbiamo individuato delle proprietà particolari di cui una relazione può godere, e abbiamo deciso una volta per tutte i nomi di tali proprietà, conviene fare qualcosa di analogo anche nel caso delle funzioni.

Definizione A.2.8 (Funzione totale) Siano X e Y due insiemi e sia f una funzione da X a Y . Allora diremo che f è una funzione totale se, per ogni $x \in X$, esiste un $y \in Y$, e quindi uno solo, tale che la coppia ordinata $\langle x, y \rangle$ appartiene ad f , vale a dire se la funzione f è definita per ogni elemento del dominio.

Definizione A.2.9 (Funzione iniettiva) Siano X e Y due insiemi e sia f una funzione da X a Y . Allora diremo che f è una funzione iniettiva se accade che elementi distinti di X hanno immagini distinte, vale a dire se ogni volta che un elemento $y \in Y$ è immagine di un qualche elemento $x \in X$ tale x è unico, vale a dire che se $\langle x_1, y \rangle \in f$ e $\langle x_2, y \rangle \in f$ allora $x_1 = x_2$.

Definizione A.2.10 (Funzione suriettiva) Siano X e Y due insiemi e sia f una funzione da X a Y . Allora diremo che f è una funzione suriettiva se accade che ogni elemento di Y è immagine di un qualche elemento di X .

Definizione A.2.11 (Funzione biettiva) Siano X e Y due insiemi e sia f una funzione da X a Y . Allora diremo che f è una funzione biettiva se f è sia iniettiva che suriettiva.

Definizione A.2.12 (Immagine diretta) Siano X e Y due insiemi e sia f una funzione da X a Y . Allora chiameremo immagine diretta

$$f(U) \equiv \{y \in Y \mid \text{esiste } x \in U \text{ tale che } y = f(x)\}$$

Definizione A.2.13 (Immagine inversa) Siano X e Y due insiemi e sia f una funzione da X ad Y . Allora chiameremo immagine inversa

$$f^{-1}(V) \equiv \{x \in X \mid \text{esiste } y \in V \text{ tale che } y = f(x)\}$$

Definizione A.2.14 (Funzione inversa) Siano X e Y due insiemi e sia f una funzione iniettiva da X ad Y . Allora chiameremo funzione inversa di f la relazione opposta f^{-} .

Definizione A.2.15 (Composizione di relazioni e di funzioni) Siano X , Y e Z tre insiemi e R sia una relazione tra X e Y e S sia una relazione tra Y e Z . Allora possiamo definire la composizione di R ed S ponendo

$$S \circ R \equiv \{\langle x, z \rangle \in X \times Z \mid \text{esiste } y \in Y \text{ tale che } \langle x, y \rangle \in R \text{ e } \langle y, z \rangle \in S\}$$

(si legge S dopo R).

Vale la pena di notare che la funzione inversa di una funzione totale può essere non totale. Tuttavia nel caso in cui la funzione f sia non solo iniettiva ma anche suriettiva abbiamo che la funzione f^{-1} è pure una funzione biettiva e $f \circ f^{-1} = \text{id}_X$ e $f^{-1} \circ f = \text{id}_Y$.

Esercizi

1. Fornire un esempio di una funzione totale, di una funzione iniettiva, di una funzione suriettiva e di una funzione biettiva.
2. Dimostrare che la composizione di relazioni è una operazione associativa, vale a dire dimostrare che se f è una funzione da X verso Y , g è una funzione da Y verso W e h è una funzione da W verso Z allora

$$(h \circ g) \circ f = h \circ (g \circ f)$$

3. Dimostrare che la relazione identica è l'identità per l'operazione di composizione.
4. Siano X e Y due insiemi e sia f una funzione da X verso Y . Dimostrare che

- $f(U_1 \cup U_2) = f(U_1) \cup f(U_2)$
- $f(U_1 \cap U_2) \subseteq f(U_1) \cap f(U_2)$
Trovare un controesempio alla inclusione mancante.
- $f^-(V_1 \cup V_2) = f^-(V_1) \cup f^-(V_2)$
- $f^-(V_1 \cap V_2) = f^-(V_1) \cap f^-(V_2)$

Costruzione dei numeri reali

A differenza dei numeri interi e dei numeri razionali che si possono sempre presentare utilizzando una quantità finita di informazione, vale a dire una coppia di numeri naturali, ogni numero reale richiede di specificare una quantità infinita di informazione, vale a dire, ad esempio, tutte le infinite cifre di una sua espansione decimale. I primi oggetti che possono racchiudere tutta questa informazione in un solo insieme sono le relazioni e le funzioni; questo è il motivo per cui useremo particolari insiemi di funzioni per rappresentare i numeri reali all'interno della teoria degli insiemi.

L'intuizione che sta alla base della particolare scelta che faremo è che possiamo specificare un numero reale r se specifichiamo tutte le sue approssimazioni decimali, vale a dire, supponendo che

$$r \equiv n, n_0 n_1 n_2 n_3 \dots$$

dove n è la parte intera di r e n_0, n_1, n_2 sono le varie cifre decimali, noi possiamo specificare la seguente *successione* di numeri razionali

$$\begin{aligned} q_0 &\equiv n, n_0 \\ q_1 &\equiv n, n_0 n_1 \\ q_2 &\equiv n, n_0 n_1 n_2 \\ &\dots \end{aligned}$$

Ora, una successione di numeri naturali altro non è che una funzione

$$r : \omega \rightarrow Q$$

Naturalmente, il problema è che non tutte le funzioni da ω in Q possono rappresentare un numero reale ma vanno bene solo le successioni che *convergono*. Tra tutte le funzioni da ω in Q dobbiamo quindi scegliere solo quelle i cui valori, da un certo punto in poi, siano tra loro tanto vicini di una qualsiasi distanza decisa a priori (formalmente possiamo esprimere questa condizione dicendo che una funzione $r : \omega \rightarrow Q$ è un numero reale se accade che, fissato un qualsiasi $\epsilon \in Q$ esiste un $M \in \omega$ tale che, per ogni $n_1, n_2 \in \omega$ che siano maggiori di M accade che $|q_{n_1} - q_{n_2}| \leq \epsilon$).

Purtroppo non siamo ancora arrivati in fondo con la nostra costruzione. Considerate infatti le due successioni seguenti:

$$\begin{aligned} r_1 &\equiv 0 \frac{1}{2} \frac{3}{4} \frac{7}{8} \dots \\ r_2 &\equiv 0 \frac{3}{2} \frac{3}{4} \frac{9}{8} \dots \end{aligned}$$

Si tratta chiaramente di due successioni diverse ma non è difficile convincersi che sono entrambi convergenti a 1. Siamo cioè nella solita situazione in cui abbiamo più modi diversi per denotare lo stesso oggetto (ad esempio avevamo molte coppie di numeri naturali che denotavano lo stesso numero intero o molte coppie di frazioni che denotavano lo stesso numero razionale). La soluzione

è quindi la solita: i numeri reali non saranno semplicemente l'insieme delle funzioni dai naturali ai razionali che convergono ma piuttosto le classi di equivalenza su tale insieme di funzioni determinate dalla relazione di equivalenza che pone equivalenti due funzioni convergenti quando convergono allo stesso valore (formalmente, le funzioni convergenti r_1 e r_2 da ω verso Q sono in relazione se accade che, fissato un qualsiasi $\epsilon \in Q$ esiste un $M \in \omega$ tale che, per ogni $n \in \omega$ maggiore di M accade che $|q_n^{r_1} - q_n^{r_2}| \leq \epsilon$).

Esercizi

1. Dimostrare che la relazione sopra definita tra funzioni convergenti è una relazione di equivalenza
2. Definire le usuali operazioni tra numeri reali.
3. Definire la relazione d'ordine sui numeri reali.
4. Definire la radice quadrata di un numero reale

Grafico di una funzione

Visto che una funzione altro non è che una speciale relazione il grafico di una funzione è semplicemente un caso speciale del grafico di una relazione. Tuttavia, visto l'ampio spettro di applicazioni, lo studio dei grafici delle funzioni costituirà tuttavia un grosso capitolo dei vostri studi futuri per cui non si ritiene necessario insistere fin da ora sull'argomento.

Esercizi

1. Si consideri la seguente funzione f dall'insieme Q dei numeri razionali verso l'insieme Z dei numeri interi

$$f \equiv \{\langle x, y \rangle \in Q \times Z \mid y \text{ è uguale ad } x \text{ o è il numero intero immediatamente minore a } x\}$$

(f si chiama la funzione *parte intera*)

Se ne tracci il grafico.

2. Si tracci il grafico della seguente funzione f dall'insieme Q verso se stesso

$$f \equiv \{\langle x, y \rangle \in Q \times Q \mid y = x + 1\}$$

A.3 Insiemi induttivi non numerici

Nei precedenti paragrafi abbiamo visto come costruire i naturali e, a partire da essi, tutti gli altri insiemi numerici. I naturali non sono tuttavia l'unico esempio di insieme che riusciamo a costruire tramite una definizione induttiva. Infatti, in questo capitolo vogliamo introdurre altri esempi di insiemi definiti tramite un processo induttivo. Essi saranno utili nel seguito dei vostri studi anche se non si tratta di insiemi numerici.

A.3.1 Liste

Il prossimo esempio di insieme induttivo che vogliamo introdurre è una semplice generalizzazione dell'insieme dei numeri naturali. Se analizziamo la definizione induttiva dei numeri naturali ci possiamo infatti rendere conto che possiamo darne una rappresentazione della seguente forma

$$0, 0I, 0II, 0III, \dots$$

vale a dire che, a partire da 0, possiamo ottenere un generico numero naturale aggiungendo un nuovo I alla destra di qualcosa che sia stato già costruito come numero naturale. Possiamo allora generalizzare questa costruzione se invece che aggiungere sempre lo stesso simbolo I aggiungiamo

un simbolo preso da un insieme fissato di simboli. Ad esempio se consideriamo l'insieme delle lettere dell'alfabeto e sostituiamo lo 0 con un simbolo bianco, che indicheremo con *nil*, otteniamo l'insieme delle *parole*, vale a dire le liste finite di lettere. I numeri naturali sono quindi le parole che si possono scrivere usando un'unica lettera.

Definizione A.3.1 (Liste su un insieme) *Sia A un insieme. Allora indicheremo con $\text{List}(A)$ l'insieme i cui elementi sono le parole ottenute utilizzando gli elementi di A come lettere, vale a dire che gli elementi di $\text{List}(A)$ sono i seguenti:*

- (*Base*) *nil è un elemento di $\text{List}(A)$;*
- (*Passo*) *Se l è un elemento di $\text{List}(A)$ e a è un elemento di A allora $l \cdot a$ è un elemento di $\text{List}(A)$.*

Visto che l'insieme delle liste nasce tramite una definizione induttiva è immediato pensare ad una generalizzazione della proprietà di induzione per dimostrare proprietà sui suoi elementi.

Definizione A.3.2 (Induzione sulle liste) *Sia A un insieme e sia $\text{List}(A)$ l'insieme delle liste sull'insieme A . Sia inoltre $P(-)$ una proprietà sugli elementi dell'insieme $\text{List}(A)$. Allora per dimostrare che $P(-)$ vale per ogni lista l basta dimostrare che vale $P(\text{nil})$ e che se vale $P(l)$ allora vale anche $P(l \cdot a)$ per ogni elemento $a \in A$.*

Esercizi

1. Sia A un insieme. Si definisca la funzione *length* che data una lista l di $\text{List}(A)$ ne fornisce la lunghezza (si noti che $\text{length}(\text{nil}) = 0$).
2. Si consideri la seguente definizione di una funzione sulle coppie di liste l_1 e l_2 di $\text{List}(A)$ il cui intento è quello di fornire la lista ottenuta concatenando l_1 e l_2

$$\begin{cases} \text{append}(m, \text{nil}) &= m \\ \text{append}(m, l \cdot a) &= \text{append}(m, l) \cdot a \end{cases}$$

Dimostrare per induzione che, per ogni coppia di liste l_1 e l_2 ,

$$\text{length}(\text{append}(l_1, l_2)) = \text{length}(l_1) + \text{length}(l_2)$$

3. Si consideri la seguente funzione sugli elementi di $\text{List}(A)$ che presa una lista l fornisce la lista opposta di l .

$$\begin{cases} \text{reverse}(\text{nil}) &= \text{nil} \\ \text{reverse}(l \cdot a) &= \text{append}(a \cdot \text{nil}, \text{reverse}(l)) \end{cases}$$

Dimostrare che, per ogni lista l ,

$$\text{length}(l) = \text{length}(\text{reverse}(l))$$

4. Definire una funzione *norep* che data una lista l di $\text{List}(A)$ fornisce una nuova lista i cui elementi sono esattamente quelli di l ma nessun *carattere* appare due volte.
5. Verificare che la seguente funzione $\text{ord}(-) : \text{List}(\omega) \rightarrow \text{List}(\omega)$ quando applicata ad una lista di numeri naturali fornisce una lista ordinata che contiene gli stessi elementi:

$$\begin{cases} \text{ord}(\text{nil}) &= \text{nil} \\ \text{ord}(l \cdot a) &= \text{insert}(a, \text{ord}(l)) \end{cases}$$

$$\begin{cases} \text{insert}(a, \text{nil}) &= \text{nil} \cdot a \\ \text{insert}(a, l \cdot b) &= \text{if } b \leq a \text{ then } l \cdot b \cdot a \text{ else } \text{insert}(a, l) \cdot b \end{cases}$$

A.3.2 Alberi

Un altro esempio di insieme induttivo, notevolmente più complesso del precedente, sono gli alberi. In queste note ci limiteremo a considerare gli *alberi binari* che risultano comunque notevoli per la ricchezza delle loro applicazioni. Come al solito dobbiamo cominciare con una definizione induttiva degli elementi dell'insieme degli alberi binari.

Definizione A.3.3 (Alberi binari) *Sia A un insieme. Allora gli elementi di $\text{BinTree}(A)$ sono i seguenti:*

- (Base) se a è un elemento di A allora $\text{leaf}(a)$ è un elemento di $\text{BinTree}(A)$;
- (Passo) se t_1 e t_2 sono due elementi di $\text{BinTree}(A)$ allora (t_1, t_2) è un elemento di $\text{BinTree}(A)$

Può essere utile osservare che in letteratura si incontrano molte altre definizioni sostanzialmente analoghe del concetto di albero binario.

Esercizi

1. Definire la funzione $\text{depth}(-)$ che dato un albero binario b ne fornisce la profondità, vale a dire la lunghezza del ramo più lungo.

A.4 Cardinalità

È immediato rendersi conto che il principio di buona fondatezza ci assicura che possiamo riconoscere come insieme anche la collezione di tutti i sottoinsiemi di un qualunque insieme.

Definizione A.4.1 (Assioma delle parti) *Sia X un insieme. Allora la collezione di tutti i sottoinsiemi di X è un insieme che chiameremo potenza dell'in-sieme X e indicheremo con $\mathcal{P}(X)$.*

Infatti, per l'assioma del sottoinsieme, ogni sottoinsieme Y di X è a sua volta un insieme e quindi la collezione $\mathcal{P}(X)$ è formata solo da elementi che abbiamo già riconosciuto come insiemi nel momento in cui essa viene formata. Bisogna tuttavia ammettere che se l'insieme X è infinito, come ad esempio ω , il processo di riconoscimento di tutti i sottoinsiemi sembra proprio essere tutto meno che effettivo!

Ora, supponiamo che X sia un insieme finito e indichiamo con $|X|$ il numero di elementi di X , che chiameremo *cardinalità* dell'insieme X . È allora facile dimostrare, usando ad esempio il principio di induzione, che se $|X| = n$ allora $|\mathcal{P}(X)| = 2^n$. Ma cosa succede nel caso X non sia un insieme con un numero finito di elementi? È facile convincersi che anche $\mathcal{P}(X)$ conterrà un insieme non finito di elementi, ma quanti sono? Sono tanti quanti gli elementi in X ? Sono di più come accade nel caso in cui X sia un insieme finito? Per rispondere a queste domande è necessario prima di tutto chiarire cosa si deva intendere in generale con l'operazione di contare il numero di elementi di un insieme.

L'idea generale è quella di rifarsi al modo di contare dei bambini che contano sulle dita, ricordandosi però che noi abbiamo mani con infinite dita! Quello che faremo per contare non sarà cioè, in prima istanza, lo scoprire quanti elementi compaiono in un insieme ma cercheremo piuttosto di capire quando due insiemi abbiano lo stesso numero di elementi. Poi, nello stesso modo di un bambino, diremo che ci sono cinque oggetti quando ci sono tanti oggetti quante sono le dita di una sola mano.

In realtà, abbiamo già predisposto tutto quello che ci serve per risolvere questo primo problema almeno nel caso degli insiemi con un numero finito di elementi. In questo caso infatti due insiemi hanno lo stesso numero di elementi esattamente quando esiste una funzione totale e biettiva dal primo insieme verso il secondo. Infatti, supponiamo di avere due insiemi finiti X e Y e supponiamo che Y abbia meno elementi di X ; allora è impossibile definire una funzione tra X e Y che sia totale e iniettiva. D'altra parte se supponiamo che X abbia meno elementi di Y è impossibile definire tra X e Y una funzione che sia suriettiva. Quindi, nel caso stiamo considerando due insiemi finiti per vedere se essi hanno lo stesso numero di elementi possiamo provare a trovare una funzione totale e biettiva dal primo al secondo. Il trucco consiste ora nell'usare lo stesso test come definizione del fatto che due insiemi, anche non finiti, hanno lo stesso numero di elementi.

Definizione A.4.2 (Equinumerosità) *Siano X e Y due insiemi. Allora diremo che X e Y sono equinumerosi se e solo se esiste una funzione totale e biettiva da X verso Y .*

In questo modo siamo riusciti ad avere una nozione di equinumerosità tra due insiemi che non richiede di parlare del numero di elementi di un insieme che è una nozione chiaramente problematica, almeno per quanto riguarda gli insiemi con un numero non finito di elementi.

A.4.1 Insiemi numerabili

Bisogna però notare subito che questa nozione di equinumerosità produce talvolta risultati inaspettati. Ad esempio, una sua conseguenza immediata è il fatto che l'insieme P dei numeri pari risulta equinumeroso con l'insieme ω dei numeri naturali visto che possiamo definire la seguente mappa **double** da ω verso P

$$\text{double}(n) \equiv 2n$$

che è chiaramente totale e biettiva. Abbiamo cioè ottenuto che l'insieme dei numeri naturali è equinumeroso con un suo sottoinsieme proprio! Naturalmente questo non dovrebbe stupirci più che tanto visto che nessuno può garantire che una nozione di equinumerosità tra due insiemi, che nel caso tali insiemi siano finiti significa che essi hanno lo stesso numero di elementi e quindi esclude che un insieme finito possa avere lo stesso numero di elementi di un suo sottoinsieme proprio, preservi per intero il suo significato quando essa viene utilizzata su insiemi che finiti non sono. Potremo invece approfittare di questa mancanza di uniformità della nozione di equinumeroso per distinguere gli insiemi finiti da quelli che finiti non sono.

Definizione A.4.3 (Insiemi finiti e infiniti) *Sia X un insieme. Allora diremo che X è in insieme infinito se esso è equinumeroso con una sua parte propria. Diremo invece che X è un insieme finito quando esso non è infinito.*

Non è difficile accorgersi che anche l'insieme dei numeri interi è equinumeroso con l'insieme dei numeri naturali. Infatti possiamo utilizzare la seguente mappa **Nat2Int**

$$\text{Nat2Int}(n) = (-1)^n \lfloor \frac{n+1}{2} \rfloor$$

dove $\lfloor q \rfloor$ rappresenta la parte intera del numero razionale q , che risulta evidentemente totale e biettiva.

Un po' più sorprendente è forse il fatto che anche l'insieme delle frazioni, tramite le quali è possibile dare un nome a tutti numeri razionali (in realtà infiniti nomi per ciascun numero razionale!), è equinumeroso con l'insieme dei numeri naturali. In generale, se trascuriamo il problema delle frazioni con denominatore nullo, possiamo pensare ad una frazione come ad una coppia ordinata di numeri naturali. Dobbiamo quindi trovare una mappa totale e biettiva dall'insieme dei numeri naturali all'insieme delle coppie ordinate di numeri naturali. In pratica è più facile risolvere il problema nell'altra direzione, vale a dire definire una mappa totale e biettiva dall'insieme delle coppie ordinate di numeri naturali verso l'insieme dei numeri naturali. È facile rendersi conto che vale il seguente lemma e quindi trovare questa mappa è completamente equivalente a trovarne una dall'insieme dei numeri naturali all'insieme delle coppie.

Lemma A.4.4 *Siano X e Y due insiemi e sia f una mappa totale e biettiva da X a Y . Allora è possibile definire una mappa totale e biettiva da Y a X .*

Dimostrazione. Si definisca la mappa g da Y a X nel modo seguente:

$$\langle y, x \rangle \in g \equiv \langle x, y \rangle \in f$$

L'idea è naturalmente che la mappa g associa ad $y \in Y$ quell'unico elemento $x \in X$ tale che $f(x) = y$. Dobbiamo allora dimostrare tre cose: prima di tutto bisogna far vedere che g è davvero una funzione, poi che si tratta di una funzione totale e infine che è una funzione biettiva.

- (buona definizione) Supponiamo che sia $\langle y, x_1 \rangle$ che $\langle y, x_2 \rangle$ siano elementi di g . Allora dobbiamo dimostrare che $x_1 = x_2$. Ma noi sappiamo che se $\langle y, x_1 \rangle \in g$ allora $\langle x_1, y \rangle \in f$ e che se $\langle y, x_2 \rangle \in g$ allora $\langle x_2, y \rangle \in f$. Ma allora, visto che f è iniettiva in quanto biettiva, $x_1 = x_2$.
- (totalità) Bisogna dimostrare che, per ogni elemento $y \in Y$ esiste un $x \in X$ tale che $\langle y, x \rangle \in g$. Ma noi sappiamo che affinché questo accada basta che, per ogni elemento $y \in Y$ esiste un $x \in X$ tale che $\langle x, y \rangle \in f$ che vale visto che f è suriettiva in quanto biettiva.
- (biattività) Dobbiamo dimostrare che g è sia iniettiva che suriettiva. Dimostrare l'iniettività significa far vedere che se $\langle y_1, x \rangle$ e $\langle y_2, x \rangle$ sono elementi di g allora $y_1 = y_2$. Ora $\langle y_1, x \rangle \in g$ implica che $\langle x, y_1 \rangle \in f$ e $\langle y_2, x \rangle \in g$ implica che $\langle x, y_2 \rangle \in f$. Ma allora otteniamo immediatamente che $y_1 = y_2$ visto che f è una funzione. Per quanto riguarda la surattività di g , dobbiamo dimostrare che per ogni $x \in X$ esiste un $y \in Y$ tale che $\langle y, x \rangle \in g$. Ora per dimostrare questo risultato basta far vedere che per ogni $x \in X$ esiste un $y \in Y$ tale che $\langle x, y \rangle \in f$ che è una immediata conseguenza della totalità di f .

Possiamo ora risolvere il nostro problema originale esibendo una mappa totale e biettiva che fornisce un numero naturale in corrispondenza di ogni coppia ordinata di numeri naturali:

$$\text{Pair2Nat}(\langle x, y \rangle) \equiv \frac{1}{2}(x + y)(x + y + 1) + x$$

Sebbene il risultato non sia immediato si può dimostrare che $\text{Pair2Nat}(-)$ gode delle proprietà richieste.

A.4.2 Insiemi più che numerabili

Dopo aver visto che i numeri razionali non sono di più che i numeri naturali può forse risultare sorprendente il fatto che non tutti gli insiemi infiniti sono equinumerosi. Infatti con un argomento dovuto a Cantor, e che useremo spesso nel seguito, possiamo far vedere che l'insieme dei sottoinsiemi dei numeri naturali non è equinumeroso con l'insieme dei numeri naturali. È ovvio che per dimostrare questo risultato basta far vedere che ogni funzione dall'insieme ω verso l'insieme $\mathcal{P}(\omega)$ non è suriettiva.

Teorema A.4.5 (Teorema di Cantor) *Ogni funzione f da ω a $\mathcal{P}(\omega)$ non è suriettiva.*

Dimostrazione. Basta far vedere che esiste un sottoinsieme di numeri naturali che non è immagine secondo f di nessun numero naturale. È ovvio che tale sottoinsieme dipende dalla funzione f che stiamo considerando. Una possibile scelta è la seguente

$$U_f \equiv \{n \in \omega \mid n \notin f(n)\}$$

Non dovrebbe essere difficile riconoscere nella definizione del sottoinsieme U_f una forte analogia con il paradosso di Russell, ma questa volta la stessa idea, la *diagonalizzazione*, viene utilizzata per definire il sottoinsieme che goda della proprietà richiesta invece che per trovare una contraddizione. Tuttavia il metodo è analogo: faremo infatti vedere che l'ipotesi che U_f sia nell'immagine di f porta a contraddizione e quindi la funzione f non può essere suriettiva. Supponiamo infatti che esista un numero naturale n tale che $U_f = f(n)$. Allora, $n \in f(n)$ se e solo se $n \in U_f$ e quindi, per definizione di U_f , se e solo se $n \notin f(n)$, cioè abbiamo ottenuto una contraddizione.

Esercizi

1. Siano X e Y due insiemi tali che $|X| = n$ e $|Y| = m$. Determinare il numero degli elementi di $X \times Y$. Quante sono le relazioni tra X e Y ? Quante sono le funzioni iniettive da X a Y ? (Difficile) Quante sono le funzioni tra X e Y ?
2. Sia X un insieme con n elementi. Allora l'insieme $\mathcal{P}(X)$ ha 2^n elementi.

3. Dimostrare che la funzione

$$\text{Nat2Int}(n) = (-1)^n \lfloor \frac{n+1}{2} \rfloor$$

che manda un numero naturale in un numero intero è totale e biettiva.

4. (Difficile) Dimostrare che la funzione

$$\text{Pair2Nat}(\langle x, y \rangle) \equiv \frac{1}{2}(x+y)(x+y+1) + x$$

che mappa una coppia ordinata di numeri naturali in un numero naturale è totale e biettiva.

5. Dimostrare che l'insieme F delle funzioni totali da ω a ω è un insieme più che numerabile. (sugg.: supponete che $\phi(-)$ sia una funzione da ω verso F e dimostrate che $\phi(-)$ non può essere suriettiva; infatti, basta definire una funzione $f_\phi(-)$ da ω a ω ponendo $f_\phi(n) = \phi(n)(n) + 1$ per ottenere una funzione totale da ω a ω tale che ...)
6. Dimostrare che l'insieme dei numeri reali compresi tra 0 e 1 è più che numerabile. (sugg.: supponete che $\phi(-)$ sia una funzione da ω verso $[0, 1]$ e dimostrate che $\phi(-)$ non può essere suriettiva; infatti, basta definire una funzione $f_\phi(-)$ da ω a ω che fornisce le varie cifre decimali di un numero reale r ponendo $f_\phi(n) = (\phi(n)[n] + 1) \bmod 10$, dove con $\phi(n)[n]$ intendiamo la n -ma cifra decimale del numero reale $\phi(n)$ e dimostrare che r non può essere nell'immagine di $\phi(-)$)

Appendice B

Espressioni con arietà

B.1 Premessa

Prima di parlare della teoria delle espressioni vera e propria richiamiamo alcune nozioni generali riguardanti il linguaggio che utilizzeremo per poter affrontare questioni inerenti *segni*, *espressioni*, *deduzioni*, *proposizioni*, *proprietà*, *verità*, *dimostrazioni* e *dimostrabilità*. Uno dei compiti fondamentali della logica è infatti quello di riuscire a dare una definizione che sia la più appropriata possibile di questi concetti.

Diciamo che un segno è una forma (sonora, gestuale, visiva, ...) riconoscibile senza ambiguità ogni volta che compare e riproducibile a piacimento, mentre una espressione è un gruppo di segni combinati secondo delle regole sintattiche, che chiameremo regole di *formazione* delle espressioni corrette.

Le espressioni infatti, affinché si possano leggere, devono essere sequenze di segni scritte in modo sensato, anche se non necessariamente dotate di significato. Vogliamo quindi astrarre dal significato delle espressioni stesse preoccupandoci della loro chiarezza e non ambiguità. Definiremo quindi un sistema di calcolo, il cosiddetto λ -calcolo tipato, che ci permetterà di decidere se due espressioni sono uguali o meno.

B.2 Il linguaggio e i suoi elementi

Introduciamo ora gli elementi che costituiscono il linguaggio simbolico della logica. Essi sono: *costanti*, *variabili*, *applicazioni* e *astrazioni*.

Le costanti sono i segni che vengono utilizzati per indicare oggetti precisi che appartengono all'insieme di cui vogliamo parlare.

Le variabili sono segni che vengono utilizzati per indicare elementi che possono essere sostituiti con altri.

Applicazioni ed astrazioni meritano invece una trattazione di maggior dettaglio.

B.2.1 Applicazione

L'applicazione è l'operazione che consiste nell'applicare una espressione, vista come una funzione, ad un'altra espressione, vista come un argomento per tale funzione. Affinchè l'espressione risultante continui ad avere senso è necessario fare attenzione al *tipo* dell'elemento a cui la funzione viene applicata. Il λ -calcolo tipato fornisce appunto le regole per l'utilizzo corretto dell'applicazione. Introduciamo quindi una nuova nozione per rappresentare il tipo degli argomenti a cui può essere applicata un'espressione: l'*arietà*. Essa viene definita in modo induttivo tramite le seguenti regole:

Definizione B.2.1 (Arietà)

$$\begin{array}{ll} (base) & 0 \text{ arietà} \\ (freccia) & \frac{\alpha \text{ arietà} \quad \beta \text{ arietà}}{\alpha \rightarrow \beta \text{ arietà}} \end{array}$$

Ad una espressione viene assegnata l'arietà $\alpha \rightarrow \beta$ se essa si applica ad oggetti di arietà α per ottenere un risultato di arietà β mentre una espressione di arietà 0 non si può applicare a nessun argomento. Ad esempio l'espressione $\sin$, ottenuta dalla costante $\sin$, può essere applicata ad un argomento x di arietà 0 ottenendo l'espressione $\sin(x)$, a sua volta di arietà 0.

Formalizziamo quanto detto introducendo le regole seguenti che permettono di formare espressioni in modo corretto.

Definizione B.2.2 (λ -espressioni (prima parte))

$$\begin{array}{ll}
 \text{(costanti)} & \frac{C : \alpha \text{ const}}{C : \alpha \text{ exp}} \\
 \text{(variabili)} & \frac{x : \alpha \text{ var}}{x : \alpha \text{ exp}} \\
 \text{(applicazione)} & \frac{b : \alpha \rightarrow \beta \text{ exp} \quad a : \alpha \text{ exp}}{b(a) : \beta \text{ exp}}
 \end{array}$$

Una costante o una variabile di arietà α è quindi una espressione della stessa arietà, mentre l'applicazione dà luogo ad una espressione solo se l'arietà dell'applicando e quella dell'argomento vanno d'accordo.

Come già accennato se una espressione di arietà $\alpha \rightarrow \beta$ viene applicata ad un argomento di arietà α otteniamo una espressione di arietà β . Ora possiamo comprendere perchè, in questo contesto, espressioni quali $\sin(\sin)$ e $x(a)$, dove x è una variabile di arietà 0, sono prive di senso: $\sin$, avendo arietà $0 \rightarrow 0$, può essere applicata solamente ad un argomento di arietà 0 e non ad un argomento di arietà $0 \rightarrow 0$, mentre x , avendo arietà 0, non si può applicare a nessun tipo di argomento.

Esempio 1

L'espressione $+$ si applica a due argomenti di arietà 0. Quindi se la applichiamo ad una espressione di arietà 0, per esempio $3 : 0 \text{ exp}$ otteniamo una espressione di arietà $0 \rightarrow 0$, cioè $+(3) : 0 \rightarrow 0 \text{ exp}$ che può nuovamente essere applicata ad una espressione di arietà 0, per esempio $x : 0 \text{ exp}$, ottenendo una espressione di arietà 0, cioè $+(3)(x) : 0 \text{ exp}$.

Esempio 2

Per il segno di integrale si può procedere in modo analogo. Consideriamo la seguente scrittura:

$$\int_a^b f(x) dx \equiv \int (a)(b)(f)$$

L'espressione $\int$ si applica a tre argomenti: a , b di arietà 0 e f di arietà $0 \rightarrow 0$. Otteniamo così l'arietà della costante $\int$:

$$\int : 0 \rightarrow (0 \rightarrow ((0 \rightarrow 0) \rightarrow 0)) \text{ const}$$

Se applichiamo $\int$ ad una espressione di arietà 0, ad esempio $4 : 0 \text{ exp}$, otteniamo l'espressione $\int(4)$ di arietà $0 \rightarrow ((0 \rightarrow 0) \rightarrow 0)$. Applicando a sua volta questa espressione ad una nuova espressione di arietà 0, per esempio $6 : 0 \text{ exp}$ otteniamo una espressione di arietà $(0 \rightarrow 0) \rightarrow 0$, vale a dire $\int(4)(6)$. Quest'ultima espressione può essere applicata solamente ad una espressione di arietà $0 \rightarrow 0$, ad esempio $\sin$, ottenendo finalmente una espressione di arietà 0, cioè $\int(4)(6)(\sin)$. La prova formale è la seguente:

$$\begin{array}{c}
 \frac{\frac{\int : 0 \rightarrow (0 \rightarrow ((0 \rightarrow 0) \rightarrow 0)) \text{ const}}{\int : 0 \rightarrow (0 \rightarrow ((0 \rightarrow 0) \rightarrow 0)) \text{ exp}} \quad \frac{4 : 0 \text{ const}}{4 : 0 \text{ exp}}}{\int(4) : 0 \rightarrow ((0 \rightarrow 0) \rightarrow 0)} \quad \frac{6 : 0 \text{ const}}{6 : 0 \text{ exp}} \\
 \frac{\int(4)(6) : (0 \rightarrow 0) \rightarrow 0 \text{ exp}}{\int(4)(6)(\sin)} \quad \frac{\sin : 0 \rightarrow 0 \text{ const}}{\sin : 0 \rightarrow 0 \text{ exp}}
 \end{array}$$

B.2.2 Astrazione

L'astrazione è l'operazione che permette di "astrarre" una variabile da una espressione trasformandola in una espressione che può essere applicata. Introduciamo la regola che la caratterizza.

Definizione B.2.3 (λ -espressioni (segue))

$$(astrazione) \quad \frac{b : \beta \text{ exp} \quad x : \alpha \text{ var}}{(\lambda x.b) : \alpha \rightarrow \beta \text{ exp}}$$

cioè, se b è una espressione di arietà β ed x è una variabile di arietà α , astraendo x da b otteniamo una espressione di arietà $\alpha \rightarrow \beta$ che *non contiene più x come variabile*.

Esempio 3

$$\frac{+(3)(x) : 0 \text{ exp} \quad x : \alpha \text{ var}}{(\lambda x. +(3)(x)) : 0 \rightarrow 0 \text{ exp}}$$

Esempio 4

$$\frac{\frac{3 : 0 \text{ const}}{3 : 0 \text{ exp}} \quad x : 0 \text{ exp}}{(\lambda x.3) : 0 \rightarrow 0 \text{ exp}}$$

cioè $(\lambda x.3)$ è una espressione che rappresenta la funzione costante che dà risultato 3 se applicata ad un qualsiasi argomento di arietà 0.

Esempio 5

$$\frac{\frac{\sin(+(x)(y)) : 0 \text{ exp} \quad x : 0 \text{ var}}{(\lambda x. \sin(+(x)(y))) : 0 \rightarrow 0 \text{ exp}} \quad y : 0 \text{ var}}{(\lambda y. (\lambda x. \sin(+(x)(y)))) : 0 \rightarrow (0 \rightarrow 0)}$$

dove le variabili x e y non appaiono in $(\lambda y. (\lambda x. \sin(+(x)(y))))$.

In generale, per indicare che una variabile x non appare in una espressione b scriveremo $x \notin \text{FV}(b)$ e diremo che $\text{FV}(b)$ indica l'insieme delle variabili libere nell'espressione b .

$\text{FV}(b)$ si può definire formalmente per induzione sulla complessità della struttura della espressione b nel modo che segue:

$$\begin{array}{lll} \text{FV}(C) & \equiv & \emptyset \quad \text{dove } C \text{ è una costante} \\ \text{FV}(x) & \equiv & \{x\} \quad \text{dove } x \text{ è una variabile} \\ \text{FV}(b(a)) & \equiv & \text{FV}(b) \cup \text{FV}(a) \\ \text{FV}((\lambda x.b)) & \equiv & \text{FV}(b) \setminus \{x\} \end{array}$$

Quindi la variabile y risulta libera in $(\lambda x. \sin(x))(y)$ mentre la variabile x non appare libera perchè tutte le sue occorrenze sono astratte. Si noti tuttavia che la variabile x appare libera in $(\lambda x. \sin(x))(x)$ visto che c'è comunque una sua occorrenza libera, anche se la sua occorrenza in $\sin(x)$ è astratta.

Nel seguito cercheremo di rendere esplicito il fatto che il significato di una espressione deve dipendere solo dalle variabili che vi appaiono libere, come accade ad esempio nel caso di un integrale il cui valore non dipende dal nome della variabile di integrazione.

In particolare, faremo spesso uso del fatto che data una espressione come $(\lambda x.b)$ la sola variabile che sicuramente *non* vi compare (libera) è x .

Abbiamo così definito gli elementi del nostro linguaggio ed introdotto le regole che costituiscono la nostra *teoria delle espressioni*. Vediamo ora quale uso possiamo farne.

B.2.3 La sostituzione

Il prossimo passo consiste nel definire un processo di calcolo per semplificare le scritture quali $(\lambda x. \sin(x))(\pi)$ oppure $(\lambda y. (\lambda x. +(x)(y)))(3)$ che rappresentano le consuete funzioni, rispettivamente ad una e a due variabili, applicate ad un argomento.

A tale scopo introduciamo l'operazione di *sostituzione*. Essa consiste nel sostituire, nell'espressione b , l'espressione a al posto delle occorrenze libere della variabile x (notazione $b[x := a]$).

Mentre la spiegazione intuitiva del suo significato è chiara, la presenza della operazione di astrazione rende la sua definizione formale non completamente banale.

Definizione B.2.4 (Sostituzione) Sia $b : \beta \text{ exp}$, $a : \alpha \text{ exp}$ e $x : \alpha \text{ var}$ allora la sostituzione della variabile x con l'espressione a nella espressione b è definita per induzione sulla complessità della struttura dell'espressione b come segue:

$$\begin{aligned}
 (\text{cost.}) \quad C[x := a] &\equiv C \\
 (\text{var.}) \quad y[x := a] &\equiv \begin{cases} a & \text{se } x \equiv y \\ y & \text{altrimenti} \end{cases} \\
 (\text{appl.}) \quad b(c)[x := a] &\equiv b[x := a](c[x := a]) \\
 (\text{astr.}) \quad (\lambda y. b)[x := a] &\equiv \begin{cases} (\lambda y. b) & \text{se } x \equiv y \\ (\lambda y. b[x := a]) & \text{se } x \not\equiv y \text{ e } y \notin \text{FV}(a) \\ (\lambda t. b[y := t])[x := a] & \text{se } x \not\equiv y, y \in \text{FV}(a) \\ & \text{e } t \text{ è una nuova va-} \\ & \text{riabile di arietà } \alpha \end{cases}
 \end{aligned}$$

La maggior parte dei passi considerati nella definizione dovrebbero essere di comprensione immediata. Qualche dubbio potrebbe forse sorgere nella comprensione del primo punto relativo alla sostituzione della variabile x con l'espressione a nel caso di una astrazione $(\lambda x. b)$, ma è chiaro che la soluzione proposta nell'algoritmo non fa altro che rendere esplicito il fatto che la variabile x non appare in $(\lambda x. b)$.

Tuttavia, per comprendere meglio la ragione per cui nell'ultimo punto della definizione dell'algoritmo di sostituzione si usa una *nuova* variabile, cioè una variabile che non compare nè in b nè in a , nel caso in cui si intenda sostituire in $(\lambda y. b)$ la variabile x con l'espressione a e $x \not\equiv y$ e $y \in \text{FV}(a)$, è forse conveniente fare un esempio. Consideriamo la seguente sostituzione

$$(\lambda x. + (x)(y))[y := x]$$

Potrebbe allora venire spontaneo pensare che il risultato della sostituzione sia:

$$(\lambda x. + (x)(x))$$

Bisogna tuttavia notare che, per la proprietà dell'astrazione, x non compare più come variabile nel risultato della sostituzione, mentre appariva chiaramente nella espressione di partenza. Il significato della scrittura assume quindi una certa importanza: $(\lambda x. + (x)(y))$ rappresenta una espressione che applicata ad un generico elemento c , della stessa arietà della variabile x , dà come risultato $+(c)(y)$. D'altra parte $(\lambda x. + (x)(x))$ rappresenta una funzione completamente diversa in quanto quando essa viene applicata ad un generico elemento c , che naturalmente deve ancora avere la stessa arietà di x , fornisce come risultato $+(c)(c)$. Pertanto la scrittura $(\lambda x. + (x)(y))[y := x]$ *deve* rappresentare una espressione che applicata ad un argomento c restituisce $c + x$, in quanto il significato inteso della sostituzione è quello di ottenere una nuova funzione in cui y viene sostituito con x ma che per il resto si comporta come la precedente.

Nella definizione dell'algoritmo di sostituzione abbiamo formalizzato quanto detto introducendo una nuova variabile non presente nelle espressioni coinvolte in modo da ottenere:

$$\begin{aligned}
 (\lambda x. + (x)(y))[y := x] &\equiv (\lambda t. + (x)(y)[x := t])[y := x] \\
 &\equiv (\lambda t. + (x)(y)[x := t][y := x]) \\
 &\equiv (\lambda t. + (t)(y)[y := x]) \\
 &\equiv (\lambda t. + (t)(x))
 \end{aligned}$$

Vogliamo ora dimostrare che l'algoritmo di sostituzione termina per qualsiasi scelta corretta delle espressioni su cui la sostituzione viene effettuata. Visto che la definizione dell'algoritmo è ricorsiva la cosa non è completamente scontata. Iniziamo osservando il seguente risultato.

Lemma B.2.5 *Le espressioni b e $b[x := t]$, dove t è una nuova variabile della stessa arietà di x , hanno la stessa complessità strutturale.*

Dimostrazione. La dimostrazione procede per induzione sulla complessità della struttura dell'espressione b .

- (Costante) $b \equiv C$. Bisogna dimostrare che $C[x := t]$ ha la stessa complessità strutturale di C . Ma questo è ovvio visto che $C[x := t] \equiv C$

- (Variabile) Guardando la definizione dell'algoritmo di sostituzione è chiaro che dobbiamo considerare due casi:
 - $b \equiv y \neq x$. Questo caso è completamente simile al caso della costante.
 - $b \equiv x$. Allora $b[x := t] \equiv t$, ma t per ipotesi è una variabile come x e quindi le espressioni b e $b[x := t]$ hanno la stessa complessità.
- (Applicazione) $b \equiv d(a)$. Supponiamo, per ipotesi induttiva, che il risultato valga per le espressioni $d[x := t]$ e $a[x := t]$, in quanto d e a sono espressioni meno complesse di b . Allora $d[x := t]$ ha la stessa complessità di d e $a[x := t]$ ha la stessa complessità di a . Dobbiamo ora dimostrare che $d(a)[x := t]$ ha la stessa complessità di $d(a)$. Ma $d(a)[x := t] \equiv d[x := t](a[x := t])$ e quindi il risultato è ovvio.
- (astrazione) $b \equiv ((z) d)$. Dobbiamo dimostrare che $((z) d)$ e $((z) d)[x := t]$ hanno la stessa complessità strutturale. Per farlo distinguiamo i seguenti tre casi in conformità alla definizione dell'algoritmo di sostituzione:
 - $z \equiv x$. Per definizione dell'algoritmo di sostituzione, otteniamo $((z) d)[x := t] \equiv ((z) d)$ e quindi $((z) d)[x := t]$ e $((z) d)$ hanno ovviamente la stessa complessità.
 - $z \not\equiv x$ e $z \not\equiv t$. Per ipotesi induttiva possiamo supporre che l'enunciato valga per l'espressione d , che è meno complessa dell'espressione $((z) d)$. Quindi d e $d[x := t]$ hanno la stessa complessità. Ora, utilizzando l'algoritmo di sostituzione otteniamo $((z) d)[x := t] \equiv ((z) d[x := t])$, pertanto $((z) d)[x := t]$ e $((z) d)$ hanno la stessa complessità.
 - $z \not\equiv x$ e $z \equiv t$. Per ipotesi induttiva possiamo supporre che l'enunciato valga per l'espressione d in quanto essa è meno complessa dell'espressione $((z) d)$, cioè possiamo supporre che d abbia la stessa complessità di $d[t := w]$, per ogni variabile w della corretta arietà. Allora, applicando nuovamente l'ipotesi induttiva all'espressione $d[t := w]$, in quanto complessa tanto quanto l'espressione d e quindi meno complessa dell'espressione $((z) d)$, otteniamo che $d[t := w][x := t]$ ha la stessa complessità di $d[t := w]$ e quindi di d . Ora, per l'algoritmo di sostituzione, $(\lambda t.d)[x := t] \equiv ((w) d[t := w][x := t])$, dove w è una nuova variabile. Quindi le espressioni $(\lambda t.d)[x := t]$ e $(\lambda t.d)$ hanno la stessa complessità strutturale.

Possiamo ora dimostrare il seguente teorema.

Teorema B.2.6 (Terminazione dell'algoritmo di sostituzione) *Supponiamo che $b : \beta \text{ exp}$, $y : \alpha \text{ var}$ e $e : \alpha \text{ exp}$. Allora l'esecuzione dell'algoritmo di sostituzione per sostituire la variabile y con l'espressione e nell'espressione b termina.*

Dimostrazione. Per dimostrare la terminazione dell'algoritmo di sostituzione si procede per induzione sulla complessità strutturale dell'espressione b .

- (Costante) $b \equiv C$. In questo caso $b[y := e] \equiv b$ e, non essendoci chiamate ricorsive, l'algoritmo termina.
- (Variabile) $b \equiv x$. Anche in questo caso non ci sono chiamate ricorsive dell'algoritmo di sostituzione, quindi la terminazione è ovvia.
- (Applicazione) $b \equiv d(a)$. Supponiamo, per ipotesi induttiva, che l'enunciato valga per le due espressioni d e a , in quanto esse sono meno complesse dell'espressione b . Quindi l'algoritmo di sostituzione termina quando viene applicato a $d(a)$.
- (Astrazione) $b \equiv (\lambda x.b')$. Vediamo allora i tre casi previsti nell'algoritmo di sostituzione.
 - $x \equiv y$. Per definizione in questo caso l'applicazione dell'algoritmo di sostituzione richiede che $(\lambda x.b')[x := e] \equiv (\lambda x.b')$. Non essendoci chiamate ricorsive la terminazione è ovvia.

- $x \neq y$ e $x \notin \text{FV}(e)$. Supponiamo, per ipotesi induttiva, che l'esecuzione dell'algoritmo di sostituzione termini quando esso viene applicato a b' in quanto b' è una espressione meno complessa dell'espressione b . Dobbiamo allora dimostrare che l'algoritmo termina quando viene applicato a $(\lambda x.b')$. Ma in questo caso $(\lambda x.b') \equiv (\lambda x.b'[y := e])$ e quindi la terminazione è ovvia.
- $x \neq y$ e $x \in \text{FV}(e)$. Per ipotesi induttiva l'algoritmo di sostituzione termina quando applicato alla espressione b' in quanto essa è meno complessa della espressione $(\lambda x.b')$. Inoltre, per il lemma precedente, b' e $b'[x := t]$ hanno la stessa complessità e possiamo quindi ri-applicare l'ipotesi induttiva ottenendo che l'esecuzione dell'algoritmo di sostituzione termina anche nel caso in cui esso sia applicato per sostituire l'espressione e al posto della variabile y in $b'[x := t]$. Ma, per definizione dell'algoritmo di sostituzione per questo caso, sappiamo che $(\lambda x.b')[y := e] \equiv (\lambda t.b'[x := t][y := e])$ e quindi ne segue la terminazione della sua esecuzione quando sia applicato a $(\lambda x.b')$.

B.2.4 Il processo di calcolo

Come abbiamo più volte ricordato le espressioni che abbiamo definito intendono rappresentare, in modo estremamente astratto, delle funzioni e i loro argomenti. Inoltre l'arietà è il più astratto dei modi per imporre dei vincoli sull'applicazione che ne rendano sensato l'uso. In questa sezione volgiamo vedere come si possano fare dei calcoli con queste funzioni. Introduciamo quindi alcune definizioni.

Definizione B.2.7 (β -riduzione) *Sia $b : \beta \text{ exp}$, $a : \alpha \text{ exp}$ e $x : \alpha \text{ var}$. Allora diciamo β -riduzione la regola seguente:*

$$(\lambda x.b)(a) \Rightarrow b[x := a]$$

La β -riduzione permette di *semplificare* una espressione della forma $(\lambda x.b)(a)$ in una espressione di forma più semplice in quanto viene eliminata una istanza di astrazione-applicazione.

Definizione B.2.8 (η -riduzione) *Sia $d(x) : \delta \text{ exp}$ una espressione tale che la variabile $x : \alpha \text{ var}$ non appaia in d . Allora diciamo η -riduzione la regola seguente:*

$$(\lambda x.d(x)) \Rightarrow d$$

La η -riduzione è un processo di calcolo che permette di ottenere l'espressione $d : \alpha \rightarrow \delta$ astraendo la variabile x dall'espressione $d(x)$. Tale regola è stata introdotto allo scopo di garantirsi la validità nel calcolo delle espressioni della condizione di *estensionalità*: se due espressioni f e g sono tali che, per ogni argomento x della arietà corretta, $f(x) = g(x)$ allora $f = g$. L'estensionalità risulta particolarmente intuitiva nel caso si consideri una nozione insiemistica delle funzioni visto che sostiene che due funzioni con lo stesso grafico coincidono, mentre la sua validità è molto più problematica quando si considerino le funzioni come programmi: in generale è infatti un problema quello di capire quando due programmi sono uguali. Ad ogni modo, per ora non siamo in grado di fornire alcuna dimostrazione formale del fatto che la η -riduzione implica l'estensionalità visto che non abbiamo ancora introdotto una teoria dell'uguaglianza tra le espressioni; lo faremo comunque in uno degli esercizi alla fine della dispensa.

Esempio 6

Ecco un esempio di applicazione della β -riduzione:

$$(\lambda x. + (x)(3))(5) \Rightarrow +(x)(3)[x := 5] \equiv +(5)(3)$$

Esempio 7

Mentre il seguente è un esempio di applicazione della η -riduzione:

$$(\lambda x. \sin(x)) \Rightarrow \sin$$

B.3 Teoria dell'uguaglianza

Nei paragrafi precedenti si è visto che è sorta la necessità di sapere quando due espressioni sono uguali. Dovrebbe essere abbastanza chiaro che l'identità delle scritture è una nozione di uguaglianza troppo stretta visto che è desiderabile che due espressioni che si riducono una all'altra siano uguali, allo stesso modo in cui desideriamo in matematica che il *risultato* dell'applicazione della funzione f all'argomento a sia uguale a $f(a)$ (ad esempio sicuramente vogliamo che $3 + 2$ sia considerato uguale a 5 anche se si tratta chiaramente di due scritture completamente diverse). Questo è il motivo per cui introduciamo le seguenti regole che caratterizzano la minima relazione di equivalenza tra espressioni indotta dalla $\beta\eta$ -riduzioni.

Definizione B.3.1 (λ -uguaglianza) *La relazione di uguaglianza tra espressioni è la minima relazione di equivalenza chiusa per le seguenti regole.*

(costante)	$\frac{C : \alpha \text{ const}}{C = C : \alpha}$
(variabile)	$\frac{x : \alpha \text{ var}}{x = x : \alpha}$
(applicazione)	$\frac{b_1 = b_2 : \alpha \rightarrow \beta \quad a_1 = a_2 : \alpha}{b_1(a_1) = b_2(a_2) : \beta}$
(ξ -uguaglianza)	$\frac{b = d : \beta \quad x : \alpha \text{ var}}{(\lambda x.b) = (\lambda x.d) : \alpha \rightarrow \beta}$
(α -uguaglianza)	$\frac{b : \beta \text{ exp} \quad x : \alpha \text{ var} \quad y : \alpha \text{ var} \quad y \notin \text{FV}(b)}{(\lambda x.b) = (\lambda y.b[x := y]) : \alpha \rightarrow \beta}$
(β -uguaglianza)	$\frac{b : \beta \text{ exp} \quad a : \alpha \text{ exp}}{(\lambda x.b)(a) = b[x := a] : \beta}$
(η -uguaglianza)	$\frac{b : \alpha \rightarrow \beta \text{ exp} \quad x : \alpha \text{ var} \quad x \notin \text{FV}(b)}{(\lambda x.b(x)) = b : \alpha \rightarrow \beta}$

Qualche commento può essere utile allo scopo di capire meglio il significato delle regole di uguaglianza. Le prime quattro regole hanno il solo scopo di garantire che l'uguaglianza rispetta le operazioni che abbiamo deciso di usare per costruire le espressioni, vale a dire costanti, variabili, applicazioni e astrazioni.

La regola di α -uguaglianza garantisce che due espressioni che differiscono solo per i nomi delle variabili astratte sono uguali. Abbiamo già detto che desideriamo la validità di questa condizione per evitare che le variabili astratte usate nel definire una funzione nel influenzino il significato. La condizione a lato della regola serve per assicurarsi di non legare con l'astrazione variabili diverse da quella voluta. Si consideri ad esempio la seguente espressione $((x) + (x)(y))$ che denota la funzione che, per ogni argomento c , fornisce $+(c)(y)$. E' chiaro che essa è diversa dalla funzione $((y) + (y)(y))$, che otterremo dalla sostituzione $+(x)(y)[y := x]$. Infatti, quest'ultima è la funzione che per ogni argomento c fornisce $+(c)(c)$.

Le ultime due regole si limitano a garantire che una espressione è comunque uguale ad una espressione a cui è possibile ridurla utilizzando una β -riduzione o una η -riduzione.

B.3.1 Decidibilità dell'uguaglianza tra λ -espressioni

Nel seguito sarà per noi essenziale poter decidere quando due espressioni sono uguali. Fortunatamente l'arietà che accompagna le espressioni, ci permette di dimostrare in modo effettivo quando due espressioni sono uguali. Diremo quindi che l'uguaglianza tra espressioni è effettivamente decidibile.

La dimostrazione di questo teorema non è completamente banale e non la inseriremo in queste note. Tuttavia è possibile avere un'idea di come funziona tale dimostrazione dopo aver introdotto la nozione di *forma normale* di una espressione.

Definizione B.3.2 (Forma normale) *Una espressione è in forma normale se non può essere applicata una β -riduzione a nessuna delle sue sotto-espressioni.*

Data una qualunque espressione e , il seguente algoritmo permette di ottenere una espressione in forma normale, che indicheremo con $\text{nf}(e)$, equivalente ad e .

$$\text{nf}(e) \equiv \begin{cases} (\lambda x. \text{nf}(e(x))) & \begin{array}{l} \text{se } e : \alpha \rightarrow \beta \text{ exp e } x \\ \text{è una nuova variabile di arietà } \alpha \end{array} \\ x(\text{nf}(a_1)) \dots (\text{nf}(a_n)) & \begin{array}{l} \text{se } e \equiv x(a_1) \dots (a_n) : 0 \text{ exp} \\ x : \alpha_1 \rightarrow (\dots (\alpha_n \rightarrow 0) \dots) \text{ var} \\ a_1 : \alpha_1 \text{ exp}, \dots, a_n : \alpha_n \text{ exp} \end{array} \\ C(\text{nf}(a_1)) \dots (\text{nf}(a_n)) & \begin{array}{l} \text{se } e \equiv C(a_1) \dots (a_n) : 0 \text{ exp} \\ C : \alpha_1 \rightarrow (\dots (\alpha_n \rightarrow 0) \dots) \text{ var} \\ a_1 : \alpha_1 \text{ exp}, \dots, a_n : \alpha_n \text{ exp} \end{array} \\ \text{nf}(d[y := a_0](a_1) \dots (a_n)) & \begin{array}{l} \text{se } e \equiv (\lambda y. d)(a_0)(a_1) \dots (a_n) : 0 \text{ exp} \\ d : \alpha_1 \rightarrow (\dots (\alpha_n \rightarrow 0) \dots) \text{ exp} \\ a_0 : \alpha_0 \text{ exp}, \dots, a_n : \alpha_n \text{ exp} \end{array} \end{cases}$$

Per provare la parziale correttezza dell'algoritmo di conversione in forma normale si dimostra il seguente teorema (si noti che per dimostrare che l'algoritmo è totalmente corretto bisognerebbe anche dimostrare che la sua esecuzione termina, ma tale dimostrazione eccede i limiti che ci siamo posti nella stesura di questi appunti).

Teorema B.3.3 (Correttezza parziale dell'algoritmo di forma normale) *Sia e un'espressione di arietà α . Allora, se l'algoritmo di forma normale termina su e , $\text{nf}(e)$ è una espressione in forma normale uguale ad e .*

Dimostrazione. Prima di tutto si noti che stiamo lavorando nell'ipotesi che l'esecuzione dell'algoritmo termini e quindi esiste un numero m di volte in cui esso deve essere ancora applicato all'espressione e per ottenere la sua forma normale $\text{nf}(e)$. È chiaro che è allora possibile dimostrare una proprietà per induzione su m .

Prima di tutto non è difficile vedere che se l'algoritmo termina quando viene applicato ad e allora il suo risultato, vale a dire $\text{nf}(e)$, è una espressione in forma normale. Consideriamo infatti i vari casi previsti dall'algoritmo:

1. Dobbiamo dimostrare che $(\lambda x. \text{nf}(e(x)))$ è una espressione in forma normale, ma questo risultato è immediato, visto che se c è una espressione in forma normale allora anche $(\lambda x. c)$ è in forma normale.
2. Dobbiamo dimostrare che $x(\text{nf}(a_1)) \dots (\text{nf}(a_n))$ è una espressione in forma normale, ma anche questo risultato è immediato, visto che se $c_1, \dots, c_n$ sono espressioni in forma normale allora anche $x(c_1) \dots (c_n)$ è in forma normale.
3. Dobbiamo dimostrare che $C(\text{nf}(a_1)) \dots (\text{nf}(a_n))$ è una espressione in forma normale, ma per ottenere questo risultato possiamo procedere esattamente come nel punto precedente.
4. Ovvio.

Notare che una dimostrazione formale richiede di procedere per induzione su m .

Anche la dimostrazione che $e = \text{nf}(e)$ si sviluppa per induzione sul numero m . Analizziamo anche in questo caso i vari possibili passi dell'algoritmo di forma normale:

1. Supponiamo che e sia una espressione di arietà $\alpha \rightarrow \beta$. Quindi secondo l'algoritmo di normalizzazione $\text{nf}(e) \equiv (\lambda x. \text{nf}(e(x)))$. Supponiamo allora per ipotesi induttiva che il risultato valga per l'espressione $e(x)$, vale a dire $e(x) = \text{nf}(e(x)) : \beta$. Quindi, usando una istanza di ξ -uguaglianza, otteniamo $(\lambda x. e(x)) = (\lambda x. \text{nf}(e(x))) : \alpha \rightarrow \beta$. Ma ora per la η -uguaglianza, otteniamo $e = (\lambda x. e(x)) : \alpha \rightarrow \beta$ (si ricordi che $x : \alpha$ è una variabile nuova e non appare quindi in e). Perciò

$$\text{nf}(e) \equiv (\lambda x. \text{nf}(e(x))) = (\lambda x. e(x)) = e : \alpha \rightarrow \beta$$

2. Supponiamo che $e \equiv x(a_1) \dots (a_n) : 0 \text{ exp}$. Quindi, secondo l'algoritmo di normalizzazione, $\text{nf}(e) \equiv x(\text{nf}(a_1)) \dots (\text{nf}(a_n))$. Supponiamo, per ipotesi induttiva, che il risultato valga per le espressioni $a_1, \dots, a_n$. Allora, per ogni $i = 1, \dots, n$, $a_i = \text{nf}(a_i)$. Utilizzando ora l'uguaglianza tra variabili e n volte l'uguaglianza tra applicazioni otteniamo

$$x(a_1) \dots (a_n) = x(\text{nf}(a_1)) \dots (\text{nf}(a_n)) : 0$$

quindi

$$\text{nf}(e) \equiv x(\text{nf}(a_1)) \dots (\text{nf}(a_n)) = x(a_1) \dots (a_n) : 0$$

3. La dimostrazione del caso in cui $e \equiv C(a_1) \dots (a_n) : 0 \text{ exp}$ è completamente analoga a quella del caso precedente.
4. Supponiamo che $e \equiv (\lambda y.d)(a_0)(a_1) \dots (a_n) : 0 \text{ exp}$. Quindi, per l'algoritmo di normalizzazione, $\text{nf}(e) \equiv \text{nf}(d[y := a_0](a_1) \dots (a_n))$. Supponiamo ora, per ipotesi induttiva, che il risultato valga per l'espressione

$$d[y := a_0](a_1) \dots (a_n)$$

Allora

$$d[y := a_0](a_1) \dots (a_n) = \text{nf}(d[y := a_0](a_1) \dots (a_n)) : 0$$

Ma, per la β -uguaglianza,

$$(\lambda y.d)(a_0)(a_1) \dots (a_n) = d[y := a_0](a_1) \dots (a_n) : 0$$

e quindi

$$\begin{aligned} \text{nf}((\lambda y.d)(a_0)(a_1) \dots (a_n)) &\equiv \text{nf}(d[y := a_0](a_1) \dots (a_n)) \\ &= d[y := a_0](a_1) \dots (a_n) \\ &= (\lambda y.d)(a_0)(a_1) \dots (a_n) \end{aligned}$$

Il teorema di forma normale stabilisce una forma *canonica* per le espressioni, cioè la forma più semplice.

La dimostrazione della decidibilità, che procede usando il teorema di forma normale, prova che due espressioni sono uguali se e solo se le loro forme normali, che esistono sempre e sono univocamente determinate, differiscono solo per il nome delle variabili astratte che, come conseguenza della α -uguaglianza possono essere sempre scelte liberamente.

Esempio 8

Consideriamo l'espressione $\sin : 0 \rightarrow 0 \text{ exp}$. Allora

$$\text{nf}(\sin) \equiv (\lambda x.\text{nf}(\sin(x))) \equiv (\lambda x.\sin(\text{nf}(x))) \equiv (\lambda x.\sin(x)) : 0 \rightarrow 0$$

Esempio 9

Consideriamo l'espressione $+(x) : 0 \rightarrow 0 \text{ exp}$. Allora

$$\begin{aligned} \text{nf}(+(x)) &\equiv (\lambda y.\text{nf}(+(x)(y))) \\ &\equiv (\lambda y. +(\text{nf}(x))(\text{nf}(y))) \\ &\equiv (\lambda y. +(x)(y)) : 0 \rightarrow 0 \end{aligned}$$

Esempio 10

Consideriamo l'espressione $\int(3)(5)(\sin) : 0 \text{ exp}$. Allora

$$\text{nf}\left(\int(3)(5)(\sin)\right) = \int(\text{nf}(3))(\text{nf}(5))(\text{nf}(\sin)) = \int(3)(5)(\lambda x.\sin(x))$$

Esempio 11

Consideriamo l'espressione $(\lambda x. + (x)(3))((\lambda x. \sin(x))(\pi)) : 0 \text{ exp}$. Allora

$$\begin{aligned}
 \text{nf}((\lambda x. + (x)(3))((\lambda x. \sin(x))(\pi))) &\equiv \text{nf}(+(x)(3)[x := (\lambda x. \sin(x))(\pi)]) \\
 &\equiv \text{nf}((\lambda x. \sin(x))(\pi))(3) \\
 &\equiv \text{nf}(+(\sin(x)[x := \pi])(3)) \\
 &\equiv \text{nf}(+(\sin(\pi))(3)) \\
 &\equiv +(\text{nf}(\sin(\pi))(\text{nf}(3))) \\
 &\equiv +(\sin(\pi))(3)
 \end{aligned}$$

Naturalmente non possiamo spingere oltre il processo di semplificazione perchè la nostra teoria è molto astratta e non sa nulla di aritmetica elementare e trigonometria.

B.4 Esercizi risolti

Esercizio 1

Siano $c : \beta \text{ exp}$ e $d : \beta \text{ exp}$ due espressioni tali che $c = d : \beta$, x una variabile di arietà α e a e b due espressioni di arietà α tali che $a = b : \alpha$. Dimostrare che

$$c[x := a] = d[x := b]$$

Soluzione. Dal fatto che $c = d : \beta$ otteniamo, utilizzando un'istanza di ξ -uguaglianza, che $(\lambda x. c) = (\lambda x. d) : \alpha \rightarrow \beta$. Quindi, essendo per ipotesi $a = b : \alpha$, per la uguaglianza tra applicazioni otteniamo $(\lambda x. c)(a) = (\lambda x. d)(b) : \beta$. Ma ora, per β -uguaglianza, $(\lambda x. c)(a) = c[x := a] : \beta$ e $(\lambda x. d)(b) = d[x := b] : \beta$ e quindi otteniamo la tesi per *transitività*.

Esercizio 2

Siano c e b due espressioni e x e y due variabili. Dimostrare che se $x \notin \text{FV}(b)$ e $x \neq y$ allora

$$c[x := a][y := b] = c[y := b][x := a[y := b]]$$

Soluzione. Per la β -uguaglianza abbiamo che $c[x := a] = (\lambda x. c)(a)$. Se effettuiamo su ambo i membri dell'uguaglianza la stessa sostituzione $[y := b]$ otteniamo (vedi esercizio precedente) $c[x := a][y := b] = (\lambda x. c)(a)[y := b]$. Ora, per la definizione di sostituzione nelle ipotesi dell'esercizio, $(\lambda x. c)(a)[y := b] \equiv (\lambda x. c[y := b])(a[y := b])$. Ma, per la β -uguaglianza, $(\lambda x. c[y := b])(a[y := b]) = c[y := b][x := a[y := b]]$ e quindi otteniamo la tesi per *transitività*.

Esercizio 3

Dimostrare che la relazione di uguaglianza tra espressioni è estensionale, vale a dire che, supponendo $f : \alpha \rightarrow \beta \text{ exp}$ e $g : \alpha \rightarrow \beta \text{ exp}$ siano due espressioni e x sia una variabile di arietà α che non compare ne in f ne in g , se $f(x) = g(x) : \beta$ allora $f = g : \alpha \rightarrow \beta$.

Soluzione. Visto che per ipotesi $f(x) = g(x) : \beta$, per la ξ -uguaglianza, otteniamo che $(\lambda x. f(x)) = (\lambda x. g(x)) : \alpha \rightarrow \beta$. Ma ora, per la η -uguaglianza che possiamo applicare visto che $x \notin \text{FV}(f)$ e $x \notin \text{FV}(g)$, otteniamo $(\lambda x. f(x)) = f : \alpha \rightarrow \beta$ e $(\lambda x. g(x)) = g : \alpha \rightarrow \beta$ e quindi la tesi per *transitività*.

Appendice C

Dittatori e Ultrafiltri

Quando si è in presenza di opinioni divergenti e bisogna comunque prendere decisioni comuni è inevitabile ricorrere a qualche forma di aggregazione delle scelte individuali a favore di una scelta collettiva. Le forme di aggregazione possono essere molte, ma quelle più usuali, almeno dalle nostre parti, passano attraverso delle votazioni regolate da una qualche *legge elettorale*.

Naturalmente, l'accettabilità sociale del risultato del voto dipende fortemente dalla legge elettorale scelta e si pone quindi la questione di quali condizioni dobbiamo richiedere su tale legge elettorale affinché la si possa considerare equa.

In questo lavoro presentiamo il famoso teorema di Arrow sulle scelte democratiche (si veda [A63, D82]), che illustra bene le difficoltà insite nella definizione di una legge elettorale equa, e ne diamo una dimostrazione basata su idee e risultati propri della logica matematica.

Abbiamo scelto di collocare in appendice i dettagli tecnici e le dimostrazioni più complesse. Il lettore interessato solamente alle idee fondamentali può trascurare tali parti, che possono però essere utili a chi voglia approfondire e capire meglio lo sviluppo dell'argomento.

C.1 Il problema

Supponiamo che una comunità sia nella condizione di dover scegliere tra diverse alternative, come succede ad esempio quando bisogna scegliere quale coalizione politica mandare al governo. Supponiamo inoltre che quando si vota ogni membro della comunità metta le varie alternative in ordine di preferenza. Il problema è allora scegliere una regola che ci consenta di stabilire, una volta note le opinioni individuali, quale debba essere la scelta della collettività.

Se le alternative sono solo due, e il numero dei votanti è dispari, la soluzione è banale: basta votare a maggioranza! Se invece le alternative sono più di due il problema è più complesso di quanto appaia a prima vista.

Vediamo infatti un esempio che ci farà capire quali problemi possono sorgere se non siamo abbastanza accorti nella scelta della legge elettorale. Supponiamo che l'insieme dei votanti sia costituito da tre individui, che chiameremo i_1 , i_2 e i_3 , e che le scelte possibili siano pure tre: A , B e C . Ogni votante esprime la sua opinione definendo un ordine totale stretto tra le tre alternative. Naturalmente vogliamo che anche l'ordine di preferenza collettivo sia un ordine totale stretto sull'insieme delle alternative.

Potrebbe allora sembrare ragionevole stabilire il seguente criterio: *per ogni coppia di alternative X e Y , se la maggioranza dei votanti preferisce X a Y allora anche nell'ordine di preferenza collettivo X deve essere preferito a Y .*

Vediamo però come questo approccio ingenuo crea dei problemi. Supponiamo infatti che la votazione sia la seguente:

$$\begin{aligned} i_1 : & A > B > C \\ i_2 : & B > C > A \\ i_3 : & C > A > B \end{aligned}$$

Notiamo che i_1 e i_3 sostengono che $A > B$. Quindi siccome due votanti su tre sostengono che $A > B$, anche nell'ordine di preferenza collettivo A dovrebbe essere preferito a B . Analogamente,

due votanti su tre sostengono che $B > C$ e altrettanti sostengono che $C > A$. Quindi l'ordine di preferenza collettivo dovrebbe soddisfare contemporaneamente:

$$A > B, B > C, C > A.$$

Ma allora $B > A$ dovrebbe seguire dato che valgono $B > C$ e $C > A$. Quindi nell'ordine di preferenza collettivo dovrebbe valere sia $A > B$ che $B > A$ e questo è assurdo.

C.2 Il Teorema di Arrow

Proviamo ora a formalizzare il problema in modo rigoroso. Siano I un insieme non vuoto di votanti e $\mathcal{A} = \{A_1, A_2, \dots, A_N\}$ un insieme di N alternative da ordinare.

Supponiamo inoltre di scrivere $\phi_{k,l}$ per dire che A_k è preferibile a A_l . Quindi un votante mette in ordine di preferenza le alternative esprimendo la sua opinione sulle formule dell'insieme $\Phi = \{\phi_{k,l} \mid 1 \leq k < l \leq N\}$. Partendo dalle formule dell'insieme Φ possiamo costruire formule più complesse utilizzando ad esempio i connettivi di negazione $\neg$ e congiunzione $\wedge$. Così la formula $\phi_{1,2} \wedge \phi_{2,3}$ dice che A_1 è preferibile a A_2 e che A_2 è preferibile a A_3 .

In realtà se conosciamo l'opinione di un votante sulle formule di Φ , conosciamo anche cosa pensa di ogni altra formula del linguaggio. Ad esempio se ritiene vere le formule $\phi_{1,2}$ e $\phi_{2,3}$, se ritiene cioè vero il fatto che A_1 sia preferibile ad A_2 e che A_2 sia preferibile ad A_3 , allora riterrà vera anche $\phi_{1,2} \wedge \phi_{2,3}$. Possiamo formalizzare l'opinione di un votante come una funzione f che va dall'insieme delle formule all'algebra di Boole delle tabelline di verità sull'insieme $\{\top, \perp\}$ con operazioni **and** e **not** definite nel modo ben noto (si veda ad esempio [BM77]). Questa funzione sarà definita così:

$$f(\psi) = \begin{cases} \top & \text{se il votante pensa che } \psi \text{ sia vera} \\ \perp & \text{se il votante pensa che } \psi \text{ sia falsa} \end{cases}$$

e dovrà soddisfare alle seguenti proprietà

$$\begin{aligned} f(\psi_1 \wedge \psi_2) &= f(\psi_1) \text{ and } f(\psi_2) \\ f(\neg\psi) &= \text{not } f(\psi) \end{aligned}$$

D'ora in poi chiameremo *valutazione* ogni funzione dall'insieme delle formule a $\{\top, \perp\}$ che goda di queste proprietà. Quindi l'opinione di un votante verrà formalizzata con una valutazione.

Quando si vota ogni votante esprime le sue scelte; l'*esito del voto* è quindi una funzione ω che associa ad ogni individuo $i \in I$ una valutazione $\omega(i)$ che formalizza le scelte del votante i . Data una qualsiasi formula ψ indicheremo con $I_{\omega,\psi} = \{i \in I \mid \omega(i)(\psi) = \top\}$ la coalizione dei votanti a favore di ψ nella votazione di esito ω .

Siamo finalmente in grado di trattare il problema che ci interessa: quello della legge elettorale. Una legge elettorale permette di determinare quale sia la decisione collettiva una volta che i votanti abbiano espresso le loro opinioni individuali. Quel che ci serve è quindi una *funzione di aggregazione* che ad ogni esito del voto associa una valutazione. Per chiarire le idee vediamo un esempio. Scegliamo un particolare votante $i_0 \in I$. Possiamo allora definire una funzione di aggregazione F nel modo seguente:

$$F(\omega) = \omega(i_0) \text{ per ogni esito del voto } \omega$$

Questa è una funzione di aggregazione, perché ad ogni esito del voto associa una valutazione. Tuttavia questa legge elettorale non è per nulla democratica. Infatti qualunque sia l'esito del voto la scelta collettiva si conformerà sempre all'opinione del votante i_0 ; insomma, il votante i_0 è un dittatore!

Il problema che ci poniamo a questo punto è quello di trovare una funzione di aggregazione che sia equa e che non crei paradossi come quello visto nel paragrafo precedente. Vediamo quindi alcune proprietà che sembra naturale richiedere ad una funzione di aggregazione F equa.

1. (*Unanimità*) Per ogni esito del voto ω e per ogni formula ψ se $I_{\omega,\psi} = I$ allora $F(\omega)(\psi) = \top$,
2. (*Monotonia*) Siano ω e ω' due esiti del voto e ψ e ψ' due formule. Allora se $I_{\omega,\psi} \subseteq I_{\omega',\psi'}$ e $F(\omega)(\psi) = \top$ deve essere anche $F(\omega')(\psi') = \top$.

La prima proprietà dice che se tutti i votanti sono d'accordo su qualcosa allora la decisione collettiva deve adeguarsi. Questa richiesta sembra decisamente ragionevole per una funzione di aggregazione equa.

Per capire meglio la seconda proprietà immaginiamo di considerare due votazioni. Supponiamo ora che la coalizione che sostiene ψ nella prima votazione sia X e che l'insieme Y dei votanti che sostengono ψ' nella seconda votazione sia un sovrainsieme di X . Allora sembra ragionevole supporre che, se nel primo caso la decisione collettiva si adegua, per quanto riguarda ψ , a ciò che pensano i votanti dell'insieme X , nel secondo caso essa debba adeguarsi, per quanto riguarda ψ' , a ciò che pensano i votanti della più ampia coalizione Y .¹

A questo punto non possiamo certo sperare di aver trovato la soluzione al nostro problema della definizione di una legge elettorale equa visto che risultano godere di *unanimità* e *monotonia* sia la funzione di aggregazione che utilizza la legge della maggioranza sia quella che si conforma all'opinione di un dittatore. Tuttavia il teorema di Arrow ci riserva una sorpresa ancora peggiore.

Teorema C.2.1 (Teorema di Arrow) *Supponiamo che l'insieme I dei votanti sia finito e che l'insieme delle alternative da ordinare abbia almeno tre elementi. Sia inoltre F una funzione di aggregazione che goda di unanimità e monotonia. Allora esiste $i_0 \in I$ tale che, per ogni esito del voto ω , $F(\omega) = \omega(i_0)$.*

Insomma, se dobbiamo ordinare almeno tre elementi allora ogni funzione di aggregazione che goda di *unanimità* e *monotonia* è dittatoriale².

C.3 Ultrafiltri e aggregazioni di preferenze

La dimostrazione che daremo del teorema di Arrow è basata sullo stretto legame che intercorre tra funzioni di aggregazione e ultrafiltri (si veda [LL94, EK09]).

Definizione C.3.1 (Ultrafiltro) *Sia I un insieme non vuoto. Un ultrafiltro U su I è una collezione di sottoinsiemi di I che gode delle seguenti proprietà, per ogni sottoinsieme X e Y di I :*

(coerenza)	$\emptyset \notin U$
(chiusura per inter.)	se $X \in U$ e $Y \in U$ allora $X \cap Y \in U$
(decidibilità)	X sta in U o il suo complemento X^C sta in U

È facile verificare che in conseguenza di *coerenza* e *decidibilità* ogni ultrafiltro U soddisfa anche

$$(totalità) \quad I \in U$$

ed è appena più laborioso vedere che, in conseguenza di *coerenza*, *chiusura per intersezione* e *decidibilità*, vale anche

$$(chiusura in su) \quad \text{se } X \in U \text{ e } X \subseteq Y \text{ allora } Y \in U$$

Tanto per fissare le idee vediamo un semplice esempio di ultrafiltro. Scelto un qualsiasi elemento i_0 di I , consideriamo la collezione U dei sottoinsiemi X di I tali che $i_0 \in X$; è allora immediato verificare che U è un ultrafiltro su I . Nel seguito ci riferiremo a questo ultrafiltro come all'*ultrafiltro principale generato da i_0* .

¹La versione originale del teorema di Arrow considera al posto della nostra condizione di *monotonia* una condizione di *indipendenza* che richiede che, per ogni coppia di esiti del voto ω e ω' e ogni coppia di formule ψ e ψ' , se $I_{\omega, \psi} = I_{\omega', \psi'}$ e $F(\omega)(\psi) = \top$ allora anche $F(\omega')(\psi') = \top$ e che sembra quindi essere una condizione più debole di quella di *monotonia*. In realtà, per quanto riguarda il teorema di Arrow, *monotonia* ed *indipendenza* sono due condizioni equivalenti e noi abbiamo preferito usare quest'ultima in quanto più vicina alla analoga condizione di *chiusura in su* sugli ultrafiltri che considereremo nel prossimo paragrafo.

²In realtà il risultato si può notevolmente generalizzare; in particolare non dipende per nulla dal fatto che siamo interessati ad ordinare gli elementi di un insieme ma vale per qualsiasi insieme di proposizioni che contenga delle proposizioni ψ_1 e ψ_2 tali che $\psi_1 \wedge \psi_2$, $\neg\psi_1 \wedge \psi_2$, $\psi_1 \wedge \neg\psi_2$ e $\neg\psi_1 \wedge \neg\psi_2$ possono essere ciascuna vera nell'opinione di qualche elettore come avviene nel nostro caso per le proposizioni $\phi_{1,2}$ e $\phi_{2,3}$ (si veda [MV11]).

Il nostro interesse per gli ultrafiltri viene dal fatto che possiamo facilmente usarli per definire funzioni di aggregazione. Infatti se U è un ultrafiltro su I possiamo definire una funzione F_U sugli esiti del voto ω nel modo seguente:

$$F_U(\omega)(\psi) = \begin{cases} \top & \text{se } I_{\omega,\psi} \in U \\ \perp & \text{altrimenti} \end{cases}$$

Questa funzione è costruita stabilendo che gli elementi dell'ultrafiltro U sono le coalizioni di votanti in grado di influenzare l'esito del voto. Utilizzando le proprietà di un ultrafiltro si può dimostrare che la funzione così definita è una funzione di aggregazione che gode di *unanimità* e *monotonia*. Infatti la dimostrazione che si tratta di una funzione di aggregazione richiede solo di fare un po' di conti (ma questi si possono anche evitare se ci si accorge che il risultato è un banale corollario del teorema di Loš, si veda [BM77]³) mentre *unanimità* e *monotonia* seguono rispettivamente dalle condizioni di *totalità* e *chiusura in su* di un ultrafiltro.

Il seguente teorema dice che in realtà vale anche il viceversa e quindi parlare di funzioni di aggregazione che soddisfano *unanimità* e *monotonia* o di ultrafiltri su I è di fatto la stessa cosa.

Teorema C.3.2 (Teorema di corrispondenza) *Sia I una popolazione che deve decidere come ordinare un insieme con almeno tre elementi e sia F una funzione di aggregazione che goda di unanimità e monotonia. Allora esiste un ultrafiltro U_F su I tale che $F = F_{U_F}$.*

Il concetto di ultrafiltro U sull'insieme dei votanti sembra quindi adatto a formalizzare l'idea di coalizione sufficiente di votanti almeno nel caso di una legge elettorale che soddisfi *unanimità* e *monotonia*.

Se pensiamo che *unanimità* e *monotonia* caratterizzino i sistemi di votazione equi e che in un sistema di votazioni equo siano gli insiemi grandi di elettori che determinano il risultato delle elezioni allora un ultrafiltro su I formalizza l'idea di sottoinsieme *grande* di I (da qui il nome *filtro* per tale collezione di sottoinsiemi di I visto che come un filtro trattiene le parti più grandi).

C.4 Teorema di Arrow e ultrafiltri

In questa ultima sezione vediamo come usare le idee della sezione precedente per dimostrare il teorema di Arrow. Il seguente risultato sarà cruciale.

Teorema C.4.1 (Teorema dell'ultrafiltro principale) *Se I è un insieme finito allora ogni ultrafiltro su I è principale.*

Veniamo ora alla dimostrazione del teorema di Arrow. Supponiamo che I sia finito e che l'insieme delle alternative da ordinare abbia almeno tre elementi distinti. Sia inoltre F una funzione di aggregazione che goda di *unanimità* e *monotonia*. Segue allora dal teorema C.3.2 che esiste un ultrafiltro U_F su I tale che $F = F_{U_F}$. Quindi, per ogni esito del voto ω e per ogni formula ψ , $F(\omega)(\psi) = \top$ se e solo se $I_{\omega,\psi} \in U_F$. Poichè I è finito, segue dal teorema C.4.1 che U_F è un ultrafiltro principale. Sia i_0 il suo generatore. Quindi un sottoinsieme X di I appartiene all'ultrafiltro U_F se e solo se $i_0 \in X$. Ma allora, mettendo insieme questo ragionamento con il precedente, otteniamo la seguente serie di equivalenze:

$$F(\omega)(\psi) = \top \iff I_{\omega,\psi} \in U_F \iff i_0 \in I_{\omega,\psi} \iff \omega(i_0)(\psi) = \top$$

Perciò ψ viene approvata nella decisione collettiva se e solo se il votante i_0 è d'accordo con ψ e quindi il dittatore del teorema di Arrow è proprio il generatore dell'ultrafiltro U_F .

³È interessante osservare che un risultato come il teorema di Arrow che da un lato ha richiesto al suo autore la stesura di un intero libro e dall'altro gli ha fruttato il premio Nobel per l'economia fosse *in nuce* già contenuto in un classico teorema della Logica Matematica!

C.5 Dimostrazione del teorema di corrispondenza

Sia I una popolazione che deve decidere come ordinare un insieme che abbia almeno tre elementi distinti e sia F una funzione di aggregazione che goda delle proprietà di *unanimità* e *monotonia*.

Sia inoltre U_F la collezione di sottoinsiemi di I della forma $I_{\omega,\psi}$ per qualche esito del voto ω e qualche formula ψ tali che $F(\omega)(\psi) = \top$.

I seguenti due lemmi dimostrano che U_F è un ultrafiltro su I tale che $F = F_{U_F}$.

Teorema C.5.1 *Siano I , F ed U_F definiti come sopra. Allora, per ogni esito del voto ω e ogni formula ψ ,*

$$F(\omega)(\psi) = \top \iff I_{\omega,\psi} \in U_F$$

Dimostrazione. Se $F(\omega)(\psi) = \top$ allora $I_{\omega,\psi} \in U_F$ vale per definizione. Viceversa, supponiamo che $I_{\omega,\psi} \in U_F$. Questo significa che esistono un esito del voto ω' e una formula ψ' tali che $F(\omega')(\psi') = \top$ e $I_{\omega,\psi} = I_{\omega',\psi'}$. Ma allora $F(\omega)(\psi) = \top$ segue per *monotonia*.

Teorema C.5.2 *Siano I , F ed U_F definiti come sopra. Allora U_F è un ultrafiltro.*

Dimostrazione. Proviamo prima di tutto che $\emptyset \notin U_F$. Supponiamo infatti che $\emptyset \in U_F$. Allora esistono un esito del voto ω e una formula ψ tali che $\emptyset = I_{\omega,\psi}$ e $F(\omega)(\psi) = \top$. Ma allora $I = I_{\omega,\neg\psi}$ e quindi $F(\omega)(\neg\psi) = \top$ segue per *unanimità*. Perciò $F(\omega)(\psi) = \perp$, visto che $F(\omega)$ è una valutazione, contro l'ipotesi che $F(\omega)(\psi) = \top \neq \perp$.

Vediamo ora che, per ogni sottoinsieme X di I , $X \in U_F$ oppure $X^C \in U_F$. A tal fine consideriamo un esito del voto ω tale che $\omega(i)(\phi_{1,2}) = \top$ se e solo se $i \in X$ (visto che sia $\phi_{1,2}$ che $\neg\phi_{1,2}$ possono essere asserite da qualche elettore, in conseguenza del teorema dell'ultrafiltro e dello stretto legame esistente tra ultrafiltri su una algebra di Boole e valutazioni (vedi [BM77]), possiamo essere sicuri che ci sono sia valutazioni che rendono vero $\phi_{1,2}$ sia valutazioni che rendono vera la sua negazione e quindi, per ogni sottoinsieme X , esiste un esito del voto come ω). Allora $X = I_{\omega,\phi_{1,2}}$ e quindi $X^C = I_{\omega,\neg\phi_{1,2}}$. Ma, dato che $F(\omega)$ è una valutazione, si ha che $F(\omega)(\phi_{1,2}) = \top$ oppure $F(\omega)(\neg\phi_{1,2}) = \top$ e quindi nel primo caso $X \in U_F$ mentre nel secondo $X^C \in U_F$.

Infine proviamo che U_F è chiuso per intersezioni. Supponiamo quindi che X, Y siano due elementi di U_F . Visto che abbiamo almeno tre elementi da ordinare e quindi ognuna tra le proposizioni $\phi_{1,2} \wedge \phi_{2,3}$, $\neg\phi_{1,2} \wedge \phi_{2,3}$, $\phi_{1,2} \wedge \neg\phi_{2,3}$ e $\neg\phi_{1,2} \wedge \neg\phi_{2,3}$ può essere asserita da qualche elettore, in virtù del teorema dell'ultrafiltro per le algebra di Boole e del già ricordato legame tra ultrafiltri e valutazioni, esiste un esito del voto ω tale che:

$$\begin{aligned} \omega(i)(\phi_{1,2} \wedge \phi_{2,3}) &= \top && \text{se e solo se } i \in X \cap Y \\ \omega(i)(\neg\phi_{1,2} \wedge \phi_{2,3}) &= \top && \text{se e solo se } i \in X^C \cap Y \\ \omega(i)(\phi_{1,2} \wedge \neg\phi_{2,3}) &= \top && \text{se e solo se } i \in X \cap Y^C \\ \omega(i)(\neg\phi_{1,2} \wedge \neg\phi_{2,3}) &= \top && \text{se e solo se } i \in X^C \cap Y^C \end{aligned}$$

Abbiamo allora che $X = I_{\omega,\phi_{1,2}}$, $Y = I_{\omega,\phi_{2,3}}$ e $X \cap Y = I_{\omega,\phi_{1,2} \wedge \phi_{2,3}}$. Poichè $X, Y \in U_F$ dal lemma precedente segue che $F(\omega)(\phi_{1,2}) = \top$ e $F(\omega)(\phi_{2,3}) = \top$. Ma allora, visto che $F(\omega)$ è una valutazione, $F(\omega)(\phi_{1,2} \wedge \phi_{2,3}) = \top$ e quindi $X \cap Y \in U_F$ per definizione.

C.6 Dimostrazione del teorema dell'ultrafiltro principale

Dobbiamo dimostrare che ogni ultrafiltro U su un insieme finito I è principale. Per dimostrare questo teorema faremo uso del seguente lemma.

Teorema C.6.1 *Sia I un insieme e U un ultrafiltro su I . Allora se $X \cup Y \in U$ allora $X \in U$ o $Y \in U$.*

Dimostrazione. Supponiamo che $X \notin U$ e $Y \notin U$. Allora $X^C, Y^C \in U$. Quindi $X^C \cap Y^C = (X \cup Y)^C \in U$. Ma allora $X \cup Y$ non può appartenere a U . Infatti se così fosse avremmo che $\emptyset = (X \cup Y) \cap (X \cup Y)^C \in U$, contro la proprietà di *coerenza* di U .

Vediamo ora che esiste un elemento $i^* \in I$ tale che $\{i^*\} \in U$. Poiché stiamo lavorando nell'ipotesi che I sia finito, possiamo scriverlo così: $I = \{i_1, \dots, i_n\}$. Dato che U è un ultrafiltro, $I = \{i_1\} \cup \{i_2, \dots, i_n\} \in U$ per *totalità*. Ma allora, per il lemma C.6.1, o $\{i_1\} \in U$ oppure $\{i_2, \dots, i_n\} \in U$. Se $\{i_1\} \in U$ abbiamo finito. Altrimenti possiamo iterare il ragionamento sull'insieme $\{i_2\} \cup \{i_3, \dots, i_n\} \in U$ e, in un numero di passi minore o uguale del numero di elementi di I , troviamo $i^* \in I$ tale che $\{i^*\} \in U$.

Sia ora X un qualsiasi sottoinsieme di I . Allora se $i^* \in X$, cioè $\{i^*\} \subseteq X$, $X \in U$ segue dalla *chiusura in su* di U . Viceversa se $i^* \notin X$, abbiamo che $i^* \in X^C$; quindi è X^C che sta in U e di conseguenza $X \notin U$ perché altrimenti $\emptyset = X \cap X^C \in U$ contro la *coerenza* di U . Ma allora abbiamo provato che $X \in U$ se e solo se $i^* \in X$, cioè U è l'ultrafiltro principale generato da i^* .

Appendice D

Correzione degli esercizi

D.1 La verità matematica

Esercizio 2.1.1.

Una possibile soluzione è costituita dalla coppia $a = \sqrt{2}$ e $b = \log_2 9$. Infatti è ben noto che $\sqrt{2}$ è un numero irrazionale (chi non ne ha mai visto una dimostrazione può facilmente costruirne una). Inoltre se supponiamo che $\log_2 9$ sia razionale, i.e. che esistano p e q tali che $\frac{p}{q} = \log_2 9$, allora $q \log_2 9 = p$ e quindi $2^p = 2^{q \log_2 9} = (2^{\log_2 9})^q = 9^q$ ma questo è assurdo. Infine è immediato verificare, utilizzando la definizione della funzione logaritmo e di banali proprietà delle potenze, che $a^b = 3$, i.e. che si tratta di un numero razionale.

Esercizio 2.2.1

Abbiamo già visto come risolvere l'esercizio nel caso $n = 1$. Let us suppose now that we know that the property holds for all the infinite sequences of k -tuples and that we want to show that it holds also for all the infinite sequence of $k + 1$ -tuples. Then we have to prove that for any infinite sequence $(n_{0_1}, n_{0_2}, \dots, n_{0_{k+1}}), (n_{1_1}, n_{1_2}, \dots, n_{1_{k+1}}) \dots$ of $k + 1$ -tuples of natural numbers there are two indexes i and j such that $i < j$, and, for any $1 \leq h \leq k + 1$, $n_{i_h} \leq n_{j_h}$. Also in this case we can reason by absurd as above and suppose that l is an infinite sequence of $k + 1$ -tuples that does not enjoy the desired property. Then, let us consider the $k + 1$ -tuple $l_0 \equiv (n_{0_1}, n_{0_2}, \dots, n_{0_{k+1}})$, namely, the first element in the sequence l ; then all the next elements must have at least one component which is smaller than the corresponding component in l_0 and hence for some $1 \leq h \leq k + 1$ there is an infinite numbers of $k + 1$ -tuples in the sequence l such that their h -th component is smaller than the element n_{0_h} . So, again by the pigeon hole principle, for some fixed $m < n_{0_h}$ there is an infinite numbers of $k + 1$ -tuples of the following shape $(n_1, n_2, \dots, n_{h-1}, m, n_{h+1}, \dots, n_{k+1})$; hence, by inductive hypothesis, the infinite sequence built with all the k -tuples $(n_1, n_2, \dots, n_{h-1}, n_{h+1}, \dots, n_{k+1})$, obtained from the infinite $k + 1$ -tuples $(n_1, n_2, \dots, n_{h-1}, m, n_{h+1}, \dots, n_{k+1})$, would satisfy the required property and hence l itself would enjoy it which is again contrary to the assumption.

D.2 L'approccio classico e quello intuizionista

Esercizio 3.1.3

Se indichiamo con π_1 il ticket per $3 + 2 = 5$ e con π_2 quello per $2 \times 2 = 4$, ottenuti semplicemente calcolando $3 + 2$ e 2×2 per verificare che i risultati siano rispettivamente 5 e 4, un ticket per la proposizione $(3 + 2 = 5) \ \& \ (2 \times 2 = 4)$ è $\langle \pi_1, \pi_2 \rangle$.

Bibliografia

- [A63] K.J. Arrow, *Social choice and individual values*, John Wiley and Sons, Inc., New York, London, Sydney 1963
- [BM77] J. L. Bell and M. Machover, *A course in mathematical logic*, Amsterdam : North-Holland Pub. Co. ; New York : Elsevier Pub. Co., 1977
- [Bir67] Birkhoff, G., *Lattice Theory*, 3rd ed. Vol. 25 of American Mathematical Society Colloquium Publications. American Mathematical Society, 1967.
- [Bolz1817] Bolzano, B., *Rein analytischer Beweis*.
- [BBJ07] Boolos G., Burgess J., Jeffrey R., *Computability and Logic*, Cambridge University Press
- [D82] G. dall'Aglia, *Decisioni di gruppo: il paradosso di Arrow*, Archimede, vol. XXXIV n. 1-2 (1982), pp 3-14
- [Gal79] Gale, D., *The game of Hex and the Brouwer Fixed-Point Theorem*, The American Mathematical Monthly, vol. 86 (10), 1979, pp. 818–827.
- [Gar59] Gardner, M., *The Scientific American Book of Mathematical Puzzels and Diversions*, Simon and Schuster, New York, 1959, pp. 73–83.
- [EK09] C. Klamler and D. Eckert, *A Simple Ultrafilter Proof For an Impossibility Theorem in Judgment Aggregation*, Economics Bulletin, 2009
- [LL94] L. Lawers and L. Van Liedekerke, *Ultraproducts and Aggregation*, Journal of Mathematical Economics, 1994
- [LeoTof07] Leonesi, S., Toffalori, C., *Matematica, miracoli e paradossi*, Bruno Mondadori editore
- [Mat1970] Y. Matiyasevich, *Enumerable sets are Diophantine*, Doklady Akademii Nauk SSSR, 191, pp. 279-282, 1970, Traduzione inglese in Soviet Mathematics. Doklady, vol. 11, no. 2, 1970
- [MV11] A. Montino and S. Valentini, *Generalizing Arrow's theorem: a logical point of view*, in corso di pubblicazione
- [Rasiowa-Sikorski 63] H. Rasiowa, R. Sikorski, *The mathematics of metamathematics*, Polish Scientific Publishers, Warsaw, 1963
- [Smullyan81] R. Smullyan, *Quale è il titolo di questo libro?* Zanichelli, Bologna, 1981
- [Smullyan85] R. Smullyan, *Donna o tigre?* Zanichelli, Bologna, 1985
- [Smullyan92] R. Smullyan, *Gödel incompleteness Theorems* Oxford University Press, Oxford, 1992
- [Takeuti 75] G. Takeuti *Proof Theory*, North Holland, 1975