

Lorenzo Peña

INTRODUCCIÓN A LAS LÓGICAS NO CLÁSICAS

México: UNAM, 1993
ISBN 968-36-3451-6.

A mi apreciado colega,
el Profesor **Emilio Terzaga**,
en testimonio de estima y amistad

Sumario

Prólogo	001
Introducción	005
1.— ¿Qué tan justificable es el monopolio docente de la lógica clásica?	
2.— Noción de verdad lógica	
3.— El carácter logográfico de las notaciones lógicas	
4.— Prerrequisitos para la lectura del presente opúsculo	
Capítulo 1. Notaciones, paréntesis, puntos	013
Esquemas y funtores	
Paréntesis y puntos	
Capítulo 2. Noción de dominio de valores de verdad	017
Noción de valor de verdad	
Valores designados, valores antidesignados	
Relaciones de orden en un dominio de valores de verdad	
Lógicas escalares y lógicas tensoriales	
Capítulo 3. Noción de tautología	020
Lógica bivalente	
Lista de algunas tautologías	
¿Cómo calcular las tautologías de una lógica tensorial?	
Capítulo 4. Estudio de varios funtores	025
Funtores de afirmación	
Funtores de negación	
Funtores de conyunción	
Funtores de disyunción	
Funtores condicionales	
Funtores bicondicionales y equivalenciales	
Sobreimplicaciones	
Escolio al Capítulo 4. Aplicaciones a la formalización de enunciados en lenguaje natural	
Capítulo 5. Ventajas de la lógica infinivalente como lógica de lo difuso	033
Matices	
Reales e hiperreales	
El sistema Ar	
Asignaciones de valores de verdad	
El sistema Ap	
Valuaciones de Ap	
Relaciones entre Ap y Ar	
La lógica Abp	
Capítulo 6. Noción de teoría, clasificación sintáctica y semántica de las teorías	042
Reglas de inferencia	
Cierre de un cúmulo con respecto a una relación	
Reglas de formación	
Noción de teoría	

Teorías axiomatizadas	
Teorías inconsistentes	
Teorías superconsistentes y paraconsistentes	
Modelo, interpretación, valuación, validez	
Completez	
Tablas de verdad y modelos	
Un modelo alternativo para A_p	
Capítulo 7. Los principios de no-contradicción y tercio excluso	051
La regla de Cornubia	
El principio de Cornubia	
El principio sintáctico de no-contradicción	
Principio de no-contradicción y paraconsistencia	
Antinomia y contradicción	
El principio semántico de no-contradicción	
El principio sintáctico de tercio excluso	
El principio semántico de tercio excluso	
Relación entre los principios de no-contradicción y de tercio excluso	
Capítulo 8. Sistemas lógicos deductivos: el sistema A_t	060
Lecturas de signos a emplear	
Reglas de formación	
Definiciones	
Regla de inferencia primitiva (única)	
Desarrollo del sistema	
Derivación de reglas de inferencia	
Demostración de teoremas	
Pruebas de otros teoremas	
Capítulo 9. Una extensión de A_t : el sistema infinivalente y tensorial A_j	093
Nuevas lecturas	
Reglas de formación	
Reglas de inferencia	
Esquemas axiomáticos	
Comentarios sobre la base axiomática de A_j	
Diferencia entre la conyunción natural ' $\wedge$ ' y la superconyunción ' $\bullet$ '	
Consideraciones sobre la pluralidad de funtores bicondicionales	
Diversos grados de certeza de los axiomas	
El principio de Heráclito	
Capítulo 10. El cálculo cuantificacional de primer orden	107
Consideraciones preliminares	
Introducción de nueva terminología	
Variación alfabética	
De la pluralidad de cálculos sentenciales a la opción entre diversos cálculos cuantificacionales	
Reglas de formación	

Esquema definicional	
Reglas de inferencia	
Esquemas axiomáticos	
Comentarios	
Capítulo 11. La lógica combinatoria	115
Capítulo 12. Modelos algebraicos	120
Álgebras cuasitransitivas	
Postulados	
Terminología suplementaria	
Productos directos	
El cálculo sentencial $\mathcal{A}_p$ y las álgebras libres asociadas con la clase de las aa.cc.tt.	
Álgebras transitivas y transitivoides	
Consideraciones finales	
Anejo Nº 1. El recurso a una lógica infinivalente en la defensa del realismo científico .	135
Anejo Nº 2. Nota sobre la noción quineana de verdad lógica	139
Bibliografía selecta comentada	141

Prólogo

Este opúsculo es fruto no sólo de una labor investigativa, sino también de la experiencia docente del autor y asimismo de años de debates, de contrastaciones argumentativas, que le han permitido comprender mejor no sólo, en general, dónde están las dificultades principales que se yerguen ante la tarea de adentrarse en el campo de las lógicas no clásicas, sino también, en particular, qué incomprensiones son más frecuentes, qué confusiones desorientan más a menudo y retardan ese adentramiento, esa exploración. Teniendo en cuenta todo ello, el autor ha querido ofrecer un manual claro, sencillo, con explicaciones bien desmenuzadas, aunque hayan de ser prolijas a veces; pero todo ello, dentro de un constreñimiento de brevedad y concisión. Quizá sea eso como la cuadratura del círculo. En cualquier caso, trátase de fijar unos objetivos, conferir grados de prelación a cada uno de ellos, y tender lo más posible a, conciliándolos cuanto se pueda, ofrecer un trabajo que equilibradamente pueda considerarse una razonable aproximación a la realización de los mismos —de cada uno en la medida en que goce de prelación en la jerarquización recién aludida. ¡Ojalá el lector, juzgando benévola-mente pero sin faltar a la verdad, pueda estimar que la presente obra constituye tal aproximación razonable a los dos objetivos arriba mencionados!

Lo que el lector tiene en sus manos es el resultado de reelaboraciones de un trabajo anterior, que en su primera versión se tituló «Apuntes introductorios a la lógica matemática elemental». Aunque el autor se sintió satisfecho por los logros docentes conseguidos mediante la utilización de ese texto poligrafiado, la ulterior práctica más arriba aludida ha aconsejado las modificaciones del mismo que han desembocado en la presente obra. Espero, a la vez, que su utilización como texto universitario de iniciación al estudio de lógicas no clásicas permita un más amplio conocimiento de estas lógicas y, junto con ello, aporte una nueva experiencia que se tendría en cuenta para ulteriores ediciones de este librito.

Según lo acabo de sugerir, espero haber aportado, con el presente trabajo, un instrumento de iniciación al estudio de lógicas no clásicas. De ahí que el título más apropiado de la obra sería el de **Introducción al estudio de** tales lógicas, más que a esas lógicas. Y es que no cabe en un trabajo de la envergadura del aquí presentado —ni probablemente en ningún otro de extensión razonable— brindar una introducción a todas las lógicas no clásicas, ni siquiera a todas aquellas que cuentan con adalides notables. Hoy es tan abundante y variado el panorama de tales lógicas que, si resulta difícil ir siguiendo los desarrollos nuevos en este campo, difícilísimo es, o imposible, compendiar ese panorama sin omitir nada de lo que juzguen valioso unos u otros investigadores destacados. Entre las muchas lógicas no clásicas cuyo estudio se ha omitido en esta Introducción figuran varios sistemas de la familia constructivista —como la lógica intuicionista de Heyting—, las lógicas relevantes y otras similares —como las conexivistas—, las lógicas cuánticas (como las no-distributivas). Todos esos sistemas han sido puestos en pie por consideraciones filosóficas dignas de atención. Merecen ser estudiados y discutidos. Pero la organización del presente manual ha hecho aconsejable no incluir en él un tratamiento de tales sistemas.

Por otro lado suscítase la cuestión de cuándo un sistema lógico es una lógica no clásica y cuándo es meramente un desarrollo, un complemento de la lógica clásica. Algunos llaman lógicas no clásicas también a la lógica modal (la que estudia operadores como ‘posiblemente’ y ‘necesariamente’) y a otras de las llamadas “intensionales”. Desde mi punto de vista, son simples extensiones de la lógica clásica, porque no entran en ningún conflicto con ésta ni en el terreno del cálculo sentencial ni siquiera en el del cálculo cuantificacional. O sea: esas lógicas “intensionales” no añaden ni quitan nada a lo que dice la lógica clásica cuando se reduce su vocabulario al del cálculo cuantificacional (o sea —tratándose de sistemas cuyo cálculo cuantificacional es el mismo de la lógica clásica—: el conjunto o cúmulo de estas partículas:

'y', 'o', 'sólo si', 'si y sólo si', 'no', 'cada', 'algún'). (En este trabajo no se establecerá diferencia ninguna entre los significados de '**cúmulo**' y de '**conjunto**'. Y —salvo en una parte del capítulo XII y último— no hará falta tampoco hacer diferencia entre cúmulos y **clases**.) En cambio —según se verá— algunos de los sistemas lógicos considerados en este librito tienen una divergencia respecto a la lógica clásica en lo tocante al 'no', porque hacen un distingo entre dos **noes**: el mero 'no' —al que asignan características **no** clásicas— y una negación fuerte, el 'no... en absoluto', que conciben como la negación de la lógica clásica (mientras que los lógicos clásicos **leen** así a su negación, sino que la leen como un mero 'no'). Esa es la razón de que, mientras no es cierto que las lógicas modales y otras así sean —por el mero hecho de contener teoremas que no contiene el cálculo cuantificacional clásico— lógicas no clásicas, sí lo sean en cambio sistemas como varios de los que se examinarán en este trabajo, a pesar de que tanto los unos como los otros **contienen** a la lógica (cuantificacional) clásica.

Una de las decisiones difíciles en la elaboración del presente manual ha consistido en la selección de una notación apropiada. Varios son, en efecto, los criterios que hacen recomendable una escritura o notación más que otra. Es doble el uso de cualquier signo gráfico: lectura y escritura. Son ventajas de un signo el hacer económica la lectura o la escritura. Mas hay varios géneros de economía al respecto. Están la condición de fácil discernibilidad visual y la de cuán parsimonioso sea el cúmulo de trazos distintivos que han de recordarse (de algún modo) para formar o reconocer esos grafemas. Este último criterio (economía **paradigmática**, referida a las **oposiciones** en terminología de la lingüística funcional) pugna con el otro (economía **sintagmática**, referida a los **contrastes**); pero también hay un conflicto entre esas dos economías —la sintagmática y la paradigmática— y otra, que es la referida al esfuerzo mayor o menor del trazado de signos. (Esto último no tiene analogía en el lenguaje hablado; es dudosísimo, en efecto, que, en general, haya emisiones bucales más fáciles de hacer —sabemos cuán relativo es eso a una educación articuladora.) Por otro lado, está el problema de cuán arbitrario sea un signo, ya que los menos arbitrarios son más fáciles de recordar. (Y hay grados en eso, así como aspectos, motivaciones **relativas** no más.) Optar por una notación u otra sin tener en cuenta ninguna de tales consideraciones es incurrir en cierta ligereza, aunque de suyo el tema no sea muy importante.

Apartarse de lo normal es imponer una carga a la memoria y capacidad de reconocimiento de los estudiantes. Pero ése es un aspecto de la cuestión, no el único. ¿Qué ha de hacerse cuando la norma estándar no es estrictamente aplicable? No queda más remedio que diseñar nuevas notaciones teniendo en cuenta los criterios precedentes. Y operando siempre en un marco pragmático: el de unas posibilidades de reproducción gráfica que no son ilimitadas.

Así pues, y habida cuenta de todo ello, planteábase como sigue el problema de seleccionar una notación apropiada en nuestro caso. Por un lado, razones de índole pragmática aconsejaban la mayor estandarización posible, esto es: reproducir en la mayor medida las notaciones más usuales. Por otro lado, sin embargo (aparte ya de objeciones que cabe formular a varias de tales notaciones por acudir al empleo de signos menos reconocibles, o menos fáciles de encontrar entre los disponibles en impresoras corrientes, cuando a menudo bien pudieran emplearse, para tales usos, otros signos que sí suelen estar disponibles), surge una dificultad especial tratándose de sistemas no clásicos, donde aparecen pluralidades desconocidas en la lógica clásica. En ésta última p.ej. hay tan sólo una negación, tan sólo un bicondicional, tan sólo una conyunción, etc. En muchos sistemas no clásicos hay varias negaciones, varios condicionales, varios bicondicionales, varias conyunciones y así sucesivamente. Entonces, si en una notación estándar para la lógica clásica se escribe "la" negación como ' $\neg$ ', ¿cómo cabrá operar al pasar a un sistema no clásico con dos o tres negaciones? Puédese acaso explotar el hecho de que esa notación estándar no es única, sino que hay otras también

estándar, como la que representa “la” negación por una tilde, ‘~’. Sin embargo, usar en un sistema no estándar ‘¬’ para una negación y ‘~’ para otra tiene el inconveniente de que con ello lo que es una inocua alternativa notacional en los simbolismos estándar pasa a ser una alternancia significativa en la notación así pergeñada. Por otro lado, si se reserva ‘¬’ (o ‘~’) para representar una de las negaciones de un sistema no clásico —escribiéndose las otras negaciones con símbolos no estándar— ello conlleva también una dificultad, a saber: que es muy posible que o bien las características de la negación clásica no sean las de esa negación no clásica que se opte por representar como ‘¬’, o bien la lectura en lengua natural que se ofrezca para ésta última no sea el mero ‘no’, al paso que los lógicos clásicos suelen leer su [única] negación, su ‘¬’, como ‘no’, sin más.

Semejantes dificultades surgen con el uso de otros símbolos estándar también, aunque a mi juicio menos agudos que con la negación. Por otro lado, hay pasajes del presente opúsculo donde se usan ciertos símbolos como **esquemas** de funtores, o sea haciendo las veces de un signo cualquiera que posea ciertas características enumeradas. A menudo es mejor usar en tales casos un signo no usual, a fin de no prejuizar una lectura determinada y de no exigir de antemano que haya de tener todos los rasgos que sea habitual asignar a un signo para el cual ya se reserva cierta lectura.

Ante esta encrucijada, el autor ha optado por estandarizar la notación empleada, pero dejando a salvo el empleo en algunos lugares cruciales de signos no estándar, donde le ha parecido que el uso de una notación más estándar motivaría confusiones dificultando la comprensión del texto. Sin embargo, es de esperar que en general el lector familiarizado con notaciones habituales rápidamente se acostumbre a esos signos no estándar, que son los menos.

Para poner punto final a este Prólogo, he de manifestar mi gratitud a cuantos me han estimulado a llevar a cabo la redacción de la presente obra con sus sugerencias, objeciones y consejos. En primer lugar, los relatores de la UNAM, quienes me han hecho llegar valiosas indicaciones. El Prof. Marcelo Vásconez Carrasco —de la Universidad de Cuenca (República del Ecuador)— ha aportado correcciones a la primera versión del trabajo arriba mencionado a partir del cual se originó el aquí ofrecido. Igualmente me han ayudado a esa tarea diversas críticas y comentarios a otros trabajos míos, como los de Newton C.A. da Costa (Universidad de São Paulo), Graham Priest, Richard Sylvan, Igor Urbas (todos ellos de Australia), Marcelo Dascal (Universidad de Tel Aviv), Diderik Batens (Universidad de Gante), Katalin Havas (Instituto de Filosofía de la Academia de Ciencias de Hungría y Universidad de Budapest), Emilio López Medina (Universidad de Granada) y Vicente Muñoz Delgado (Universidad Pontificia de Salamanca).

Introducción

§1.— ¿Qué tan justificable es el monopolio docente de la lógica clásica?

Muchos manuales de lógica matemática (casi todos desgraciadamente) incurren en el defecto de exponer un sistema particular de lógica, sin tomarse siquiera la molestia de advertir al lector de que están haciendo tal cosa. El sistema expuesto es casi siempre, bajo alguna presentación particular, la llamada lógica clásica, a la que el autor este libro prefiere llamar ‘lógica bivalente verifuncional’, (LBV), pues es la lógica en la que se admiten dos únicos valores de verdad (verdadero y falso) y que trata a cada signo de su cálculo sentencial (‘y’ y ‘no’ pueden ser los únicos signos que aparezcan como primitivos en tal cálculo) de modo estrictamente verifuncional (tales nociones se explicarán en el texto mismo de este opúsculo).

Quizá tal preferencia por LBV sea tan respetable como la preferencia por otros sistemas de lógica. Lo que no es respetable es esa falta de claridad (o de sinceridad) que estriba en exponer un sistema particular de lógica como si se tratara de “la” lógica. Y no vale como remedio (ni aun como remiendo) el adjuntar al final del libro un apéndice, indicando la existencia de “otros” sistemas de lógica —sobreentendiéndose, o insinuándose, que se trata de meros juegos formales sin aplicación a la realidad o al genuino discurso “científico”.

Tales procedimientos son tan poco recomendables en lógica como lo serían otros similares en cualquier disciplina filosófica, y aun en muchas disciplinas no filosóficas. Se cubriría de ridículo quien, exponiendo la metafísica, p.ej., presentara al comienzo una serie de axiomas y reglas de inferencia y procediera a aplicar éstos, obteniendo así un cierto número de teoremas, condescendiendo luego, en un apéndice o anexo de su obra, a enunciar que lo que ha expuesto es “la” metafísica, pero que la metafísica es la metafísica clásica (quizá la aristotélica —bajo su propia interpretación, ciertamente—, presentada en una particular axiomatización), si bien hay “también” metafísicas no clásicas, puros juegos. Verdad es que a nadie se le ha ocurrido escribir, de conformidad con una pauta semejante, un manual de metafísica.

Claro está, los autores de esos manuales de lógica creen que, en ésta, la confrontación de sistemas alternativos se plantea de un modo radicalmente diverso, puesto que habría verdades irrefutables más allá de toda polémica. Quizá haya alguna verdad así (el principio de identidad, tal vez), pero ciertamente hay muchísimo en la LBV que no está, ni de lejos, más allá de toda controversia; hay en dicha lógica muchas tesis y reglas que han sido objeto de controversia.

Lo más controvertible que hay en la LBV es el principio de Cornubia, según el cual de una antinomia, cualquiera que sea, se sigue cualquier afirmación ($p \vee \neg p$ sólo si q , —donde ‘ q ’ es cualquier enunciado [el empleo de estos símbolos especiales, ‘ $\neg$ ’ y ‘ $\vee$ ’, llamados ‘**esquinas**’, será explicado en el primer acápite del Capítulo I]). Ese principio será dilucidado ampliamente en este estudio. Y veremos que —desde determinados horizontes de intelección, como el del autor de estas páginas— tal principio es aceptable tan sólo cuando el ‘no’ se lee como ‘no es cierto **en absoluto** que’; pero debe ser rechazado, en cambio, tal principio, cuando se aplique a un simple y mero ‘no’ a secas.

De aceptarse el principio de Cornubia **para la negación simple** o natural (o sea: para el mero ‘no’), resultaría que serían absurdas e ilógicas todas las doctrinas que han defendido la contradictorialidad de lo real, como —en el plano filosófico— las de Heráclito, Platón (en el *Parménides* y el *Sofista*), Enesidemo, Plotino, Proclo, el *Corpus Dionysianum*, Mario Victorino, Escoto Eriúgena, Nicolás de Cusa, Giordano Bruno, Jacob Boehme, Robert Fludd, Hegel, el materialismo dialéctico de Engels y Lenin, el energetismo de Stéphane Lupasco y Marc Beigbeder; y, en otros planos, las de San Juan de la Cruz, Sta. Teresa de Jesús, Sor Juana Inés de la Cruz, Petrarca, Du Bellay, Francisco de Quevedo, Melville, G. Manley Hopkins, Gus-

tavo Adolfo Bécquer, Guadalupe Amor y tantos otros. Los escritos de todos esos autores contienen antinomias. Y, si de una antinomia —cualquiera que sea— se desprende cualquier afirmación, entonces de cada una de esas doctrinas se desprende cualquier afirmación, por más absurda que sea. Habría, además, que añadir a la lista de doctrinas que incluyen contradicciones determinadas concepciones de la actual ciencia de la naturaleza —como la que afirma el carácter a la vez corpuscular y ondulatorio de la luz.

Conviene, empero, puntualizar que lo que al autor de este trabajo le parece errado no es el contenido de la LBV, sino el uso que de la misma se hace, como si toda la lógica sentencial pudiera reducirse al mero par de signos conformado por la conyunción y por la negación “clásica” —que no es la negación simple de la lengua natural, sino la **supernegación** que se expresa al decir ‘es enteramente falso que’. Por eso, los sistemas lógicos descubiertos y defendidos por el autor contienen **todas** las tautologías de la lógica clásica, y **todas** sus reglas de inferencia, pero sólo para ese par de funtores; hay, en tales sistemas, otros funtores primitivos e irreducibles a los dos mencionados; y uno de esos **otros** funtores adicionales, que no pueden tratarse **dentro** de la LBV, sino sólo en una lógica más amplia y abarcadora, es la negación simple, el mero ‘no’.

Querer encerrar, de grado o por fuerza, todo el cálculo verifuncional —todo el tratamiento de los funtores sentenciales de la lengua natural— en el diminuto marco conformado por los dos funtores clásicos (la supernegación y la conyunción simple) podrá ser justificable desde el ángulo de un reduccionismo simplificador a ultranza —postura respetable pero que a muchos nos parece muy poco convincente—; mas no lo es desde otras perspectivas.

Y, aun suponiendo que los que así pensamos estemos equivocados, no cabe ignorar los argumentos que se han venido ofreciendo, durante los últimos cuatro o cinco lustros principalmente, a favor de lógicas no-clásicas. Pujantes y vigorosas corrientes de la investigación lógico-matemática han logrado, en los últimos decenios, mostrar, concluyentemente, la viabilidad y plausibilidad de lógicas no-clásicas, de modo que quienes, haciendo tabla rasa de esos recientes descubrimientos de la investigación lógico-matemática, siguen hablando de la LBV como “la” lógica, sin tomarse siquiera la molestia de justificar razonadamente su punto de vista, dan, con ello, pruebas de una voluntaria ceguera.

No trato con esto de incentivar al escepticismo. No es lícito acusar a un autor de ser escéptico o de echar leña al fuego del posible escepticismo de sus lectores por el hecho de que reconozca, franca y honradamente, que el sistema que presenta es sólo **un** sistema y no el único —ni muchísimo menos—; que hay otras concepciones alternativas, las cuales no son aberraciones ni desvaríos, sino pareceres defendibles, que tienen sus partidarios, personas cultas y honorables a quienes sería calumnioso presentar como carentes de agudeza.

Aquello de que se trata es, ni más ni menos, reconocer que el propio sistema se ha originado a partir de un determinado horizonte previo de intelección; que está enmarcado en unas determinadas posiciones últimas de valor e inspirado, pues, en unas presuposiciones determinadas, que ciertamente son discutibles pero que al autor le parecen, empero, portadoras de una evidencia sobresaliente y **más** dignas de consideración y aceptación que cuantos **peros** han sido suscitados contra ellas.

Por todas esas razones, en este opúsculo efectuaremos principalmente un estudio comparativo de diversos sistemas; y sólo después de ese estudio se hará la exposición y desarrollo de determinados sistemas lógicos, elaborados y propuestos por el autor.

En filosofía, como en la vida, cada uno debe escoger su propio camino. Eso mismo ocurre en lógica. Que cada cual escoja su lógica, según sus opciones metafísicas fundamentales. La lógica por la que opta el autor de este opúsculo está —como su filosofía toda— inserta

en la perspectiva de un gradualismo contradictorio —que rechaza toda desnivelación categorial. Pero, naturalmente, a muchos les parecerán preferibles otras vías.

Intentaré aquí dilucidar esa pluralidad de caminos de tal modo que se pueda hacer más lúcida la opción personal de cada uno. Y el autor se considerará satisfecho si sus lectores emprenden efectivamente, cada uno de ellos, su propio camino, si abrazan su propia opción personal, y si a ello ha contribuido en algo la dilucidación comparativa que se presenta en estas páginas.

Propónese, pues, el presente librito brindar a estudiantes universitarios de carreras de letras —y sobre todo de filosofía— un manual introductorio a la lógica matemática elemental que escape a los procedimientos dogmáticos —tan indubitables y caros para los adeptos de un clasicismo intransigente que ni siquiera se conoce a sí mismo, pues ignora que es una opción—, procedimientos que presentan como “la” lógica un sistema particular de lógica, en vez de ir presentando desde el comienzo abanicos de sistemas lógicos alternativos para, a tenor de criterios filosóficos, ir optando por unos u otros. Y, en segundo lugar, allanarles el camino a quienes estén familiarizados con la lógica clásica para explorar ese otro campo, mucho más variado y complejo, de los sistemas de lógica no clásicos, mostrando cómo algunos de ellos poseen capacidades —mayores que las de la lógica clásica, que son modestas— para dar cuenta no sólo de la corrección de una serie de raciocinios usuales, sino también de la incorrección de ciertos pasos inferenciales que a veces se profieren a la ligera, por falta de adecuados distinguos —p.ej. por desconocimiento de la pluralidad de negaciones.

§2.— **Noción de verdad lógica**

Para definir qué es una verdad lógica empezamos por definir lo que es una **ocurrencia** de una expresión en otra, y qué es una **ocurrencia esencial** de una expresión en un enunciado. (Obsérvese que la definición que sigue es meramente propedéutica, provisional, ya que está afectada por cierta inexactitud, la cual vendrá corregida en el Anejo N° 2 que figura al final del opúsculo.)

Ocurrencia. Dícese que una expresión ‘e’ tiene una ocurrencia en otra expresión cualquiera ‘u’ ssi [si, y sólo si,] una instancia de ‘e’ aparece en cada instancia de ‘u’. (Llámase instancia de una expresión a cada resultado particular de escribir o proferir tal expresión. Así, p.ej., la expresión ‘manzana’ tiene tantas instancias como casos hay en que esa palabra se escribe o profiere. Nótese que, en adelante, ‘si, y sólo si,’ vendrá expresado abreviadamente como ‘ssi’.)

Veamos esa noción de **ocurrencia** con algunos ejemplos. La expresión ‘occidental’ tiene una ocurrencia en la expresión ‘Europa occidental’; ésta, a su vez, tiene una ocurrencia en la expresión ‘Europa occidental sufre una terrible crisis económica’. (Esta última expresión es un enunciado; todo enunciado es una expresión, mas no toda expresión es un enunciado). Es obvio que, si una expresión u tiene una ocurrencia en una expresión u¹, y si u¹ tiene una ocurrencia en otra expresión u², entonces u tiene una ocurrencia en u². (Por eso, ‘occidental’ tiene una ocurrencia en ‘Europa occidental sufre una terrible crisis económica’.)

Nótese bien que, para saber si una expresión tiene o no una ocurrencia en otra expresión, hay que analizar bien a ésta última. Si, al analizar a una expresión u, vemos que tiene dos constituyentes, u¹ y u², entonces no podemos decir que haya en u una ocurrencia de **otra** expresión conformada por alguna expresión u³ que tenga una ocurrencia en u¹ y por otra expresión que tenga una ocurrencia en u². Así, p.ej., en el enunciado ‘el filósofo cuyo más famoso alumno fue Alejandro Magno fundó el Liceo en Atenas’, no hay ninguna ocurrencia de ‘Alejandro Magno fundó el Liceo’; ese enunciado es un caso concreto de u; en u, hay una

ocurrencia de u^1 —que será ‘el filósofo cuyo más famoso alumno fue Alejandro Magno’— y de u^2 —que será ‘fundó el Liceo en Atenas’.

Aclarado qué es una ocurrencia de una expresión en otra, veamos ahora qué es una ocurrencia esencial de una expresión en un enunciado.

Diremos que una expresión ‘ e ’ tiene una ocurrencia esencial en el enunciado ‘ p ’ ssi ‘ e ’ tiene una ocurrencia en ‘ p ’ y, además, **suponiendo** que ‘ p ’ sea —en uno u otro grado— verdadero, la sustitución de **esa** ocurrencia de ‘ e ’ en ‘ p ’ por una ocurrencia de otra expresión puede dar como resultado un enunciado enteramente falso. Así, p.ej., sea ‘ p ’ ‘Averroes es un filósofo sumamente adicto a la doctrina de Aristóteles’. En tal enunciado las expresiones ‘Averroes’, ‘doctrina’, ‘Aristóteles’ tienen ocurrencias esenciales. (Suponemos que el enunciado es verdadero, como de hecho lo es.) Si sustituimos ‘Averroes’ por ‘Nicolás de Cusa’, el resultado es un enunciado enteramente falso. (El Cusano no era, **en absoluto**, sumamente adicto a la doctrina de Aristóteles.) Si sustituimos ‘doctrina’ por ‘esclava’, también obtenemos como resultado una oración enteramente falsa; y lo mismo si sustituimos ‘Aristóteles’ por ‘Platón’, p.ej.

Ahora ya podemos definir qué es una verdad de lógica: un enunciado ‘ p ’ es una verdad de lógica ssi cumple dos condiciones:

- 1) ‘ p ’ es verdadero;
- 2) las únicas expresiones que tienen ocurrencias esenciales en ‘ p ’ son (algunos de entre) los signos: ‘y’, ‘o’, ‘no es cierto que’, ‘es enteramente cierto que’, ‘no sólo... sino que también’, ‘es en todos los aspectos cierto que’, ‘es un sí es no cierto que’, ‘es tan cierto que... como que...’, ‘si... entonces’, ‘algún ente’, ‘es’, etc.

¿Qué se quiere decir al escribir ese ‘etc.’? Que la lista no está cerrada. No hay cómo determinar, de una vez por todas, qué expresiones son aceptables en la lista de expresiones que son las únicas en poder tener ocurrencias esenciales en verdades lógicas. Ello lleva consigo que la noción de verdad lógica ha de conservar un margen de indeterminación. Nada se opone a que se llame ‘lógico’ a un sistema en el que, además de las palabras indicadas, también otras entren en juego como palabras que tendrán ocurrencias esenciales en las verdades del sistema (p.ej.: las palabras ‘antes de que’, ‘desde que’, ‘hasta que’; las palabras ‘es un hecho necesario que’, ‘es un hecho posible que’, etc.).

En cualquier caso, aquí llamaremos ‘lógico’, desde luego, a todo sistema de enunciados en los que sólo las expresiones de la lista indicada tengan ocurrencias esenciales.

Diremos que todas las expresiones de la lista indicada constituyen el **vocabulario** de la lógica (similaramente, expresiones como ‘planta’ constituyen el vocabulario de la botánica; expresiones como ‘Morelos’, ‘Santa Anna’, ‘Juárez’, ‘Zapata’, etc. constituyen el vocabulario de la Historia de México, etc.)

Nótese que, si un enunciado es una verdad de lógica, entonces cada resultado de sustituir en él las ocurrencias de una expresión cualquiera que **no** pertenezca al vocabulario de la lógica por ocurrencias respectivas de otras expresiones seguirá siendo una verdad de lógica, siempre y cuando se respeten dos condiciones:

- 1ª) Esa sustitución ha de hacerse de tal modo que el resultado sea también una oración sintácticamente bien formada. P.ej., no se puede sustituir en ‘Laura tiene celos de Toribia’, la ocurrencia que hay de ‘Toribia’ por una ocurrencia de ‘así’, pues el resultado no sería una oración.
- 2ª) La sustitución ha de hacerse **uniformemente**. Ello quiere decir que, si ‘ p ’ es una verdad de lógica en la que aparece una expresión ‘ e ’ que **no** forma parte del vocabulario lógico, y ‘ e ’ tiene, p.ej., dos ocurrencias en ‘ p ’, entonces la sustitución de ‘ e ’ por otra expresión ‘ w ’

cualquiera ha de hacerse de tal modo que cada una de las dos ocurrencias de 'e' en 'p' sea reemplazada por una ocurrencia respectiva de 'w'. Así, p.ej., no es una verdad de lógica 'Si Maximino tiene lumbago y si Gisberta tiene dolor de muelas, entonces Teótimo se va a divorciar'.

Veamos algunos ejemplos de verdades lógicas: 'Conrado ama a Cleta o Conrado no ama a Cleta'. Es obvio que ese enunciado es una verdad de lógica (si es que es una verdad; los intuicionistas podrían negarlo, puesto que rechazan el principio de tercio excluso). Aunque sustituyamos una palabra de ese enunciado —salvo 'o' y 'no'— por otra, el resultado seguirá siendo una verdad de lógica (supuesto siempre un sistema en el que el principio de tercio excluso sea válido, y dándose por supuesto que la sustitución habrá de hacerse de tal modo que el resultado de la misma siga siendo una fórmula sintácticamente bien formada).

He aquí otro ejemplo de verdad lógica: 'Si Felipe está cocinando, [entonces] Felipe está cocinando'. Hagamos cuantas sustituciones **uniformes** queramos en ese enunciado, con tal de que mantengamos la expresión 'si..., entonces'. El resultado seguirá siendo verdad.

Para concluir este párrafo, conviene precisar unos pocos puntos más. No es original el intento que aquí se lleva a cabo de representar simbólicamente fragmentos —lo más amplios posible— de la lengua natural. Insértase ese intento en lo mucho que al respecto se lleva haciendo en los últimos lustros en la filosofía analítica. No sería acertado hablar aquí de una "traducción", pues —según se verá en el cap. 1 de esta obra y según lo ha puesto de relieve Quine en trabajos ampliamente conocidos— el mejor modo de ver las notaciones simbólicas es el de concebirlas como representaciones **esquemáticas** de mensajes en lengua natural —en lengua natural regimentada, si se quiere, pero que no por ello deja de pertenecer al ámbito de la lengua natural o hablada en general. En verdad, el fundador de la lógica matemática, Gottlob Frege se proponía ya, al articular el primer sistema de lógica matemática en una notación simbólica, más que acuñar mediante ésta última un lenguaje diverso del natural, elaborar, para representar por escrito ciertos fragmentos del discurso en lengua natural, una **escritura** (de ahí el término que utiliza: 'Begriffsschrift', o sea: 'conceptografía'). Y, si bien habla a menudo de un lenguaje artificial para expresar en él el saber matemático, puede eso deber entenderse con las paráfrasis adecuadas a tenor de esa noción fregeana de la escritura conceptográfica.

Cabe prever un cierto asombro de algunos retardatarios ante el planteamiento de la tarea de representar simbólicamente fragmentos del habla en lengua natural —mediante letras esquemáticas y gracias a lecturas apropiadas. Sin embargo —cabe recalcar— no es ahí donde estriba la originalidad del presente trabajo, pues no en vano ha existido la obra de Richard Montague, la de Cresswell, la de Fitch, la de Sommers y la de tantos otros que, con fortuna diversa, han emprendido esa tarea. Y en lo tocante al recurso a lógicas no clásicas con esa finalidad cabe citar la obra de Lofti Zadeh, de Goguen, de G. Lakoff, para no mencionar a los relevantistas, conexivistas y conceptivistas, todos los cuales sostienen que el discurso en lengua natural es informalizable con los magros recursos de la lógica clásica.

Es anhelo compartido en esa comunidad científica el alcanzar esa representación formalizada tanto de los discursos correctos como de los sofísticos, justamente para poder someterlos a patrones de enjuiciamiento rigurosos y poder detectar así los sofismas. Sólo que, mal que le pese al dogmático y ciego absolutismo de algún clasicista intransigente, qué discursos sean sofísticos y cuáles no lo sean no es asunto que quepa decidir de una vez por todas e independientemente de qué código de reglas de inferencia se esté tomando como patrón, explícita o implícitamente. El *modus ponens* (e.d. la regla que autoriza a concluir 'q' de 'p' y 'Si p, q') parecerá a muchos infrangible; y, si bien quien esto escribe también lo considera válido (mas no indubitavelmente tal), algunos lógicos lo han rechazado, no dándole cabida en los sistemas que han puesto en pie, de suerte que —de ser válidos tales sistemas como adecuados patrones de inferencia correcta— serían sofísticos ciertos argumentos que sólo

utilizaran como regla de deducción ésa del *modus ponens*. Por el contrario, son lógicamente no-válidos (y, por ende, sofísticos) muchos razonamientos bastante banales en los que figuran construcciones comparativas, si el criterio de lo lógicamente válido es la lógica clásica: no se sigue —según esa lógica, que es aquella cuya incuestionabilidad es ciegamente asumida por los clasicistas recién aludidos— de que Laponia sea fría y de que Siberia sea más fría que Laponia que Siberia es fría: eso es según la lógica clásica un sofisma (si sofisma es una secuencia de enunciados $\langle p^1, \dots, p^n, q \rangle$ tales que no es una regla de inferencia derivable [en el sistema lógico que uno adopte] la regla $\{p^1, \dots, p^n\} \vdash q$). (Obsérvese, por cierto, que aquí representamos una inferencia de una conclusión $\lceil q \rceil$ a partir de las premisas $\lceil p^1 \rceil, \lceil p^2 \rceil, \dots, \lceil p^n \rceil$ como $\{p^1, \dots, p^n\} \vdash q$. En cambio —y en aras de un tratamiento más uniformemente general de las relaciones— en el capítulo VI la representaremos como $\langle p^1, p^n \rangle \vdash q$ —aunque haciendo la salvedad de que, cuando una regla autoriza tal inferencia, igualmente autoriza una inferencia cualquiera que únicamente difiera de ésa por permutaciones en la secuencia de premisas. Con esa cláusula, ambas presentaciones son equivalentes, pues ambas hacen que se extraiga la conclusión de la [única] combinación n -aria de las n premisas. Por otra parte, la misma representación sirve para una inferencia y para una regla que la autorice, puesto que ésta última viene concebida como el cúmulo de todas las inferencias de cierta índole —índole que se exhibe mediante los signos pertinentes que, en cada caso, figuren en la representación gráfica de las premisas y de la conclusión.)

Para aclarar mejor el final del párrafo precedente —y adelantándonos aquí a explicaciones ulteriores— cabe apuntar esto. Se usan las llaves, $\{ \}$, para indicar un **conjunto** (o un cúmulo de cosas), de suerte que $\{x, z, u, v\}$ denotará p.ej. un cúmulo que abarque tan sólo a x , a z , a u , a v , y a nada más. Los ángulos, $\langle \rangle$, se usan para denotar a **n -tuplos ordenados**. El signo $\vdash$ úsase para indicar la inferibilidad del enunciado escrito a la derecha a partir del cúmulo de enunciados a la izquierda de ese signo. (Normalmente, empero, suprimense las llaves a la izquierda de $\vdash$ sin desmedro por ello de la claridad, pues están sobreentendidas.)

§3.— El carácter logográfico de las notaciones lógicas

Conviene reflexionar sobre las definiciones precedentes, ya que dista de ser baladí el problema en ellas involucrado. ¿Qué es en verdad una notación simbólica? Hay dos maneras de entenderla. Para unos es un sistema representativo (semántico) independiente del lenguaje, en el cual los signos hacen [directamente] las veces de cosas, de relaciones entre cosas, o al menos de “ideas” o “conceptos” (sea eso lo que fuere, si es que lo hay). Para otros, cada sistema de notación simbólica en lógica y en matemáticas es una escritura de mensajes en lengua natural; los signos de tal notación hacen las veces de elementos de un lenguaje hablado. A la primera concepción cabe denominarla **semasiográfica**, a la segunda **glotográfica** (pues convencional y corrientemente se llama ‘semasiográfico’ a un sistema gráfico del primer tipo, y ‘glotográfico’ a una escritura de una lengua hablada). La discrepancia está en principio clara. Pero es difusa, admitiendo grados. Es semasiográfico un sistema gráfico cuando admite muy diversas “traducciones” a una misma lengua natural —a cualquier lengua natural. Es glotográfico cuando hay una lengua natural con respecto a la cual cabe una sola lectura o “traducción”. Con arreglo a ese criterio, se suele alegar que $\langle 2^2+52 \rangle 3$ es un mensaje no glotográfico, porque cabe traducirlo de maneras diversas a diversos idiomas y, en cada uno, con mensajes diversos.

Ahora bien, hay escrituras glotográficas que no responden a esa condición de univocidad, por la presencia de abreviaturas, simplificaciones y ambigüedades (piénsese p.ej. en la ambigüedad resultante en escrituras taquigráficas, o en la de ciertos mensajes cifrados etc.). Además, de los sistemas glotográficos, unos —los logográficos— estriban en que cada signo gráfico represente una unidad significativa del idioma —puede ser un monema, pero puede

ser una palabra o una locución cualquiera. Un sistema logográfico puede representar por una misma ristra de grafemas a muy diversos mensajes de la lengua hablada.

¿No hay diferencia? Sí, pero de grado. En un sistema glotográfico es menor la variedad de lecturas o traducciones que en uno semasiográfico; al menos hay un orden lineal más o menos igual —entre el mensaje gráfico total y su lectura o “traducción”— en lo tocante a la ordenación de oraciones que lo forman, mientras que en una representación semasiográfica no sucede así. Pero también eso es cuestión de grado, evidentemente.

Quienes conciben a las notaciones simbólicas de la lógica como “lenguajes” propios —o sea: sistemas semasiográficos, con su propia sintaxis— pueden jactarse de emancipar a esas notaciones de toda sumisión al lenguaje natural, con sus defectos reales o supuestos. Pero esa postura suscita grandes dificultades epistemológicas, toda vez que está entonces por mostrar: (1º) qué relevancia tiene lo escrito en esa notación para el saber que se expresa en lengua natural; (2º) cómo ha sido posible pasar del pensamiento expresado en lengua natural al que se plasma en la notación en cuestión (porque, o bien ésta es capaz sólo de expresar pensamientos también expresables en lengua natural, sólo que con mayor ambigüedad, o bien no es así en absoluto: pero entonces, ¿cómo se ha accedido a esos pensamientos nuevos?); (3º) cómo puede uno “estar” en esa notación, o sea cómo se integra y se engrana lo en ella expresado en y con el cúmulo de nuestros otros pensamientos.

Mucho más verosímil parece, pues —porque nos ahorra esos quebraderos de cabeza y porque es más simple y clara—, la concepción de las notaciones como escrituras. Esa es la concepción de Quine y también la aquí postulada. Las notaciones son como taquigrafías, que tienen varias lecturas, y que son reutilizables de una lengua a otra, dentro de ciertos límites. La necesidad de esas notaciones para la lógica es meramente pragmática; para decirlo sin pretensiones de rigor, son muletas que ayudan a que el pensamiento principalmente oral y auditivo venga ayudado por el pensamiento visual.

Todo lo que puede hacerse en lógica matemática con el recurso de notaciones simbólicas puede hacerse también sin él. (Puede. Pero, ¿podemos efectivamente nosotros —con un poder a tenor de nuestras limitadísimas capacidades?) Las notaciones simplemente permiten poner de relieve qué vocabulario es esencial o pertinente para ciertas reflexiones intelectuales. Reemplazan por letras **esquemáticas** (en el sentido que se explicará en el capítulo I de este trabajo) los segmentos de un mensaje que no sean esenciales —aquellos que no contengan ocurrencias de **vocablos** de la disciplina con la que se esté trabajando, e.d. ocurrencias de expresiones que no aparezcan en teoremas de esa disciplina con ocurrencias esenciales—; y representan con signos compactamente dibujables a los vocablos de esa disciplina, al paso que en la lengua natural tales vocablos pueden ser locuciones complejas, o discontinuas, o signos suprasegmentales (prosódicos —entonación, p.ej.—, o bien orden de palabras) etc.

Entendida así una notación lógica o matemática está claro por qué la lógica ni es el estudio de un delimitado terreno, de una zona especial, de cierto ámbito de entidades aparte, ni tampoco es un mero juego o una mera “apofántica” formal que meramente se limite a “prescribir” condiciones “formales” de significatividad —las cuales (en la concepción aquí criticada) vendrían empero condicionadas por las cláusulas de traducción a la lengua natural), o cosa así. No, las verdades lógicas son verdades sobre cualesquiera cosas, sin excepción. Verdades sobre los tomates, las alubias, los alquileres, los protones, los diccionarios de malgache, los granos de arena, las declaraciones de amor o las películas. Son verdades lógicas (o **de lógica**) sólo todas aquellas que se expresan por enunciados de un idioma cualquiera en los cuales aparezcan con ocurrencias esenciales únicamente vocablos “lógicos”; e.d. vocablos que figuran en una lista ofrecida por quienes se llaman lógicos como constituyendo el vocabulario de su disciplina. Hay diversas listas tales, y enconados desacuerdos. Unos meten más palabras, otros menos.

Todos parecen estar de acuerdo en, como máximo, estos vocablos: ‘no’, ‘y’, ‘o’, ‘si ... entonces’, ‘algún’, ‘cada’ y otros parónimos (según algunos, sinónimos).

Vienen luego las discusiones de qué hace que sean ésas las ‘constantes lógicas’. Y hay opiniones para todos los gustos. Es uno de los muchos temas debatidos en filosofía de la lógica. Una respuesta simple (y, *cæteris paribus*, ¿no son preferibles las respuestas simples?) es que los entoides significados por esos vocablos son “generales” en el sentido de que entran en cualesquiera “combinaciones” o relaciones de otros entes, sean los que fueren, al paso que relaciones como la de **estar contiguo a**, o la de **ser de mayor volumen que**, no poseen esa generalidad (la justicia y la democracia no están relacionadas por ellas, p.ej.).

Con todos los problemas que encierra esa noción de generalidad, parece preferible atenerse a ella mientras no se ofrezca algo mucho más claro. Y hasta ahora nada más claro se ha ofrecido.

§4.— Prerrequisitos para la lectura del presente opúsculo

Según quedó ya apuntado en el §1 de esta Introducción, la redacción del presente trabajo se ha llevado a cabo a fin de: por un lado, ofrecer al lector estudioso un instrumento para adentrarse por sí mismo en la lógica matemática con un espíritu no dogmático, sino pluralista, abierto al cuestionamiento, a la crítica, a la confrontación argumentativa entre diversas teorías lógicas que vienen aquí presentadas cual opciones alternativas, en un abanico; y, por otro lado, facilitar a quienes ya poseen conocimientos básicos de lógica matemática —sólo que centrados exclusivamente en la lógica clásica— un medio sencillo, fácil y entretenido de acceder a las lógicas **no** clásicas y de percatarse de que son lógicas con todas las de la ley, genuinas alternativas a la lógica clásica.

Metas son ésas, sí, difíciles de alcanzar. Sobre todo cuando, aunque el autor hubiera preferido que el libro sirviera más para la primera de ellas, se da cuenta de lo probable que es que venga usado principalmente para alcanzar la segunda, es decir para que estudiantes ya algo familiarizados con la lógica clásica se adentren en el ámbito de las lógicas no clásicas.

En cualquier caso, trátase de lo uno o de lo otro, la obra permite adentrarse en el terreno que cubre sin necesidad de ningún conocimiento previo, gracias a una redacción que ha tratado de eliminar todo lo que se podría dar por sabido, explicitando lo más posible y desmenuzando las explicaciones. Aun así, espérase del lector un manejo previo, siquiera mínimo, de lo más elemental en el estudio introductorio a la lógica matemática, a saber: un conocimiento de las tablas de verdad bivalentes **verticales** —aquellas en que dos letras esquemáticas están colocadas una al lado de otra y, entre ellas, un signo como la disyunción, la conjunción, el condicional o el bicondicional; un conocimiento así suele venir proporcionado por las más elementales iniciaciones a la lógica en la enseñanza secundaria. Igualmente se espera que el lector posea una familiaridad con las nociones más elementales de teoría ingenua de conjuntos. Ningún otro conocimiento técnico especial es menester para entender el presente trabajo. Sólo que, eso sí, hay diversas maneras, y diversos grados, en la comprensión de un escrito; y naturalmente aquellos lectores que estén más familiarizados con la lógica clásica podrán —pasando por alto algunas explicaciones pormenorizadas que para ellos serán ya ociosas, por lo obvias— concentrarse mejor en los temas más avanzados. (Cabe advertir, empero, que el cap. 12 —y último— es de lectura más difícil que el resto, aunque tampoco requiere especiales conocimientos previos.)

Capítulo 1

Notaciones, paréntesis, puntos

Esquemas y funtores

Vamos a explicar aquí algunos signos que se utilizarán en el resto de esta exposición. Aclaremos, ante todo, el uso que hacemos de las letras 'p', 'q', 'r', 's', 'p¹', 'q¹'..., 'p²' etc. Esas letras serán usadas como letras **esquemáticas**. Ello significa que, al escribir una o varias de esas letras afectadas por funtores monádicos o diádicos, no se estará escribiendo ninguna fórmula, ninguna oración, sino un esquema oracional. Un esquema oracional no es ni verdadero ni falso; para que algo pueda tener un valor de verdad, es menester que sea una oración. Y un esquema **no** es una oración. Un esquema es, simplemente, un sucedáneo, una representación formal de la forma de un número infinito de oraciones. Al escribir un esquema, al decir que es una verdad lógica, estamos empleando un modo impropio de hablar, que puede y debe ser corregido —a costa de la brevedad—, diciendo que cualquier oración que tenga la forma de ese esquema es una verdad lógica.

Sea, p.ej., el esquema siguiente: $\lceil p \vee Np \rceil$. Diremos que eso es una verdad lógica. En realidad, lo que queremos decir —al hablar así, de manera impropia, pero útil, por lo breve—, es que es una verdad lógica cualquier resultado de sustituir —uniformemente—, en ese esquema, la letra 'p' por una oración cualquiera. Es decir, que son verdades lógicas las oraciones:

Almudena está enamorada o Almudena no está enamorada

Menem ganó las elecciones o Menem no ganó las elecciones

Cuando se habla de algo o alguien, se le da un nombre. Cuando se habla de una expresión, se le da un nombre a la expresión; y, para dárselo, se usa —comúnmente— el procedimiento de encerrar la expresión entre comillas **simples**. Pero, a veces, queremos hablar, no sobre una oración en particular, sino sobre una oración **cualquiera** que tenga la forma de un esquema dado: entonces, evidentemente, no podemos dar un nombre a "la" oración, puesto que no hablamos de ninguna oración en particular, sino de **una oración cualquiera** de esa forma. Para hablar de esa oración cualquiera, acuñamos un **pseudonombre**, encerrando el esquema entre esquinas (*corners*), e.d. los signos ' $\lceil$ ' a la izquierda —al comienzo del esquema— y ' $\rceil$ ' a la derecha —al final del esquema—, según lo hemos hecho unas pocas líneas más atrás. Así, decimos: $\lceil p \vee Np \rceil$ es una verdad. Lo que con ello se quiere decir es que es verdad una oración cualquiera de esa forma (cualquier oración de la forma '... o no es cierto que...'). En algunos casos, informalmente, pueden reemplazarse las esquinas por comillas dobles, especialmente cuando ello suceda en una formulación en román paladino donde se estén brindando **lecturas** de los signos que constituyen la notación simbólica utilizada.

Aclarado esto, expliquemos ahora el uso de ciertos signos.

El signo ' $\supset$ ' ha de leerse 'si... entonces' o 'sólo si'; así $\lceil p \supset q \rceil$ se leerá: 'Si p, entonces q', o bien «p sólo si q»; y también «El hecho de que p entraña (o conlleva) el hecho de que q».

El signo 'N' ha de leerse 'No es cierto que', o 'Es falso que'. Así $\lceil Np \rceil$ se leerá «No es cierto que p» o «Es falso que p».

El signo ' $\neg$ ' se leerá 'Es del todo falso que' (y sus sinónimas 'Es enteramente falso que', 'Es totalmente falso que', etc.).

El signo 'H' ha de leerse 'Es enteramente cierto que'

El signo 'L' ha de leerse 'Es, hasta cierto punto por lo menos, verdad que', o bien 'Es, en uno u otro grado, verdad que', o bien 'Es más o menos cierto que'.

El signo ' $\vee$ ' se leerá 'o'. El signo ' $\wedge$ ' se leerá 'y'. El signo '&' se leerá 'y sobre todo' ($\lceil p \& q \rceil$ se leerá «Siendo verdad que p, [lo es que] q»).

$\lceil p \rightarrow q \rceil$ se leerá de uno de los siguientes modos: «El hecho de que p implica el hecho de que q », «El hecho de que p es a lo sumo tan cierto como el hecho de que q ».

$\lceil p \sqcap q \rceil$ se leerá: «El hecho de que p es cierto en la misma medida en que lo es el hecho de que q »; «El hecho de que p equivale al hecho de que q ».

$\lceil Bp \rceil$ se leerá: «Es afirmable con verdad que p », o bien «Es en todos los aspectos cierto que p », o bien «Es genuinamente cierto que p », o bien «Es realmente cierto que p ».

$\lceil p \sqsubset q \rceil$ se leerá: «El hecho de que p es menos cierto que el hecho de que q »; «El hecho de que q es más cierto que el hecho de que p ».

$\lceil Pp \rceil$ se leerá: «Es más bien cierto que p ».

$\lceil Sp \rceil$ se leerá: «Es verdadero y falso a la vez que p » y también «No es ni verdadero ni falso que p ».

Una fórmula puede ser atómica o molecular. Es atómica ssi no hay en ella otra fórmula que sí misma. P.ej., 'Eusebia es emperatriz'. Es molecular si hay en ella alguna ocurrencia de **otra** fórmula, p.ej. 'Eusebia no es emperatriz' o sea: 'No es cierto que Eusebia sea emperatriz'. En esa fórmula hay una ocurrencia de 'Eusebia es emperatriz'.

Si una fórmula $\lceil p \rceil$ es molecular, entonces cada fórmula que tenga una ocurrencia en $\lceil p \rceil$ será llamada una subfórmula de $\lceil p \rceil$.

Antes de proseguir, impónense dos observaciones incidentales. La primera es que el uso, en los sistemas lógicos que se van a presentar en este opúsculo, de operadores como los recién mencionados —con esas respectivas lecturas— no significa pretensión de exhaustividad o ambición de ofrecer, gracias a ellos, formalización de cualquier discurso en lengua natural. No, trátase tan sólo de dar un paso adelante con respecto a la lógica clásica, cuyo vocabulario es sumamente pobre y desconoce toda matización alética o veritativa, todo lo que sea un **más** o un **menos**. Sin duda tendemos (pero asintóticamente no más) al desideratum de una formalización exhaustiva de la lengua natural. Mas de hecho constituye un avance en esa dirección, importante, sí, pero modesto, la inclusión de esos pocos operadores, desconocidos en la lógica clásica (la negación simple, 'N', el operador 'L', la implicación ' $\rightarrow$ ', la sobreimplicación ' $\sqsupset$ ' etc.). Sucede empero que gracias a pasos adelante como ése va resultando viable la formalización de fragmentos de la lengua natural que son suficientemente amplios para determinados propósitos —p.ej. que son [muchos de] los que (acaso en versiones un tanto regimentadas) figuran en los contextos de elocución del quehacer investigativo.

La segunda observación que conviene hacer aquí es que, en los ejemplos aducidos de expresiones que contienen lo que, en lengua natural, corresponde a varios de los operadores lógicos mencionados, úsase el adjetivo 'cierto' —según es usual en español— como equivalente, idiomáticamente apropiado, de 'verdadero' (que en esos contextos no se suele usar); no es, pues, un uso en el cual 'cierto' signifique lo mismo que 'dotado de certeza', e.d. **seguro**. En otros idiomas no existe un uso así de los adjetivos por los que se suele traducir 'cierto': en francés se dice 'Il est très vrai' (y no 'Il est très certain'); en inglés 'It is very true' (y no 'It is very certain'); etc. Los grados involucrados en nuestro uso idiomático de 'cierto' en esos contextos son grados de verdad, no grados de certeza.

Paréntesis y puntos

Vamos ahora a explicar el uso de los paréntesis y puntos.

Cada ocurrencia de un functor **monádico** en una oración o esquema afectará a la subfórmula más corta que siga a tal ocurrencia. Cuando se quiera que afecte a alguna subfórmula más amplia que la fórmula (o letra esquemática) que la sigue inmediatamente, entonces se escribirá, inmediatamente después de la ocurrencia en cuestión, un paréntesis izquierdo; y, al terminar

la fórmula afectada por dicha ocurrencia del functor monádico en cuestión, se cerrará el paréntesis. Así, p.ej., tenemos $\neg(N(p \supset q) \vee r)$. Una fórmula así —una instancia sustitutiva de ese esquema— difiere de [una instancia sustitutiva de] $\neg(Np \supset q \vee r)$. En este último esquema, la ocurrencia del functor 'N' sólo afecta a ' p '; en el anterior, la ocurrencia de 'N' afecta a ' $p \supset q$ '.

Veamos ahora lo que ocurre con cada ocurrencia de un functor diádico. Cada ocurrencia de un functor **diádico** se ha de hallar entre una fórmula a su izquierda y otra a su derecha, afectando a ambas (o sea, haciendo corresponder a ese par de fórmulas otra fórmula. Así, en ' $p \supset q$ ', el functor ' $\supset$ ' hace corresponder al par de fórmulas conformado por ' p ' y ' q ', **en ese orden**, otra fórmula, a saber: ' $p \supset q$ '). Diremos, pues, que cada ocurrencia de un functor diádico tiene un miembro izquierdo y un miembro derecho. Pues bien: el miembro izquierdo de cada ocurrencia de un functor diádico será toda la parte de la fórmula que haya desde el comienzo de la misma (si bien ha de entenderse esto con una restricción que viene explicitada en el último párrafo de este capítulo). Y su miembro derecho será siempre la subfórmula más corta que siga a la ocurrencia en cuestión del functor diádico. Así ' $p \vee q \supset r \wedge s$ ' habrá de analizarse como sigue: la primera ocurrencia de un functor diádico que encontramos es la de ' $\vee$ '; el miembro izquierdo de esa ocurrencia será ' p '; su miembro derecho será ' q '. Luego encontramos una ocurrencia de ' $\supset$ ': su miembro izquierdo será ' $p \vee q$ '; su miembro derecho será ' r '. Luego encontramos una ocurrencia de ' $\wedge$ ': su miembro izquierdo será ' $p \vee q \supset r$ '; su miembro derecho será ' s '.

¿Cómo hacer, entonces, para que una ocurrencia de un functor diádico tenga como miembro derecho una subfórmula más grande que la más pequeña subfórmula que siga a tal ocurrencia? Podemos distinguir dos casos:

1) Se quiere que el miembro derecho sea todo lo que sigue, hasta el final de la fórmula total; entonces se escribe un punto inmediatamente después de la ocurrencia en cuestión. Así en ' $p \supset q \vee r$ ', el miembro derecho de ' $\supset$ ' es ' $q \vee r$ '; igualmente, en ' $p \wedge q \supset s$ ', el miembro derecho de la ocurrencia de ' $\wedge$ ' es ' $q \supset s$ '.

2) Se quiere que el miembro derecho no sea ni la más pequeña subfórmula que sigue a la ocurrencia del functor diádico en cuestión, ni tampoco todo el resto de la fórmula total, sino algo intermedio; entonces, hay que acudir a los paréntesis, abriendo un paréntesis izquierdo inmediatamente después de la ocurrencia del functor en cuestión, y cerrando un paréntesis derecho inmediatamente después del final de la subfórmula que sea el miembro derecho de la ocurrencia en cuestión del functor diádico. Así en ' $p \supset (q \vee r) \wedge s$ ' tenemos que el miembro izquierdo de ' $\supset$ ' es ' p ', y su miembro derecho es ' $q \vee r$ '; (el miembro izquierdo de ' $\wedge$ ' será, obviamente, ' $p \supset (q \vee r)$ ' y su miembro derecho será ' s ').

Téngase en cuenta, por último, que nunca puede saltarse dentro de un paréntesis a fuera de él, ni a la inversa tampoco: una ocurrencia de un functor dentro de un paréntesis no puede afectar a nada fuera del paréntesis. Y, por otro lado, **dentro** del paréntesis rigen las mismas normas que para una fórmula independiente. Así, p.ej., en ' $p \supset (q \vee r \wedge s) \wedge p$ ', la ocurrencia que hay de ' $\vee$ ' tiene como miembro derecho a ' $r \wedge s$ '.

Resumiendo —y para dar una pauta—: cuando se vea un punto escrito inmediatamente después de una ocurrencia de un functor diádico, se sabe que su miembro derecho es todo el resto de la fórmula total. Si, al recorrer una fórmula total, encontramos una primera ocurrencia de un functor diádico inmediatamente seguida de un punto, esa ocurrencia es la ocurrencia principal de la fórmula total; si encontramos luego una segunda ocurrencia que sea también inmediatamente seguida de un punto, esa ocurrencia es la ocurrencia principal del miembro derecho de la fórmula total; si encontramos una tercera ocurrencia también inmediatamente seguida de un punto, es que esta ocurrencia es la ocurrencia principal **del miembro derecho**

del miembro derecho de la fórmula total. (Naturalmente, el miembro derecho del miembro derecho de una fórmula **no** es lo mismo en absoluto que el miembro derecho de la fórmula.)

Ahora ya sólo resta exponer algunos ejemplos, y dejar al lector, como ejercicio, el elaborar otros similares y más complicados. (Para ilustrar el procedimiento y ver su correspondencia con un sistema notacional que use profusión de paréntesis, encerraremos entre paréntesis todo esquema no atómico, salvo la fórmula total):

$\lceil p \vee (r \wedge s) \wedge q \vee p \supset s \rceil$	abr	$\lceil ((p \vee (r \wedge s)) \wedge q) \vee (p \supset s) \rceil$
$\lceil p \vee r \wedge s \wedge q \vee p \supset s \rceil$	abr	$\lceil p \vee (((r \wedge s) \wedge q) \vee p) \supset s \rceil$
$\lceil p \vee r \wedge s \wedge (q \vee p) \supset s \rceil$	abr	$\lceil (((p \vee r) \wedge s) \wedge (q \vee p)) \supset s \rceil$
$\lceil p \vee (r \wedge s \wedge q \vee p) \supset s \wedge p' \rceil$	abr	$\lceil (p \vee ((r \wedge s) \wedge (q \vee p))) \supset (s \wedge p') \rceil$
$\lceil p \wedge N(q \supset s) \wedge p' \supset q \supset r \rceil$	abr	$\lceil ((p \wedge N(q \supset s)) \wedge p') \supset (q \supset r) \rceil$

(Naturalmente, la regla dada más arriba —en el párrafo tercero de este acápite— según la cual el miembro izquierdo de una ocurrencia de un functor diádico es toda la parte de la fórmula que se halla a su izquierda ha de entenderse ahora con esta restricción: toda la parte de la fórmula a su izquierda **desde la última ocurrencia de un functor diádico inmediatamente seguido de un punto**).

Capítulo 2

Noción de dominio de valores de verdad

Noción de valor de verdad

Un valor de verdad es un ente que se hace corresponder a una oración al clasificar a ésta con respecto a la verdad o falsedad. Según cuáles sean los valores de verdad que se reconocen y según las relaciones que se admitan entre ellos, variará la lógica que uno acepte.

La lógica bivalente sólo acepta dos únicos valores de verdad: lo lisa y llanamente verdadero (que se escribe '1') y lo lisa y llanamente falso (que se escribe '0').

La lógica clásica —la que axiomatizaron Frege y Russell— es una lógica bivalente verifuncional, o sea una lógica bivalente en la cual tanto el signo 'y' —formalmente representado como ' $\wedge$ '— como el signo 'no es verdad en absoluto' (o su sinónimo 'es enteramente falso que') —formalmente representado como ' $\neg$ '— son, ambos, **functores**, el primero diádico y el segundo monádico.

Un **functor** es un signo tal que, colocado, ya delante de una fórmula, ya entre dos fórmulas, da por resultado **otra** fórmula, siempre y cuando el valor de verdad de la fórmula resultante dependa del (o los) valor(es) de verdad de la(s) fórmula(s) de que se parte. Cuando el functor es un signo que se coloca simplemente delante de una fórmula, es un functor **monádico**. Cuando se coloca entre dos fórmulas, es diádico.

(Nótese lo siguiente: se solía leer el único functor monádico de esa lógica clásica como un mero 'no'; pero esa lectura puede parecer incorrecta, desde el punto de vista filosófico del autor de estas páginas —y desde muchos otros puntos de vista.)

Valores designados, valores antidesignados

Con respecto a un dominio de valores de verdad interesa saber, no sólo cuántos son, sino también: cuáles son los valores designados; cuáles son los valores antidesignados; y qué relación(es) de orden hay entre los valores de verdad. Ahora definiremos tales nociones.

Por valor designado se entiende un valor tal que, si es el valor de verdad de una oración dada, ' p ', entonces ' p ' es afirmable (o sea: toda persona al tanto de las cosas hará bien en afirmar ' p ').

Por valor de verdad antidesignado se entiende un valor tal que, si es el valor de verdad de una oración ' p ', entonces ' p ' es negable (o sea: una persona al tanto de las cosas hará bien en afirmar **la negación (simple) de ' p '**; nótese que decimos 'la negación', no la supernegación; la negación de una fórmula cualquiera ' q ' es ' $\neg q$ ' —o sea ' Nq '—, mientras que su supernegación es «No es cierto **en absoluto** que q » —o sea ' $\neg\neg q$ ').

Hablando de modo menos riguroso, un valor de verdad es designado ssi es verdadero; es antidesignado ssi es falso.

Pues bien, en la lógica bivalente sólo hay dos valores: uno de ellos designado, y el otro antidesignado; no hay en ella ningún valor que sea, a la vez, designado y antidesignado; ni tampoco hay, en ella, valor alguno que no sea ni designado ni antidesignado.

Se llama **lógica multivalente** a la que reconoce un dominio de valores de verdad conformado por **más de** dos valores de verdad. En ciertas lógicas multivalentes hay valores que no son ni designados ni antidesignados; en otras hay valores que son, a la vez, designados y antidesignados; y, por último, hay también lógicas en las que hay valores designados y antidesignados y valores ni designados ni antidesignados —a la vez, por supuesto, que algún valor sólo designado y algún valor sólo antidesignado.

Relaciones de orden en un dominio de valores de verdad

Veamos ahora qué se entiende por relación de orden definida sobre un dominio de valores de verdad.

Una relación (diádica) puede verse como un cúmulo de pares ordenados. Así, la relación **padre-de** es un cúmulo que abarca, entre muchos otros, a los pares $\langle \text{Luis XIII}, \text{Luis XIV} \rangle$, $\langle \text{Fernando VII}, \text{Isabel II} \rangle$, $\langle \text{Isaac}, \text{Jacob} \rangle$, $\langle \text{David}, \text{Salomón} \rangle$, etc. (Con mayor vigor, sin embargo, vendrá definida una relación n -ádica más abajo como un cúmulo de pares $\langle \langle x^1, \dots, x^n \rangle, z \rangle$.)

Una relación está definida sobre un dominio ssi se define para miembros de tal dominio. Así, la mencionada relación de paternidad sólo se define para el dominio de los animales pluricelulares sexuados. Otro ejemplo de relación definida sobre un dominio es la de **sucesor-de** definida sobre el cúmulo de los números enteros.

Una relación es reflexiva ssi cada miembro del dominio en que está definida está ligado a sí mismo por tal relación.

Una relación es antisimétrica ssi, en el caso de que un miembro x del dominio sobre el que está definida tenga tal relación con respecto a un miembro z , y éste tenga esa relación con respecto a x , entonces $x=z$.

Una relación es transitiva ssi, en el caso de que un miembro, x , del dominio tenga tal relación con respecto a un miembro, u , y éste la tenga con respecto a un miembro, z , entonces x tiene tal relación con respecto a z .

La relación de paternidad no es reflexiva; la relación de ser-de-la-misma-estatura-que sí lo es; la relación de ser-a-lo-sumo-tan-grande-como (o sea: ser más pequeño o igual que) —definida sobre el dominio de los números reales— es antisimétrica; la relación de amistad no es antisimétrica; la relación de ser-abuelo-de no es transitiva; la de ser-antepasado-de sí es transitiva.

Pues bien, una relación de orden es una relación que sea reflexiva, antisimétrica y transitiva. (Una relación de preorden es una relación reflexiva y transitiva, que sea antisimétrica o no). Una relación de orden definida sobre un dominio es llamada **orden total** (o también **orden lineal** u orden conexo) ssi, dados dos miembros cualesquiera de tal dominio, o bien el primero tiene tal relación con respecto al segundo, o bien el segundo la tiene con respecto al primero. Una relación de orden que no sea total es llamada relación de **orden parcial**.

Lógicas escalares y lógicas tensoriales

Hasta ahora hemos visto únicamente dominios de valores de verdad **escalares**, o sea dominios en los que cada valor de verdad se expresa por un solo constituyente. Pero también hay dominios de valores de verdad **tensoriales** (llamados asimismo **lógicas-producto**), en los que cada valor de verdad tiene varios constituyentes. Esos constituyentes de un valor de verdad son sus componentes.

Así, p.ej., podemos construir un dominio de cuatro valores de verdad tensoriales a partir de los dos valores de verdad de la lógica clásica, combinándolos de dos en dos, como sigue: **11, 10, 01, 00**.

Toda lógica tensorial se construye a partir de una (o varias) lógica(s) escalar(es). (Aquí no consideraremos más que a las lógicas tensoriales construidas a partir de una sola lógica escalar, es decir: una sola en **cada** caso.)

En una lógica tensorial, lo más sensato es considerar como valores designados sólo aquellos todos cuyos componentes son designados en la lógica escalar de la que se parte; y considerar como valores antidesignados sólo aquellos cuyos componentes son antidesignados en la lógica

de la que se parte. Así, en la lógica tensorial arriba mencionada sólo será designado el valor **11**; sólo será antidesignado el valor **00**.

En una lógica así, el orden a definir mediante el signo ' $\leq$ ' será el siguiente:

$$\mathbf{00} \leq \mathbf{00} \quad \mathbf{00} \leq \mathbf{11} \quad \mathbf{10} \leq \mathbf{10}$$

$$\mathbf{00} \leq \mathbf{01} \quad \mathbf{01} \leq \mathbf{01} \quad \mathbf{10} \leq \mathbf{11}$$

$$\mathbf{00} \leq \mathbf{10} \quad \mathbf{01} \leq \mathbf{11} \quad \mathbf{11} \leq \mathbf{11}$$

Como se ve, esa relación de orden no es conexa: dados los valores **01** y **10**, ni el primero tiene la relación indicada con respecto al segundo, ni la tiene tampoco el segundo con respecto al primero.

Capítulo 3

Noción de tautología

Una tautología con respecto a una lógica dada es una fórmula que, cualesquiera que sean los valores de verdad de sus subfórmulas atómicas, tiene siempre, forzosamente, un valor designado.

Para saber si una fórmula es una tautología, primero tenemos que saber cómo calcular —en una lógica verifuncional— el valor de fórmulas más complejas a partir de sus subfórmulas atómicas. A este respecto, se construyen tablas de verdad. Veamos ejemplos de tablas de verdad en la lógica bivalente, en una lógica trivalente, y en otra pentavalente.

Lógica bivalente

(Único valor designado: 1; valor antidesignado: 0)

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

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

p	$\neg p$
1	0
0	1

$\supset$	1	0
1	1	0
0	1	1

Lógica trivalente A_3

(Valores designados: 1 y $\frac{1}{2}$; valores antidesignados: $\frac{1}{2}$ y 0)

$\wedge$	1	$\frac{1}{2}$	0
1	1	$\frac{1}{2}$	0
$\frac{1}{2}$	$\frac{1}{2}$	$\frac{1}{2}$	0
0	0	0	0

$\vee$	1	$\frac{1}{2}$	0
1	1	1	1
$\frac{1}{2}$	1	$\frac{1}{2}$	$\frac{1}{2}$
0	1	$\frac{1}{2}$	0

$\rightarrow$	1	$\frac{1}{2}$	0
1	$\frac{1}{2}$	0	0
$\frac{1}{2}$	$\frac{1}{2}$	$\frac{1}{2}$	0
0	$\frac{1}{2}$	$\frac{1}{2}$	$\frac{1}{2}$

$\supset$	1	$\frac{1}{2}$	0
1	1	$\frac{1}{2}$	0
$\frac{1}{2}$	1	$\frac{1}{2}$	0
0	1	1	1

I	1	0
1	$\frac{1}{2}$	0
$\frac{1}{2}$	0	$\frac{1}{2}$
0	0	0

p	S_p	H_p	$\neg p$	Np
1	0	1	0	0
$\frac{1}{2}$	$\frac{1}{2}$	0	0	$\frac{1}{2}$
0	0	0	1	1

(En esta lógica cabe concebir a **1** como lo totalmente verdadero; a **0** como lo totalmente falso; a $\frac{1}{2}$ como lo a la vez verdadero y falso.)

Lógica pentavalente A_5

(Valores designados: 4,3,2,1; Valores antidesignados: 3,2,1,0)

$\wedge$	4	3	2	1	0	$\vee$	4	3	2	1	0	$\mid$	4	3	2	1	0	
4	4	3	2	1	0	4	4	4	4	4	4	4	2	0	0	0	0	
3	3	3	2	1	0	3	4	3	3	3	3	3	0	2	0	0	0	
2	2	2	2	1	0	2	4	3	2	2	2	2	0	0	2	0	0	
1	1	1	1	1	0	1	4	3	2	1	1	1	0	0	0	2	0	
0	0	0	0	0	0	0	4	3	2	1	0	0	0	0	0	0	2	
$\supset$	4	3	2	1	0	$\rightarrow$	4	3	2	1	0	p	$\neg p$	Hp	Lp	Pp	Np	Sp
4	4	3	2	1	0	4	2	0	0	0	0	4	0	4	4	4	0	0
3	4	3	2	1	0	3	2	2	0	0	0	3	0	0	4	3	1	1
2	4	3	2	1	0	2	2	2	2	0	0	2	0	0	4	2	2	2
1	4	3	2	1	0	1	2	2	2	2	0	1	0	0	4	0	3	1
0	4	4	4	4	4	0	2	2	2	2	2	0	4	0	0	0	4	0

(En esta lógica cabe concebir a **4** como lo totalmente verdadero; a **3** como lo verdadero y falso pero más verdadero que falso; a **2** como lo igualmente falso que verdadero; a **1** como lo verdadero y falso pero más falso que verdadero, y a **0** como lo totalmente falso.)

Lógicas tensoriales

Para construir tablas de verdad en lógicas-producto (lógicas tensoriales), téngase presente lo siguiente: en cada lógica-producto hay que tener en cuenta: 1º, cuál es la lógica escalar de la que se parte; 2º, cuántos componentes constituyen cada valor de verdad. Pues bien, para construir una tabla de verdad se tendrá en cuenta que el i -ésimo componente del resultado (para cualquier $i \geq 1$) será una función del (o los) i -ésimo(s) componente(s) de la(s) fórmula(s) afectada(s) por el functor, calculado(s) según se hace en la lógica escalar a partir de la cual se ha construido la lógica-producto en cuestión. Así tendremos, en el caso más simple (una lógica-producto de dos componentes construida a partir de la lógica bivalente): (Valor designado: **11**; antidesignado: **00**):

$\wedge$	11	10	01	00	$\vee$	11	10	01	00	$\supset$	11	10	01	00	p	$\neg p$
11	11	10	01	00	11	11	11	11	11	11	11	10	01	00	11	00
10	10	10	00	00	10	11	10	11	10	10	11	11	01	01	10	01
01	01	00	01	00	01	11	11	01	01	01	11	10	11	10	01	10
00	00	00	00	00	00	11	10	01	00	00	11	11	11	11	00	11

Veamos ahora una lógica-producto de dos componentes contruidos a partir de la lógica trivalente. (Valores designados: **11**, **1½**, **½1**, **½½**; Valores antidesignados: **½0**, **0½**, **00**, **½½**.)

$\wedge$	11	1½	10	½1	½½	½0	01	0½	00
11	11	1½	10	½1	½½	½0	01	0½	00
1½	1½	1½	10	½½	½½	½0	0½	0½	00
10	10	10	10	½0	½0	½0	00	00	00
½1	½1	½½	½0	½1	½½	½0	01	0½	00
½½	½½	½½	½0	½½	½½	½0	0½	0½	00
½0	½0	½0	½0	½0	½0	½0	00	00	00
01	01	0½	00	01	0½	00	01	0½	00
0½	0½	0½	00	0½	0½	00	0½	0½	00
00	00	00	00	00	00	00	00	00	00

Añadiremos el functor 'B', con la siguiente tabla de verdad:

p	Bp
11	11
10	00
01	00
00	00

Queda como ejercicio para el lector construir tablas en esta lógica para los demás funtores previamente definidos en la lógica trivalente de partida.

Lista de algunas tautologías

Dejamos como ejercicio para el lector la comprobación de estas tautologías en la lógica trivalente A_3

$$p \rightarrow q \wedge (q \rightarrow p) \vdash p \vdash q$$

$$p \supset \neg p \supset q$$

$$\neg p \supset p \supset q$$

$$p \wedge \neg p \supset q$$

$$p \supset q \supset p \supset p$$

$$p \supset q \supset \neg q \supset \neg p$$

$$p \wedge q \vee r \vdash p \vee r \wedge q \vee r$$

$$\neg p \supset \neg q \supset q \supset p$$

$$p \vee q \wedge r \vdash p \wedge r \vee q \wedge r$$

$$p \wedge p \vdash p$$

$$p \vee p \vdash p$$

$$p \wedge q \vee p \vdash p$$

$$p \vee q \wedge p \vdash p$$

$$p \vdash q \vdash p$$

$$p \rightarrow (q \wedge r) \vdash p \rightarrow q \wedge p \rightarrow r$$

$$\neg Lp \vdash Fp$$

$$p \vdash q \rightarrow p \vdash r \vdash q \vdash r$$

$$H(LNp \vee p)$$

$$N(p \vee q) \vdash Np \wedge Nq$$

$$\neg(p \wedge q) \vdash \neg p \vee \neg q$$

$$\neg(p \vee q) \vdash \neg p \wedge \neg q$$

$$NNp \vdash p$$

$$p \rightarrow Lp$$

$$Lp \supset p$$

$$p \vee q \vee \neg p$$

$$p \vee q \vee Np$$

$$Np \vee q \vee Hp$$

$$p \vee q \vdash q \vee p$$

$$p \vdash p \vdash N(p \vdash p)$$

$$p \rightarrow q \vee q \rightarrow p$$

$$Np \supset p \vdash p$$

$$Np \vdash q \vdash p \vdash Nq$$

$$\neg Sp \vdash \neg p \vee Hp$$

$$p \rightarrow Np \rightarrow Np$$

$$NLp \vdash \neg p$$

$$N \neg p \vdash Lp$$

$$p \supset q \vee q \supset r$$

$$p \supset q \supset p$$

$$p \supset (q \wedge r) \vdash p \supset q \wedge p \supset r$$

$$p \vdash q \rightarrow p \wedge r \vdash q \wedge r$$

$$p \supset q \wedge (q \supset p) \supset Lp \vdash Lq$$

$$\neg Np \vdash Hp$$

$H(p \wedge q)I.Hp \wedge Hq$	$p \wedge qI.q \wedge p$	$LLpILp$
$L(p \wedge q)I.Lp \wedge Lp$	$p \wedge (q \wedge r)I.p \wedge .q \wedge r$	$N \neg pI \neg \neg p$
$HLpILp$	$p \vee (q \vee r)I.p \vee .q \vee r$	$N(p \wedge q)I.Np \vee Nq$
$L \neg pIFp$	$p \wedge qI(p \vee q)I.plq$	$H(\neg p \vee Lp)$
$H(Lp \vee Np)$	$p \rightarrow qI.Nq \rightarrow Np$	$HHpIHp$
$H(p \vee q)I.Hp \vee Hq$	$p \supset NpINp$	$p \rightarrow q \supset .p \supset q$
$p \supset (q \vee r)I.p \supset q \vee .p \supset r$	$p \vee qIpl.q \rightarrow p$	$plq \supset .p \supset q \wedge .q \supset p$
$p \wedge q \supset rI.p \supset r \vee .q \supset r$	$p \wedge NpISp$	$plNpl.plplp$
$p \vee q \supset rI.p \supset r \wedge .q \supset r$	$SSpISp$	$p \vee qI p \vee .p \vee qI p$
$p \wedge q \supset rI.p \supset .q \supset r$	$p \supset .q \rightarrow Lp$	$Hp \rightarrow p$
$p \supset (q \supset r)I.q \supset .p \supset r$	$p \rightarrow (q \vee r)I.p \rightarrow q \vee .p \rightarrow r$	$plqI.NpINq$
$p \supset (q \supset r) \supset .p \supset q \supset .p \supset r$	$plq \rightarrow .p \vee rI.q \vee r$	$SpISNp$
$H \neg pI \neg p$	$p \supset q \wedge (q \supset p) \supset . \neg pI \neg q$	$L(p \vee q)I.Lp \vee Lq$
$Np \rightarrow p \rightarrow p$	$HNpI \neg p$	$p \vee q \rightarrow rI.p \rightarrow r \wedge .q \rightarrow r$
	$p \wedge q \rightarrow rI.p \rightarrow r \vee .q \rightarrow r$	$LHpIHp$

¿Son tautologías los esquemas siguientes?

Compruebe el lector, para cada uno de los siguientes esquemas, si es o no una tautología de A_3

$p \rightarrow Hp$	$S(plp)I.plp$
$N(p \rightarrow q)$	$p \supset q \supset .p \rightarrow q$
$p \rightarrow .q \rightarrow p$	$N(p \rightarrow Np)$
$p \supset q \supset .Nq \supset Np$	$p \supset .Np \supset q$
$Np \supset .p \supset q$	$p \supset Nq \supset .q \supset Np$

¿Cómo calcular las tautologías de una lógica tensorial?

Con respecto a las lógicas-producto (lógicas tensoriales) es obvio que una fórmula dada, $\lceil p \rceil$, es una tautología de una lógica-producto $\mathfrak{L}$ cuya lógica escalar de base es $\mathfrak{L}$ ssi $\lceil p \rceil$ es una tautología de $\mathfrak{L}$. En efecto: para que $\lceil p \rceil$ sea una tautología tiene que tener uniformemente valores designados. Supongamos que $\lceil p \rceil$ es una tautología de $\mathfrak{L}$: entonces cada uno de los componentes de cualquiera de los valores de verdad que pueda tomar $\lceil p \rceil$ será designado en $\mathfrak{L}$; y, por consiguiente cada valor posible de $\lceil p \rceil$ será designado en $\mathfrak{L}$.

Supongamos, por otra parte, que $\lceil p \rceil$ es una tautología en $\mathfrak{L}$: entonces todo valor posible de verdad de $\lceil p \rceil$ será designado en $\mathfrak{L}$, esto es: será un valor todos cuyos componentes serán designados. Pero, entonces, quiere decirse que, en cada uno de los puestos, cualquiera que sea la combinación de componentes de los átomos (o sea de sus valores de verdad en $\mathfrak{L}$), el resultado será un componente designado (o sea: un valor de verdad designado de $\mathfrak{L}$).

Por consiguiente, mientras no se introduzcan funtores propios de las lógicas-producto —como el functor ‘B’—, no hay razón para calcular independientemente las tautologías de la lógica-producto; basta con calcular —más simplemente— las de la lógica escalar de base respectiva.

Calcúlese ahora, en una lógica A_3^2 (o sea: una lógica de dos componentes, cuya lógica escalar de base sea A_3) si las fórmulas siguientes (las enumeradas en la columna derecha) son tautologías, definiéndose el functor ‘B’ mediante la siguiente indicación: $/Bp/ = /p/$ si $/p/$ no contiene ningún 0; y en caso contrario $/Bp/ = \langle 00 \rangle$.

$Bp \rightarrow p$	$Bp \rightarrow BBp$	$Bp \vee BNBp$
$Bp \supset . Bplp$	$B(p \supset q) \rightarrow . Bp \supset Bq$	$Bp \wedge Bql . B(p \wedge q)$
$Bp \supset p$	$LBplBLp$	$Bp \vee BqlB(p \vee q)$
$p \supset Bp$	$Bp \vee B \neg p$	$Bp \vee Bq \rightarrow B(p \vee q)$
$Bp \vee B \neg Bp$	$B(p \rightarrow q) \rightarrow . Bp \rightarrow Bq$	

Pondremos fin a este acápite enunciando la noción de **cuasi-tautología**: tomenos un sistema trivalente A_3 , pero quitemos el estatuto de valor designado al valor $\frac{1}{2}$. Llamemos al resultado A_1 . Pues bien, una fórmula $\lceil p \rceil$ será una cuasi-tautología de A_1 ssi $\lceil p \rceil$ es una tautología de A_3 . Una fórmula de A_1 puede ser una cuasi-tautología de ese sistema sin ser una tautología del **mismo** (aun siendo una tautología de A_3). Pero es interesante constatar lo siguiente: el resultado de prefijar a cualquier cuasi-tautología de A_1 una ocurrencia del functor 'L' da por resultado una tautología de A_1 .

Capítulo 4

Estudio de varios funtores

Vamos a estudiar en este capítulo las condiciones que debe reunir un functor dado —cualquiera que sea el dominio escogido de valores de verdad— para poder ser clasificado en un grupo particular de funtores —monádicos o diádicos.

En toda esta exposición, el resultado de encerrar entre barras verticales a una fórmula designará el valor de verdad de la misma. También daremos por supuesto que el cúmulo de valores de verdad contiene un elemento mínimo 0 (antidesignado y no designado) y un elemento máximo 1 (designado y no antidesignado), tales que, para todo $\lceil p \rceil$, $\lceil p \rceil \leq 1$, y $0 \leq \lceil p \rceil$, siendo $\leq$ la relación de orden —total o parcial— definida sobre el dominio de valores de verdad que se haya escogido.

Funtores de afirmación

Un functor de afirmación es cualquier functor monádico $\$$ tal que, para cualquier $\lceil p \rceil$:

- 1) $\lceil \$p \rceil$ es designado sólo si $\lceil p \rceil$ es designado.
- 2) Hay algún $\lceil q \rceil$ tal que $\lceil \$q \rceil$ es designado.

Un functor de **sobreafirmación** es un functor de afirmación $\dagger$ tal que:

Hay algún $\lceil q \rceil$ tal que $\lceil q \rceil$ es designado, pero $\lceil \dagger q \rceil$ no es designado.

Un functor de afirmación se llamará functor **asertivo** si cumple la condición siguiente (para cualquier $\lceil p \rceil$):

Si $\lceil p \rceil$ es designado, $\lceil \$p \rceil$ es designado.

En A_3 el functor 'L' es un functor asertivo, mientras que 'H' es sobreafirmativo.

Un tipo particular de funtores de afirmación lo constituyen los que podemos llamar **catafánticos**: un functor monádico $\blacktriangledown$ es catafántico ssi, para cualquier $\lceil p \rceil$ y $\lceil q \rceil$ (y siendo ' $\wedge$ ' un functor de conjunción natural, y ' $\vee$ ' un functor de disyunción natural):

- 1) $\lceil \blacktriangledown p \rceil \leq \lceil \blacktriangledown q \rceil$ sólo si o bien $\lceil p \rceil \leq \lceil q \rceil$, o bien $\lceil \blacktriangledown p \rceil \leq 0$
- 2) $\lceil p \rceil \leq \lceil q \rceil$ sólo si $\lceil \blacktriangledown p \rceil \leq \lceil \blacktriangledown q \rceil$
- 3) $\lceil \blacktriangledown (p \wedge q) \rceil = \lceil \blacktriangledown p \wedge \blacktriangledown q \rceil$
- 4) $\lceil \blacktriangledown (p \vee q) \rceil = \lceil \blacktriangledown p \vee \blacktriangledown q \rceil$

Un functor puede ser asertivo y catafántico a la vez; pero puede ser catafántico sin ser asertivo (siendo sobreafirmativo); y puede ser asertivo sin ser catafántico. Así, p.ej., el functor 'L' no es catafántico, puesto que supongamos que $\lceil p \rceil = 1$ y $\lceil q \rceil = \frac{1}{2}$. Entonces:

$\lceil Lp \rceil \leq \lceil Lq \rceil$, pero —contrariamente a la condición (1)— ni $\lceil Lp \rceil \leq 0$, ni $\lceil p \rceil \leq \lceil q \rceil$

En realidad, dentro de una lógica trivalente no se ve la utilidad de esa noción de functor catafántico. Para que tal noción sea útil, hay que ascender a lógicas más ricas y, sobre todo, a las lógicas infinivalentes —de las que se hablará más adelante. Un ejemplo interesante de functor catafántico es el functor 'P' que definimos más arriba para una lógica pentavalente. ('H' es también un functor catafántico.)

Otro subconjunto interesante de los funtores de afirmación es el que abarca a los funtores internos. Un functor $\blacktriangle$ de afirmación es interno ssi, para todo $\lceil p \rceil$:

$$\lceil \blacktriangle \blacktriangle p \rceil = \lceil \blacktriangle p \rceil$$

Son internos todos los funtores de afirmación de los que hemos hablado aquí. Pero hay interesantes funtores de afirmación que no son internos. Véase un ejemplo: el functor 'X' definido como sigue para una lógica pentavalente y otra heptavalente (y que cabe leer 'Es muy cierto que'):

p	Xp	p	Xp
4	4	6	6
3	2	5	4
2	1	4	3
1	1	3	2
0	0	2	1
		1	1
		0	0

(Mucho mayor interés cobra esa noción en una lógica infinivalente).

Funtores de negación

Un functor de negación es un functor monádico ' $\sim$ ' tal que, para todo ' p ' :

- 01) Al menos uno de entre $/p/$ y $/\sim p/$ o bien es designado o bien no es antidesignado.
- 02) Al menos uno de entre $/p/$ y $/\sim p/$ o bien es antidesignado o bien no es designado.
- 03) $/p/$ es designado ssi $/\sim p/$ es antidesignado.
- 04) Si $/\sim p/$ es designado, entonces $/p/$ es antidesignado.
- 05) $/p/=0$ ssi $/\sim p/=1$.
- 06) Si $/p/=1$, entonces $/\sim p/=0$.
- 07) Si $/\sim p/=0$, entonces $/p/$ es designado.

Un functor de negación $\sim$ es una negación natural (o simple) ssi, para todo ' p ' :

- 08) Si $/p/$ es antidesignado, entonces $/\sim p/$ es designado.
- 09) $/p/=/\sim\sim p/$
- 10) Si $/\sim p/=0$, entonces $/p/=1$.

Obviamente, ' N ' es una negación simple o natural, tanto en las lógicas escalares como en las lógicas-producto.

Un functor $\sim$ de negación es una supernegación ssi hay algún ' p ' para el cual no cumple ninguna de las condiciones (08) a (10), pero, en cambio, cumple las tres siguientes, para cualquier ' p ' :

- 11) A lo sumo uno de entre $/p/$ y $/\sim p/$ es designado.
- 12) Si $/p/$ es designado, entonces $/\sim p/ \leq 0$.
- 13) Si $/\sim p/$ es designado, entonces $/p/ \leq 0$ y $1 \leq /p/$.

Funtores de conjunción

Un functor diádico, '&', es una conjunción ssi, para cualesquiera ' p ', ' q ' y ' r ' :

- 01) $/p\&q\&r/ = /p\&q\&r/$
- 02) $/p\&q/$ es designado ssi tanto $/p/$ como $/q/$ son designados.
- 03) Si o $/p/ \leq 0$ o $/q/ \leq 0$, entonces $/p\&q/ \leq 0$.
- 04) Para cualquier functor de negación, $\sim$, $/p\&\sim p/$ es antidesignado.
- 05) $/p\&p/ = /p/$.

Esta última condición es acaso demasiado fuerte. Hay conyunciones (o “superconyunciones”) que no la cumplen. Podría venir reemplazada por ésta otra:

05b) $/p \& q/ \leq q$.

Dentro de las conyunciones, definimos ahora una conyunción natural como cualquier functor ‘ $\wedge$ ’ de conyunción tal que, para cualquier ‘ p ’, ‘ q ’ y ‘ r ’:

06) $/p \wedge q/ = /q \wedge p/$

07) Si $1 \leq /p \wedge q/$, entonces tanto $1 \leq /p/$ como $1 \leq /q/$.

08) Si $/p/ = 1$, entonces $/p \wedge q/ = /q/$.

Un interesante functor de conyunción que no es natural es la conyunción ‘ $\&$ ’ definida como sigue en un sistema trivalente y en otro pentavalente:

$\&$	1	$\frac{1}{2}$	0
1	1	$\frac{1}{2}$	0
$\frac{1}{2}$	1	$\frac{1}{2}$	0
0	0	0	0

$\&$	4	3	2	1	0
4	4	3	2	1	0
3	4	3	2	1	0
2	4	3	2	1	0
1	4	3	2	1	0
0	0	0	0	0	0

Es evidente que un functor como ‘ $p \& q$ ’ en uno de esos sistemas equivale a ‘ $Lp \wedge q$ ’.

Functores de disyunción

Un functor diádico ‘ $\vee$ ’ es un functor de disyunción ssi, para cualquier ‘ p ’, ‘ q ’ y ‘ r ’:

01) $/p/ = /p \vee p/$

02) $/p \vee q \vee r/ = /p \vee q \vee r/$

03) Si ‘ $\wedge$ ’ es una conyunción natural, entonces: $/q \wedge r \vee p/ = /q \vee p \wedge r \vee p/$; y $/p \vee q \wedge r/ = /p \vee q \wedge p \vee r/$

04) Si ‘ $\wedge$ ’ es una conyunción natural entonces: $/q \wedge p \vee r \wedge p/ \leq /q \vee p/$ y $/p \wedge q \vee p \wedge r/ \leq /p \wedge q \vee r/$

05) $/p \vee q/$ es antidesignado ssi tanto $/p/$ como $/q/$ son antidesignados.

06) O bien, si $/p/ = 0$, entonces $/p \vee q/ = /q/$; o bien, si $/q/ = 0$, entonces $/p \vee q/ = /p/$

07) Si o bien $/p/ = 1$, o bien $/q/ = 1$, entonces $/p \vee q/ = 1$

08) Para cada functor de negación ‘ $\sim$ ’, $/\sim p \vee p/$ es designado

09) Si ‘ $\neg$ ’ es un functor de supernegación, entonces $/\neg p \vee \neg \neg p/ = 1$

10) Si ‘ $\neg$ ’ es un functor de supernegación, y si tanto $/\neg p/$ como $/p \vee q/$ son designados, también es lo es $/q/$.

Una disyunción ‘ $\vee$ ’ es natural si, además, satisface, para cualquier ‘ p ’, ‘ q ’ y ‘ r ’ las condiciones siguientes (siendo ‘ $\wedge$ ’ una conyunción natural y ‘ N ’ una negación natural).

11) $/p \vee q/ = /q \vee p/$

12) $/q \vee r \wedge p/ \leq /q \wedge p \vee r \wedge p/$

13) Si $/p/$ es designado $/p \vee q/$ es designado

14) Si $/p \vee q/ = 0$, entonces $/p/ = 0$ y $/q/ = 0$

15) $/p \vee q/ = /N(Np \wedge Np)/$

Un interesante functor de disyunción no natural es el siguiente (definido en A_3):

$\vee$	1	$\frac{1}{2}$	0
1	1	1	1
$\frac{1}{2}$	1	$\frac{1}{2}$	0
0	1	$\frac{1}{2}$	0

Se puede definir como sigue ' $\vee$ ': ' $p \vee q$ ' abr ' $H p \vee q$ '. También se puede definir de estos otros modos alternativos: ' $p \vee q$ ' abr ' $N p \supset q$ ' ' $p \vee q$ ' abr ' $N(N p \& N q)$ '

Una oración ' $p \vee q$ ' podría leerse: "Es enteramente cierto que p a menos que sea cierto que q».

Es interesante constatar que, tomando como primitivo a ' $\vee$ ' se puede también definir ' $\supset$ ' y '&' como sigue: ' $p \supset q$ ' abr ' $N p \vee q$ ' ' $p \& q$ ' abr ' $N(N p \vee N q)$ '

Functores condicionales

Un functor condicional es un functor diádico $\supset$ tal que, para cualesquiera ' p ', ' q ' y ' r ' (siendo ' $\vee$ ', ' $\wedge$ ', ' N ', ' $\neg$ ', respectivamente, funtores de disyunción natural, conyunción natural, negación natural y supernegación):

- 01) $/p \supset p/$ es designado
- 02) Si $/p \supset q/$ y $/q \supset r/$ son designados, $/p \supset r/$ es designado
- 03) Hay algunos ' s ' y ' r ' tales que $/s \supset r/$ es designado, pero $/r \supset s/$ no lo es
- 04) Si $/p/$ y $/p \supset q/$ son designados, también lo es $/q/$
- 05) $/p \supset p \supset q \supset q/$ es designado
- 06) $/p \supset (q \supset r) \supset . p \supset q \supset . p \supset r/$ es designado
- 07) $/p \supset . q \wedge r/ = /p \supset q \wedge . p \supset r/$
- 08) $/p \supset . q \vee r/ = /p \supset q \vee . p \supset r/$
- 09) $/p \vee q \supset r/ = /p \supset r \wedge . q \supset r/$
- 10) $/p \wedge q \supset r/ = /p \supset r \vee . q \supset r/$
- 11) $/p \supset N p \supset N p/$ es designado
- 12) $/p \supset \neg p \supset \neg p/$ es designado
- 13) $/p \supset q \supset . p \wedge r \supset q/$ es designado
- 14) $/p \supset q \supset . p \supset . q \vee r/$ es designado
- 15) Si $/q/ \leq 0$ y $/p \supset q/$ es designado, entonces $/p/ \leq 0$
- 16) Si $1 \leq /q/$, entonces $/p \supset q/$ es designado
- 17) Hay algún functor de negación $\sim$ tal que $/p \supset q/$ es designado ssi $/\sim q \supset \sim p/$ es designado.

Llamaremos condicional simple o **mero condicional** a cualquier functor condicional ' $\supset$ ' tal que, para cualquier ' p ', ' q ' y ' r ':

- 18) $/p \supset q \vee . q \supset r/$ es designado
- 19) $/p \supset p \supset q/ = /q/$
- 20) $/p \supset . p \supset q/ = /p \supset q/$
- 21) $/p \supset . q \supset p/$ es designado
- 22) $/p \supset . q \supset r/ = /p \wedge q \supset r/ = /q \supset . p \supset r/$
- 23) $/p \supset q/ = / \neg p \vee q/$
- 24) $/ \neg q \supset \neg p/$ es designado ssi también lo es $/p \supset q/$

25) $/p \supset p \supset q \supset q/$ es designado

26) $/p \supset q \supset p \supset p/$ es designado

Es evidente que el functor ' $\supset$ ', con las tablas de verdad que se adjudicaron al mismo en A_3 y A_5 , es un condicional simple o mero condicional.

Un functor condicional ' $\rightarrow$ ' es un functor implicativo ssi ' $\rightarrow$ ' **no** es un condicional simple y, en cambio, ' $\rightarrow$ ' satisface las condiciones siguientes (para cualesquiera ' p ' y ' q ')

27) $/p \rightarrow q/ = /Nq \rightarrow Np/$

28) $/p \rightarrow q/$ es designado ssi $/p \wedge q/ = /p/$

29) $/p \rightarrow q/$ y $/q \rightarrow p/$ son ambos designados a la vez ssi $/p/ = /q/$

30) Si $/p \rightarrow q/$ es designado, entonces $/p \rightarrow Nq/$ es antidesignado

(De (30) se desprende que $/N(p \rightarrow Np)/$ ha de ser designado).

Funtores bicondicionales y equivalenciales

Un functor bicondicional es un functor diádico $\leftrightarrow$ tal que hay algún functor condicional ' $\supset$ ' tal que, para cualquier ' p ' y ' q ' (siendo ' $\wedge$ ' una conyunción natural):

$/p \leftrightarrow q/ = /p \supset q \wedge q \supset p/$

Por consiguiente, cada functor bicondicional debe satisfacer las condiciones siguientes, entre otras:

01) $/p \leftrightarrow q/$ es designado ssi $/q \leftrightarrow p/$ es designado

02) Si $/p \leftrightarrow q/$ es designado, entonces, o bien tanto $/p/$ como $/q/$ son designados, o bien ni $/p/$ ni $/q/$ son designados.

03) $/p \leftrightarrow p/$ es designado

04) Si $/p \leftrightarrow q/$ y $/q \leftrightarrow r/$ son designados, también lo es $/p \leftrightarrow r/$

05) Hay algún functor de negación $\sim$ tal que $/p \leftrightarrow q/$ es designado ssi $/\sim p \leftrightarrow \sim q/$ es designado

06) Si $/p \leftrightarrow q/$ es designado, también lo es $/p \wedge r \leftrightarrow q \wedge r/$

07) Si $/p \leftrightarrow q/$ es designado, también lo es $/p \vee r \leftrightarrow q \vee r/$

Un functor bicondicional, ' $\equiv$ ', será llamado **mero bicondicional** o bicondicional simple ssi cumple las condiciones siguientes:

08) $/p \equiv q \equiv r/ = /p \equiv q \equiv r/$

09) Si tanto $/p/$ como $/q/$ son designados, $/p \equiv q/$ es designado

Un functor que es un mero bicondicional es el functor ' $\equiv$ ' con las siguientes tablas de verdad en A_3 y A_5

$\equiv$	1	$\frac{1}{2}$	0
1	1	$\frac{1}{2}$	0
$\frac{1}{2}$	$\frac{1}{2}$	$\frac{1}{2}$	0
0	0	0	1

$\equiv$	4	3	2	1	0
4	4	3	2	1	0
3	3	3	2	1	0
2	2	2	2	1	0
1	1	1	1	1	0
0	0	0	0	0	4

Ese functor puede ser definido así: ' $p \equiv q$ ' abr ' $p \supset q \wedge q \supset p$ '

Ahora veamos qué es un functor equivalencial.

Es un functor bicondicional, I , tal que:

10) Si $/p/q/$ es designado, entonces $/p/ = /q/$

Por consiguiente, si $/p/q/$ es designado, entonces, si $\lceil s \rceil$ es el resultado de reemplazar, en una fórmula **cualquiera** $\lceil s \rceil$, una ocurrencia de $\lceil p \rceil$ por otra de $\lceil q \rceil$, resulta que $/s/s'/$ es designado.

La importancia de esta conclusión es enorme. En ella consiste el principio verifuncional de extensionalidad.

Sobreimplicaciones

Una sobreimplicación es un functor diádico $\blacktriangleleft$ tal que —siempre para cualesquiera $\lceil p \rceil$, $\lceil q \rceil$ y $\lceil r \rceil$ — (siendo 'N' la negación natural, ' $\neg$ ' la supernegación, ' $\rightarrow$ ' functor implicativo y ' $\equiv$ ' functor equivalencial):

01) $/p\blacktriangleleft p/ \leq 0$

02) Si $/p\blacktriangleleft q/$ es designado, $/q\blacktriangleleft p/ \leq 0$

03) Si $/p\blacktriangleleft q/$ es designado, $/q/$ es designado

04) Si $/p\blacktriangleleft q/$ es designado, $/p/$ es antidesignado

05) Si $/p\blacktriangleleft q/$ y $/q\blacktriangleleft r/$ son designados, $/p\blacktriangleleft r/$ es designado

06) $/p\blacktriangleleft q/ = /p \rightarrow q \wedge \neg(q \rightarrow p)/$

07) $/p\blacktriangleleft q \vee q\blacktriangleleft p \vee p/q/$ es designado

08) $/p\blacktriangleleft q/ = /Nq\blacktriangleleft Np/$

09) Si $/p/$ no es designado y $/q/$ es designado, entonces $/p\blacktriangleleft q/ \neq 0$

10) Si $/q/$ no es antidesignado y $/p/$ es antidesignado, entonces $/p\blacktriangleleft q/ \neq 0$

Un functor sobreimplicativo es un functor que ha de leerse como un comparativo de inferioridad. Escribiremos un functor sobreimplicativo así ' $\backslash$ ', y lo definiremos en A_3 y en A_5 mediante las siguientes tablas de verdad:

$\backslash$	1	$\frac{1}{2}$	0
1	0	0	0
$\frac{1}{2}$	$\frac{1}{2}$	0	0
0	$\frac{1}{2}$	$\frac{1}{2}$	1

$\backslash$	4	3	2	1	0
4	0	0	0	0	0
3	2	0	0	0	0
2	2	2	0	0	0
1	2	2	2	0	0
0	2	2	2	2	0

Que el hecho de que q sea más verdadero que el hecho de que p significa lo siguiente: el hecho de que p implica el hecho de que q, pero es del todo falso que el hecho de que q implique el hecho de que p.

Escolio al Capítulo 4

Aplicaciones a la formalización de enunciados en lenguaje natural

Según se anunció ya al final del §2 de la Introducción del presente opúsculo, planteamientos como el aquí ofrecido brindan maneras —no forzosamente únicas— de formalizar enunciados, de la más amplia gama posible, pertenecientes a la lengua natural. Pero sólo como tendencia asintótica cabe aspirar a una formalización exhaustiva. De hecho, esa formalización plantea difíciles problemas, en los que no podemos entrar aquí. (Varios de los trabajos citados en la

bibliografía del presente opúsculo se refieren a temas estrechamente relacionados con dicha formalización.) Lo que sí cabe, no obstante, es hacer unas indicaciones al respecto.

1º.— Hay que tener en cuenta que, en la estructura de superficie de la lengua natural, los funtores monádicos, en vez de aparecer prefijados a los enunciados que afectan, aparecen incrustados en el interior de tales enunciados. Así, en vez de decirse ‘Es, hasta cierto punto por lo menos, verdad que Evagro es atlético’ se dice, más simplemente, ‘Evagro es, hasta cierto punto por lo menos, atlético’. Similarmente, en vez de ‘Es muy cierto que Amalio es glotón’, se dice ‘Amalio es muy glotón’.

2º.— De manera semejante, las implicaciones, equivalencias y comparativos sufren, en la estructura de superficie de la lengua natural, un proceso de incrustación parcial en una u otra de las dos fórmulas que vinculen —asociado, a menudo, a otros procesos más complicados, como la elipsis. Así, en vez de decirse ‘Es menos cierto que Telesforo es ambicioso que [no] que Cornificio es ambicioso’ se dirá ‘Telesforo es menos ambicioso que Cornificio’ (y nótese —dicho sea entre paréntesis— el carácter expletivo de ese ‘no’ que hemos encerrado, por ser opcional, entre corchetes); y, en vez de decirse, ‘El hecho de que Edelmira haya tenido suerte implica el hecho de que Rita haya tenido suerte’ —o su sinónimo ‘Edelmira ha tenido suerte a lo sumo en la medida en que Rita ha tenido suerte’— se dirá, más simplemente, ‘Edelmira ha tenido a lo sumo tanta suerte como Rita’. Y, en vez de decirse ‘El hecho de que Úrsula se lleve mal con su marido es cierto en la misma medida en que lo es el hecho de que Engracia se lleva mal con su marido’ se dirá ‘Úrsula se lleva tan mal con su marido como Engracia’.

Aunque tales ilustraciones pueden parecer obvias, y ser tildadas de perogrulladas, advierta el lector, sin embargo, que en verdad distan de serlo. Rechazan en efecto tales equivalencias la mayoría de los tratadistas de temas afines a la relación entre la lengua natural y diversos sistemas formales. Con respecto a la primera equivalencia, aducen que los grados de posesión de una propiedad no tienen por qué repercutir en sendos grados de verdad de hechos consistentes, cada uno de ellos, en la posesión de la propiedad en cuestión por cierto ente. Conque —concluyen— el que Evagro sea atlético en tal o cual medida no entraña que sea verdad en esa medida que Evagro es atlético. Según ellos o bien la verdad misma no tiene grados [en absoluto] o, si sí los tiene, no son los mismos que los de posesión de una propiedad por un ente; p.ej. —alegan algunos de esos autores—, aunque Evagro sea bastante atlético puede que no sea verdad en absoluto que Evagro es atlético, ya que para que sea verdad, lo que se dice verdad, que Evagro es atlético será menester que Evagro tenga la propiedad de ser atlético en un grado altísimo —normalmente esos autores exigen un grado total, del 100%.

El error en tales alegaciones estriba en confundir **ser verdad** con **ser completamente verdad** o con **ser verdad en un grado altísimo**. Una cosa es que, por razones pragmáticas, sea impropio o inoportuno en los más contextos usuales el proferir un mensaje como ‘Evagro es fuerte’ si lo proferido no alcanza cierto umbral de verdad —que puede que sea (mas no forzosamente siempre) el del 50%, p.ej.—; otra cosa es que sólo sean verdaderas —a secas— las aseveraciones que alcancen ese umbral de verdad (como si sólo fueran blandas las cosas que tengan al menos tal grado de blandura).

Sin embargo, aunque sean erradas esas alegaciones contra una equivalencia que a primera vista parece tan obvia —y que, a juicio de quien esto escribe, es verdadera—, la mera existencia de tales alegaciones revela que el asunto es más intrincado de lo que pudiera parecer. Sin duda quienes enuncian esas alegaciones son llevados a esa postura por un prejuicio filosófico muy en el espíritu del clasicismo: la verdad estaría exenta de grados o al menos no sufriría tantos grados, tantas gradaciones y fluctuaciones como las posesiones de unas u otras propiedades por unos u otros entes. Respetable como es esa idea, no es, tampoco ella, algo tan sin

vuelta de hoja como se lo figuran quienes la aducen como si fuera un argumento incontrovertible. En cualquier caso, queda en pie la observación general que vienen a ilustrar estas consideraciones: que, siendo plausibles —como lo son— esas equivalencias que he propuesto más arriba, no son indubitables ni incuestionables.

Similarmente —ya para concluir este inciso— cabe mencionar que muchos estudiosos rechazan la equivalencia entre ‘Es menos cierto [e.d. verdadero] que Norberto es joven que [no] que Pepe es joven’ y ‘Norberto es menos joven que Pepe’. Aducen argumentos similares: que no hay grados de verdad o que no son tantos cuantos sean los grados de juventud; que, p.ej., para que sea verdadera una aseveración del enunciado ‘Pepe es joven’ es necesario y suficiente que Pepe sea joven en, al menos, tal medida, al paso que, para que sea más joven que Norberto, basta con que posea en mayor medida esta propiedad, la juventud. (Eso cuando se admiten grados de posesión de propiedades, cosa que muchos otros tratadistas no hacen; para ellos ‘Norberto es menos joven que Pepe’ es un enunciado verdadero, no por grado de posesión de una propiedad, la juventud, por Pepe que sea mayor que el grado en que la posea Norberto, sino por otras razones: p.ej. porque en cierto contexto pertinente sea verdad ‘Norberto no es joven pero Pepe sí’, sin que en eso entren grados para nada.) A mi juicio están muy equivocados quienes ven así las cosas. Sufren los efectos de un desconocimiento de los grados de verdad, y eso los lleva a rehuir la manera más simple, clara y, a primera vista, plausible de tratar los comparativos (que es la articulable en torno a la equivalencia aquí propuesta). Pero el mero hecho de que esos estudiosos —que son la gran mayoría— no vean las cosas igual que el autor de este trabajo revela cuán complejo es el problema de las relaciones entre la lengua natural y unos u otros sistemas formales.

Capítulo 5

Ventajas de la lógica infinivalente como lógica de lo difuso

Desde 1965, aproximadamente, se han llevado a cabo importantes investigaciones lógico-matemáticas acerca de lo difuso, y se han elaborado, en particular, las teorías de conjuntos difusos (una de las cuales es el sistema *Abj*).

Esas teorías han articulado sistemática y rigurosamente la intuición de que existe lo difuso en cuanto al grado de pertenencia a determinadas clases —o sea, de poseer determinadas propiedades—; y de que, con respecto a ciertos cúmulos, los grados de pertenencia son infinitos, e incluso innumerables. Así, p.ej., sea un punto en una línea, *x*; los grados de verdad de '*z* está próximo a *x*' —donde la variable '*z*' tiene como campo de variación el cúmulo de los puntos de la línea— serán infinitos, conformando un **continuo** en ambas direcciones. Asimismo, parece que hay grados infinitos en '*x* es rico', '*x* es culto', '*x* es malo', '*x* es ruidoso', '*x* es grande', '*x* es blando', '*x* está caliente', etc., etc.

Un tratamiento adecuado de las construcciones comparativas en lengua natural tampoco parece posible sin admitir grados infinitos de verdad. Parece grotesco que haya un número **finito**, *n*, de posibilidades, tal que, para cualquier propiedad —p.ej. la belleza—, si hay una secuencia de *n* entes tales que cada uno de ellos es bello en medida mayor que el anterior, entonces el último es bello en un grado de 100%, o sea: en un grado irrebable. Se puede conjeturar que tal situación es absolutamente falsa, y que hay **más de** *n* entes cada uno de los cuales es más bello que el anterior; y ello para cualquier *n* finito. (Ello no excluye que pueda haber algún ente 100% bello; sólo excluye que sea obligatorio que, comenzando por un grado de belleza, y subiendo, peldaño por peldaño, *n* grados, se alcance ese grado absoluto del 100%.)

Sin entrar aquí —porque no es éste el lugar apropiado para ello— en detenidas consideraciones filosóficas al respecto, sí se impone una aclaración, siquiera mínima, de lo que está involucrado en el tratamiento lógico-matemático de lo difuso.

Hay que establecer una diferencia rigurosa entre el que sea difuso un conjunto (o cúmulo de cosas) y el que resulte indecible o incierta (insegura) la pertenencia a él de tal o cual cosa, o de muchas. Lo que hace difuso un cúmulo como el de los móviles rápidos no es que no sepamos qué decir en ciertos casos, si algo es rápido o no, sino el hecho de que muchas cosas poseen rapidez sólo hasta cierto punto, en vez de o bien carecer por completo de tal propiedad o bien poseerla plenamente. (Quizá nada es 100% rápido, ni siquiera la luz.)

No sabemos qué decir en muchos casos sin que el cúmulo de que se trate en tales casos haya de ser, por eso, difuso. No sabemos qué decir sobre si Stalin pensó o no en la guerra de Troya durante la última semana de su vida. Eso no hace difusa a la propiedad de pensar en dicha guerra durante la última semana de la vida de uno. No sabemos qué decir, si que todas las migraciones que llegaron a América antes de Colón cruzaron por el estrecho de Behring o que no fue así. No es eso lo que hace difusa a la propiedad de cruzar por el estrecho de Behring.

No sabemos tampoco qué decir, si que todo número perfecto es par o non. No sabemos si hay algún número perfecto non (aunque probablemente no los hay, ningún matemático —al parecer— tiene certeza al respecto). No por ello va a ser difuso el cúmulo de los números perfectos.

En cambio, en casos donde sí sabemos qué decir nos las habemos con propiedades difusas porque las cosas de que se trata las poseen —según quedó apuntado más arriba— en un grado no total. Sabemos que Alto Volta es un país árido; lo es hasta cierto punto y, por ende, lo es (**regla de apencamiento**). Pero, como no lo es totalmente (es menos árido que Mauritania,

p.ej., aunque tampoco Mauritania es totalmente árida), tenemos un caso típico de una verdad parcial o gradual no más. La aridez, pues, es una propiedad difusa.

Y no es que sean difusas aquellas propiedades en las que surge indecisión, o incertidumbre, o perplejidad sin que se sufra ninguna falta de información. Aparte de que no está tan clara esa noción de información, en el caso del ejemplo matemático recién aducido no es seguramente información lo que falta, sino capacidad de cálculo suficiente, o posesión de algún algoritmo apropiado, o superación de las limitaciones de nuestros sistemas formales, o algo así. Además, en muchísimos casos de propiedades difusas no se produce perplejidad ni incertidumbre (pertinente para el caso) ni indecisión (que no sea debida a falta de información). Casi todas las propiedades que atribuimos a las cosas en el quehacer científico y en la conversación cotidiana son, en efecto, difusas, sin que ello signifique que estamos constantemente sumidos en perplejidades, en indecisiones sobre qué decir. Cuando lo estamos suele ser más bien por falta de información. No nos quedamos suspensos, mudos ni perplejos al aplicar nociones geográficas (difusas, si las hay), como las de ser (una región) litoral, cálida, húmeda, fértil, elevada, montuosa, accidentada, rica en recursos mineros, y así sucesivamente. Lo que pasa es que ¡hay tantos, tantos grados en todo eso!

Otra confusión frecuente es la que se da entre lo difuso y lo probable. Mucha gente dice cosas como que está de más una teoría de conjuntos difusos, o una lógica de lo difuso, porque ya está ahí la teoría de probabilidades. Pero es eso desacertado. El cálculo de probabilidades tiene características muy diversas de las de una lógica de lo difuso. Y, sobre todo, son diferentes las nociones que maneja. Hay varias concepciones de la probabilidad, pero, sea cual fuere la que uno adopte, es seguro que nada tiene que ver —en general— el grado de probabilidad de una conjetura con el grado de verdad de aquello sobre lo que versa tal conjetura, o tal afirmación. Una cosa es qué tan probable sea o deje de ser [la suposición de] que Teng Xiaoping vaya a hacerse budista antes de morir; otra es cuán verdadero o real sea su hacerse budista (si es que hay grados de conversión a una religión —como seguramente los hay—, según cuánta confianza tenga uno en sus doctrinas y tradiciones, cuán motivado y cuán emocionalmente afectado esté por ellas etc.). Puede que eso sea poquísimo probable pero que, si se produce, sea un hecho muy real (o sea: muy verdadero).

Una última confusión que conviene también disipar para que resalte la significación de los tratamientos lógicos de lo difuso es la que consiste en confinar los discursos en que figuren términos susceptibles de un tratamiento difuso a un habla coloquial, no teórica y, sobre todo, no científica. Algunos aducen que “la ciencia” no usa términos así y que, por ende, una lógica para uso científico no tiene necesidad ninguna de ocuparse de lo difuso. Frente a tal alegato hay muchas consideraciones pertinentes, pero limitémonos a cinco de ellas, bastante escuetas.

En primer lugar, ese paradigma científico que proscribía términos que denoten propiedades difusas es un paradigma superado desde los últimos años 60. La enorme fertilidad de las teorías de conjuntos difusos se ha demostrado al revelarse capaces de nuevos y sugerentes planteamientos en muy diversas disciplinas científicas, desde la medicina hasta la geografía, pasando por la lingüística, la economía y otras ciencias sociales, sin descuidar diversos campos matemáticos.

En segundo lugar, el futuro dirá cuánto puedan aprovecharse de las lógicas de lo difuso también otras ciencias, incluida la física —que se halla hoy en estado de grave crisis y mutación. De hecho, dadas una serie de paradojas tales que los físicos no saben bien cómo articular sus teorías para librarse de ellas, no parece excesivamente arriesgado conjeturar que las lógicas de lo difuso puedan contribuir a encontrar soluciones. Los descubrimientos lógicos llevan en eso tal vez la delantera —igual que las geometrías no euclídeas de Riemann y Lobachevski llevaron la delantera en el siglo pasado respecto a sus aplicaciones físicas, que se perfilaron

a raíz de la invención por Einstein de la teoría de la relatividad. (¿Hubiera sido correcto desechar aquellas geometrías porque todavía no se había encontrado para ellas una aplicación física?)

En tercer lugar, cualesquiera que sean o dejen de ser las aplicaciones de las nuevas lógicas en campos de investigación pertenecientes a ciencias humanas o ciencias de la naturaleza —según suelen denominarse—, tienen importantes motivaciones y aplicaciones filosóficas (en teoría del conocimiento, en metafísica, en filosofía de la naturaleza, de la mente, del lenguaje, en ética, en filosofía de la historia, en filosofía política; varias de tales motivaciones y aplicaciones han sido desarrolladas por el autor de estas líneas en diversos trabajos, alguno de los cuales viene mencionado en la bibliografía al final de este opúsculo).

En cuarto lugar, hay un fuerte argumento transcendental en contra de que se dé un corte o salto entre el pensamiento humano precientífico y el científico: si el primero está infestado por masivo error —o, peor, por el recurso a algo que sería ilógico, como lo sería el empleo de nociones difusas si “la lógica” no se ocupara más que de nociones exentas por completo de rasgos difusos—, entonces lo más probable sería que, habiendo salido de ese pensamiento precientífico y recibiendo de él sus primeras nociones, sus problemas y sus métodos de partida, el pensamiento científico carezca, también él, de correspondencia con la realidad.

En quinto y último lugar, el tenor mismo de muchas teorías científicas hace dudoso que puedan interpretarse de manera realista más que concibiendo como difusas a muchas de las nociones en ellas involucradas. Así p.ej. suscítanse a menudo discusiones sobre los deslizamientos sin fricción, los gases perfectos, y otros entes así, que parecen entelequias, o postulaciones ideales. Interpretadas gradualísticamente, esas nociones permiten que la ciencia verse sobre la realidad —y no sobre un cielo ideal divorciado del mundo real—, con lo que se explica tanto el tránsito del pensamiento precientífico al científico como, a la inversa, las aplicaciones de éste. Podemos concebir que la física no se ocupa de deslizamientos que sean **enteramente** sin fricción, sino de aquello de lo que dice ocuparse, de deslizamientos sin fricción (tales, pues, que, en uno u otro grado, no tienen fricción): en la medida en que un deslizamiento es sin fricción, tiene tales o cuales características: tal sería el sentido de los enunciados científicos en cuestión. (Esta quinta consideración vendrá más ampliamente expuesta en el Anejo Nº 1 del presente opúsculo.)

Con el maximalismo clásico del **todo o nada** están excluidas esas vías de interpretación, que han empezado recientemente a posibilitar las lógicas de lo difuso.

Matices

En las lógicas finivalentes que hemos visto, hemos podido introducir un cierto número de funtores diádicos y monádicos. Pero nos hemos encontrado con serias dificultades a la hora de ampliar y precisar los **matices** de la aserción; no parece que haya en las lógicas finivalentes posibilidad de expresar adecuadamente ‘Es un sí es no cierto que’, p.ej. En cuanto a ‘Es muy cierto que’ (el functor ‘X’), su tratamiento en las lógicas finivalentes era tosco, ya que teníamos, p.ej., en la lógica pentavalente lo siguiente:

p	XP
4	4
3	2
2	1
1	1
0	0

Parece obvio que, si $/p/$ es el valor máximo, $/Xp/$ debe también ser ese valor; y que si $/p/$ es el mínimo, $/Xp/$ ha de ser el mínimo; y que, en los casos intermedios, $/Xp/$ debe ser inferior a $/p/$ (salvo cuando $/p/$ fuera sólo infinitesimalmente verdadero o sólo infinitesimalmente falso; pero —por ahora— tales matices son inexpresables, pues estamos en lógicas finivalentes).

Lo que parece inaceptable es que a dos valores diversos de $\lceil p \rceil$ corresponda un sólo y único valor de $\lceil Xp \rceil$.

Un tratamiento adecuado de un functor como 'X' (=‘Es muy cierto que’) lo encontramos sólo en una lógica infinivalente.

Reales e hiperreales

La mayor parte de las lógicas infinivalentes propuestas hasta ahora se construían asociando a cada valor de verdad un número real en el intervalo $[0, 1]$. Las propuestas durante los últimos años por el autor de estas páginas añaden una nueva complicación: en vez de asociar tan sólo a los valores de verdad los números reales de ese intervalo, se asociarán a ellos los números **hiperreales** (siempre dentro de tal intervalo, eso sí). Se entiende por **hiperreales** lo siguiente:

Un hiperreal es el resultado de dejar tal cual a un número real, o de aumentarlo o disminuirlo **infinitesimalmente**. Quiere ello decir que, en el intervalo abierto $]0, 1[$ —o sea: excluidos de él 0 y 1—, a cada número real le corresponderán tres hiperreales: 1º) el número real mismo; 2º) el resultado de darle un incremento infinitesimal; 3º) el resultado de disminuirlo infinitesimalmente.

En el intervalo **cerrado** $[0, 1]$ habrá, además, los cuatro siguientes hiperreales: 0; el resultado de aumentar infinitesimalmente a 0; el resultado de disminuir infinitesimalmente a 1; y 1 mismo.

Este sistema nos permite expresar matices ricos y variados, y dar un tratamiento adecuado a los comparativos.

Pero, antes de exponer un sistema de lógica infinivalente con hiperreales, veamos otro, más simple, sólo con reales. A este último lo llamaremos *Ar*.

El sistema *Ar*

Tomamos como valores de verdad todos los números reales en el intervalo $[0, 1]$. Todos los valores excepto 0 son designados; y todos, excepto 1, antidesignados. Introducimos los funtores primitivos siguientes: 'N', '¬', '∧', 'I', '•', con las lecturas siguientes:

$\lceil Np \rceil$: «No es cierto que p »

$\lceil \neg p \rceil$: «Es enteramente falso que p »

$\lceil p \wedge q \rceil$: « p y q »

$\lceil p \bullet q \rceil$: «No sólo p , sino que también q »

$\lceil plq \rceil$: «que p es cierto en la misma medida en que lo es que q »

Asignaciones de valores de verdad:

$/Np/ = 2^{\log_x 2}$, si $x = /p/$		$/p \wedge q/ = \min(/p/, /q/)$	$/plq/ =$	$1/2$, si $/p/ = /q/$
$/\neg p/ =$	0, si $/p/ \neq 0$	$/p \bullet q/ = /p \times q/$		0, en caso contrario
	1, en caso contrario			

En este sistema, *Ar*, podemos definir los siguientes funtores:

$\lceil p \vee q \rceil$ abr $\lceil N(Np \wedge Nq) \rceil$	$\lceil Hp \rceil$ abr $\lceil \neg Np \rceil$	$\lceil Xp \rceil$ abr $\lceil p \bullet p \rceil$
$\lceil Sp \rceil$ abr $\lceil p \wedge Np \rceil$	$\lceil Lp \rceil$ abr $\lceil \neg \neg p \rceil$	$\lceil p \supset q \rceil$ abr $\lceil \neg p \vee q \rceil$
$\lceil p \rightarrow q \rceil$ abr $\lceil p \wedge qlp \rceil$	$\lceil p \setminus q \rceil$ abr $\lceil p \rightarrow q \wedge \neg(q \rightarrow p) \rceil$	$\lceil p \equiv q \rceil$ abr $\lceil p \supset q \wedge q \supset p \rceil$

Se puede comprobar, ahora, que las siguientes son algunas tautologías de este sistema:

$\lceil p \setminus p \setminus q \setminus p \rceil$	$\lceil p \wedge q \bullet r \setminus r \bullet p \wedge r \bullet q \rceil$	$\lceil Xp \rightarrow Xql. p \rightarrow q \rceil$
$\lceil p \supset Xp \rceil$	$\lceil Xp \setminus p \vee Hp \vee \neg p \rceil$	$\lceil Xp \rightarrow p \rceil$
$\lceil p \bullet q \bullet r \setminus p \bullet q \bullet r \rceil$	$\lceil p \bullet ql(p \bullet r) \wedge p \supset q \setminus qlr \rceil$	

Y son también tautologías de Ar todas las tautologías de A_3 que expresamos en la lista expuesta en un capítulo anterior de este trabajo.

Pero hay también tautologías de A_3 —que no figuraban en la susodicha lista— que no son tautologías de Ar . P.ej.: $p \setminus p \setminus p \vee H \vee \neg p$ $p \setminus q \wedge (q \setminus r) \rightarrow Hr$ $p \setminus p \setminus p \setminus Sp$ $p \setminus q \wedge (q \setminus r) \rightarrow \neg p$

Igualmente, en A_5 son tautologías las siguientes fórmulas que **no** son tautologías de Ar : $p \setminus q \wedge (q \setminus r \wedge r \setminus s \wedge s \setminus p) \rightarrow Hp$ $p \setminus q \wedge (q \setminus r \wedge r \setminus s \wedge s \setminus p) \rightarrow \neg p$

Antes de pasar al punto siguiente, procede hacer una consideración sobre la asignación de valores de verdad con respecto a la negación simple, 'N'. ¿Por qué no escoger otra asignación más sencilla, a saber: $/Np/ = 1 - /p/$? En ambos casos se llega a resultados parecidos. Si $/p/ = 1/2$, $/Np/ = 1/2$, cualquiera que sea la asignación escogida de entre esas dos. Similarmente, y también para cualquiera de esas dos asignaciones, se tendrá que $/p/ \wedge /Np/$ ssi $/p/ \wedge 1/2$. Además, muchos resultados son iguales en ambos casos. P.ej. la opción por una de esas dos alternativas, en vez de la otra, no afecta al carácter tautológico de esquemas como $\lceil NNp \setminus p \rceil$, $\lceil p \wedge ql N(Np \vee Nq) \rceil$, etc.

Hay, sin embargo, ciertos esquemas que sólo son tautológicos si a 'N' se le da la asignación que hemos escogido, en vez de la más sencilla. P.ej., defínase $\lceil Kp \rceil$ como $\lceil NXNp \rceil$. Entonces sólo la asignación escogida, y no la otra, hace que sea tautológico el esquema: $\lceil XKp \setminus p \rceil$.

Ahora bien, 'K' puede leerse como 'Es (al menos) un poco cierto [=verdadero] que', al paso 'X' se leería 'Es muy cierto [=verdadero] que'. Entonces $\lceil XKp \rceil$ significa que es al menos un poco verdadero el ser muy cierto que p; y eso normalmente se entendería como equivaliendo a enunciar que p, sin más. Si $/p/ = 1/2$, p.ej., tendremos: (1) escogiendo la asignación por la que hemos optado, $/Kp/ = (1/2)^{1/2}$, o sea la raíz cuadrada de $1/2$ (e.d. 0'70711); (2) con la otra asignación más sencilla, $/Kp/ = 3/4$. Luego con la segunda asignación $/XKp/ \neq /p/$, mientras que con la primera evidentemente $/XKp/ = /p/ = /KXp/$.

El sistema Ap

Es mucho más complicado exponer las asignaciones de valores de verdad en Ap que en Ar . Los signos primitivos de Ap son los mismos que los de Ar , pero añadiendo una **constante sentencial**, a saber 'a', que se lee: 'lo infinitesimalmente real' o 'lo infinitesimalmente verdadero'.

En Ap se consideran valores de verdad todos los hiperreales en el intervalo $[0, 1]$. Llamemos números aléticos a todos esos valores de verdad o hiperreales. Dicho de otro modo, el cúmulo

de los números aléticos es engendrado, a partir del intervalo cerrado de los números reales $[0, 1]$ por los operadores monádicos m y n , definidos por los postulados siguientes:

Cualquier número real en el intervalo $[0, 1]$ es un número alético;

Si u es un número real en el intervalo $[0, 1[$, entonces mu es un número alético, y $mu \neq u$;

Si u es un número real en el intervalo $]0, 1]$, entonces nu es un número alético, y $nu \neq u$;

$mmu = u$ (para todo u); $nnu = nu$ (para todo u);

Si $u \neq 0$, entonces $nmu = nu$; Si $u \neq 1$, entonces $mnu = mu$;

$n0 = 0$; $m1 = 1$; $nm0 = m0$; $mn1 = n1$

Definamos ahora una relación de orden conexo (o sea: de orden lineal, o total) sobre el cúmulo de los números aléticos, a saber: $\leq$, mediante los postulados siguientes:

$u \leq mu$	O bien $u \leq u'$, o bien $u' \leq nu$, o bien $u = mu'$
$nu \leq u$	$mu \leq u'$, si u y u' son, ambos, números reales, y u es inferior o igual a u'
$u \leq u$	$u \leq u^1$ y $u^1 \leq u^2$ sólo si $u \leq u^2$

Con tal de que cumplan esos postulados, esos entes, nu y mu —donde u es un número real en el intervalo correspondiente— pueden ser lo que sean. P.ej., pueden ser así: para cada número real r tal que $r \leq 1$, $nr = \{r, 0\}$ y $mr = \{r, 1\}$. O sea, mr y nr serán sendos dúos de números; con tal —eso sí— de que entienda que $\{r, 0\} \leq r \leq \{r, 1\}$, y que, si r' es otro real tal que $r' < r$, entonces $mr' < nr$. Este es un uso extendido de ' $\leq$ ' porque tal como normalmente se entiende esta relación de orden no se afirmaría $r \leq r'$ más que si tanto r como r' son números reales.

Otro modo alternativo de entender a esos nuevos números —los hiperreales— es utilizando el análisis no estándar de Robinson (cada vez más ampliamente estudiado, aplicado y profesado por un gran número de matemáticos, dadas sus cualidades de fecundidad, elegancia y claridad en la fundamentación de los cálculos infinitesimal e integral); y, dentro de eso, hay varias maneras de proceder. Una sería ésta: el análisis no estándar acepta que, si r es un número real, hay una infinidad de hiperreales mayores que r que están infinitamente próximos a r y tales que no sólo para cada uno de ellos, h , hay otro, h' , más próximo a r , sino que para esos dos hiperreales, h, h' , hay otro entre ellos; similarmente hay una infinidad de reales menores que r pero por lo demás con las mismas características recién apuntadas. Tomemos para cada real r la clase de los hiperreales mayores que r pero infinitamente próximos a r , y llamémosla mr ; y similarmente nr será la clase de hiperreales menores que r pero infinitamente próximos a r . Extendemos la relación usual $\leq$ de esta manera: (1) $nr \leq r \leq mr$; (2) si $r^1 < r^2$, ambos reales, entonces $mr^1 < nr^2$. (Otra manera de proceder vendrá esbozada en el penúltimo acápite del Capítulo 6.)

Introduzcamos ahora dos operadores monádicos y tres operadores diádicos definidos sobre A , es decir: sobre **el cúmulo de los números aléticos**.

$\#u =$	$2^{\log u^2}$, si u es un número real
	$n\#u^1$, si u^1 es un número real tal que $mu^1=u$
	$m\#u^1$, si u^1 es un número real tal que $nu^1=u$
$@u =$	0, si $u \neq 0$
	1, si $u=0$
$\text{maxim}(u^1, u^2) =$	u^1 , si $u^2 \leq u^1$
	u^2 en caso contrario
$\text{minim}(u^1, u^2) =$	u^1 , si $u^1 \leq u^2$
	u^2 , en caso contrario
$u \times u' =$	el producto aritmético de u con u' , si u y u' son ambos reales;
	0, si o bien $u = 0$, o bien $u' = 0$;
	$m0$, si uno de los dos números u y u' es $m0$ y el otro es $\neq 0$;
	u' , si $u = 1$;
	u , si $u' = 1$;
	nu' , si $u = n1$;
	nu , si $u' = n1$;
	$n(u^1 \times u^2)$, si u^1, u^2 son números reales ambos $\neq 0$ y tales que tiene lugar una de las cinco situaciones siguientes: (a) $u^1 = u$ y $nu^2 = u'$ (b) $nu^1 = u$ y $nu^2 = u'$ (c) $nu^1 = u$ y $u^2 = u'$ (d) $nu^1 = u$ y $mu^2 = u'$ (e) $mu^1 = u$ y $nu^2 = u'$
	$m(u^1 \times u^2)$, si y^1 y u^2 son, ambos, números reales $\neq 0$ y $\neq 1$ y tales que tiene lugar una de las situaciones siguientes: (a) $mu^1 = u$ y $u^2 = u'$ (b) $u^1 = u$ y $mu^2 = u'$ (c) $mu^1 = u$ y $mu^2 = u'$

Puede comprobarse fácilmente la siguiente ley: para todo número alético u hay un número real u' , y sólo uno, tal que, o bien $u=u'$, o bien $u=nu'$, o bien $u=mu'$. Además, si U es un cúmulo cualquiera de números aléticos (o sea: un subconjunto cualquiera de A), entonces tanto $\inf(U)$ como $\sup(U)$ son, ambos, definidos y únicos. (El elemento ínfimo sobre un subconjunto S de un cúmulo C ordenado por una relación de orden, $\leq$, es el mayor miembro de C de entre los

minorantes de S ; y el elemento supremo sobre S es el menor miembro de C de entre los mayorantes de S . Serán **mayorantes** [respectivamente, **minorantes**] de S cuantos miembros de C , x , sean tales que cualquier elemento $z \in S$ es tal que $z \leq x$ [respectivamente, es tal que $x \leq z$]. Esas nociones conjuntuales pueden aclararse consultando cualquier manual escolar de teoría ingenua de conjuntos.)

Un número alético es **designado** ssi es diferente de 0. Un número alético es **antidesignado** ssi es diferente de 1.

Valuaciones de Ap

Una valuación de Ap , v , será una función que tenga como su dominio (o campo de argumentos) al cúmulo de fórmulas sintácticamente bien formadas (fbfs) de Ap y como campo de valores un subconjunto de A (o sea: un cúmulo de números aléticos) —siendo, por consiguiente, A su **contradominio** [para mayores precisiones terminológicas al respecto vide infra, el apartado «Modelo, interpretación, valuación, validez» del capítulo 6]—, con tal de que, para cualesquiera $\ulcorner p \urcorner$ y $\ulcorner q \urcorner$, cumpla las condiciones siguientes:

$$v(p \wedge q) = \text{minim}(v(p), v(q))$$

$$v(p \vee q) = \text{maxim}(v(p), v(q))$$

$$v(p \bullet q) = v(p) \times v(q)$$

$$v(Np) = \#v(p)$$

$$v(\neg p) = @v(p)$$

$$v(a) = m0$$

$$v(plq) =: \frac{1}{2} \text{ si } v(p) = v(q); \text{ y } 0 \text{ en caso contrario}$$

Como se ve, el único signo **primitivo** que hay en Ap y que no existía en Ar es la constante sentencial 'a' que puede leerse como 'lo infinitesimalmente real', o 'lo infinitesimalmente verdadero' o 'lo un sí es no es real' o 'lo un sí es no verdadero'. Esa importantísima constante sentencial nos permite definir, gracias a ella, numerosos funtores monádicos y diádicos, de una significación transcendental y decisiva. Su gigantesca importancia, sin embargo, sólo puede apreciarse al abordarse el cálculo cuantificacional y la teoría de conjuntos, cosa que se llevará en cabo en los capítulos finales del presente opúsculo. Pero digamos, ya desde ahora, lo siguiente: un matiz tan importante, en nuestro hablar cotidiano, como 'un sí es no' (expresado, p.ej., en la oración 'Zósimo es un sí es no distraído', o su equivalente 'Es un sí es no cierto que Zósimo es distraído') no puede formalizarse en Ar , pero sí puede formalizarse en Ap . Escribimos 'Es un sí es no cierto que' como ' Y ', definiendo así tal functor: ' $\ulcorner Yp \urcorner$ abr ' $\ulcorner pla \wedge p \urcorner$ '

Similarmente, podemos expresar en Ap el matiz afirmativo 'Es un tanto cierto que', entendiendo por tal lo siguiente: 'Es cierto que... y es del todo falso que sea un sí es no cierto que...'. Representaremos 'Es un tanto cierto que' como ' f ', definiéndolo así: ' $\ulcorner fp \urcorner$ abr ' $\ulcorner \neg Yp \wedge p \urcorner$ '

Relaciones entre Ap y Ar

La relación entre Ap y Ar es la siguiente: **no toda** tautología de Ar es una tautología de Ap . P.ej. ' $\ulcorner p \bullet q \urcorner (\ulcorner p \bullet r \urcorner \wedge p \supset q \urcorner)$ ' —lo que podríamos llamar el principio de cancelación— no es una tautología de Ap ; porque sea ' $\ulcorner p \urcorner = 'a'$ '; entonces la fórmula será: ' $\ulcorner a \bullet q \urcorner (\ulcorner a \bullet r \urcorner \wedge a \supset q \urcorner)$ '. Pero sea $/q/ = \frac{1}{2}$; y sea $/r/ = \frac{1}{3}$; entonces tendremos, en virtud de las asignaciones estipuladas, que la prótasis de esa fórmula tendrá como valor $m0$, o sea lo infinitesimalmente verdadero, que es un valor **designado**; la apódosis tendrá como valor 0, que no es designado. Pero si $/s/ = m0$ y $/s'/ = 0$, entonces $/s \supset s'/ = 0$. (En este caso ' $\ulcorner s \urcorner$ sería ' $\ulcorner a \bullet q \urcorner (\ulcorner a \bullet r \urcorner \wedge a \urcorner$ ' y ' $\ulcorner s' \urcorner$ sería ' $\ulcorner q \urcorner$ '). Luego la fórmula total tendrá, en este caso, como valor de verdad 0. Luego no es una tautología. Y, por tanto, tampoco lo es el esquema del que, por instanciación, se ha sacado tal fórmula (a saber: ' $\ulcorner p \bullet q \urcorner (\ulcorner p \bullet r \urcorner \wedge p \supset q \urcorner)$ ').

La lógica *Abp*

La lógica *Abp* es una lógica tensorial con un número infinito de componentes construida a partir de *Ap* como lógica de base. Dicho de otro modo: cada valor de verdad de *Abp* (o **tensor alético** de *Abp*) es una secuencia infinita de números aléticos. Se puede considerar como tensores aléticos designados de *Abp* aquellos que sólo están conformados por números aléticos designados de *Ap*.

Sea v una valuación de *Ap*. Para cada fórmula de *Ap*, $\lceil p \rceil$, $v(p)$ será un número alético. Ahora definiremos una valuación μ , de *Abp* como sigue.

Para cada tensor alético, s , s_i será el i^o componente de s ; o sea s será $\langle s_1, s_2, s_3, \dots \rangle$. Si $\blacktriangle$ es un functor monádico, $v(\blacktriangle p)$ será también un número alético, sea o no el mismo que $v(p)$.

Definimos μ así: para cada fórmula $\lceil p \rceil$ de *Ap*, para cada functor monádico $\blacktriangle$ y para cada valuación v de *Ap*, $\mu(\blacktriangle p)_i = v(\blacktriangle p)$ ssi $\mu(p)_i = v(p)$.

Similarmente, para cada functor diádico $\perp$ y cualesquiera fórmulas $\lceil p \rceil$, $\lceil q \rceil$ de *Ap*, $\mu(p \perp q)_i = v(p \perp q)$ (siendo v una valuación cualquiera de *Ap*) ssi $\mu(p)_i = v(p)$ y $\mu(q)_i = v(q)$.

Además de contener todos los signos de *Ap*, *Abp* va a tener un functor más, monádico, a saber 'B', definido así: para cada valuación de *Abp*, μ , y cada fórmula de *Abp*, $\lceil p \rceil$, $\mu(Bp)_i = \mu(p)_i$ si $\mu(p)$ es una secuencia en la que no hay ningún 0; en caso contrario, $\mu(p)_i = 0$ para cada i .

El sentido que cabe dar a 'B' es algo así como: 'En todos los aspectos', 'Desde todos los puntos de vista [objetivamente fundados]'. Si —según parece plausible— es afirmable con verdad sólo todo lo que es verdadero en todos los aspectos, entonces alternativamente cabe leer a 'B' como 'Es afirmable con verdad que'.

En virtud de las relaciones ya anteriormente indicadas entre las lógicas tensoriales (o lógicas-producto) y sus **respectivas** lógicas escalares de base, los teoremas de *Abp* escritos sólo con los signos de *Ap* son exactamente los mismos que los teoremas de *Ap*.

El gran interés de una lógica tensorial como *Abp* estriba en que con ayuda de la misma cabe ver el mundo de manera más compleja y sutil. En vez de pensarse que un enunciado ha de tener un solo y único grado de verdad, o bien una ausencia absoluta de cualquier grado de verdad, cabe ahora pensar que sea verdadero en unos aspectos, pero en otros no (incluso quizá no **en absoluto**), o más en unos que en otros. Tal vez un cuerpo esté formado por entes que son ondas y partículas a la vez, pero acaso sean en unos aspectos más ondas que partículas, siendo, en cambio, en otros aspectos más partículas que ondas.

Capítulo 6

Noción de teoría, clasificación sintáctica y semántica de las teorías

Reglas de inferencia

Empezaremos por definir la noción de regla de inferencia. Una regla de inferencia es una autorización para extraer, a partir de cierto tipo de premisas, conclusiones de un determinado tipo. Dicho de otro modo: una regla de inferencia es un procedimiento, estipulado como lícito, en virtud del cual, dadas determinadas premisas, cabe aseverar determinadas conclusiones.

Normalmente, para indicar cómo funciona una regla de inferencia se escriben —separadas entre sí por comas—, a la izquierda los esquemas de las premisas, a la derecha el esquema de la conclusión a sacar, y en el centro el signo sintáctico (o, si se quiere, metalingüístico) ' $\vdash$ '; o bien, se escriben las premisas encima y, debajo, separada de las premisas por una línea horizontal recta, la conclusión.

Ejemplos simples de reglas de inferencia son las siguientes:

$$p, q \vdash p \wedge q$$
$$p \vdash p \vee q$$
$$p, p \supset q \vdash q$$

(La primera se llama **regla de adjunción**; la segunda, **regla de adición**; la tercera, ***modus ponens***.

Es obvio que una regla de inferencia no es una fórmula; una regla de inferencia no tiene valor de verdad, por tanto. Como sí cabe considerar a una regla de inferencia es como una relación entre fórmulas —una relación de engendramiento si se quiere.

Cierre de un cúmulo con respecto a una relación

Veamos ahora, en segundo lugar, qué se significa al decir que un cúmulo está **cerrado** con respecto a una relación; quiere decirse con ello lo siguiente.

Concibamos de momento una relación R como un cúmulo o conjunto de pares cada uno de los cuales tiene como miembro izquierdo una secuencia ordenada de miembros de un cúmulo dado A —el **dominio** de la relación—, y como miembro derecho un objeto. [Para mayores precisiones terminológicas, vide infra, el apartado «Modelo, interpretación, valuación, validez» de este mismo capítulo.] Se dirá que A es cerrado con respecto a R ssi cada par ordenado perteneciente a R cuyo componente izquierdo conste sólo de miembros de A es tal que el objeto que es su componente derecho ha de ser también miembro de A . Si R es la relación de engendrar, y A es la especie humana, es obvio que cada par perteneciente a la relación tendrá como miembro izquierdo una pareja de individuos humanos de diverso sexo, y como miembro derecho un objeto que ha de ser, forzosamente, un individuo humano. Por tanto, A está cerrado con respecto a R .

Pues bien: concibamos cada **regla de inferencia** como una relación tal que cada uno de sus miembros es tal que su miembro izquierdo es una secuencia de fórmulas —llamadas premisas—, mientras que su miembro derecho es una fórmula. (Nótese que una secuencia de fórmulas puede tener un solo componente, o sea: estar conformada por una sola fórmula.) Por lo demás, postulamos que una secuencia de premisas es equivalente a cualquier permutación de la misma, o sea que carece de importancia el orden dentro de la secuencia. Lo cual significa que, si ' p ' es inferible de la secuencia $\langle \text{'q'}^1, \dots, \text{'q'}^n \rangle$, es también inferible de cualquier resultado de efectuar sucesivas permutaciones en esa secuencia. (Nótese, empero, que esta postulación, por razonable que sea o parezca, no es indiscutible, y de hecho hay sistemas lógicos en los cuales no rige sin restricciones.)

Se dirá, pues, que un cúmulo A de fórmulas **está cerrado con respecto** a una regla de inferencia ssi toda conclusión que pueda extraerse, mediante dicha regla de inferencia, a partir de premisas que sean fórmulas pertenecientes a A es también una fórmula perteneciente a A .

Reglas de formación

Se llama **regla de formación** (sintáctica) a una regla que permite engendrar inscripciones (o prolocuciones). Se llama **fbf** (**fórmula sintácticamente bien formada**) con respecto a un cúmulo de reglas de formación a cualquier fórmula que se ha engendrado sólo mediante aplicaciones de una o varias de esas reglas de formación.

Noción de teoría

Pues bien, tras esos preliminares, ¡definamos qué es una teoría! De manera informal —en seguida pasaremos a una definición más rigurosa— podemos ver una teoría como un cúmulo $\mathfrak{B}$ de fórmulas engendradas todas según determinadas reglas de formación, siempre y cuando $\mathfrak{B}$ esté cerrado con respecto a determinadas reglas de inferencia. Las fórmulas pertenecientes a la teoría se llaman sus **teoremas**; las reglas de inferencia con respecto a las cuales la teoría es cerrada son las reglas de inferencia de la teoría; y las reglas de formación con arreglo a las cuales las fórmulas de la teoría se han engendrado se llaman reglas de formación de la teoría. Por otro lado, se llaman **fbfs** de la teoría todas las fórmulas sintácticamente bien formadas con respecto al cúmulo de reglas de formación de la teoría.

Si una regla de formación de una teoría $\mathfrak{B}$ estipula que, en el caso de que $\ulcorner p \urcorner$ y $\ulcorner q \urcorner$ sean fbfs de $\mathfrak{B}$, $\ulcorner \Delta p \urcorner$ y $\ulcorner p \mathfrak{A} q \urcorner$ son también fbfs de $\mathfrak{B}$ (siendo $\ulcorner \Delta \urcorner$ y $\ulcorner \mathfrak{A} \urcorner$ signos cualesquiera, respectivamente monádico y diádico), entonces se dirá que $\mathfrak{B}$ contiene esos signos $\ulcorner \Delta \urcorner$ y $\ulcorner \mathfrak{A} \urcorner$, e.d. que éstos forman parte de su **vocabulario**.

Más rigurosamente expresado. Sea V un cúmulo de signos y $\mathfrak{R}$ un cúmulo de reglas de formación a partir de elementos de V . Entonces $\mathfrak{S}$ será un cúmulo de fbfs. Dicho de otro modo, $\mathfrak{S}$ será aquel subconjunto de rstras finitas de elementos de V (o sea de secuencias finitas cuyos componentes sean miembros de V) tal que cada miembro de $\mathfrak{S}$ esté constituido en virtud de una regla perteneciente a $\mathfrak{R}$. Y cada una de tales reglas puede representarse como un par ordenado $\langle A, \alpha \rangle$ donde A es una serie de rstras de miembros de $\mathfrak{S}$ y α es una rstra de miembros de $\mathfrak{S}$; con lo cual la regla es una relación entre una serie de rstras (que sean fórmulas) y una fórmula (tiene, pues, el sentido de que, cuando sean fórmulas todos los componentes de A , será también una fórmula la rstra α).

Así caracterizado un cúmulo de fórmulas (o sea fbfs) $\mathfrak{S}$ (a partir de un vocabulario, V , y un cúmulo de reglas de formación), una teoría será [representada como] un par ordenado $\langle \mathfrak{B}, \mathfrak{R} \rangle$, donde $\mathfrak{B}$ es un subconjunto de $\mathfrak{S}$ cerrado con respecto a cada una de las reglas de inferencia pertenecientes a $\mathfrak{R}$, siendo $\mathfrak{R}$ un cúmulo de reglas de inferencia, siempre que, además, se cumpla la condición de que $\mathfrak{B}$ esté cerrado con respecto a cada miembro de $\mathfrak{R}$.

Nótese que, una vez que hemos dado la definición más rigurosa, cabe apreciar que no basta que un cúmulo de fórmulas $\mathfrak{B}$ esté cerrado con respecto a cierta regla de inferencia, R , para que esa regla de inferencia R haya de pertenecer al cúmulo de reglas, $\mathfrak{R}$, tal que definamos cierta teoría como la teoría $\langle \mathfrak{B}, \mathfrak{R} \rangle$. Puede que $\mathfrak{R}$ no abarque a esa regla, R . Porque es posible que, si bien el añadir R a ese conjunto de reglas, $\mathfrak{R}$, no haría aumentar el cúmulo de teoremas de esa teoría (en este caso, $\mathfrak{B}$), así y todo ese añadir R provocaría cambios si se expandiera el cúmulo de axiomas (la noción de axioma se va a explicar en seguida); puede, por consiguiente, que, de incrementarse por postulación la clase de teoremas, el añadido de R causara un incremento más amplio que el que resultara sin ese añadido. En una situación así, una regla R que no venga abarcada por el conjunto $\mathfrak{R}$ de reglas de inferencia que sirvan

para definir a la teoría $\langle \mathfrak{J}, \mathfrak{R} \rangle$ pero que pueda agregarse a $\mathfrak{R}$ sin por ello expandir el cúmulo de teoremas, $\mathfrak{J}$, se llamará una **regla sistémica** (o también una regla **admisible**) en esa teoría.

Nótese lo siguiente: una teoría puede no coincidir (y normalmente no coincidirá) con el cúmulo de sus fbfs. Dicho de otro modo: si bien toda fórmula perteneciente a la teoría (o sea, todo teorema de la teoría) ha de ser una fbf de la teoría, en cambio no toda fbf de la teoría ha de ser un teorema suyo (o sea, no toda fbf de la teoría tiene necesariamente que ser una fórmula **perteneciente a la teoría**). Si toda fbf de una teoría es un teorema de dicha teoría, la teoría se llama **delicuescente**. Si, por el contrario, no toda fbf de la teoría es un teorema de la teoría, la teoría se llama **sólida o coherente**. (Se llama también, a veces, a una teoría sólida: teoría **absolutamente consistente**; aquí prescindimos de esa terminología para evitar confusiones).

Teorías axiomatizadas

Se llama axioma a un teorema de una teoría que está postulado como tal, sin necesidad, pues, de que sea obtenido a partir de otros mediante la aplicación de reglas de inferencia.

Un esquema axiomático es un esquema —en el sentido definido anteriormente en este trabajo— tal que es un axioma cada instancia del mismo (cada resultado de sustituir por fbfs las letras esquemáticas que lo conforman).

Una teoría $\mathfrak{J}$ está axiomatizada ssi tiene un número enumerable —finito o infinito— de axiomas, tal que hay algún procedimiento **efectivo** para decidir, tras un número finito de pasos, si una fbf cualquiera de $\mathfrak{J}$ es o no un axioma de $\mathfrak{J}$. Una teoría axiomatizada puede tener un número infinito de axiomas (todas las instancias de determinados esquemas axiomáticos, p.ej.).

Teorías inconsistentes

Vamos ahora a dar una caracterización sintáctica de qué es una teoría **simplemente** inconsistente (a la que llamamos **inconsistente a secas**). Expresándonos primero intuitivamente y sin gran rigor, diremos que una teoría es consistente ssi no contiene nunca un par de teoremas tales que el uno sea una negación del otro.

Definamos esas nociones ahora con todo rigor.

Si una teoría $\mathfrak{J}$ es tal que hay una regla de inferencia de $\mathfrak{J}$ que permite reemplazar siempre, en cualquier fórmula $\ulcorner r \urcorner$, una ocurrencia de $\ulcorner p \urcorner$ en $\ulcorner r \urcorner$ por una ocurrencia respecto de $\ulcorner q \urcorner$, entonces se dirá que $\ulcorner p \urcorner$ y $\ulcorner q \urcorner$ son **reemplazables** en $\mathfrak{J}$.

En primer lugar, se define qué es una teoría simplemente inconsistente **con respecto a** un functor de negación 'N': una teoría $\mathfrak{J}$ es llamada simplemente inconsistente con respecto a un functor monádico (de negación) 'N' ssi $\mathfrak{J}$ contiene el functor monádico 'N' y también dos funtores diádicos (respectivamente de conjunción y disyunción): ' $\wedge$ ' y ' $\vee$ ', tales que, para cualesquiera $\ulcorner p \urcorner$, $\ulcorner q \urcorner$ y $\ulcorner r \urcorner$ que sean fbfs de $\mathfrak{J}$:

- (i) Si $\ulcorner p \wedge q \urcorner$ es un teorema de $\mathfrak{J}$, también lo es $\ulcorner p \urcorner$;
- (ii) Si $\ulcorner p \urcorner$ o $\ulcorner q \urcorner$ son teoremas de $\mathfrak{J}$, también lo es $\ulcorner p \vee q \urcorner$;
- (iii) $\ulcorner p \vee q \vee r \urcorner$ y $\ulcorner q \vee r \vee p \urcorner$ son reemplazables;
- (iv) $\ulcorner p \urcorner$, $\ulcorner p \wedge p \urcorner$ y $\ulcorner p \vee p \urcorner$ son reemplazables;
- (v) $\ulcorner p \vee q \wedge r \urcorner$ y $\ulcorner p \wedge r \vee q \wedge r \urcorner$ son reemplazables;
- (vi) El functor 'N' posee las características siguientes:
 - (1) $\ulcorner p \vee Np \urcorner$ es un teorema de $\mathfrak{J}$;
 - (2) $\ulcorner N(p \wedge Np) \urcorner$ es un teorema de $\mathfrak{J}$;
 - (3) $\ulcorner p \urcorner$ y $\ulcorner NNp \urcorner$ son reemplazables;
 - (4) $\ulcorner N(p \wedge q) \urcorner$ y $\ulcorner Np \vee Nq \urcorner$ son reemplazables;

(5) $\neg N(p \vee q)$ y $\neg Np \wedge \neg Nq$ son reemplazables;

(vii) Hay algún $\neg s$ tal que tanto $\neg s$ como $\neg Ns$ son teoremas de $\mathfrak{B}$.

Una teoría $\mathfrak{B}$ es llamada **(simplemente) inconsistente** ssi $\mathfrak{B}$ es inconsistente con respecto a algún functor monádico (de negación) de $\mathfrak{B}$.

(Hay teorías que también se llaman inconsistentes, en un sentido más amplio de la palabra, pese a que sólo cumplen algunas de entre las condiciones (vi.1) a (vi.5).)

Una teoría $\mathfrak{B}$ es contradictoria ssi es inconsistente y, además, cumple la condición siguiente:

(viii) Si $\neg p$ y $\neg q$ son teoremas de $\mathfrak{B}$, también lo es $\neg p \wedge \neg q$.

Así pues, cada teoría contradictoria contiene algún teorema de la forma $\neg p \wedge \neg p$. Una teoría no inconsistente es consistente.

Teorías superconsistentes y paraconsistentes

Ahora vamos a ver qué se entiende por teoría **superconsistente**. Definimos, en primer lugar, qué es una **extensión** de una teoría. Una teoría $\mathfrak{B}'$ es una extensión de otra teoría $\mathfrak{B}$ ssi todo teorema de $\mathfrak{B}$ es también un teorema de $\mathfrak{B}'$.

Nótese lo siguiente: una teoría $\mathfrak{B}'$ puede ser extensión de otra $\mathfrak{B}$ sin que cada regla de inferencia de $\mathfrak{B}$ sea una regla de inferencia de $\mathfrak{B}'$. En efecto: supongamos que, en el acervo de reglas de inferencia de $\mathfrak{B}$ figura una regla de R que es prescindible —o sea: sus efectos pueden lograrse usando sólo las otras reglas—; supongamos ahora que $\mathfrak{B}'$ incluye un nuevo teorema $\neg p$, pero no incluye las conclusiones que se derivarían de aplicar R a $\neg p$ juntamente con los otros teoremas de $\mathfrak{B}$ —que son también, obviamente, teoremas de $\mathfrak{B}'$.

Definimos ahora una **extensión recia** de una teoría como sigue: una teoría $\mathfrak{B}'$ es una extensión recia de una teoría $\mathfrak{B}$ ssi $\mathfrak{B}'$ es una extensión de $\mathfrak{B}$ y cada regla de inferencia de $\mathfrak{B}$ es también una regla de inferencia de $\mathfrak{B}'$.

Una teoría $\mathfrak{B}$ es **superconsistente** ssi toda extensión recia de $\mathfrak{B}$ que sea inconsistente con respecto a algún functor de negación de $\mathfrak{B}$ es una teoría delicuescente.

Una teoría es paraconsistente ssi no es superconsistente.

Modelo, interpretación, valuación, validez

Un dominio es un cúmulo de objetos sobre el que se han definido ciertas relaciones. Más exactamente expresado: llamamos **dominio** a un par ordenado $D = \langle C, \mathfrak{R} \rangle$ donde C es un cúmulo de cosas cualesquiera y $\mathfrak{R}$ es un cúmulo de relaciones definidas sobre C .

Decimos que una relación n -ádica, R , está **definida sobre** un cúmulo C ssi cada miembro de R es un par ordenado cuyo componente izquierdo es una secuencia de n miembros de C . Una **función** n -aria definida sobre un cúmulo o conjunto C es una relación $[n+1]$ -ádica, I , tal que: (1º) I está definida sobre C ; (2º) cada secuencia ordenada de n miembros de C pertenece a un solo miembro de I . Por ende, si I es una función n -aria, y si un par ordenado $\langle \alpha, \beta \rangle$ pertenece a I , entonces no hay ningún otro par ordenado perteneciente a I $\langle \alpha, \gamma \rangle$ (donde $\gamma \neq \beta$).

Si I es una función n -aria, cada elemento que sea componente derecho de un miembro de I será un **valor** de I ; o sea, si —siendo I una función n -aria— $\langle \langle x^1, \dots, x^n \rangle, z \rangle$ pertenece a I , diremos que z es el **valor** (alternativamente llamado también la **imagen**) de I para $\langle x^1, \dots, x^n \rangle$, y lo denotaremos así: $I(x^1, \dots, x^n)$. Si I es una función n -aria definida sobre un cúmulo C , el conjunto de entes que sean componentes derechos de uno u otro miembro de I es el **campo de valores** de I . Ese campo de valores puede tener una intersección no vacía con C y puede también no tenerla. Al campo de valores llámasele a menudo también **contradominio** de la función; pero, más propiamente, cabe llamar **contradominio** de una función a un cúmulo D de entes tal que esté definida la función como teniendo un campo de valores incluido en

D. [Y, todavía con mayor rigor —según lo haremos cuatro párrafos más abajo— cabe llamar **contradominio** de una función ϕ a un par $\langle D, S \rangle$ tal que D incluya al campo de valores de la función ϕ y S sea un cúmulo de relaciones definidas sobre D .] Así nótese una función n -aria $C \rightarrow D$ como una función n -aria definida sobre C y cuyo campo de valores está incluido en D (siendo, pues, D un contradominio de la misma).

Si ϕ es una función unaria, es una relación diádica. Ahora bien, una relación diádica definida sobre C es tal que cada miembro izquierdo de un miembro de ϕ es una secuencia unitaria, e.d. una “serie” $\langle x \rangle$, donde x (su único miembro) pertenece a C . No habiendo confusión podemos identificar esa función unaria ϕ con un cúmulo de pares ordenados que sea igual que ϕ sólo que en cada uno de sus miembros el componente izquierdo $\langle x \rangle$ venga reemplazado por el correspondiente x —e.d. por el componente único de ese componente izquierdo. Así, si una función que es ϕ es un cúmulo de pares, cada uno de los cuales es (para cierto x , cierto z) $\langle \langle x \rangle, z \rangle$, podemos (cuando no haya confusión) reemplazar a ϕ por la correspondiente clase, ψ , de pares $\langle x, z \rangle$, y llamar a ψ una función unaria (como si fuera ϕ). En vez de **función unaria**, podemos hablar de **función monádica**, o de **función a secas**. (Las **funciones** por antonomasia son, pues, las unarias, hasta el punto de que a ellas exclusivamente nos referimos al hablar de **funciones**, salvo que se añada un adjetivo que precise la aridad.) Si I es una función [unaria], entonces llamaremos **argumento** de I a cada elemento x tal que $\langle \langle x \rangle, z \rangle$ (o, hablando laxamente, $\langle x, z \rangle$) —para cierto ente z — sea miembro de I .

Sin dar lugar a confusión podemos, hablando de un dominio $D = \langle C, \mathfrak{R} \rangle$, llamar **dominio** al propio C , o a sus miembros **miembros del dominio**. (Si $\mathfrak{R}$ es una función definida sobre C , a C , e.d. al cúmulo de esos argumentos de la función $\mathfrak{R}$, lo llamamos **campo de argumentos** de la función; pero también, alternativamente, **dominio** de la función.)

Incluido en [el componente izquierdo de] un dominio $\langle C, \mathfrak{R} \rangle$ puede haber un subconjunto de C al que llamaremos el cúmulo de **elementos aléticos** de D . Ese subconjunto puede ser no-propio (o sea, el mismo que C) o bien no abarcar a todos los miembros de C .

Una **interpretación** de una teoría $\mathfrak{Z}$ será una función [monádica] ϕ tal que: (1º) el campo de valores de ϕ es un cúmulo C incluido en un conjunto C' tal que existe el dominio $\langle C', \mathfrak{R} \rangle$, que es llamado **contradominio** (o dominio de valores) de ϕ ; (2º) cada miembro de ϕ tendrá como componente izquierdo a una fbf de $\mathfrak{Z}$ (o, más rigurosamente hablando, a la secuencia unitaria de tal fbf); (3º) cada functor monádico de $\mathfrak{Z}$, ϕ , es tal que hay una relación diádica, ϑ , perteneciente a $\mathfrak{R}$ (siendo $\langle C, \mathfrak{R} \rangle$ el contradominio de ϕ) tal que, para cada fbf de $\mathfrak{Z}$, $\ulcorner p \urcorner$, sucede esto: $\langle \langle \phi(p) \rangle, \phi(\phi p) \rangle \in \vartheta$; (4º) cada functor diádico de $\mathfrak{Z}$, $\ulcorner \mathfrak{I} \urcorner$, es tal que existe una relación triádica, δ , perteneciente a $\mathfrak{R}$ tal que para cualesquiera dos fbfs de $\mathfrak{Z}$, $\ulcorner p \urcorner$, $\ulcorner q \urcorner$, sucede esto: $\langle \langle \phi(p), \phi(q) \rangle, \phi(p \mathfrak{I} q) \rangle \in \delta$.

Si $\mathfrak{Z}$ es una teoría, ϕ es una interpretación de $\mathfrak{Z}$ cuyo contradominio es el dominio $D = \langle C, \mathfrak{R} \rangle$, entonces llamamos a D un **modelo** de $\mathfrak{Z}$ con respecto a la interpretación ϕ siempre y cuando un subconjunto de D sea llamado el cúmulo de **elementos aléticos** de D y cada fbf de $\mathfrak{Z}$, $\ulcorner p \urcorner$, es tal que $\phi(p)$ es un elemento alético. Para los fines del presente capítulo (puesto que estamos ocupándonos no más del cálculo sentencial, y todavía no del cuantificacional), consideraremos modelos todos cuyos miembros sean sus respectivos elementos aléticos.

Dentro de la clase de elementos aléticos de un dominio M , se delimita un subconjunto de la misma: el que abarca a los elementos aléticos **designados**. Un dominio es **trivial** ssi todos sus elementos aléticos son designados.

Con respecto a una teoría dada se puede escoger un modelo particular o una clase particular de modelos, llamándose a ese o esos modelo(s), **modelos propios** de la teoría.

Una vez determinados cuáles son los modelos propios de una teoría $\mathfrak{Z}$ se estipulan determinadas condiciones que debe cumplir una interpretación de $\mathfrak{Z}$ en cualquiera de esos modelos para denominarse una **valuación** de la teoría.

Ssi cada valuación de la teoría $\mathfrak{Z}$ envía al argumento $\ulcorner p \urcorner$ —siendo $\ulcorner p \urcorner$ una fbf de $\mathfrak{Z}$ — sobre un valor que sea un elemento alético designado de uno de los modelos propios de $\mathfrak{Z}$, se dirá que $\ulcorner p \urcorner$ es **válida** [con respecto a esa clase de modelos propios de $\mathfrak{Z}$].

Se dirá también —de modo más restringido— que una fbf $\ulcorner p \urcorner$ de una teoría $\mathfrak{Z}$ es **válida con respecto a** un modelo M de $\mathfrak{Z}$ ssi cada valuación de $\mathfrak{Z}$ que sea una interpretación de $\mathfrak{Z}$ en M hace corresponder a $\ulcorner p \urcorner$ un elemento alético designado de M .

Y se llama **válida** a una regla de inferencia, $\mathfrak{R}$, de $\mathfrak{Z}$ ssi cada valuación v de $\mathfrak{Z}$ es tal que, para cada aplicación de dicha regla (o sea, para cada par $\langle \alpha, \beta \rangle$ perteneciente a $\mathfrak{R}$, donde α es una serie de fbfs de $\mathfrak{Z}$ y β es una fbf de $\mathfrak{Z}$), si v hace corresponder a cada premisa un elemento alético designado, entonces v hace también corresponder a la conclusión un elemento alético designado.

Un modelo M es **idóneo** para una teoría $\mathfrak{Z}$, con respecto a las valuaciones de $\mathfrak{Z}$ en M , ssi éstas están definidas de tal modo que:

- 1º.— Para todo teorema $\ulcorner p \urcorner$ de $\mathfrak{Z}$, cada valuación de $\mathfrak{Z}$ en M haga corresponder a $\ulcorner p \urcorner$ un elemento alético designado de M ;
- 2º.— Ninguna valuación de $\mathfrak{Z}$ en M haga corresponder a cada una de las fbfs de $\mathfrak{Z}$ un elemento alético designado.

Conviene notar que, si $\mathfrak{Z}$ es una teoría axiomatizada, entonces es condición necesaria y suficiente para que un dominio M sea **modelo idóneo** de $\mathfrak{Z}$ que se cumplan las dos condiciones siguientes:

- 1) Para cada axioma $\ulcorner p \urcorner$ de $\mathfrak{Z}$, cada valuación de $\mathfrak{Z}$ en M asigna a $\ulcorner p \urcorner$ un elemento alético designado de M ;
- 2) Cada aplicación de cada regla de inferencia de $\mathfrak{Z}$ es tal que, dada una valuación cualquiera de $\mathfrak{Z}$ en M , si tal valuación hace corresponder a cada premisa de esa aplicación un elemento alético designado, también hace corresponder a la conclusión un elemento alético designado.

Una teoría $\mathfrak{Z}$ es **robusta** ssi existe un cúmulo de modelos propios de $\mathfrak{Z}$ cada uno de los cuales es un modelo idóneo de $\mathfrak{Z}$, o sea: ssi cada uno de los teoremas de $\mathfrak{Z}$ es una fórmula válida de $\mathfrak{Z}$ (con respecto a ese cúmulo de modelos propios, habiéndose definido las valuaciones con sujeción a las dos condiciones recién enumeradas). Es obvio que, una vez conocido algún modelo idóneo de una teoría $\mathfrak{Z}$, siempre es posible definir los modelos propios de $\mathfrak{Z}$ de tal modo que tan sólo sean modelos propios de $\mathfrak{Z}$ aquellos que sean modelos idóneos de $\mathfrak{Z}$. Por ello basta con que una teoría tenga algún modelo idóneo para que sea robusta. Y es fácil demostrar que una teoría es robusta ssi es sólida.

(Una observación terminológica parentética: a menudo se denomina **modelo** de una teoría a lo que aquí se ha llamado **modelo idóneo**; o a un par ordenado formado por un modelo idóneo más una valuación; o incluso a la valuación misma, con respecto a un dominio que sea lo que aquí hemos llamado ‘modelo idóneo’.)

Completez

Una teoría $\mathfrak{Z}$ se llama completa ssi hay alguna clase Θ de modelos idóneos de $\mathfrak{Z}$ tal que es un **teorema** de $\mathfrak{Z}$ cada fbf de $\mathfrak{Z}$ válida con respecto a todos los miembros de Θ .

Si $\mathfrak{Z}$ es una teoría completa y M es un modelo idóneo de $\mathfrak{Z}$ con la susodicha característica (es decir: si cada fbf de $\mathfrak{Z}$ que sea válida con respecto a M es un teorema de $\mathfrak{Z}$), entonces M es un **modelo característico** de $\mathfrak{Z}$.

La noción de completez es menos importante que las de robustez y solidez. Una teoría cuya robustez quede indemostrada quizá todavía no habrá probado su utilidad; pero hay teorías dotadas de enorme interés y utilidad sin ser completas.

Es **fuertemente característico** un modelo $M = \langle C, \mathfrak{R} \rangle$ de una teoría $\mathfrak{Z}$ ssi es modelo característico de $\mathfrak{Z}$ y, además, cumple esta condición adicional: para cada relación R perteneciente a $\mathfrak{R}$ o resultante de las pertenecientes a $\mathfrak{R}$ por composición relacional —producto relativo— hay un functor $\mathfrak{F}$ que forma parte del vocabulario de $\mathfrak{Z}$ —o es definible mediante ese vocabulario— tal que: o bien (1) R es una relación diádica y $\mathfrak{F}$ es un functor monádico y entonces, para cada $x \in C$, cada fbf $\ulcorner p \urcorner$ de $\mathfrak{Z}$ y cada valuación v de $\mathfrak{Z}$ en M , $v(\mathfrak{F}p) = x$ ssi $\langle \langle v(p), x \rangle \in R \rangle$; o bien (2) R es una relación triádica y $\mathfrak{F}$ es un functor diádico y entonces, para cualesquiera fbfs $\ulcorner p \urcorner$, $\ulcorner q \urcorner$ de $\mathfrak{Z}$, para cualquier miembro x de C y cualquier valuación v de $\mathfrak{Z}$ en M , sucede esto: $v(\mathfrak{F}pq) = x$ ssi $\langle \langle v(p), v(q), x \rangle \in R \rangle$.

Una teoría es **fuertemente completa** ssi tiene algún modelo que sea **fuertemente característico** de ella. Que una teoría $\mathfrak{Z}$ sea fuertemente completa quiere decir que cualquier fórmula —de la teoría que sea— válida en cierto modelo característico de $\mathfrak{Z}$, siendo una fórmula expresable en el vocabulario de $\mathfrak{Z}$, es un teorema de $\mathfrak{Z}$.

La noción de completez fuerte sólo tiene interés para los cálculos sentenciales cuyos modelos característicos son finivalentes.

Tablas de verdad y modelos

En los capítulos precedentes, al hablar de una lógica, se ha entendido un cúmulo de fórmulas válidas **en un modelo** —a saber, en **un conjunto de valores de verdad**— según ciertas valuaciones, que eran las tablas de verdad. (En el caso de las lógicas infinivalentes no había exactamente tablas de verdad, sino asignaciones de verdad).

Pero, tal como hemos venido empleando la expresión hasta ahora, una lógica (o un sistema lógico) era una teoría no axiomatizada, sin reglas de inferencia, ya que se tomaban directamente como teoremas de la misma todas sus fórmulas válidas. Una teoría cuyos teoremas se determinen de antemano de ese modo es, por definición, una teoría completa. Y una teoría definida de ese modo es una teoría **semánticamente definida**. En los capítulos 8 y siguientes nos interesaremos, en cambio, por teorías lógicas axiomatizadas.

Si una teoría axiomatizada es completa, cualquiera de las dos representaciones de la misma —la axiomatizada y la que **define** los teoremas como las fórmulas válidas equivaldrá a la otra.

Pero, a partir del Capítulo 9 centraremos principalmente nuestra atención en sistemas cuya completez únicamente podrá reconocerse mediante las técnicas de hallazgo de modelos algebraicos que exploraremos en el último capítulo. Y, por ello, las teorías lógicas de que hablaremos serán **subteorías** de otras estudiadas en capítulos precedentes de este opúsculo, sin coincidir con ellas. (Técnicamente hablando: éstas últimas son extensiones no-conservativas de las que se estudiarán entonces; una teoría $\mathfrak{Z}'$ es una extensión conservativa de una teoría $\mathfrak{Z}$ ssi $\mathfrak{Z}'$ es una extensión de $\mathfrak{Z}$ y cada teorema de $\mathfrak{Z}'$ que esté formulado **sólo** mediante el vocabulario de $\mathfrak{Z}$ es un teorema de $\mathfrak{Z}$.)

Un modelo alternativo para $\mathcal{A}_p$

El lector puede ahora, utilizando la noción de modelo idóneo expuesta en el presente capítulo, comprobar que constituye un modelo idóneo para el sistema lógico $\mathcal{A}_p$ presentado en el capítulo anterior el dominio A que se va a indicar a continuación.

Partimos del cúmulo de los números reales no negativos, R . Sea $D' = R \cup \{\infty\}$ (o sea: el conjunto que sólo abarca a ∞ y a todo miembro de R , donde ∞ es un número infinito) uno cualquiera de los números enteros mayores que cualquier número obtenible a partir del 0 por

la operación de ir a su sucesor; obsérvese que la existencia de números enteros así es reconocida en el análisis no-estándar de Robinson. Sea D = el cúmulo que abarca a cada miembro de D' , a cada resultado de **adicionar** un infinitésimo, α , a un miembro cualquiera de R y también a cada resultado de **restar** ese infinitésimo, α , de un miembro cualquiera de D' que no sea 0. (α será $1/\infty$.) Definimos ahora los elementos y operadores siguientes.

$nx =$	$z+\alpha$ si $z \in R$ es tal que $z=x$ o $x=z+\alpha$	$mx =$	$z-\alpha$ si $z \in R$ es tal que $z=x$ o $x=z+\alpha$	$Nx =$	∞ , si $x=0$
	x en caso contrario		x en caso contrario		0 , si $x=\infty$
$x \downarrow z = \max(Nx, Nz)$		$Hx =$	0 , si $x=0$		$1/x$ si $x \neq 0$ pero $x \in R$
$x \bullet z = x \bullet z = x \times z$, si $x, z \in R$			∞ , en caso contrario		mNz si $x=nz$
$nx \bullet z = z \bullet nx = n(x \bullet z)$					nNz si $x=mz$
$x \bullet \infty = \infty \bullet x = \infty$		$x \rightarrow z =$	1 , si $x \geq z$		
$mx \bullet z = z \bullet mx = m(x \bullet z)$ si $z \neq nz$			∞ , en caso contrario		

Son **designados** todos los miembros de D salvo ∞ . Definamos ahora A como el dominio $\langle D, \mathfrak{R} \rangle$ donde $\mathfrak{R}$ es el siguiente cúmulo de relaciones: $\{r^1, r^2, r^3, r^4, r^5\}$, siendo r^1 el cúmulo de pares ordenados $\langle \langle x \rangle, Hx \rangle$ (para cada $x \in D$); r^2 = el cúmulo de pares ordenados $\langle \langle x \rangle, nx \rangle$ donde $x \in D$; sea ahora un cúmulo cualquiera de pares $\langle B, x \rangle$ donde $B = \langle c, d \rangle$, siendo $x, c, d \in D$; entonces: (1) r^3 = el cúmulo de tales pares $\langle B, x \rangle$ tal que $x=c \downarrow d$; (2) r^4 = el cúmulo de tales pares tal que $x=c \bullet d$; (3) r^5 = el cúmulo de tales pares tal que $x=c \rightarrow d$.

Tócale al lector como ejercicio calcular cómo A es un modelo idóneo de Ap . (Pauta: percatarse del isomorfismo entre A y el dominio de valores de verdad estudiado en el capítulo precedente para Ap : para cada número real perteneciente a A , x , hay un solo miembro de ese dominio que es 2^{-x} . Búsquese entonces qué operadores, de los definibles con miembros de $\mathfrak{R}$, corresponden a n , a m , a $\#$, a $@$, a $\times$; y qué condiciones han de satisfacer las interpretaciones del conjunto de fbfs de Ap en el dominio A para ser valuaciones de Ap , tales que, en virtud de ellas, no sólo A sea un modelo idóneo de Ap , sino que además A sea un modelo fuertemente característico de Ap .)

Un modelo no fuertemente característico de la LBV

Al final del Capítulo 3 estudiamos las lógicas-producto o lógicas tensoriales. Tomemos una de ellas, a saber aquella en la que cada valor de verdad es un par ordenado, $\langle x, z \rangle$ donde cada uno de entre x, z es o bien $=0$ o bien $=1$. Sea $\langle 1, 1 \rangle$ el único elemento designado en ese conjunto de pares ordenados, o sea en ese cúmulo de valores de verdad. Sobre ese conjunto de valores de verdad definimos estas dos relaciones: r^1 = el cúmulo de pares ordenados $\langle \langle u, u' \rangle, v \rangle$ tales que $v_i=1$ ssi no sólo $u_i=0$ sino que también $u'_i=0$; r^2 = el cúmulo de pares ordenados $\langle \langle u \rangle, v \rangle$ donde $v=u$ si $u=\langle 1, 1 \rangle$ y, en caso contrario, $v=\langle 0, 0 \rangle$.

Ese dominio —llamémoslo **D**— es un modelo característico de la lógica clásica, e.d. de la LBV. En efecto es fácil encontrar qué condiciones deben cumplir las valuaciones de LBV en **D**, a saber: (1ª) cada valuación v será tal que para cualesquiera fbfs $\lceil p \rceil, \lceil q \rceil, v(\neg p) = u$ ssi $\langle \langle v(p) \rangle, v(p), u \rangle \in r^1$; (2ª) $v(p \wedge q) = u$ ssi $\langle \langle v(\neg p), v(\neg q) \rangle, u \rangle \in r^1$. Con ello se comprueba cómodamente que cada fórmula válida de LBV con respecto a ese modelo es un teorema

de LBV y viceversa. Sin embargo, hay una relación en ese modelo, a saber r^2 , a la que no corresponde ningún functor en LBV. En efecto:

$$r^2 = \{\langle\langle 0, 1 \rangle\rangle, \langle 0, 0 \rangle\rangle, \langle\langle 1, 0 \rangle\rangle, \langle 0, 0 \rangle\rangle, \langle\langle 0, 0 \rangle\rangle, \langle 0, 0 \rangle\rangle, \langle\langle 1, 1 \rangle\rangle, \langle 1, 1 \rangle\rangle\}$$

Supongamos que hubiera un functor monádico de LBV, $\oint$, tal que, siendo v una valuación cualquiera de LBV en $\mathbf{D}$, para cada fbf, $\ulcorner p \urcorner$, $v(\oint p) = u$ ssi $\langle\langle v(p), u \rangle\rangle \in r^2$.

Es fácil comprobar que eso no sucede, que ningún functor definible en LBV es un functor $\oint$ con esa característica. De hecho los únicos modelos fuertemente característicos de LBV son modelos isomórficos al modelo $\langle\{0, 1\}, \{r^1\}\rangle$, donde r^1 es el cúmulo de pares ordenados $\langle\langle u, v \rangle, v' \rangle$ tales que $v' = 1$ ssi no sólo $u = 0$ sino que también $v = 0$. Todo modelo característico de la LBV es un **álgebra booleana**. (Vide al respecto el cap. 12, y último, del presente opúsculo.) Pero no toda álgebra booleana es un modelo **fuertemente** característico de dicha lógica. Si expandiéramos la lógica clásica con un functor monádico 'B' —léible como 'Es afirmable con verdad que' o 'Es verdad en todos los aspectos que'— con tal de que a una exposición axiomática de la lógica clásica, de entre las varias que hay, se le añadieran, junto con él, ciertos axiomas y reglas de inferencia (la regla $p \vdash Bp$; y los dos esquemas axiomáticos: (1) $\ulcorner B \neg (p \wedge \neg q) \wedge Bp \supset Bq \urcorner$; y (2) $\ulcorner \neg B \neg p \supset B \neg B \neg p \urcorner$), entonces, sí, el resultado sería una lógica que tendría como modelo fuertemente característico al recién considerado dominio $\mathbf{D}$.

Desde el punto de vista filosófico lo anterior quiere decir que todo dominio de valores de verdad característico de la lógica clásica, siendo **booleano**, excluye de sí la presencia de **grados**, mas no forzosamente excluye de sí los **aspectos**. Sólo que, si comprende aspectos, entonces ni siquiera es un modelo fuertemente característico de LBV. La LBV es, pues, una teoría de la cual son característicos sólo modelos donde, por tener en ellos vigencia un principio de maximalidad (todo o nada), no hay matices o grados, y de la cual, además, son fuertemente característicos sólo modelos donde no caben aspectos, e.d. donde no cabe que un hecho suceda o sea existente sólo desde ciertos puntos de vista **objetivos** o con relación a ciertos lados de las cosas.

En el otro polo, una lógica como *Abp* tan sólo admite modelos que tengan tanto infinitos grados como una pluralidad de aspectos.

Capítulo 7

Los principios de no-contradicción y tercio excluso

Hemos visto que no existe una lógica única, sino una pluralidad de lógicas. Escogerá cada pensador una lógica determinada según cuáles sean sus concepciones ontológicas fundamentales y posiciones últimas de valor —en suma, todo su cosmorama.

Pero lo que es interesante, desde el ámbito de una pulcra investigación lógico-matemática, es indagar las propiedades de los diversos sistemas desde el ángulo de las tautologías y reglas de inferencia que contienen.

Uno de los puntos importantes, a ese respecto, consiste en saber qué pasa en los diversos sistemas con los principios de no contradicción y de tercio excluso, dada la importancia filosófica de ambos.

La regla de Cornubia

Como ya sabemos, una lógica **superconsistente** es aquella tal que, si se le añade una inconsistencia simple (o sea un par de fórmulas una de las cuales sea **una** negación de la otra), el resultado es un sistema delicuescente.

Todo sistema superconsistente es, pues, un sistema sólido que contiene la siguiente regla de inferencia (para cualquier functor de negación ' $\sim$ ' del sistema): $p, \sim p \vdash q$

A esa regla de inferencia la llamaremos *regla de Cornubia*, pues aparece en un escrito lógico medieval, cuyo autor fue, al parecer, Juan de Cornubia, aunque durante tiempo se atribuyó por error al Doctor Sutil Juan Duns Escoto —por lo cual a menudo la llaman 'regla de Escoto'.

En cambio, un sistema S es **paraconsistente** sólo si es sólido y no sucede en absoluto que S contenga la regla de Cornubia para cada functor de negación de S —aunque, desde luego, puede contenerla para algún functor de negación de S .

Pero, ¿qué se significa al decir que un sistema contiene una regla de inferencia determinada? La respuesta es obvia cuando se trata de un sistema axiomatizado, en el que se estipulan explícitamente reglas de inferencia que permiten extraer teoremas a partir de otros teoremas —y, por ende, teoremas a partir, en último término, de los axiomas. Mas cuando un sistema no está axiomatizado y no se postulan en él reglas de inferencia sino que se lo define semánticamente —o sea, diciendo que son teoremas del sistema todas las fórmulas válidas del mismo con arreglo a determinado(s) modelo(s) y valuaciones—, entonces es menester precisar qué se entiende, con respecto a tal sistema, al decir que contiene o deja de contener una regla de inferencia.

Lo que se puede significar al decir de un sistema semánticamente definido que contiene una regla de inferencia puede ser, en primer lugar, que tal regla de inferencia es sana o válida en el sistema, o sea: que, si las premisas tienen valores de verdad designados, la conclusión también tendrá un valor de verdad designado.

En un sistema superconsistente es imposible que un enunciado —cualquiera que sea— y una negación del mismo tengan ambos, a la vez, valores designados.

Tal es el caso, obviamente, en la lógica bivalente. Si $/p/=1$, su negación tendrá el valor 0, y viceversa. Y 0 no es designado. Por ello, la regla de Cornubia es una regla sana en la lógica bivalente.

Pero esa regla de inferencia es también una regla válida en otros sistemas, como A_1 (un sistema trivalente en el que sólo es designado el valor 1). El lector comprobará por sí mismo que, en tal sistema, no es tampoco posible que un enunciado y su negación tengan, ambos, valores de verdad designados.

En cambio, ocurre algo muy distinto en los sistemas A_3 , Ar , Ap , Abp . Ninguno de ellos contiene la regla de Cornubia. Esto se ve, del modo más sencillo e inmediato, en A_3 : si $/p/=1/2$, $/Np/=1/2$; supongamos que $/q/=0$; es evidente que del par de premisas $\lceil p \rceil$ y $\lceil Np \rceil$ no cabe extraer la conclusión $\lceil q \rceil$, pues los valores de verdad de las premisas son designados, mientras que el valor de verdad de la conclusión no es designado. La regla de Cornubia no es válida en esos sistemas.

Pero podemos también entender algo diferente cuando decimos, de un sistema semánticamente definido, que contiene una regla de inferencia: podemos querer decir que esa regla es **dura** en el sistema, o sea: que es una regla de inferencia que, además de ser válida, es tal que la conclusión tiene como valor 0 (el valor mínimo) sólo si alguna de las premisas tiene también como valor 0.

Pues bien, la regla de Cornubia es una regla de inferencia dura de la lógica bivalente, pero no es una regla de inferencia dura del sistema A_1 .

Hay, pues, sistemas que contienen la regla de Cornubia como regla válida, pero nunca como regla dura. Pues, si un sistema contiene la regla de Cornubia como regla dura, quiere ello decir que, puesto que $\lceil q \rceil$ es —en la enunciación de dicha regla— cualquier oración —y, por tanto, es posible que $/q/=0$ —, entonces, dadas un par de oraciones tales que una de ellas sea negación de la otra —o sea: tales que una de ellas sea $\lceil p \rceil$ y la otra $\lceil \sim p \rceil$ —, la siguiente situación se da: O bien $/p/=0$, o bien $/\sim p/=0$.

No obstante, no basta que la regla de Cornubia no sea regla de inferencia dura en un sistema para que el mismo sea paraconsistente.

Así, veamos lo que ocurre en L_3 (la lógica trivalente de Łukasiewicz, el primer sistema multivalente de lógica que fue construido). En ese sistema tenemos tres valores: 1, $1/2$, 0. El único valor designado es 1. Las tablas de verdad para la conjunción ($\wedge$), para la negación (N) y para la disyunción ($\vee$) son las mismas que las tablas respectivas de A_1 . Pero el condicional de L_3 ($\Rightarrow$) tiene la tabla siguiente:

$\Rightarrow$	1	$1/2$	0
1	1	$1/2$	0
$1/2$	1	1	$1/2$
0	1	1	1

Es evidente que en L_3 la regla de Cornubia no es una regla de inferencia dura, aunque sí es una regla de inferencia válida. Pero otras reglas de inferencia válidas de L_3 son *modus ponens* y *adjunción*, o sea:

MP: $p \Rightarrow q, p \vdash q$

Adj: $p, q \vdash p \wedge q$

Pues bien, supongamos que, para alguna constante sentencial s , añadimos a L_3 esa constante sentencial y su negación. En L_3 no tenemos como teorema (o sea como tautología): $\lceil p \wedge Np \Rightarrow q \rceil$ —si lo tuviéramos, se concluiría inmediatamente $\lceil q \rceil$ de las premisas $\lceil s \rceil$ y $\lceil Ns \rceil$, mediante las dos reglas de Adj y MP. Pero, así y todo, tenemos en L_3 el teorema siguiente: $\lceil p \wedge Np \Rightarrow p \wedge Np \Rightarrow q \rceil$

Supongamos, pues, que añadimos a los teoremas de L_3 los dos siguientes: $\lceil s \rceil$, $\lceil Ns \rceil$. Por aplicación de Adj, obtendremos: $\lceil s \wedge Ns \rceil$. Por instanciación en el esquema teorematóico más arriba indicado tendremos:

$\lceil s \wedge Ns \rceil \supset \lceil s \rceil$. Por MP tendremos: $\lceil s \wedge Ns \rceil \supset \lceil q \rceil$. Y, nuevamente por MP: $\lceil q \rceil$. (Conque la regla de Cornubia es derivable en $\mathcal{L}_3$.)

El principio de Cornubia

El principio de Cornubia corresponde a la regla de idéntica denominación. Se enuncia así: «Si p y no- p , entonces q » (siendo ‘no’ cualquier functor de negación). ¿Qué relación hay entre el principio de Cornubia y la regla de Cornubia? Ello depende de cuáles sean las propiedades del functor ‘si... entonces’, en un sistema dado.

Si un sistema S contiene la regla de adjunción y la regla de MP entonces contiene el principio de Cornubia sólo si también contiene la regla de Cornubia.

En efecto: supongamos que S contiene la regla de adjunción y la regla de MP, o sea:

$p \supset q, p \vdash q$

y también el principio de Cornubia:

$p \wedge \neg p \supset q$

Entonces, por aplicación de la regla de adjunción, se obtiene la prótasis del principio de Cornubia a partir del par de premisas $\lceil p \rceil$ y $\lceil \neg p \rceil$ (cualquiera que sea $\lceil p \rceil$); y luego, por aplicación de la regla MP, se obtiene $\lceil q \rceil$.

Pero un sistema puede contener la regla de Cornubia sin contener el principio de Cornubia (tal es el caso de $\mathcal{L}_3$). Y puede también un sistema contener el principio de Cornubia pero no la regla de Cornubia, siempre y cuando tal sistema carezca, ya sea de la regla de adjunción, ya sea de MP. (Hay quienes no consideran correcta la regla de adjunción.) El inconveniente, sin embargo, de tales sistemas es que es difícil admitir que un functor diádico sea un condicional si no es válida para él la regla de MP; y que es difícil admitir que **no** haya **ningún** functor condicional ‘ $\supset$ ’ tal que, si $p \vdash q$, entonces $\lceil p \supset q \rceil$ sea un teorema.

Por ello, todo sistema que sea, por lo demás, satisfactorio y contenga como teorema el principio de Cornubia debe contener la regla de Cornubia. Y, como tal regla es inadmisibles desde el ángulo de las posiciones que admiten la contradictorialidad de lo real, el principio será también inadmisibles desde ese ángulo.

Antes de pasar al epígrafe siguiente, conviene precisar que, en la formulación tanto del principio como de la regla de Cornubia, sólo hemos tenido en cuenta la formulación más fuerte de los mismos, a saber: que sean válidos para **cualquier** negación. Pero un sistema puede ser paraconsistente —o incluso contradictorio— y, sin embargo, contener para **algún** functor de negación (de negación fuerte o **super**negación) tanto el principio como la regla de Cornubia. Tal es el caso, en lo tocante al functor ‘ $\neg$ ’, de los sistemas paraconsistentes A_3 , A_5 , Ar , Ap , Abp .

El principio sintáctico de no-contradicción

Podemos llamar **antinomia** a cualquier fórmula del tipo $\lceil p \wedge Np \rceil$. (También se suele llamar a esas fórmulas **contradicciones**. Sólo que, más comúnmente, se llama hoy **contradicción**, en la lógica actual, a cualquier fórmula de un sistema semánticamente definido que tome uniformemente un valor antidesignado, cualesquiera que sean los valores de sus fórmulas atómicas.)

El principio sintáctico de no contradicción es una fórmula o, más exactamente, alguna de entre un abanico de fórmulas.

Si ‘ $\sim$ ’ es un functor de negación y ‘ $\&$ ’ un functor de conjunción, **un** principio de no-contradicción es el esquema: $\lceil \sim(p \& \sim p) \rceil$. O sea el principio es la negación de cualquier antinomia.

Ahora bien, teniendo en cuenta que un sistema puede contener varios funtores de negación y también varios funtores de conjunción, ¿cuándo cabrá decir si el sistema posee o no el principio de no-contradicción? La respuesta es que, en vez de hablar del principio de no-contradicción, hay que hablar de principios de no-contradicción. Un sistema puede no contener, p.ej., un principio semejante para determinados funtores pero sí contenerlo para otros funtores.

Sea, p.ej., el sistema A_1 . En él el esquema $\lceil N(p \wedge Np) \rceil$ no es teorematizado (no es una tautología). Pero el esquema $\lceil \neg(p \wedge \neg p) \rceil$ sí es una tautología.

En cambio, los sistemas A_3 , Ar , Ap , Abp contienen, todos ellos, el principio de no-contradicción para **cada** functor de negación y para **cada** functor de conjunción. Así, en Ap las fórmulas siguientes son tautologías (definiendo así '&': $\lceil p \& q \rceil$ abr $\lceil Lp \wedge q \rceil$):

$N(p \& Np)$ $\neg(p \& \neg p)$ $N(p \bullet Np)$ $\neg(p \bullet \neg p)$ $N(p \wedge Np)$ $\neg(p \wedge \neg p)$

Cabe seguramente pensar que sólo son satisfactorios y plausibles los sistemas que contengan el principio de no-contradicción para cada functor de negación y para cada functor de conjunción. Un sistema que satisfaga tal requisito será llamado **eunómico**.

Principio de no-contradicción y paraconsistencia

¿Qué relación hay entre la posesión por un sistema del o los principio(s) de no-contradicción y el hecho de que tal sistema sea o deje de ser paraconsistente?

Un sistema puede ser, no ya paraconsistente, sino incluso contradictorio, siendo, con todo, un sistema eunómico. (P.ej.: A_3 , Ar , Ap , Abp .)

Por otro lado, un sistema puede carecer no ya de algún principio de no-contradicción, sino de cualquier principio de no-contradicción, y ser, empero, un sistema superconsistente. Tal es el caso del sistema $\perp_3$, como ya sabemos.

Por supuesto, un sistema puede también ser paraconsistente sin ser eunómico (el sistema A_1 , p.ej.); y también ser eunómico y superconsistente (la lógica bivalente, p.ej.)

Mas —se dirá— ¿qué sentido tiene que un sistema contradictorio contenga el principio de no-contradicción? En efecto: si equiparamos (como se suele hacer —y correctamente a juicio del autor de este trabajo) «no-p» con «Es falso que p», entonces lo que nos viene a decir el principio de no-contradicción (entendiendo aquí por tal $\lceil N(p \wedge Np) \rceil$ —o sea: tomando, como functor de negación, la negación natural, y, como functor de conjunción, la conjunción natural) es que, cualquiera que sea $\lceil p \rceil$, es falso que p-y-no-p. Pero un sistema contradictorio reconoce alguna antinomia, o sea: alguna fórmula del tipo «p-y-no-p»; si la reconoce, es que la tiene por verdadera; mas hemos dicho que, si el sistema admite el principio de no-contradicción, sostendrá que cada antinomia es falsa; así pues, un sistema que reconozca como verdadera una antinomia y que contenga el principio de no-contradicción vendrá a sostener que hay al menos una fórmula verdadera y falsa.

Así es, en efecto. Mas ello es legítimo desde el ángulo del partidario de la contradictorialidad de lo real.

Al afirmar que, para al menos un $\lceil p \rceil$, la antinomia $\lceil p \wedge Np \rceil$ es cierta, viene a afirmar el contradictorialista que $\lceil p \rceil$ es verdadero y $\lceil Np \rceil$ es también verdadero, o sea: que $\lceil p \rceil$ es verdadero y falso a la vez. Ello es forzosamente así si su sistema admite la conmutatividad de la conjunción ($\lceil p \wedge q \rceil = \lceil q \wedge p \rceil$) y la regla de simplificación, a saber: $p \wedge q \vdash p$

Por consiguiente, el contradictorialista admite, precisamente, que hay alguna oración que es, a la vez, verdadera y falsa —y esa admisión es lo que lo caracteriza como contradictorialista. Y, por ello, no tiene escrúpulo en aceptar igualmente, con respecto a otra fórmula más compleja, que esa otra fórmula es, también ella, a la vez verdadera y falsa.

Recapitulando: el contradictorialista acepta, para algún $\lceil p \rceil$, la antinomia $\lceil p \wedge Np \rceil$; si también acepta el principio de no-contradicción, aceptará $\lceil N(p \wedge Np) \rceil$. O sea: considera a la antinomia $\lceil p \wedge Np \rceil$ a la vez como verdadera y como falsa, al haber afirmado $\lceil p \wedge Np \rceil$ (y por ello —en virtud de la conmutatividad de la conyunción y de la regla de simplificación— tanto $\lceil p \rceil$ como $\lceil Np \rceil$), ya había sostenido con respecto a una fórmula —a saber: $\lceil p \rceil$ — que es, a la vez, verdadera y falsa. Por consiguiente, como el contradictorialista acepta que se dan fórmulas a la vez verdaderas y falsas, nada le impide aceptar que una antinomia sea, a la vez, verdadera y falsa.

Es más: supongamos que el sistema del contradictorialista en cuestión contiene la regla de adjunción:

$p, q \vdash p \wedge q$

Tal regla es plausibilísima y de uso corriente, y es validada por la gran mayoría de los sistemas de lógica. (Los sistemas que tienen esa regla son los **copulativos**.)

Pues bien, de ser así, si el sistema contradictorial en cuestión sostiene, a la vez, una antinomia $\lceil s \wedge Ns \rceil$ y el principio de no contradicción $\lceil N(p \wedge Np) \rceil$, el contradictorialista sostendrá la verdad de todas las fórmulas siguientes:

- 1) $\lceil s \wedge Ns \rceil$
- 2) $\lceil N(s \wedge Ns) \rceil$
- 3) $\lceil s \wedge Ns \wedge N(s \wedge Ns) \rceil$
- 4) $\lceil N(s \wedge Ns \wedge N(s \wedge Ns)) \rceil$
- 5) $\lceil s \wedge Ns \wedge N(s \wedge Ns) \wedge N(s \wedge Ns \wedge N(s \wedge Ns)) \rceil$
- 6) $\lceil N(s \wedge Ns \wedge N(s \wedge Ns) \wedge N(s \wedge Ns \wedge N(s \wedge Ns))) \rceil$
- 7) $\lceil s \wedge Ns \wedge N(s \wedge Ns) \wedge N(s \wedge Ns \wedge N(s \wedge Ns)) \wedge N(s \wedge Ns \wedge N(s \wedge Ns) \wedge N(s \wedge Ns \wedge N(s \wedge Ns))) \rceil$

etc. etc. Se pueden abreviar esas fórmulas definiendo así un functor monádico 'S':

$\lceil Sp \rceil$ abr $\lceil p \wedge Np \rceil$

(Queda como ejercicio para el lector abreviar las anteriores fórmulas de ese modo, dándoles así mayor perspicuidad.)

Pero si algo puede objetarse a las fórmulas antinómicas de esa lista —que están indicadas por números de orden nones— (o a la inconsistencia simple formada por cada par de fórmulas de la lista tal que la primera tiene un número de orden non y la segunda es la fórmula con número de orden par que la sigue inmediatamente), eso mismo podía objetarse, ya de entrada, contra la antinomia inicial $\lceil s \wedge Ns \rceil$ (y contra la inconsistencia simple: $\lceil s \rceil$, $\lceil Ns \rceil$, que se deriva de ella por la conmutatividad de la conyunción más la regla de simplificación).

Lo único, pues, que el contradictorialista debe evitar y rehuir en su sistema es la regla de Cornubia **con respecto al** functor de negación 'N' tal que haya en su sistema una antinomia $\lceil s \wedge Ns \rceil$.

Antinomia y contradicción

Como ya hemos dicho, se entiende normalmente por contradicción, en la lógica actual, cualquier fórmula que —cualesquiera que sean los valores de verdad de sus subfórmulas atómicas— tiene siempre, forzosamente, un valor de verdad antidesignado.

Si suponemos que —según las normas estipuladas para la negación natural en el Capítulo 4— $\lceil p \rceil$ es antidesignado ssi $\lceil Np \rceil$ es designado —cualquiera que sea $\lceil p \rceil$ —, concluiremos que una contradicción es una fórmula tal que su negación simple o natural tiene siempre, forzosamente, un valor de verdad designado; o sea: una fórmula cuya negación —natural— es una tautología. Y como la negación natural es **involutiva** (o sea: cumple la condición $\lceil NNp \rceil = \lceil p \rceil$), podemos concluir que una fórmula es una contradicción ssi es la negación —simple

o natural— de una tautología (o, lo que es lo mismo —por la involutividad de la negación natural— si su negación —simple o natural— es una tautología).

Pues bien, un sistema que contenga el principio de no-contradicción es un sistema en el que cada antinomia es una contradicción. Cabe preguntarse: en un sistema semejante ¿es una antinomia cada contradicción? No, pero, dada una contradicción, se puede siempre —mediante la regla de adjunción— engendrar —a partir de ella y de la tautología cuya negación es— **otra** contradicción que sea una antinomia.

Y, por otro lado, nada impide considerar como contradictorias por antonomasia a aquellas contradicciones —de un sistema en el que el principio de no-contradicción sea una tautología— que sean antinomias.

En un sistema de lógica contradictoria hay siempre **tautologías contradictorias** —o sea: contradicciones tautológicas.

El principio semántico de no-contradicción

Por **principio semántico de no-contradicción** se pueden entender varias cosas:

- (1) Para cada fórmula $\lceil p \rceil$, o bien $\lceil p \rceil$ es antidesignado, o bien $\lceil \sim p \rceil$ es antidesignado (siendo $\sim$ cualquier functor de negación).
- (2) Para cada fórmula $\lceil p \rceil$, o bien $\lceil p \rceil = 0$, o bien $\lceil \sim p \rceil = 0$ (siendo $\sim$ cualquier functor de negación).
- (3) Para cada fórmula $\lceil p \rceil$, no sucede que tanto $\lceil p \rceil$ como $\lceil \sim p \rceil$ sean ambos designados (siendo $\sim$ cualquier functor de negación).
- (4) Para cada fórmula $\lceil p \rceil$, o bien $\lceil p \rceil \neq 1$, o bien $\lceil \sim p \rceil \neq 1$ (siendo $\sim$ cualquier functor de negación).
- (5) Cada fórmula $\lceil p \rceil$ es tal que $\lceil p \rceil$ es designado sólo si $\lceil \sim p \rceil$ es antidesignado (siendo $\sim$ cualquier functor de negación).
- (6) Cada fórmula $\lceil p \rceil$ es tal que $\lceil \sim p \rceil$ es designado sólo si $\lceil p \rceil$ es antidesignado (siendo $\sim$ cualquier functor de negación).
- (7) Ninguna fórmula $\lceil p \rceil$ es tal que $\lceil p \rceil$ sea, a la vez, designado y antidesignado.

Esos siete principios son de desigual valor y aceptabilidad. Por otra parte, es interesante considerar reformulaciones de los seis primeros en las que, en vez de hablarse de **cualquier** functor de negación, sólo se hable de **algún** functor de negación —y, concretamente, de la **supernegación**. A esas reformulaciones las llamaremos: (1'), (2'), (3'), (4'), (5') y (6').

Los principios (1) [y, por tanto, también (1')], (2'), (3'), (4) [y, por tanto, también (4')], (5) y (6) son, todos, válidos para los sistemas A_3 , A_5 , A_r y A_p , que son —todos ellos— sistemas eunómicos paraconsistentes.

Pero, si pasamos a un sistema tensorial —como A_{bp} —, entonces ni (1), ni (1'), ni (2') son principios válidos. Ni son tampoco, en general, principios plausibles, ya que —¡justamente!— las lógicas tensoriales gozan de un enorme atractivo. Así, p.ej., el principio (1) nos obliga a que cualquier enunciado $\lceil p \rceil$ sea negable si no lo es su negación (o sea: a menos que $\lceil p \rceil$ sea afirmable). Mas es muy posible que una oración no sea ni afirmable ni negable, pues tal vez en unos aspectos sea —poco o mucho— verdadera, mientras que, en otros aspectos, es cabalmente falsa, lo cual nos impediría tanto afirmarla como negarla.

En A_{bp} es válido, en cambio, el principio (4) (y (4')), así como también (3'), (5), (6) (y (6')). El principio (3) está rechazado en cualquier sistema contradictorio. El principio (7) es rechazado en cualquier sistema contradictorio verifuncional.

En cuanto a la aplicación de esos principios a sistemas no contradictorios, la abordaremos ya sólo escuetísimamente, con tres ejemplos: la lógica bivalente, L_3 y A_1 .

La lógica bivalente (verifuncional) satisface las siete condiciones. $\mathcal{L}_3$ no satisface ni (1), ni (1'), ni (2), ni (2'), pero satisface (3), (4), (5), (6) y (7). A_1 no satisface ni (1), ni (2) —pero sí (2')—, ni (6) —pero sí el resultado de restringir (6) precisamente a la **supernegación**—, aunque sí satisface (3), (4), (5) y (7).

El principio sintáctico de tercio excluso

El principio sintáctico de tercio excluso es el esquema: $\lceil p \vee \sim p \rceil$ (donde ' $\vee$ ' es el functor de disyunción natural y ' $\sim$ ' es un functor de negación cualquiera).

En un sistema puede ser válido el principio de tercio excluso para algún functor de negación, pero no para cualquier functor de negación; más concretamente, puede ser válido para la negación natural sin serlo para la supernegación.

Así, p.ej., sea el sistema pentavalente siguiente. Los valores de verdad son: 4, 3, 2, 1 y 0. Sólo 4, 3 y 2 son designados, mientras que 2, 1 y 0 son antidesignados. Tendremos los funtores definidos por las mismas tablas de verdad que en A_5 . En tal sistema (compruébelo el lector) $\lceil p \vee \sim p \rceil$ no es una tautología, aunque $\lceil p \vee Np \rceil$ sí lo es.

Si un sistema eunómico satisface el principio de tercio excluso para cada functor de negación, entonces tal sistema será llamado **ortonómico**.

Los sistemas A_3 , A_5 , Ar , Ap , Abp son, todos ellos, ortonómicos. En todos ellos son tautologías tanto $\lceil p \vee \sim p \rceil$ como $\lceil p \vee Np \rceil$.

La lógica bivalente es también un sistema ortonómico. En cambio no son ortonómicos ni A_1 , ni $\mathcal{L}_3$, ni el sistema pentavalente recién bosquejado. (Los dos primeros no contienen **ningún** principio de tercio excluso).

Antes de concluir, señalemos que algunos sistemas multivalentes, si bien carecen del principio sintáctico de tercio excluso (o sea: no hay en ellos ninguna tautología del tipo $\lceil p \vee \sim p \rceil$), contienen, no obstante, como tautologías principios de cuarto excluso, de quinto excluso, etc. (en general, cada una de esas lógicas es tal que, si es una lógica n-valente, contiene un principio de $(n+1)^\circ$ excluso). La desventaja de **esas** lógicas estriba, no en que posean tales principios, sino en que tan sólo contienen esos principios —o sea: en que carecen del principio de tercio excluso. En cambio, un sistema como A_3 contiene, a la vez, el principio de tercio excluso y un principio de cuarto excluso, que es el siguiente: $\lceil Hp \vee \sim p \vee Sp \rceil$

Y los sistemas Ar , Ap y Abp contienen, para cada n finito igual o superior a 2, un principio de $n+1$ excluso, del tipo: $\lceil p_1 \vee p_2 \vee \dots \vee p_n \rceil$, siendo cada ' p_i ' una fórmula que consiste en una secuencia de funtores monádicos seguida de ' p ', y tal que cada fórmula del tipo $\lceil p_1 \wedge p_j \rceil$ —si $i \neq j$ — es una supercontradicción.

El principio semántico de tercio excluso

Por **principio semántico de tercio excluso** podemos entender varias cosas:

- (1) Para cada fórmula ' p ', o bien $\lceil p \rceil$ es designado o bien $\lceil \sim p \rceil$ es designado (siendo ' $\sim$ ' cualquier functor de negación).
- (2) Para cada fórmula ' p ', o bien $\lceil p \rceil = 1$, o bien $\lceil \sim p \rceil = 1$ (siendo ' $\sim$ ' cualquier functor de negación).
- (3) Para cada fórmula ' p ', no sucede que tanto $\lceil p \rceil$ como $\lceil \sim p \rceil$ sean ambos antidesignados (siendo ' $\sim$ ' cualquier functor de negación).
- (4) Para cada fórmula ' p ', o bien $\lceil p \rceil \neq 0$, o bien $\lceil \sim p \rceil \neq 0$ (siendo ' $\sim$ ' cualquier functor de negación).
- (5) Cada fórmula ' p ' es tal que $\lceil p \rceil$ es antidesignado sólo si $\lceil \sim p \rceil$ es designado (siendo ' $\sim$ ' cualquier functor de negación).
- (6) Cada fórmula ' p ' es tal que $\lceil \sim p \rceil$ es antidesignado sólo si $\lceil p \rceil$ es designado (siendo ' $\sim$ ' cualquier functor de negación).

(7) Cada fórmula $\lceil p \rceil$ es tal que, $/p/$ es designado y/o $/p/$ es antidesignado.

Llamaremos a los resultados de sustituir en los seis primeros de esos siete principios ‘cualquier functor de negación’ por ‘algún functor de negación’, respectivamente: (1’), (2’), (3’), (4’), (5’) y (6’).

Los sistemas escalares A_3 , A_5 , Ar y Ap satisfacen los principios (1), (4), (5’), (6) y (7). (Esos sistemas son, todos ellos, ortonómicos y contradictorios.)

El sistema tensorial Abp no satisface (1) (ni (1’)) ni (7); pero satisface (4), (5’) y (6’).

Cierto es que podemos introducir en esos sistemas un functor monádico ‘—’, definido así: $\lceil \text{—}p \rceil$ abr $\lceil L N p \rceil$

Pero ese functor monádico, ‘—’, no respetaría todas las condiciones que impusimos en el Capítulo 4 para los funtores de negación (no respetaría íntegramente la condición (03), ya que $/p/$ pudiera ser $\frac{1}{2}$, p.ej., —o, en Abp , la secuencia $\langle \frac{1}{2}, \frac{1}{2}, \frac{1}{2}, \dots \rangle$ — que es un valor designado, sin que $/\text{—}p/$ fuera un valor antidesignado, ya que, en ese caso, $/\text{—}p/ = 1$).

Podemos, empero, flexibilizando un poquito las normas para la negación, considerar que el functor ‘—’ es un functor de negación. Si así lo hacemos, entonces los sistemas escalares A_3 , A_5 , Ar y Ap satisfacen también las condiciones (2’), (3’).

Podemos también debilitar las condiciones (1), (2’), (3’) y (7), transformándolas en (1²), (2²), (3²) y (7’) como sigue. Notaremos con el símbolo $/p/_{\text{i}}$ el i -ésimo componente de un valor de verdad. (En una lógica escalar, el i -ésimo componente de un valor de verdad será su único componente, o sea: ese mismo valor de verdad en cuestión.) Y llamaremos componente designado de un valor de verdad a un componente que es un valor designado en la lógica escalar de base, a partir de la cual se ha constituido la lógica tensorial en cuestión. (O —dicho de otro modo—: llamaremos designado a un componente j ssi es designado el valor alético uniformemente constituido por j —repetido tantas veces cuantos componentes tenga un valor de verdad en la lógica tensorial en cuestión.) Luego, en las formulaciones de (1), (2’), (3’) y (7) sustituimos $/p/$ por $/p/_{\text{i}}$, e igualmente $/\text{—}p/$ por $/\text{—}p/_{\text{i}}$, y así hemos formulado (1²), (2²), (3²) y (7’).

Cabe, entonces, decir que un sistema tensorial como Abp satisface (1²), (2²), (3²) y (7’). (Notemos, entre paréntesis, que, mediante ese procedimiento, las versiones (1) y (2’) del principio **semántico** de no contradicción —que estudiamos dos epígrafes atrás—, las cuales eran válidas en A_3 , A_5 , Ar y Ap , pero no en Abp , pasan a ser —reformulados como se acaba de indicar— válidas también con respecto a Abp .)

Y, para terminar, y de la manera más escueta, indiquemos cómo se comportan ciertos sistemas con respecto a las siete versiones del principio semántico de tercio excluso.

La lógica bivalente respeta las siete condiciones.

L_3 no respeta ni (1), ni (1’), ni (2), ni (2’), ni (7); pero respeta (3), (4), (5) y (6’).

A_1 no respeta (1), ni (1’), ni (2) (aunque sí (2’) si también en A_1 introducimos definicionalmente el functor ‘—’, definido como para A_3 , considerándolo un functor de negación) ni (6) (aunque sí (6’)) ni (7). Pero satisface (3), (4) y (5’).

Relación entre los principios de no-contradicción y de tercio excluso

No entraremos aquí en el debate filosófico en torno al principio (sintáctico) de tercio excluso. Al autor de este estudio —como quizá a la mayoría de los locutores de la lengua natural— tal principio le parece un principio obvio, si los hay.

Lo que aquí nos interesa es mostrar que los principios de no-contradicción y tercio excluso son equivalentes, si se aceptan otros dos principios, a saber:

la ley de De Morgan: $\lceil p \vee q \rceil \sim (\sim p \wedge \sim q)$ (o, expresada semánticamente: $\lceil p \vee q \rceil = \lceil \sim(\sim p \wedge \sim q) \rceil$;

y la ley involutiva de la negación (simple): $\lceil p \vdash \neg\neg p \rceil$ (o expresada semánticamente: $\lceil p \rceil = \lceil \neg\neg p \rceil$).

En virtud de esas dos leyes y de la sustituibilidad de los equivalentes, se verá que los principios (sintácticos) de no-contradicción y de tercio excluso son equivalentes entre sí.

Veamos cuál es la relación que se da entre las versiones de esos dos principios para el functor de negación fuerte o supernegación ($\neg$), versiones que podemos llamar, respectivamente, principio reforzado de tercio excluso y principio débil de no-contradicción —o principio de no-supercontradicción—; esa denominación se debe a lo siguiente: el principio reforzado de tercio excluso — $\lceil p \vee \neg p \rceil$ — **implica** el principio simple o normal de tercio excluso — $\lceil p \vee Np \rceil$ — sin ser implicado por él; en cambio, el principio de no-supercontradicción (o principio débil de no-contradicción) es implicado por el principio simple o normal de no-contradicción, pero no implica a éste último.

Tales principios **no** son equivalentes entre sí, pero sí están ligados —en los sistemas A_3 , A_5 , Ar , Ap , Abp — por un bicondicional, formando la siguiente tautología de esos sistemas: $\lceil p \vee \neg p \equiv \neg(p \wedge \neg p) \rceil$

Ello se prueba de modo similar —pero no idéntico—, en virtud de las siguientes fórmulas bicondicionales válidas —que son tautologías de todos esos sistemas—: $\lceil p \vee q \equiv \neg(\neg p \wedge \neg q) \rceil$
 $\lceil p \equiv \neg\neg p \rceil$

(En **cada uno** de esos sistemas, una fórmula bicondicional es tautológica —¡no se olvide!— sólo si, en el caso de que el miembro derecho sea designado, también lo es el izquierdo; y, en el caso de que el miembro izquierdo sea designado, también lo es el derecho.)

De todo ello se desprende que, si hay motivos para negar el principio de tercio excluso, también los hay para negar el principio de no-contradicción. Pero es muy común pensar que hay situaciones que ni se dan ni dejan de darse; se dice, contestando a la pregunta de si Grecia es un país industrializado, que ni sí ni no. Pero, como acabamos de ver, ‘ni sí ni no’ —o sea: ‘no (sí o no)’—, en virtud de una ley de De Morgan— equivale a ‘sí y no’, o sea a una antinomia (que es la negación de una instancia del principio de no contradicción).

Por otro lado, aun quien no acepte alguno de los principios involucrados en nuestra prueba —una de las leyes de De Morgan y la ley involutiva de la negación simple— deberá, empero, aceptar una inconsistencia simple siempre que acepte: 1º) el principio de tercio excluso; 2º) la negación de alguna instancia de ese principio (como que el Chad ni es ni deja de ser un país árabe). Además, si alguien acepta una inconsistencia simple y también la regla de adjunción ($p, q \vdash p \wedge q$) entonces no podrá por menos de aceptar una antinomia. Pero la regla de adjunción es una de las reglas cuyo rechazo parece más sobrecogedor desde un punto de vista previo o externo a la construcción de sistemas formales.

La conclusión es que quienquiera que acepte situaciones difusas —situaciones que ni tienen lugar ni dejan de tener lugar— debe, para ser consecuente, aceptar un sistema contradictorio.

Capítulo 8

Sistemas lógicos deductivos: el sistema At

Ya vimos anteriormente que una teoría lógica puede definirse de dos maneras:

- 1) semánticamente (como el conjunto de tautologías que se engendran en un dominio de valores de verdad en virtud: 1º, del otorgamiento de la calidad de valor designado a determinados valores pertenecientes al dominio; y 2º, de ciertas asignaciones de valores de verdad);
- 2) sintácticamente: estableciendo un conjunto enumerable de axiomas, más determinada(s) regla(s) de inferencia, mediante la(s) cual(es) se pueden obtener otros teoremas que no son axiomas.

En este capítulo vamos a estudiar sistemas definidos sintácticamente. A esos sistemas los llamaremos: sistemas lógicos deductivos. Lo que aquí nos interesa, pues, es la deducción de teoremas a partir de otros teoremas —y, en última instancia, a partir de axiomas—, mediante reglas de inferencia; y también estudiaremos cómo se pueden derivar ciertas reglas de inferencia a partir de otras —y, en última instancia, a partir de reglas de inferencia primitivas, con la ayuda de teoremas.

At es un sistema trivalente, si bien no es un sistema completo; el conjunto de los teoremas de At es un subconjunto **propio** del conjunto de las tautologías de A_3 . (Podemos transformar At en un conjunto completo —tal que el conjunto de sus teoremas coincida con el conjunto de las tautologías de A_3 — sustituyendo el miembro conjuntivo izquierdo del axioma A_{13} de At —vide infra— por la fórmula $\lceil q \downarrow q \downarrow p \rceil$. Pero esa transformación no ofrece demasiado interés, toda vez que, de hacerla, ya no tendríamos el resultado apetecido de que el sistema Ap —del que se hablará en seguida— sea una extensión de At).

Emplearemos siempre las letras 'p', 'q', 'r' etc. como letras esquemáticas.

Lecturas de signos a emplear

Conviene recordar, ante todo, las lecturas de los signos a emplear —unos como primitivos, otros como definidos— en la exposición de At y Ap . En cada caso, escribimos a la izquierda la notación simbólica, y a la derecha su lectura en lengua natural.

$\lceil Np \rceil$: «No sucede que p» $\approx$ «Sucedé que no-p» $\approx$ «No es cierto que p» $\approx$ «Es falso que p» $\approx$ «No-p»;

$\lceil \neg p \rceil$: «Es enteramente falso que p» $\approx$ «No es cierto en absoluto que p» $\approx$ «Es cien por cien falso que p» $\approx$ «Es totalmente falso que p» $\approx$ «Es plenamente falso que p» $\approx$ «No sucede en absoluto que p» $\approx$ «Es enteramente cierto que no-p», etc.;

$\lceil Sp \rceil$: «No es ni cierto ni falso que p» $\approx$ «Es cierto y falso a la vez que p» $\approx$ «Sucedé y no sucede que p» $\approx$ «Sucedé que p y que no-p» $\approx$ «Ni sucede ni deja de suceder que p».

$\lceil Lp \rceil$: «Es más o menos cierto que p» $\approx$ «Es, por lo menos hasta cierto punto, verdad que p» $\approx$ «Sucedé, por lo menos en alguna medida, que p»;

$\lceil Yp \rceil$: «Es un sí es no cierto que p» $\approx$ «Es infinitesimalmente cierto que p»;

$\lceil fp \rceil$: «Es más que infinitesimalmente cierto que p» $\approx$ «Es un tanto cierto que p»;

$\lceil np \rceil$: «Es supercierto que p»;

$\lceil mp \rceil$: «Viene a ser cierto que p»;

$\lceil p \wedge q \rceil$: «p y q»;

$\lceil p \vee q \rceil$: «p o q» $\approx$ «p y/o q»;

$\lceil p \bullet q \rceil$: «No sólo p, sino que también q» $\approx$ «p así como también q»;

$\lceil plq \rceil$: «p en la misma medida en que q» $\approx$ «Es cierto que p en la misma medida en que lo es que q» $\approx$ «El hecho de que p equivale al hecho de que q»;

- $\lceil p \rightarrow q \rceil$: «Sucede que p a lo sumo en la medida en que sucede que q» $\approx$ «El hecho de que p implica el hecho de que q» $\approx$ «[Por lo menos] en tanto en cuanto es [o sea] verdad que p, q»;
- $\lceil p \setminus q \rceil$: «Es menos cierto que p que (que) q» $\approx$ «Es más cierto que q que (que) p»;
- $\lceil p \supset q \rceil$: «p sólo si q» $\approx$ «Sucede que p con tal de que también suceda que q» $\approx$ «Si p, entonces q» $\approx$ «El hecho de que p entraña el hecho de que q»;
- $\lceil p \equiv q \rceil$: «p ssi q» $\approx$ «El hecho de que p y el hecho de que q se entrañan mutuamente»;
- $\lceil Hp \rceil$: «Es enteramente cierto que p» $\approx$ «Es ciento por ciento verdad que p» $\approx$ «Es plenamente (cabalmente, completamente) cierto que p»;
- $\lceil p \& q \rceil$: «Sucediendo que p, q» $\approx$ «p y, sobre todo, q»;
- $\lceil p \vee q \rceil$: «Es cabalmente cierto que p, a menos que q»;
- $\lceil \neg p \rceil$: «Es más o menos falso que p» $\approx$ «No es enteramente cierto que p»;
- '0': 'Lo absolutamente falso' $\approx$ 'Lo absolutamente irreal';
- '1': 'Lo absolutamente verdadero' $\approx$ 'Lo absolutamente real';
- 'a': 'Lo infinitesimalmente real' $\approx$ 'Lo infinitesimalmente verdadero' $\approx$ 'El grado mínimo ($\approx$ ínfimo) de verdad (de realidad)';
- ' $\frac{1}{2}$ ': 'Lo igualmente verdadero que falso' $\approx$ 'Lo tan real como irreal'.

EL SISTEMA *At*

Reglas de formación

- 1.— Si $\lceil p \rceil$ y $\lceil q \rceil$ son fbfs, también lo son $\lceil \neg p \rceil$, $\lceil Np \rceil$, $\lceil p \wedge q \rceil$ y $\lceil plq \rceil$.
- 2.— 'a' es una bbf.

Definiciones

(Ya se sabe que una definición es una mera abreviación: la expresión que se halla entre esquinas a la izquierda abrevia a la que se halla —también entre esquinas— a la derecha, estando unidas ambas por el signo sintáctico o metalingüístico 'abr'):

df02 $\lceil p \vee q \rceil$ abr $\lceil N(Np \wedge Nq) \rceil$	df03 $\lceil p \supset q \rceil$ abr $\lceil \neg p \vee q \rceil$
df04 $\lceil p \rightarrow q \rceil$ abr $\lceil p \wedge qlp \rceil$	df05 '0' abr $\lceil \neg a \wedge a \rceil$
df06 '1' abr $\lceil N0 \rceil$	df07 $\lceil p \equiv q \rceil$ abr $\lceil p \supset q \wedge q \supset p \rceil$
df08 $\lceil Lp \rceil$ abr $\lceil \neg \neg p \rceil$	df09 $\lceil Sp \rceil$ abr $\lceil p \wedge Np \rceil$
df10 $\lceil Hp \rceil$ abr $\lceil \neg Np \rceil$	df11 $\lceil p \& q \rceil$ abr $\lceil N(p \supset Nq) \rceil$
df12 $\lceil p \setminus q \rceil$ abr $\lceil N(Np \& Nq) \rceil$	df14 ' $\frac{1}{2}$ ' abr $\lceil ala \rceil$
df13 $\lceil \neg p \rceil$ abr $\lceil LNp \rceil$	

Esquemas axiomáticos de *At*

- A01 $plq \supset \neg pl \neg q$
 A02 $plq \supset rlql.plr$
 A03 $p \wedge lp$
 A04 $p \vee q \wedge plp$

- A05 $plq \supset p \wedge r \wedge sl.q \wedge r \wedge s$
 A06 $q \vee r \wedge pl.p \wedge q \vee p \wedge r$
 A07 $\neg p \vee \neg ql \neg (p \wedge q)$

A08 $\neg p \wedge \neg q \vdash (p \vee q)$

A09 $N \neg p \vdash \neg p$

A10 $p \vdash q \supset p$

A11 $p \vdash Nq \vdash Np$

A12 $p \wedge q \supset p$

A13 $p \vdash p \wedge (p \vee q) \supset p$

A14 $Lp \vee q \supset Np$

Regla de inferencia primitiva (única)

rinf01: $p \supset q, p \vdash q$

Desarrollo del sistema

La base de un sistema deductivo, como *At*, está constituida por sus axiomas y sus reglas de inferencia primitivas. Se desarrolla el sistema demostrando teoremas (no todos los teoremas son axiomas, ni de lejos), y derivando reglas de inferencia no primitivas.

En esos procesos de demostración y derivación seguiremos las normas que se indican a continuación:

1º.— Cada prueba o derivación consta de varias líneas. A la izquierda de cada fórmula expuesta en una línea —y separado de ella— figura, encerrado entre paréntesis, un número entero positivo que se considerará como **un nombre de la fórmula** en cuestión. (En esos números sólo figurarán los guarismos del 2 al 9 por las razones abajo indicadas —en el punto 4º.)

2º.— En cada línea, a la derecha de la fórmula que en ella se asevera —y separados de ella—, figuran los nombres de los teoremas, reglas de inferencia y/o fórmulas previamente probadas —éstas últimas, probadas dentro de la misma prueba o derivación— (y, asimismo, de las hipótesis de que se parte, en el caso de que se trate, no de una prueba de un teorema, sino de una derivación de una regla de inferencia) que justifican la aserción de la fórmula expuesta en esa línea.

3º.— Si al nombre de una fórmula expuesta en una línea (o sea: al número entero positivo encerrado entre paréntesis que figura en la línea, a la izquierda de la fórmula en cuestión) se le quita el par de paréntesis, el resultado es una **abreviación** de la fórmula por él nombrada. (Una abreviación de una fórmula es algo diferente del nombre de la fórmula.) En cambio, los nombres de teoremas (cada uno de los cuales consta de una 'A' seguida de un número entero positivo) se utilizarán también como abreviaciones de los teoremas que ellos nombran —para mayor simplicidad. Todas esas abreviaciones —tanto las de líneas previamente probadas o deducidas dentro de la misma prueba o derivación, como las de teoremas previamente probados en el desarrollo del sistema— podrán aparecer en líneas ulteriores, como subfórmulas de la fórmula total que figure en tal línea ulterior. (Pero las abreviaciones de líneas de una prueba sólo podrán ser utilizadas en líneas ulteriores de la misma prueba.)

4º.— Como '1' y '0' son dos signos de los sistemas *At* y *Ap* (dos constantes sentenciales definidas), se ha preferido —para evitar confusiones— que esos dos guarismos no figuren en ninguna abreviación de fórmula alguna expuesta en una línea de una prueba o derivación, ni, por consiguiente, tampoco en el nombre de la misma.

5º.— Presentamos separadamente la derivación de reglas de inferencia no primitivas y la demostración de teoremas. Ahora bien, en la derivación de muchas reglas de inferencia no primitivas se presuponen, como **ya** demostrados, ciertos teoremas; y en la demostración de muchos teoremas se presuponen, como **ya** derivadas, reglas de inferencia no primitivas. Pero el lector podrá observar cuidadosamente, al toparse —en la derivación de una regla de

inferencia— con una referencia justificatoria a un teorema que no sea un axioma, que la regla de inferencia que se está derivando no ha sido utilizada **ni** en la prueba del aludido teorema, **ni** tampoco —¡claro está!— en la prueba de ningún otro teorema anterior al mismo; y podrá asimismo comprobar, al toparse —en la prueba de un teorema— con una referencia justificatoria a una regla de inferencia no primitiva, que, **ni** en la derivación de la misma **ni** en la de ninguna otra regla de inferencia anterior, se ha aducido el teorema que se está demostrando. (Si se hubiera hecho alguna de esas dos cosas, se hubiera incurrido en una viciosa circularidad, que invalidaría la prueba o derivación.)

6º.— Cada fórmula es, o bien una fórmula atómica, o bien una fórmula que comienza por una ocurrencia de un functor monádico, o bien una fórmula que tiene como functor principal una ocurrencia de un functor diádico; en este último caso, la fórmula tendrá un miembro derecho y otro izquierdo. Sea un número cualquiera, '3', p.ej., una abreviación de tal fórmula; entonces 'σ3' será una abreviación de su miembro izquierdo (la sigma se usa con relación a la palabra griega 'σκούτον', 'izquierdo'); al paso que 'δ3' será una abreviación de su miembro derecho (la delta viene usada por relación a 'δεξιόν', 'derecho'). También podrá escribirse 'σσ3', 'σδ3', 'δσ3', 'σσδ3', y así sucesivamente, con significados claros y obvios.

Derivación de reglas de inferencia

rinf11 $plq \vdash qlp$

Derivación:

hip: plq

qlp hip, A104, rinf01

rinf12 $plq \vdash p \supset q$

Derivación:

hip: plq

(2) qlp

hip, rinf11

$p \supset q$

(2), A10, rinf01

rinf13 $p, plq \vdash q$

(Derivación: rinf12, rinf 01)

rinf14 $p, qlp \vdash q$

(Derivación: rinf11, rinf13)

rinf15 $plq, qlr \vdash plr$

Derivación:

hip.2^a qlr

hip.1^a plq

(2) $rlql.plr$

hip.1^a, A02, rinf01

(3) $rlq \supset .plr$

(2), rinf12

(4) rlq

hip.2^a, rinf11

plr

(3), (4), rinf01

rinf16 $plq \vdash p \wedge rl.q \wedge r$

Derivación:

hip: plq

(2) $p \wedge r \wedge 1l.q \wedge .r \wedge 1$

hip, A05

(3) $p \wedge rl.q \wedge .r \wedge 1$

(2), A118, rinf11, rinf15

(4) $q \wedge r \wedge 1l.q \wedge .r \wedge 1$

A05, A101, rinf01

(5) $q \wedge (r \wedge 1)l.q \wedge r \wedge 1$

(4), rinf11

(6) $q \wedge r \wedge 1l.q \wedge r$	A118	
(7) $q \wedge (r \wedge 1)l.q \wedge r$	(5), (6), rinf15	
$p \wedge rl.q \wedge r$	(3), (7), rinf15	
rinf17 $plq \vdash p \vee rl.q \vee r$		
Derivación:		
hip: plq		
(2) $Np \vdash Nq$	hip, A109, rinf13	
(3) $Np \wedge Nrl.Nq \wedge Nr$	(2), rinf16	
$p \vee rl.q \vee r$	(3), A110, df02, rinf01	
rinf18 $plq \vdash q \wedge rl.p \wedge r$	(Derivación: rinf11, rinf16)	
rinf19 $plq \vdash q \vee rl.p \vee r$	(Derivación: rinf11, rinf17)	
rinf20 $plq \vdash r \wedge pl.r \wedge q$		
Derivación:		
hip: plq		
(2) $p \wedge rl.q \wedge r$	hip, rinf16	
(3) $q \wedge rl.p \wedge r$	hip, rinf18	
(4) $q \wedge rl.r \wedge q$	A121	
(5) $r \wedge pl.p \wedge r$	A121	
(6) $r \wedge pl.q \wedge r$	(5), (2), rinf15	
$r \wedge pl.r \wedge q$	(6), (4), rinf15	
rinf21 $plq \vdash r \vee pl.r \vee q$		
(Derivación similar a la de rinf20, utilizando A122 en vez de A121, rinf17 en vez de rinf16, y rinf19 en vez de rinf18)		
rinf22 $p, q \vdash p \wedge q$		
Derivación:		
hip1 ^a p		
hip2 ^a q		
(2) $q \supset.p \wedge q$	hip1 ^a , A135, rinf01	
$p \wedge q$	(2), hip2 ^a , rinf01	
rinf23 $plq \vdash p \supset rl.q \supset r$		
Derivación:		
hip. plq		
(2) $\neg pl \neg q$	hip, A01, rinf01	
(3) $\neg p \vee rl.\neg q \vee r$	(2), rinf17	
$p \supset rl.q \supset r$	(3), df03	
rinf24 $plp^1, p^1lp^2, p^2lp^3, \dots, p^{n-1}lp^n \vdash plp^n$		
Derivación:		
hip1 ^a plp^1		
hip2 ^a p^1lp^2		
hip3 ^a p^2lp^3		

.	
.	
hipn ^a	p ⁿ⁻¹ lp ⁿ
(2)	plp ² hip.1 ^a , hip.2 ^a , rinf15
(3)	plp ³ (2), hip.3 ^a , rinf15
.	
.	
.	
(n)	plp ⁿ (n-1), hip.n ^a , rinf15

rinf25 plq $\vdash$ rlr¹ (si 'r' es una fórmula en la que 'p' figura afectado sólo por
ocurrencias de 'I', '¬', 'N', '∧', mientras que 'r¹' sólo difiere de 'r' por el reemplazamiento
de un número finito cualquiera de ocurrencias de 'p' en 'r' por otras tantas ocurrencias
respectivas de 'p')

La derivación se efectúa por **Inducción matemática**. La inducción matemática es un procedimiento que se funda en el principio siguiente (principio de inducción matemática). Supongamos que se quiere probar que un teorema vale para un número cualquiera de entes que satisfagan determinada condición. Para probarlo basta con demostrar:

- 1º) que, para al menos un ente que satisfaga la condición en cuestión, el teorema es correcto;
- 2º) que, si es correcto para un número n —cualquiera que sea n— de entes que satisfagan la condición, también valdrá para n+1.

Ahora bien, derivar una regla de inferencia es probar un teorema sintáctico —o, si se quiere, metalingüístico— que dice: si una o varias fórmulas de tal o cual tipo son dadas como premisas, entonces otra fórmula de determinado tipo es obtenida como conclusión.

La regla rinf25 nos dice: si una fórmula del tipo 'plq' es una premisa, entonces una fórmula del tipo 'rlr¹' es obtenible de ella como conclusión correcta (siempre y cuando 'r' y 'r¹' sean como se indica en la explicación añadida a la regla).

Vamos a probar, en primer lugar, que la regla vale para el caso en que 'r¹' sólo difiera de 'r' en la sustitución de una sola ocurrencia de 'p' por otra de 'q'; y luego que, si la regla vale para n sustituciones, entonces vale también para n+1 sustituciones.

Como los únicos funtores que hemos introducido como primitivos en At son 'I', '¬', 'N' y '∧', si la regla de inferencia en cuestión vale para ellos, valdrá para cualquier fórmula engendrada de conformidad con nuestras dos reglas de formación.

Vamos a derivar la regla por partes. Primero la probaremos para el caso de que 'r¹' sólo difiera de 'r' por el reemplazamiento de **una** sola ocurrencia de 'p' en 'r' por una ocurrencia respectiva de 'q'.

Pero también aquí iremos por partes, y acudiremos a una inducción matemática particular incrustada dentro de la inducción matemática general en que consiste toda la derivación.

Esta inducción matemática **particular** estriba en lo siguiente: Se prueba, primero, que la regla vale para el caso de que 'p' esté afectada en 'r' por una sola ocurrencia de uno de los cuatro funtores primitivos. (Una fórmula 'p' está afectada por una ocurrencia de un functor '✱' ssi 'p' es miembro —derecho o izquierdo— de un miembro —derecho o izquierdo— de un miembro —derecho o izquierdo—... de la mencionada ocurrencia de '✱' —o sea: de una fórmula cuyo functor principal es dicha ocurrencia de '✱'; y diremos que una fórmula cualquiera

precedida inmediatamente por una ocurrencia de un functor monádico es miembro derecho de tal ocurrencia.)

En segundo lugar, se prueba que, si la regla de inferencia —restringida a una sola sustitución de $\ulcorner p \urcorner$ por $\ulcorner q \urcorner$ — vale para el caso de que $\ulcorner p \urcorner$ esté afectado en $\ulcorner r \urcorner$ por n ocurrencias de cualesquiera de los cuatro funtores primitivos, entonces también vale para el caso de que $\ulcorner p \urcorner$ esté afectado en $\ulcorner r \urcorner$ por $n+1$ ocurrencias de cualesquiera de esos cuatro funtores primitivos.

Así, pues, empecemos haciendo la derivación para el primer caso del primer caso:

hip. $\text{pl}q$

(2) $\text{NplN}q$	hip., A110, rinf01
(3) $\neg \text{pl} \neg q$	hip., A01, rinf01
(4) $\text{p} \wedge \text{sl} . q \wedge \text{s}$	hip., rinf16
(5) $\text{s} \wedge \text{pl} . \text{s} \wedge q$	hip., rinf20
(6) $\text{slql} . \text{pls}$	hip., A02, rinf01
(7) $\text{plsl} . \text{slp}$	A103
(8) $\text{slpl} . \text{slq}$	(6), (7), rinf15, rinf11
(23) $\text{plsl} . \text{slq}$	(6), rinf11
(24) $\text{slql} . \text{qls}$	A103
(25) $\text{plsl} . \text{qls}$	(23), (24), rinf15

Así pues, hemos probado lo siguiente: de la premisa $\ulcorner \text{pl}q \urcorner$ se desprenden las conclusiones $\ulcorner \text{NplN}q \urcorner$ (2); $\ulcorner \neg \text{pl} \neg q \urcorner$ (3); $\ulcorner \text{p} \wedge \text{sl} . q \wedge \text{s} \urcorner$ (4); $\ulcorner \text{s} \wedge \text{pl} . \text{s} \wedge q \urcorner$ (5); $\ulcorner \text{plsl} . \text{qls} \urcorner$ (25); $\ulcorner \text{slpl} . \text{slq} \urcorner$ (8). Ahora bien, esas seis fórmulas son **todas** las fórmulas del tipo $\ulcorner \text{rlr}^1 \urcorner$ en las que $\ulcorner r^1 \urcorner$ difiere de $\ulcorner r \urcorner$ por la sustitución de una sola ocurrencia de $\ulcorner p \urcorner$ por una sola ocurrencia de $\ulcorner q \urcorner$, y estando en cada caso afectado por una sola ocurrencia de, respectivamente, 'N', '¬', '∧' y 'l' —y, en los dos últimos casos, estando afectado, ya como miembro derecho, ya como miembro izquierdo. Por tanto, para una sola ocurrencia de $\ulcorner p \urcorner$ en $\ulcorner r \urcorner$, y para el caso de que $\ulcorner p \urcorner$ esté afectado en $\ulcorner r \urcorner$ por una sola ocurrencia de alguno de esos cuatro funtores, la regla de inferencia es válida.

Veamos ahora cómo se generaliza, suponiendo siempre que la sustitución se efectúe sobre una sola ocurrencia de $\ulcorner p \urcorner$ en $\ulcorner r \urcorner$. Lo que ahora hay que probar es que, si la regla vale —siempre para una sola sustitución de $\ulcorner p \urcorner$ por $\ulcorner q \urcorner$ — cuando $\ulcorner r \urcorner$ contiene n ocurrencias de cualesquiera de esos cuatro funtores que estén afectando a $\ulcorner p \urcorner$, también vale entonces para cuando $\ulcorner r \urcorner$ contiene, **además**, una ocurrencia suplementaria de '¬', o de 'N', o de '∧', o de 'l'.

Supongamos, pues, que se ha probado ya que la regla es válida para n ocurrencias de cada uno de esos funtores, o sea que, de la hipótesis, se ha deducido (26), a saber:

(26) rlr^1 (siendo $\ulcorner r^1 \urcorner$ el resultado de reemplazar una ocurrencia de $\ulcorner p \urcorner$ por otra de $\ulcorner q \urcorner$, y estando afectado $\ulcorner p \urcorner$ en $\ulcorner r \urcorner$ por n ocurrencias de cualesquiera de esos cuatro funtores).

Ahora se trata de deducir (27) a partir de (26):

(27) sls^1 (siendo $\ulcorner s \urcorner$ una de estas fórmulas: $\ulcorner \text{Nr} \urcorner$, $\ulcorner \neg r \urcorner$, $\ulcorner \text{p}^1 \wedge r \urcorner$, $\ulcorner r \wedge \text{p}^1 \urcorner$, $\ulcorner \text{p}^1 \text{lr} \urcorner$, $\ulcorner \text{rlp}^1 \urcorner$; y siendo $\ulcorner s^1 \urcorner$ el resultado de reemplazar $\ulcorner r \urcorner$ por $\ulcorner r^1 \urcorner$ en $\ulcorner s \urcorner$ —o sea: el resultado de reemplazar **una** ocurrencia de $\ulcorner p \urcorner$ en $\ulcorner s \urcorner$ por una ocurrencia de $\ulcorner q \urcorner$).

Pues bien, se pasa de (26) a (27) del mismo modo que de la hipótesis originaria se pasaba a (2), (3), (4), (5), (25) y (8).

Así pues, de la premisa $\ulcorner \text{pl}q \urcorner$ hemos deducido (suponiendo que la fórmula $\ulcorner r^1 \urcorner$ difiere de $\ulcorner r \urcorner$ tan sólo por la sustitución de una ocurrencia de $\ulcorner p \urcorner$ en $\ulcorner r \urcorner$ por otra de $\ulcorner q \urcorner$) que:

1º) Si $\ulcorner p \urcorner$ está afectado en $\ulcorner r \urcorner$ por una sola ocurrencia de cualquiera de los cuatro funtores primitivos, podemos concluir $\ulcorner rlr^1 \urcorner$;

2º) Si concluir $\ulcorner rlr^1 \urcorner$ está justificado cuando $\ulcorner p \urcorner$ está afectado en $\ulcorner r \urcorner$ por n ocurrencias de cualesquiera de esos cuatro funtores primitivos, entonces el concluir $\ulcorner rlr^1 \urcorner$ estará también justificado cuando $\ulcorner p \urcorner$ esté afectado en $\ulcorner r \urcorner$ por $n+1$ ocurrencias de cualesquiera de esos cuatro funtores primitivos.

Por consiguiente —y en virtud del principio ya explicado de la inducción matemática—, el tránsito de la premisa $\ulcorner plq \urcorner$ a la conclusión $\ulcorner rlr^1 \urcorner$ está siempre justificado, siendo $\ulcorner r \urcorner$ una fórmula cualquiera, con tal de que $\ulcorner r^1 \urcorner$ sólo difiera de $\ulcorner r \urcorner$ por la sustitución de **una** ocurrencia de $\ulcorner p \urcorner$ en $\ulcorner r \urcorner$ por una ocurrencia respectiva de $\ulcorner q \urcorner$.

Nos falta ahora probar el segundo paso de la inducción matemática general en que consiste nuestra derivación. Suponemos el antecedente (y seguimos suponiendo la hipótesis originaria, o sea: la premisa $\ulcorner plq \urcorner$); ese antecedente será (28):

(28) rlr^1 (siendo $\ulcorner r \urcorner$ una fórmula cualquiera, y difiriendo $\ulcorner r^1 \urcorner$ de $\ulcorner r \urcorner$ por la sustitución de n ocurrencias de $\ulcorner p \urcorner$ en $\ulcorner r \urcorner$ por n ocurrencias respectivas de $\ulcorner q \urcorner$).

Queremos probar que también será válida la regla :

$plq \vdash sls^1$

(O sea: queremos deducir $\ulcorner sls^1 \urcorner$, ya que seguimos suponiendo la hipótesis $\ulcorner plq \urcorner$.) La fórmula $\ulcorner s^1 \urcorner$ diferirá de $\ulcorner s \urcorner$ por la sustitución de $n+1$ ocurrencias de $\ulcorner p \urcorner$ en $\ulcorner s \urcorner$ por otras tantas ocurrencias respectivas de $\ulcorner q \urcorner$.

Sea $\ulcorner s^2 \urcorner$ el resultado de sustituir una ocurrencia de $\ulcorner q \urcorner$ en $\ulcorner s^1 \urcorner$ por una de $\ulcorner p \urcorner$; por lo cual $\ulcorner s^2 \urcorner$ será también un resultado de sustituir n ocurrencias de $\ulcorner p \urcorner$ en $\ulcorner s \urcorner$ por n ocurrencias respectivas de $\ulcorner q \urcorner$. Ya sabemos (en virtud de (28)) que:

(29) sls^2

Pero $\ulcorner s^1 \urcorner$ sólo difiere de $\ulcorner s^2 \urcorner$ por la sustitución de **una** ocurrencia de $\ulcorner p \urcorner$ por una ocurrencia respectiva de $\ulcorner q \urcorner$. Luego —en virtud de la hipótesis originaria ($\ulcorner plq \urcorner$) y del primer paso, ya demostrado de toda la derivación:

(32) s^2ls^1

(33) sls^1 (29), (32), rinf15

Por consiguiente, de (28) se extrae (33); o sea: si la regla vale para el caso de que $\ulcorner r^1 \urcorner$ sólo difiere de $\ulcorner r \urcorner$ por la sustitución de n ocurrencias de $\ulcorner p \urcorner$ por n ocurrencias respectivas de $\ulcorner q \urcorner$, entonces también vale para el caso de que $\ulcorner r^1 \urcorner$ difiera de $\ulcorner r \urcorner$ por la sustitución de $n+1$ ocurrencias de $\ulcorner p \urcorner$ por $n+1$ ocurrencias respectivas de $\ulcorner q \urcorner$. Y con ello ha quedado demostrado el segundo paso de la inducción matemática general.

Con lo cual queda concluida la inducción matemática general que habíamos abordado. Es decir, queda derivada la regla rinf25.

rinf26 $p \supset q, q \supset r \vdash p \supset r$

Derivación:

hip1ª $p \supset q$

hip2ª $q \supset r$

(2) $p \supset q \supset . q \supset r \supset . p \supset r$

A152

(3) $q \supset r \supset . p \supset r$

(2), hip.1ª, rinf01

$p \supset r$

(3), hip.2ª, rinf01

rinf27 $p \supset p^1, p^1 \supset p^2, \dots, p^{n-1} \supset p^n \vdash p \supset p^n$

Derivación:

hip1^a $p \supset p^1$

hip2^a $p^1 \supset p^2$

.

.

.

hipn^a $p^{n-1} \supset p^n$

(2) $p \supset p^2$

hip.1^a, hip.2^a, rinf26

.

.

.

(n—1) $p \supset p^{n-1}$

hip.(n—1)^a, (n—2), rinf26

$p \supset p^n$

(n—1), hip.n^a, rinf26

rinf28 $p \supset q \supset r, r \supset r^1, r^1 \supset r^2, \dots, r^{n-1} \supset r^n \vdash p \supset q \supset r^n$

Derivación:

hip1^a $p \supset q \supset r$

hip2^a $r \supset r^1$

hip3^a $r^1 \supset r^2$

.

.

.

hip(n+1)^a $r^{n-1} \supset r^n$

(2) $r \supset r^n$

hip.2^a... hip.(n+1)^a, rinf27

(3) $q \supset r \supset q \supset r^n$

A153, (2), rinf01

(4) $p \supset (q \supset r) \supset p \supset q \supset r^n$

(3), A153, rinf01

$p \supset q \supset r^n$

(4), hip.1^a, rinf01

rinf29 $p \supset r^1, r^1 \supset r^2, \dots, r^{n-1} \supset r^n, p \supset q \supset q^1 \supset r \vdash p \supset q \supset q^1 \supset r^n$

Derivación: A161, rinf27

rinf30 $r \supset r^1, r^1 \supset r^2, \dots, r^{n-1} \supset r^n, p \supset q \supset q^1 \supset q^2 \supset r \vdash p \supset q \supset q^1 \supset q^2 \supset r^n$

Derivación: A162, rinf27

Demostración de teoremas

A101 plp

Prueba:

(2) $p \wedge plp \supset plpl.p \wedge plp$

A02

(3) $plpl.p \wedge plp$

(2), A03, rinf01

(4) $p \wedge plp \supset plp$

(3), A10. rinf01

plp

(4), A03, rinf01

A102 $p \text{INN} p$

Prueba:

(2) $p \text{INN} p \text{I}. p \text{I} \text{N} p$

A11

(3) $p \text{I} \text{N} p \text{C}. p \text{I} \text{N} p$

(2), A10, rinf01

(4) $p \text{I} \text{N} p$

A101

$p \text{I} \text{N} p$

(3), (4), rinf01

A103 $p \text{I} q \text{I}. q \text{I} p$

Prueba:

(2) $q \text{I} q \supset. p \text{I} q \text{I}. q \text{I} p$

A02

(3) $q \text{I} q$

A101

A103

(2), (3), rinf01

A104 $p \text{I} q \supset. q \text{I} p$

Prueba:

(2) $q \text{I} p \text{I}. p \text{I} q$

A103

(3) $q \text{I} p \text{I} (p \text{I} q) \supset. p \text{I} q \supset. q \text{I} p$

A10

A104

(2), (3), rinf01

A105 $p \text{I} q \text{I}. \text{NN} q \text{I} p$

Prueba:

(2) $\text{NN} q \text{I} q$

A104, rinf01, A102

A105

A02, rinf01, (2)

A106 $p \text{I} q \text{I}. p \text{I} \text{NN} q$

Prueba:

(2) $\text{NN} q \text{I} p \text{I}. p \text{I} \text{NN} q$

A103

A106

A105, (2), rinf15

A107 $p \text{I} \text{NN} q \text{I}. p \text{I} q$

Prueba: A106, rinf11

A108 $p \text{I} \text{N} q \text{I}. p \text{I} q$

Prueba:

(2) $p \text{I} \text{NN} q \text{I} (p \text{I} q) \supset. p \text{I} \text{N} q \text{I} (p \text{I} q) \text{I}. p \text{I} \text{NN} q \text{I}. p \text{I} \text{N} q$

A02

(3) $\delta 2$

(2), A107, rinf01

(4) $p \text{I} \text{NN} q \text{I} (p \text{I} \text{N} q) \supset. p \text{I} \text{N} q \text{I}. p \text{I} q$

(3), A10

$\delta 4$

rinf01, A11, (4)

A109 $p \text{I} q \text{I}. p \text{I} \text{N} q$ Prueba: A108, rinf11

A110 $p \text{I} q \supset. p \text{I} \text{N} q$ Prueba: A109, rinf12

A111 $p \text{I} \text{N} q \supset. p \text{I} q$ Prueba: A109, A10, rinf01

A112 $p \vee p \text{I} p$

Prueba:

(2) $p \wedge p \text{I} \text{N} p$

A03

(3) $\text{N}(p \wedge p) \text{I} \text{NN} p$

A110, (2), rinf01

(4) $p \vee p \text{I} \text{NN} p$

(3), df02

A112

(4), A107, rinf13

A113 $p \supset p$ Prueba: A101, rinf12

A114 $\neg p \vee p$ Prueba: A113, df03

A115 $\neg 0$

Prueba:

(2) $\neg \neg a \vee \neg a$
 $\neg 0$

A114

(2), A07, rinf13, df05

A116 $\neg p \supset . q \rightarrow Np$ Prueba: A14, df08, df03

A117 $\neg p \supset . q \wedge Nplq$ Prueba: A116, df04

A118 $p \wedge 1lp$ Prueba: A117, A115, df06, rinf01

A119 $p \vee 0lp$

Prueba:

(2) $Np \wedge 1Np$

A118

(3) $Np \wedge N0INp$

(2), df06

(4) $N(Np \wedge N0)INNp$
 $p \vee 0lp$

(3), A110, rinf01

(4), df02, A107, rinf13

A120 $q \vee q \wedge pl.q \wedge p$

Prueba:

(2) $q \vee qlq$
 $q \vee q \wedge pl.q \wedge p$

A112

(2), rinf16

A121 $q \wedge pl.p \wedge q$

Prueba:

(2) $q \wedge pl.q \vee q \wedge p$

A120, rinf11

(3) $q \vee q \wedge pl.p \wedge q \vee . p \wedge q$

A06

(4) $p \wedge q \vee (p \wedge q)l.p \wedge q$

A112

(5) $q \vee q \wedge pl.p \wedge q$

(3), (4), rinf15

A121

(2), (5), rinf15

A122 $q \vee pl.p \vee q$

Prueba:

(2) $Nq \wedge . Npl.Np \wedge Nq$

A121

A122

(2), A110, rinf01, df02

A123 $plq \supset . p \wedge rl.q \wedge r$

Prueba:

(2) $r \wedge 1lr$

A118

(3) $q \wedge (r \wedge)l.q \wedge r$

(2), rinf20

(4) $q \wedge rl.q \wedge . r \wedge 1$

(3), rinf11

(5) $p \wedge r \wedge 1l(q \wedge . r \wedge 1)l.q \wedge rl.p \wedge r \wedge 1$

(4), A02, rinf01

(6) $\neg(plq) \vee \sigma 5$	A05, df03
(7) $p \wedge r l. p \wedge r \wedge 1$	A118, rinf11
(8) $q \wedge r l.(p \wedge r \wedge 1) l. p \wedge r l. q \wedge r$	(7), A02, rinf01
(9) $\sigma 5 l \delta 8$	(5), (8), rinf15
(22) $\neg(plq) \vee \delta 8$	(9), (6), rinf21, rinf13
A123	(22), df03
A124 $plq \supset r \wedge pl. r \wedge q$	
Prueba:	
(2) $r \wedge pl(r \wedge q) l. q \wedge r l. r \wedge p$	A121, A02, rinf01
(3) $q \wedge r l(r \wedge p) l. p \wedge r l. q \wedge r$	similarmente
(4) $r \wedge pl(r \wedge q) l. p \wedge r l. q \wedge r$	(2), (3), rinf15
(5) $\neg(plq) \vee \sigma 4 l. \neg(plq) \vee \delta 4$	(4), rinf21
A124	(5), rinf14, A123, df03
A125 $plq \supset p \vee r l. q \vee r$	
Prueba:	
(2) $Np \wedge Nr l(Nq \wedge Nr) l. p \vee r l. q \vee r$	A109, df02
(3) $\neg(Np l Nq) l \neg(plq)$	A108, A01. rinf01
(4) $\neg(Np l Nq) \vee \delta 2$	(2), rinf21, A123, df02, rinf13
A125	(3), (4), rinf19, rinf14
A126 $plq \supset r \vee pl. r \vee q$	
Prueba: a partir de A125, similar a la de A124 a partir de A123 (utilizando A122 en vez de A121)	
A127 $plq \supset r \supset pl. r \supset q$	Prueba: A126, df02
A128 $p \wedge q \wedge r l. p \wedge q \wedge r$	Prueba: A05, rinf01, A101
A129 $p \wedge (q \wedge r) l. p \wedge q \wedge r$	Prueba: A128, rinf11
A130 $Np \vee Nq l N(p \wedge q)$	
Prueba:	
(2) $Np \vee Nq l N(NNp \wedge NNq)$	A101, df02
(3) $p \wedge NNq l. NNp \wedge NNq$	A102, rinf16
(4) $p \wedge q l. p \wedge NNq$	A102, rinf20
(5) $N(p \wedge q) l N(NNp \wedge NNq)$	(3), (4), rinf15, A110, rinf01
A130	(2), (5), rinf11, rinf15
A131 $p \wedge q l N(Np \vee Nq)$	Prueba: A130, A102, rinf11, A110, rinf01, rinf15
A132 $Np \wedge Nq l N(p \vee q)$	
Prueba:	
(2) $NNp \vee NNq l. p \vee NNq$	A102, rinf11, rinf17
(3) $p \vee NNq l. p \vee q$	A102, rinf11, rinf21
(4) $N\sigma 2 l N\delta 3$	(2), (3), rinf15, A110, rinf01
A132	A131, (4), rinf15

A133 $p \vee q \vee r \vdash p \vee q \vee r$

Prueba:

- (2) $Np \wedge Nq \wedge Nr \vdash Np \wedge Nq \wedge Nr$
 - (3) $\sigma 2 \vdash N(p \vee q) \wedge Nr$
 - (4) $N(p \vee q) \wedge Nr \vdash Np \wedge Nq \wedge Nr$
 - (5) $\delta 2 \vdash Np \wedge N(q \vee r)$
 - (6) $\sigma 4 \vdash \delta 5$
 - (7) $N(N(p \vee q) \wedge Nr) \vdash N(Np \wedge N(q \vee r))$
- A133

- A128
- A132, rinf16
- (2), (3), rinf11, rinf15
- A132, rinf20
- (4), (5), rinf15
- (6), A110, rinf01
- (7), df02

A134 $p \wedge q \vee p \vdash p$

Prueba:

- (2) $Np \vee Nq \wedge Np \vdash Np$
 - (3) $N(p \wedge q) \wedge Np \vdash \sigma 2$
 - (4) $N(p \wedge q) \wedge Np \vdash Np$
 - (5) $N(N(p \wedge q) \wedge Np) \vdash NNp$
 - (6) $p \wedge q \vee p \vdash NNp$
- $p \wedge q \vee p \vdash p$

- A04
- A130, rinf18
- (3), (2), rinf15
- (4), A110, rinf01
- (5), df02
- (6), A102, rinf11, rinf15

A135 $p \supset q \supset p \wedge q$

Prueba:

- (2) $\neg(p \wedge q) \vee p \wedge q$
 - (3) $\neg p \vee \neg q \vee p \wedge q$
 - (4) $\neg p \vee \neg q \vee p \wedge q$
 - (5) $p \supset \neg q \vee p \wedge q$
- A135

- A114
- (2), A07, rinf17, rinf14
- (3), A133, rinf13
- (4), df03
- (5), df03

A136 $p \wedge (q \vee r) \vdash p \wedge q \vee p \wedge r$

Prueba:

- (2) $p \wedge (q \vee r) \vdash q \vee r \wedge p$
 - (3) $\delta 2 \vdash p \wedge q \vee p \wedge r$
- $\sigma 2 \vdash \delta 3$

- A121
- A06
- (2), (3), rinf15

A137 $p \vee (q \wedge r) \vdash p \vee q \wedge p \vee r$

Prueba:

- (2) $Np \wedge (Nq \vee Nr) \vdash Np \wedge Nq \vee Np \wedge Nr$
- (3) $N(q \wedge r) \vdash Nq \vee Nr$
- (4) $Np \wedge N(q \wedge r) \vdash \sigma 2$
- (5) $Np \wedge N(q \wedge r) \vdash \delta 2$
- (6) $N(p \vee q) \vee (Np \wedge Nr) \vdash \delta 2$
- (7) $N(p \vee q) \vee N(p \vee r) \vdash \sigma 6$

- A136
- A130, rinf11
- (3), rinf20
- (4), (2), rinf15
- A132, rinf19
- A132, rinf11, rinf21

(8) $\sigma 7 \delta 2$	(7), (6), rinf15
(9) $\delta 2 \iota \sigma 7$	(8), rinf11
(22) $\sigma 5 \iota \sigma 7$	(5), (9), rinf15
(23) $N \sigma 5 \iota N \sigma 7$	(22), A110, rinf01
(24) $\delta 2 3 \iota . p \vee q \wedge . p \vee r$	A131, rinf11
(25) $\sigma 2 3 \iota \delta 2 4$	(23), (24), rinf15
A137	(25), df02
A138 $\neg(p \wedge \neg p)$	
Prueba:	
(2) $\neg \neg p \vee \neg p$	A114
(3) $\neg p \vee \neg \neg p$	(2), A122, rinf13
$\neg(p \wedge \neg p)$	(3), A07, rinf13
A139 $p \vee \iota 1 \iota 1$	
Prueba:	
(2) $p \wedge \iota 1 \vee \iota 1 \iota 1$	A134, A121, rinf17, rinf15
(3) $p \wedge \iota 1 \vee \iota 1 \iota . p \vee \iota$	A118, rinf17
A139	(3), rinf11, (2), rinf15
A140 $p \wedge 0 \iota 0$	
Prueba:	
(2) $N(N p \vee \iota 1) \iota N 1$	A139, A110, rinf01
(3) $p \wedge 0 \iota N N 0$	(2), df06, A131, rinf15
A140	A107, (3), rinf13
A141 $p \wedge q \supset q$	Prueba: A12, A121, rinf23, rinf13
A142 $p \supset . p \vee q$	
Prueba:	
(2) $p \vee q \wedge p \supset . p \vee q$	A12
$p \supset . p \vee q$	(2), A04, rinf23, rinf13
A143 $\neg p \vee p \vee q$	Prueba: A142, A114, rinf01
A144 $p \supset . q \supset p$	
Prueba:	
(2) $p \supset . p \vee \neg q$	A142
(3) $\neg p \vee . p \vee \neg q$	(2), df03
(4) $\neg p \vee . \neg q \vee p$	A122, rinf21, (3), rinf13
A144	(4), df03
A145 $p \supset (p \supset q) \iota . p \supset q$	
Prueba:	
(2) $\neg p \vee (\neg p \vee q) \iota . \neg p \vee \neg p \vee q$	A133, rinf11
(3) $\neg p \vee \neg p \iota \neg p$	A112
(4) $\delta 2 \iota . \neg p \vee q$	(3), rinf17
(5) $\sigma 2 \iota . \neg p \vee q$	(2), (4), rinf15
A145	(5), df03

A146 $p \supset (q \supset r) \vdash q \supset p \supset r$

Prueba:

- | | |
|--|------------------|
| (2) $\neg p \vee (\neg q \vee r) \vdash \neg p \vee \neg q \vee r$ | A133, rinf11 |
| (3) $\neg p \vee \neg q \vee r \vdash \neg q \vee \neg p \vee r$ | A122, rinf17 |
| (4) $\neg q \vee \neg p \vee r \vdash \neg q \vee \neg p \vee r$ | A133 |
| (5) $\sigma 2 \delta 3$ | (2), (3), rinf15 |
| (6) $\sigma 2 \delta 4$ | (5), (4), rinf15 |
| A146 | (6), df03 |

A147 $p \supset (q \supset r) \supset q \supset p \supset r$

Prueba: A146, rinf12)

A148 $p \wedge q \supset r \vdash p \vee r \vee q \supset r$

Prueba:

- | | |
|--|------------------------------------|
| (2) $\neg(p \wedge q) \vdash \neg p \vee \neg q$ | A07, rinf11 |
| (3) $\neg(p \wedge q) \vee r \vdash \neg p \vee \neg q \vee r$ | (2), rinf17 |
| (4) $\neg p \vee \neg q \vee r \vdash \neg p \vee \neg q \vee r$ | A133 |
| (5) $\neg q \vee r \vdash \neg q \vee r \vee r$ | A112, rinf11, rinf21, A133, rinf15 |
| (6) $\neg p \vee (\neg q \vee r) \vdash \neg p \vee \neg q \vee r \vee r$ | (5), rinf21 |
| (7) $\neg p \vee (\neg q \vee r) \vdash \neg p \vee r \vee \neg q \vee r$ | A122, rinf21, rinf15, (6) |
| (8) $\neg p \vee (r \vee \neg q \vee r) \vdash \neg p \vee r \vee \neg q \vee r$ | A133, rinf11 |
| (9) $\neg(p \wedge q) \vee r \vdash \neg p \vee r \vee \neg q \vee r$ | (3), (4), (7), (8), rinf24 |
| A148 | (9), df03 |

De aquí en adelante haremos un uso implícito de la regla rinf24 juntamente con rinf11. Si hemos demostrado, en una línea, una fórmula $\ulcorner plq \urcorner$ y si tenemos como teoremas (o sea: como instancias de algún esquema teorematizado —axiomático o no—) las fórmulas $\ulcorner qlr \urcorner$, $\ulcorner rlr^1 \urcorner$, ..., $\ulcorner r^{n-1}lr^n \urcorner$ (o lo que —en virtud de rinf11— equivale a lo mismo, a saber: $\ulcorner rlq \urcorner$, etc. —o sea, la inversa de alguna de esas fórmulas equivalenciales—), entonces podremos escribir:

(m) plq

lr

lr^1

lr^2

.

.

.

lr^n

Además, cuando así se haga, lo que nombrará el número de orden que, encerrado entre paréntesis, se halla a la izquierda de la primera línea de esa cadena de líneas será la fórmula equivalencial cuyo miembro izquierdo será el de la primera línea, y cuyo miembro derecho será el de la última línea (o sea —en el caso supuesto que se acaba de citar— '(m)' será un nombre de la fórmula $\ulcorner plr^n \urcorner$). Y, haciendo un uso implícito de rinf11, aducir una equivalencia $\ulcorner plq \urcorner$ será igual que aducir $\ulcorner qlp \urcorner$ —la equivalencia inversa.

A149 $p \vee q \supset r \vdash p \supset r \wedge q \supset r$

Prueba:

- | | |
|--|-------------|
| (2) $\neg(p \vee q) \vee r \vdash \neg p \wedge \neg q \vee r$ | A08, rinf17 |
|--|-------------|

(3) $\neg p \wedge \neg q \vee r \vdash r \vee \neg p \wedge \neg q$	A122
$\quad \quad \quad \vdash r \vee \neg p \wedge r \vee \neg q$	A137
$\quad \quad \quad \vdash r \vee \neg p \wedge r \vee \neg q \vee r$	A122, rinf20, rinf16
A149	(2), (3), rinf24

A150 $p \supset (q \wedge r) \vdash p \supset q \wedge p \supset r$

Prueba:

(2) $\neg p \vee (q \wedge r) \vdash \neg p \vee q \wedge \neg p \vee r$	A137
A150	(2), df03

A151 $p \supset (q \vee r) \vdash p \supset q \vee p \supset r$

Prueba:

(2) $\neg p \vee (q \vee r) \vdash \neg p \vee \neg p \vee q \vee r$	A112, rinf19
$\quad \quad \quad \vdash \neg p \vee \neg p \vee q \vee r$	A133
$\quad \quad \quad \vdash \neg p \vee \neg p \vee q \vee r$	A133, rinf21
$\quad \quad \quad \vdash \neg p \vee r \vee \neg p \vee q$	A122, rinf21
$\quad \quad \quad \vdash \neg p \vee r \vee \neg p \vee q$	A133
$\quad \quad \quad \vdash \neg p \vee q \vee \neg p \vee r$	A122
A151	(2), df03

En la demostración de los siguientes teoremas, ya no haremos, en la justificación de cada prueba, mención de las reglas de inferencia siguientes: rinf01, rinf11, rinf12, rinf13, rinf14, rinf22, rinf24, rinf25. Con ello quedará más despejado —menos sobrecargado— el conjunto de referencias que justifican cada prueba y, así, quedará más claramente puesto de relieve el *neruus probandi* de la misma.

A152 $p \supset q \supset q \supset r \vdash p \supset r$

Prueba:

(2) $\neg \neg p \vee \neg p \vee r \vdash \neg \neg q \wedge \neg r$	A143
(3) $\neg \neg q \vee \neg q \vee \neg p \vee r$	id.
(4) $\neg r \vee r \vdash \neg p \vee \neg q$	id.
(5) $\neg \neg p \vee (\neg \neg q \wedge \neg r) \vee \neg p \vee r$	(2), A133, A122
(6) $\neg p \vee r \vee \neg q \vee \neg \neg q$	(3), A122
(7) $\neg p \vee r \vee \neg q \vee \neg r$	(4), A133, A122
(8) $6 \wedge 7$	(6), (7)
(9) $\neg q \vee (\neg \neg q \wedge \neg r) \vee \neg p \vee r$	A122, (8), A133, A137
(22) $\neg \neg p \vee (\neg \neg q \wedge \neg r \vee \neg p \vee r) \wedge \neg q \vee \neg \neg q \wedge \neg r \vee \neg p \vee r$	(5), (9)
(23) $\neg \neg p \wedge \neg q \vee \neg \neg q \wedge \neg r \vee \neg p \vee r$	(22), A137, A122, A121
(24) $\neg (\neg p \vee q) \vee \neg \neg q \wedge \neg r \vee \neg p \vee r$	(23), A08
(25) $\neg (\neg p \vee q) \vee \neg (\neg q \vee r) \vee \neg p \vee r$	(24), A08
A152	(25), df03

A153 $p \supset q \supset r \supset p \supset r \supset p$

Prueba:

(2) $r \supset p \supset p \supset q \supset r \supset q$	A152
A153	(2), A147

De ahora en adelante, haciendo un uso implícito de rinf27, podremos escribir, si son teoremas ya demostrados las fórmulas $\lceil q \supset r \rceil$, $\lceil r \supset s \rceil$, $\lceil s \supset s^1 \rceil$, $\lceil s^1 \supset s^2 \rceil$, $\lceil s^2 \supset s^3 \rceil$, ..., $\lceil s^{n-1} \supset s^n \rceil$, y si, en una línea, se ha probado $\lceil p \supset q \rceil$, lo siguiente:

(m) $p \supset q$

$\supset r$

$\supset s$

$\supset s^1$

$\supset s^2$

$\supset s^3$

.

.

.

$\supset s^n$

Y en lo sucesivo, $\lceil (m) \rceil$ será un nombre de $\lceil p \supset s^n \rceil$.

Otro procedimiento que utilizaremos será el siguiente. Si tenemos un teorema $\lceil p \supset q \rceil$ y si se demuestra en una línea $\lceil p \rceil$, entonces escribiremos —haciendo un uso implícito de rinf01—:

(m) $p \supset q$

Y, en adelante, $\lceil (m) \rceil$ será un nombre de la fórmula $\lceil q \rceil$. Generalizando el procedimiento, supongamos que se ha probado: $\lceil p \rceil$, $\lceil p^1 \rceil$, $\lceil p^2 \rceil$... $\lceil p^{n-1} \rceil$; y supongamos también que se prueba:

$p \supset p^1 \supset p^2 \supset \dots \supset p^{n-1} \supset p^n$

Haciendo n pasos del procedimiento abreviatorio susodicho, tendríamos:

$p \supset [p^1 \supset [p^2 \supset [p^3 \supset \dots \supset [p^{n-1} \supset p^n]]]]$

Pues bien, abreviaremos la sucesión de esos n pasos, escribiendo:

$p \supset p^1 \supset p^2 \supset p^3 \supset \dots p^{n-1} \supset p^n$

Y, en lo sucesivo, $\lceil (m) \rceil$ será el nombre de la fórmula $\lceil p^n \rceil$.

Estos procedimientos se combinarán con el anterior, fundiéndose en uno solo cuando sea conveniente.

A154 $p \supset q \supset r \supset q \supset r \wedge p$

Prueba:

- | | |
|--|------------------------|
| (2) $\neg p \vee p \vee \neg q \vee \neg r$ | A143 |
| (3) $\neg r \vee r \vee \neg p \vee \neg q$ | A143 |
| (4) $\neg \neg q \vee \neg q \vee \neg p \vee r \wedge p$ | A143 |
| (5) $3 \wedge 2$ | (3), (2) |
| (6) $\neg r \vee r \vee \neg p \vee \neg q \wedge \neg p \vee p \vee \neg q \vee \neg r$ | (5), A133 |
| (7) $\neg p \vee \neg q \vee \neg r \vee r \wedge \neg p \vee \neg q \vee \neg r \vee p$ | (6), A122, A133 |
| (8) $\neg p \vee \neg q \vee \neg r \vee r \wedge p$ | (7), A137 |
| (9) $\neg p \vee \neg q \vee (r \wedge p) \vee \neg r$ | (8), A122, A133 |
| (22) $\neg p \vee \neg q \vee (r \wedge p) \vee \neg \neg q$ | (4), A122, A133 |
| (23) $9 \wedge 22$ | (9), (22) |
| (24) $\neg p \vee \neg \neg q \wedge \neg r \vee \neg q \vee r \wedge p$ | (23), A137, A122, A133 |

(25) $p \supset . \neg(\neg q \vee r) \vee . \neg q \vee . r \wedge p$	(24), df03, A08
$\supset . \neg q \vee r \supset . \neg q \vee . r \wedge p$	df03
$\supset . q \supset r \supset . q \supset . r \wedge p$	df03
A155 $p \supset (q \supset r) \supset . p \supset q \supset . p \supset r$	
Prueba:	
(2) $\neg \neg p \vee \neg p \vee r$	A143
(3) $\neg p \vee \neg q \vee r \supset . \neg p \vee \neg q \vee r$	A113
(4) $2 \supset . 3 \supset .] \neg p \vee \neg q \vee r \supset . \neg p \vee \neg q \vee r \wedge . \neg \neg p \vee \neg p \vee r$	A154
$\supset . \neg \neg p \wedge \neg q \vee . \neg p \vee r$	A122, A133. A137
$\supset . \neg(\neg p \vee q) \vee . \neg p \vee r$	A08
$\supset . p \supset q \supset . p \supset r$	df03
A155	(4), A133, df03
A156 $q \supset (p \supset r) \supset . p \supset q \supset . p \supset r$	Prueba: A155, A146

Metateorema de la deducción

Lo que vamos a probar ahora no es **ni** un teorema de *At*, ni **tampoco** una regla de inferencia de *At*, sino un teorema sintáctico (o, si se quiere, metalingüístico) **acerca de** *At*, a saber:

Si $p \vdash q$ es una regla de inferencia de *At*, entonces $\lceil p \supset q \rceil$ es un teorema de *At*.

Que $p \vdash q$ significa que es educible de la premisa $\lceil p \rceil$ hasta la conclusión $\lceil q \rceil$.

En *At* no hay más que una sola regla de inferencia primitiva: rinf01 (o sea la regla de *modus ponens*). Pero, como la base axiomática de *At* está expuesta por modo de **esquemas axiomáticos** (y el desarrollo de *At* se está efectuando por modo de **esquemas teorematizados**), es obvio que, siempre que tengamos un esquema teorematizado, tenemos como teorema cualquier resultado de sustituir uniformemente todas las ocurrencias de cada una de las letras esquemáticas que en el mismo figuran por ocurrencias de fbfs cualesquiera.

Pero que haya una deducción $p \vdash q$ puede significar dos cosas: o bien, 1º, que haya una prueba o demostración dentro de *At*, de $\lceil p \rceil$ a $\lceil q \rceil$ (y, en ese caso, $\lceil p \rceil$ y $\lceil q \rceil$ son dos esquemas teorematizados de *At*, y sólo contienen letras esquemáticas); o bien, 2º, que, en una extensión de *At*, se deduzca de la premisa $\lceil p \rceil$ la conclusión $\lceil q \rceil$ (y, en este segundo caso, $\lceil p \rceil$ podrá ser una genuina fórmula, y no un mero esquema). En este último caso, evidentemente, ninguna regla de inferencia de *At* autoriza una sustitución de esa(s) fbf(s) que sean subfórmulas de la premisa $\lceil p \rceil$ por otra(s) fbf(s). Así, supongamos que aplicando la rinf22, en una extensión de *At* en la que tengamos las constantes sentenciales $\lceil s \rceil$ (que signifique lo mismo que 'Anatole France fue galardonado con el Premio Nobel'), $\lceil s^1 \rceil$ (que signifique lo mismo que 'Juan Ramón Jiménez fue galardonado con el Premio Nobel') y $\lceil r \rceil$ (que signifique lo mismo que 'Antonio Machado fue galardonado con el Premio Nobel'), y teniendo las dos primeras de esas fórmulas como premisas, hacemos la deducción siguiente: (como caso concreto de rinf22):

$s, s^1 \vdash s \wedge s^1$

Pero, naturalmente, no podremos nunca concluir de esas dos premisas, la conclusión $\lceil s \wedge r \rceil$. No podemos sustituir $\lceil s^1 \rceil$ por $\lceil r \rceil$, porque $\lceil s^1 \rceil$ **no** es ahí un esquema teorematizado, sino una fbf (y *At* no contiene ninguna regla de sustitución de fbfs por otras, ni tan siquiera dentro de determinados límites).

Por consiguiente, para que $p \vdash q$ tenga lugar (para que $\lceil q \rceil$ se deduzca de $\lceil p \rceil$), ello debe, en última instancia, ser posible tan sólo en virtud de la regla rinf01, y de los axiomas de *At*, gracias a los cuales se engendran, a partir de rinf01, otras reglas de inferencia derivadas; mas

esas otras reglas de inferencia son **prescindibles**, ya que su empleo es un mero expediente para **abreviar** las pruebas, que pueden obtenerse con el uso **exclusivo** de rinf01, siempre y cuando se aduzcan, en cada caso, cuantos teoremas de *At* han sido aducidos en las derivaciones de esas reglas de inferencia no primitivas.

Que haya una deducción de $\lceil q \rceil$ a partir de $\lceil p \rceil$ (o sea que tenga lugar $p \vdash q$) quiere decir que hay una serie de fórmulas, $\lceil r_1 \rceil, \lceil r_2 \rceil, \lceil r_3 \rceil \dots \lceil r_n \rceil$ tal que $\lceil r_1 \rceil = \lceil p \rceil$, y $\lceil r_n \rceil = \lceil q \rceil$, y que, en esa serie, cada $\lceil r_i \rceil$ sea deducible de los anteriores **más** los teoremas de *At* **más** rinf01. Ello quiere decir que, para cada i tal que $i \neq 1$, debe haber dos fórmulas (o esquemas), $\lceil r_j \rceil$ y $\lceil r_k \rceil$ tales que tanto j como k sean menores que i , y $\lceil r_j \rceil$ sea $\lceil r_k \rceil \supset r_i$. Recapitulando, tendremos, en las líneas que conforman la deducción desde $\lceil p \rceil$ hasta $\lceil q \rceil$:

(j) $r_k \supset r_i$

(k) r_k

Pero, de ahí, aplicando el teorema A144 de *At* podemos proseguir la deducción como sigue:

(j ¹) $r_k \supset r_i \supset . p \supset . r_k \supset r_i$	A144
(j ²) $p \supset . r_k \supset r_i$	(j ¹), (j), rinf01
(j ³) $p \supset r_k \supset . p \supset r_i$	A155, (j ²), rinf01
(k ¹) $r_k \supset . p \supset r_k$	A144
(k ²) $p \supset r_k$	(k ¹), (k), rinf01
(k ³) $p \supset r_i$	(j ³), (k ²), rinf01

La conclusión que obtenemos en (k³) es, justamente, $\lceil p \supset r_i \rceil$. Pero esa conclusión vale para cualquier i , y, por tanto, para **cualquier** r_i , en la serie que va de r_1 a r_n o sea: vale también para r_n es decir: para $\lceil q \rceil$.

Luego, como caso particular de (k³) tenemos: $\lceil p \supset q \rceil$. Con ello queda probado el **metateorema de la deducción**. En adelante llamaremos a ese metateorema '(MD)'.

Pruebas de otros teoremas

A157 $p \wedge q \supset r \vdash . p \supset . q \supset r$

Prueba:

(2) $\neg(p \wedge q) \vee r \vdash . \neg p \vee \neg q \vee r$	A07
$\vdash . \neg p \vee . \neg q \vee r$	A133
A157	(2), df03

A158 $p \supset (q \supset r) \vdash . p \supset r \vee . q \supset r$

Prueba: A157, A148

A159 $p \wedge q \supset r \vdash . p \supset q \supset . p \supset r$

Prueba: A155, A157

A160 $r \supset s \vdash . p \supset (q \supset r) \supset . p \supset . q \supset s$

Prueba:

(2) $q \supset r \supset (q \supset s) \supset . p \supset (q \supset r) \supset . p \supset . q \supset s$	A153
(3) $r \supset s \vdash . q \supset r \supset . q \supset s$	A153
$\vdash . p \supset (q \supset r) \supset . p \supset . q \supset s$	(2)

A161 $r \supset s \vdash . p \supset (q \supset . q' \supset r) \supset . p \supset . q \supset . q' \supset s$

Prueba:

(2) $r \supset s \vdash . q' \supset r \supset . q' \supset s$	A153
(3) $q' \supset r \supset (q' \supset s) \supset . q \supset (q' \supset r) \supset . q \supset . q' \supset s$	A153

(4) $\delta \supset p \supset (q \supset q^1 \supset r) \supset p \supset q \supset q^1 \supset s$

A153

A161

(2), (3), (4), rinf27

A162 $\supset s \supset p \supset (q \supset q^1 \supset q^2 \supset r) \supset p \supset q \supset q^1 \supset q^2 \supset s$ A163 $\supset s \supset p \supset (q \supset q^1 \supset q^2 \supset q^3 \supset r) \supset p \supset q \supset q^1 \supset q^2 \supset q^3 \supset s$

(Las pruebas de A162 y A163 son similares a la de A161)

Por inducción matemática se puede generalizar la secuencia de teoremas A161, A162, A163... y la secuencia de reglas rinf28, rinf29, rinf30... Y, en virtud de ellos, podemos, en adelante, siempre que tengamos como teoremas o fórmulas ya probadas $\ulcorner r \supset r^1 \urcorner \dots \ulcorner r^{n-1} \supset r^n \urcorner$, y que tengamos una fórmula probada del tipo:

 $p \supset p^1 \supset p^2 \supset \dots \supset p^n \supset r$

escribir:

 $p \supset p^1 \supset p^2 \supset \dots \supset p^n \supset r$ $\supset r^1$ $\supset r^2$ $\supset r^{n-1}$ $\supset r^n$

Por otro lado, como, cada vez que tenemos un teorema $\ulcorner plq \urcorner$, tenemos —por rinf11— el teorema $\ulcorner qlp \urcorner$ y, por consiguiente —en virtud de A10 y rinf01— el teorema $\ulcorner p \supset q \urcorner$, utilizaremos también el procedimiento siguiente (siendo (n) un teorema o una fórmula previamente probada):

(m) $nl.jp$ Y, en adelante, $\ulcorner (m) \urcorner$ nombrará a $\ulcorner p \urcorner$.

Por otro lado, en virtud de rinf 25, cada vez que se haya probado previamente $\ulcorner plq \urcorner$ y que tengamos, en una fórmula cualquiera, una ocurrencia de 'l' seguida de una ocurrencia de $\ulcorner r \urcorner$, siendo $\ulcorner r \urcorner$ una fórmula cualquiera que contenga m ocurrencias de $\ulcorner p \urcorner$, podemos escribir:

(n) $\dots lr \dots$ $lr^1 \dots$

siendo $\ulcorner r^1 \urcorner$ el resultado de reemplazar una o varias de esas m ocurrencias de $\ulcorner p \urcorner$ en $\ulcorner r \urcorner$ por ocurrencias respectivas de $\ulcorner q \urcorner$. Y, en lo sucesivo, $\ulcorner (n) \urcorner$ nombrará a:

 $\dots lr^1 \dots$ A164 $p \supset q \supset p \wedge r \supset q$

Prueba:

(2) $p \wedge r \supset p \supset q \supset p \wedge r \supset q$

A152

A164

(2), A12

A165 $p \supset q \supset r \wedge p \supset q$

Prueba: A152, A141

A166 $p \supset q \supset p \supset r \supset p \supset q \wedge r$

Prueba:

(2) $p \supset q \wedge (p \supset r) \supset p \supset q \wedge r$

A150

A166

A157, (2)

A167 $p \supset q \supset p \supset p \wedge q$

Prueba: A113, A166

A168 $p \supset q \supset . p \wedge r \supset . q \wedge r$

Prueba:

- (2) $p \wedge r \supset q \supset . p \wedge r \supset . p \wedge r \wedge q$
 $\supset . q \wedge r$

A168

A167

A121, A128, A12

(2), A164, A153

A169 $p \supset . \neg p \supset q$

Prueba:

- (2) $\neg \neg p \vee \neg p \vee q$
(3) $\neg p \vee . \neg \neg p \vee q$

A169

A143

A122, A133

(3), df03

A170 $\neg p \supset . p \supset q$

Prueba: A169, A147

A171 $p \wedge \neg p \supset q$

Prueba: A169, A157

A172 $p \supset q \supset . \neg q \supset \neg p$

Prueba:

- (2) $\neg \neg p \vee \neg p \vee \neg q$
(3) $\neg \neg q \vee \neg q \vee \neg p$
(4) $2 \wedge 3$
(5) $\neg \neg p \wedge \neg q \vee . \neg \neg q \vee \neg p$
(6) $\neg (\neg p \vee q) \vee . \neg \neg q \vee \neg p$

A172

A143

A143

(2), (3)

(4), A122, A133, A137

(5), A08

(6), df03

A173 $p \supset . q \vee p$

Prueba: A142, A122

A174 $p \supset Lp$

Prueba:

- (2) $\neg \neg p \vee \neg p$
(3) $\neg p \vee \neg \neg p$
 $p \supset Lp$

A114

(2), A122

(3), df03, df08

A175 $p \supset q \supset . r \wedge p \supset . r \wedge q$

Prueba: A168, A121

A176 $p \supset q \supset . p \vee r \supset . q \vee r$

Prueba:

- (2) $\neg \neg p \vee \neg p \vee q \vee r$
(3) $\neg r \vee r \vee . \neg \neg p \vee q$
(4) $\neg \neg p \vee q \vee r \vee \neg p$
(5) $\neg \neg p \vee q \vee r \vee \neg r$
(6) $\neg p \wedge \neg r \vee . \neg \neg p \vee q \vee r$
(7) $\neg q \vee q \vee . \neg p \wedge \neg r \vee r$
(8) $\neg p \wedge \neg r \vee . \neg q \vee q \vee r$
(9) $\neg p \wedge \neg r \vee . \neg \neg p \vee (q \vee r) \wedge . \neg q \vee q \vee r$
(22) $\neg p \wedge \neg r \vee . \neg \neg p \wedge \neg q \vee q \vee r$
(23) $\neg \neg p \wedge \neg q \vee . \neg p \wedge \neg r \vee q \vee r$
(24) $\neg (\neg p \vee q) \vee . \neg (p \vee r) \vee q \vee r$

A176

A143

A143

(2), A122, A133

(3), A122, A133

(4), (5), A137, A122, A133

A143

(7), A122, A133

(6), (8), A137

(9), A122, A137

(22), A122, A133

A08, (23)

(24), df03

A177 $p \supset q \supset . r \vee p \supset . r \vee q$

Prueba: A176, A122

A178 $p \vee q \supset . r \supset . r \wedge p \vee . r \wedge q$

Prueba:

 $p \vee q \supset . r \supset . p \vee q$

A144

 $\supset . r \wedge . p \vee q$

A167

 $\supset . r \wedge p \vee . r \wedge q$

A136

A179 $p \supset q \supset . r \supset s \supset . p \wedge r \supset . q \wedge s$

Prueba:

(2) $\neg \neg p \vee \neg p \vee . \neg \neg r \wedge \neg s \vee \neg r \vee . q \wedge s$

A143

(3) $\neg q \vee q \vee . \neg \neg r \wedge \neg s \vee \neg p \vee \neg r$

A143

(4) $\neg \neg r \vee \neg r \vee . \neg p \vee s \vee \neg q$

A143

(5) $\neg s \vee s \vee . \neg q \vee \neg p \vee \neg r$

A143

(6) $\neg q \vee (\neg \neg r \wedge \neg s) \vee \neg p \vee \neg r \vee s$

(4), (5), A122, A133, A137

(7) $\neg q \vee (\neg \neg r \wedge \neg s) \vee \neg p \vee \neg r \vee . q \wedge s$

(3), (6), id.

(8) $\neg \neg p \wedge \neg q \vee . \neg \neg r \wedge \neg s \vee . \neg p \vee \neg r \vee . q \wedge s$

(2), (7), id.

(9) $\neg (\neg p \vee q) \vee . \neg (\neg r \vee s) \vee . \neg (p \wedge r) \vee . q \wedge s$

(8), A07, A08

A179

(9), df03

A180 $Lp \supset p$

Prueba:

(2) $NN \neg p \vee p$

A114, A102

(3) $N \neg \neg p \vee p$

(2), A09

(4) $\neg \neg \neg p \vee p$

(3), A09

 $Lp \supset p$

(4), df03, df08

A181 $r \supset s \supset . p \supset (s \supset q) \supset . p \supset . r \supset q$

Prueba:

 $r \supset s \supset . s \supset q \supset . r \supset q$

A152

 $\supset . p \supset (s \supset q) \supset . p \supset . r \supset q$

A153

A182 $p \supset \neg q \supset . q \supset \neg p$

Prueba:

 $p \supset \neg q \supset . \neg \neg q \supset \neg p$

A172

 $\supset . q \supset \neg p$

A174, A153, df08

A183 $p \supset (q \wedge r) \supset . p \supset q$

Prueba:

(2) $q \wedge r \supset q$

A12

 $p \supset (q \wedge r) \supset . p \supset q$

A153, (2)

A184 $p \supset (q \wedge r) \supset . p \supset r$ A185 $p \supset (p \wedge q) \supset . p \supset q$ A186 $p \supset (q \wedge p) \supset . p \supset q$

A187 $p \vee q \wedge \neg q \supset p$

Prueba:

- | | |
|--|------------------|
| (2) $p \vee q \wedge \neg q \vdash p \wedge \neg q \vee q \wedge \neg q$ | A136, A121 |
| (3) $q \wedge \neg q \supset p$ | A171 |
| (4) $p \wedge \neg q \supset p$ | A12 |
| (5) $4 \wedge 3 \vdash \delta 2 \supset p$ | A149, A136, A121 |
| $\sigma 2 \supset p$ | (5), (2) |

A188 $p \supset (q \vee r) \supset \neg (p \wedge r) \supset p \supset q$

Prueba:

- | | |
|--|------------------|
| (2) $p \supset (q \vee r) \supset p \supset p \wedge q \vee r$ | A167 |
| $\supset p \wedge q \vee p \wedge r$ | A136 |
| (3) $\neg (p \wedge r) \supset 2 \supset \sigma 2 \wedge \neg (p \wedge r)$ | A154 |
| (4) $2 \supset \neg (p \wedge r) \supset \sigma 2 \wedge \neg (p \wedge r)$ | (3), A147 |
| (5) $\sigma 2 \supset \neg (p \wedge r) \supset \delta 2 \wedge \neg (p \wedge r)$ | (4), A146 |
| $\supset p \supset \delta \delta 2 \wedge \neg (p \wedge r)$ | A121, A154, A157 |
| $\supset p \wedge q$ | A187 |
| $\supset q$ | A141 |

A189 $p \supset (q \vee r) \supset r \supset p \supset p \supset q$ Prueba: A188, df03, A07, A122

A190 $p \wedge q \supset r \supset p' \supset p \wedge (q' \supset q) \supset p' \wedge q' \supset r$

Prueba:

- | | |
|--|------------------|
| (2) $p' \wedge q' \supset (p \wedge q) \supset p \wedge q \supset r \supset p' \wedge q' \supset r$ | A152 |
| (3) $p' \supset p \wedge (q' \supset q) \supset p' \wedge q' \supset p \wedge q$ | A179, A157 |
| (4) $p' \supset p \wedge (q' \supset q) \supset p \wedge q \supset r \supset p' \wedge q' \supset r$ | (3), (2), rinf26 |
| A190 | (4), A147 |

A191 $p \supset p \supset q \supset q$

Prueba:

- | | |
|--|----------------------------|
| (2) $\neg \neg p \vee \neg p \vee q$ | A143 |
| (3) $\neg q \vee q \vee \neg p$ | A143 |
| (4) $\neg p \vee \neg \neg p \wedge \neg q \vee q$ | (2), (3), A122, A133, A137 |
| (5) $\neg p \vee \neg (\neg p \vee q) \vee q$ | (4), A08 |
| $p \supset p \supset q \supset q$ | (5), df03 |

A192 $p \supset q \vee q \supset r$

Prueba:

- | | |
|--|-----------------|
| (2) $\neg q \vee q \vee r \vee \neg p$ | A143 |
| (3) $\neg p \vee q \vee \neg q \vee r$ | (2), A122, A133 |
| A192 | (3), df03 |

A193 $q \supset p \supset (p \supset q \supset p) \supset p \supset q \supset q \supset p \supset q$

Prueba:

- | | |
|---|-----------|
| (2) $p \supset p \supset q \supset q$ | A191 |
| (3) $p \supset q \supset p \supset p \supset q \supset p \supset q \supset q$ | (2), A153 |

(4) $q \supset p \supset (p \supset q \supset p) \supset . q \supset p \supset . p \supset q \supset . p \supset q \supset q$

(3), A153

 $\supset . q \supset p \supset . p \supset q \supset q$

A145

 $\supset . p \supset q \supset . q \supset p \supset q$

A147

A194 $\neg q \supset . p \supset q \supset \neg p$ Prueba: A172, A146A195 $\neg(p \supset q) \supset . p \wedge \neg q$

Prueba: df03, A08, df08, A180, A168

A196 $p \supset q \supset p \supset p$

Prueba:

(2) $\neg \neg \neg p \vdash \neg \neg p$

A09

 $\vdash \neg \neg \neg p$

A09

 $\vdash p$

A102

(3) $\neg p \vee p \vee . \neg q \wedge \neg p$

A143

(4) $\neg \neg \neg p \vee p \vee . \neg q \wedge \neg p$

(3), (2)

(5) $\neg \neg \neg p \wedge \neg p \vee (\neg q \wedge \neg p) \vee p$

(3), (4), A137, A122, A133

(6) $\neg \neg \neg p \vee \neg q \wedge \neg p \vee p$

(5), A121, A136

(7) $\neg(\neg \neg p \wedge \neg q) \wedge \neg p \vee p$

A07, (6)

(8) $\neg \neg(\neg p \vee q) \wedge \neg p \vee p$

A08, (7)

(9) $\neg(\neg(\neg p \vee q) \vee p) \vee p$

A08, (8)

(22) $\neg(\neg(p \supset q) \vee p) \vee p$

(9), df03

(23) $\neg(p \supset q \supset p) \vee p$

(22), df03

 $p \supset q \supset p \supset p$

(23), df03

A197 $H p \vee N p$

Prueba: A113, df03, df10

A198 $p \rightarrow q \supset . p \supset q$

Prueba:

 $p \wedge q \vdash p \supset . p \wedge q$

A10

 $\supset . p \supset q$

A185, df04

A199 $p \supset 0 \vdash p$

Prueba: A119, df03

A200 $p \vdash q . r \vdash r$ (donde $\ulcorner r \urcorner$ sólo difiere de $\ulcorner r \urcorner$ por el reemplazamiento de n ocurrencias de $\ulcorner p \urcorner$ en $\ulcorner r \urcorner$ por n ocurrencias respectivas de $\ulcorner q \urcorner$)

Prueba: (MD), rinf25

A201 $p \vee q \wedge (p \supset r \wedge . q \supset s) \supset . r \vee s$

Prueba:

(2) $r \supset . r \vee s$

A142

(3) $s \supset . r \vee s$

A173

(4) $p \supset r . p \supset . r \vee s$

A153, (2)

(5) $4 \supset .] p \supset r \wedge (q \supset s) \supset \delta 4$

A164

(6) $q \supset s \supset . q \supset . r \vee s$

A153, (3)

(7) $6 \supset .] p \supset r \wedge (q \supset s) \supset \delta 6$

A165

(8) $p \supset r \wedge (q \supset s) \supset . \delta 4 . \delta 6$

(5), (7), A150

 $\supset . p \vee q \supset . r \vee s$

A149

(9) $p \supset r \wedge (q \supset s) \wedge (p \vee q) \supset . r \vee s$	(8), A157
A201	(9), A121
A201/2 $p \vee q \wedge (p \supset s \wedge . q \supset s) \supset s$	Prueba: A201, A112
A201/3 $p \vee q \wedge (p \supset r) \supset . r \vee q$	A201/4 $p \vee q \wedge (q \supset r) \supset . p \vee r$
A202 $p \supset q \supset . q \supset p \supset . p \equiv q$	Prueba: A135, df07
A203 $p \vee q \wedge (q \vee r) \supset . p \vee r$	
Prueba:	
(2) $p \vee q \supset . p \vee r \vee . q \vee r$	A200
$\supset . q \vee r \supset . p \vee r$	A10
A203	(2), A157
A204 $p \vee q \wedge (p \vee r) \supset . q \vee r$	Prueba: A203, A103
A205 $p \rightarrow q \wedge (q \rightarrow p) \supset . p \vee q$	
Prueba:	
(2) $p \wedge q \vee p \wedge (p \wedge q \vee q) \supset . p \vee q$	A204
A205	(2), A121, df04
A206 $p \vee q \supset . p \rightarrow q \wedge . q \rightarrow p$	
Prueba:	
(2) $p \vee q \supset . p \wedge p \vee . p \wedge q$	A200
$\supset . p \vee . p \wedge q$	A03
$\supset . p \wedge q \vee p$	A103
(3) $p \vee q \supset . q \wedge p \vee q$	similarmente
$2 \wedge 3 \supset]A206$	A150, df04
A207 $p \rightarrow q \wedge (q \rightarrow p) \equiv . p \vee q$	Prueba: A205, A206, A202
A208 $p \rightarrow q \equiv . p \vee q \vee p$	
Prueba:	
(2) $p \wedge q \vee p \supset . p \wedge q \vee q \vee . p \vee q$	A200
$\supset . q \vee . p \vee q$	A134, A121
$\supset . p \vee q \vee q$	A103
(3) $p \vee q \vee q \supset . p \vee q \wedge p \vee . q \wedge p$	A200
$\supset . p \vee . q \wedge p$	A04
$\supset . p \vee . p \wedge q$	A121
$\supset . p \wedge q \vee p$	A103
$2 \supset . 3 \supset .]A208$	A202, df04
A209 $p \rightarrow q \equiv . Nq \rightarrow Np$	
Prueba:	
(2) $p \wedge q \vee p \supset . p \vee q \vee q$	A208, A12, df07, df04
$\supset . Np \wedge Nq \vee Nq$	A108, A132
$\supset . Nq \wedge Np \vee Nq$	A121
(3) $Nq \wedge Np \vee Nq \supset . NNp \wedge NNq \vee NNp$	similarmente
$\supset . p \wedge q \vee p$	A102
$2 \supset . 3 \supset .]A209$	A202, df04

A210 $\neg p \supset p \rightarrow q$ Prueba: A14, df08, df03, A209

A211 $p \wedge \neg p \vdash q \wedge \neg q$

Prueba:

(2) A138 $\supset$.] $p \wedge \neg p \rightarrow q \wedge \neg q$

A210

(3) $q \wedge \neg q \rightarrow p \wedge \neg p$

similarmemente

2 $\wedge$ 3 $\supset$.] A211

A205

A212 $p \wedge \neg p \vdash 0$

Prueba: A211, A121, df05

A213 $p \wedge \neg p \wedge q \vdash 0$

Prueba: A212, A140

A214 $p \supset \neg p \rightarrow q$

Prueba:

(2) $\neg \neg p \supset \neg p \rightarrow q$

A210

A214

(2), A174, rinf26

A215 $p \supset \neg p \vdash 0$

Prueba:

(2) $p \supset \neg p \rightarrow 0$

A214

(3) $0 \rightarrow \neg p$

A140, A121, df04

(4) $p \supset 0 \rightarrow \neg p$

(3), A144

$p \supset \delta 2 \wedge \delta 4$

(2), (4), A150

$\supset \neg p \vdash 0$

A205

A216 $\neg p \supset p \vdash 0$

Prueba similar, a partir de A210 en vez de A214.

A217 $p \supset p \supset p \vdash p$

Prueba:

(2) $\neg p \vee p \wedge p \vdash p$

A04, A122

(3) $p \rightarrow \neg p \vee p$

(2), df04, A121

(4) $p \rightarrow p \supset p$

(3), df03

(5) $p \supset \neg p \vdash 0$

A215

$\supset \neg p \vee p \vdash 0 \vee p$

A200

$\vdash p$

A119, A122

A217

(5), df03

A218 $p \vee q \supset \neg p \supset q$

Prueba:

(2) $p \vee q \wedge \neg p \vdash q$

A187, A122

A218

A157, (2)

A partir de este momento, y en las restantes pruebas de teoremas, ya no mencionaremos expresamente, como instancia justificatoria, el empleo de ninguno de los axiomas —salvo A13 y A14—; ni tampoco de los teoremas siguientes: A101, A102, A103, A112, A113, A114, A115, del A118 al A123 —ambos inclusive—, del A130 al A153 —ambos inclusive—, A155, A157, del A168 al A177 —ambos inclusive—, A180, A182, A183, A194, del A200 al A218 —ambos inclusive—; ni de ninguna de las definiciones de df02 a df10 —ambas inclusive—; ni de ninguna de las reglas de inferencia hasta ahora derivadas.

A250 $plp' \wedge (rlr') \supset q$ no difiere de $\ulcorner q \urcorner$ más que por el reemplazamiento de n ocurrencias de $\ulcorner p \urcorner$ en $\ulcorner q \urcorner$ por n ocurrencias respectivas de $\ulcorner p' \urcorner$, y/o de m ocurrencias de $\ulcorner r \urcorner$ en $\ulcorner q \urcorner$ por m ocurrencias respectivas de $\ulcorner r' \urcorner$.

Prueba (Sea $\ulcorner q \urcorner$ el resultado de reemplazar en $\ulcorner q \urcorner$ esas n ocurrencias de $\ulcorner p \urcorner$ por n ocurrencias respectivas de $\ulcorner p' \urcorner$):

(2) $plp' \supset qlq'$

(3) $rlr' \supset q'ls$

$\sigma 2 \wedge \sigma 3 \supset \delta 2 \wedge \delta 3$

$\supset q'ls$

A250/2 $plp' \vee (rlr') \supset q'ls$ es el resultado de reemplazar en $\ulcorner q \urcorner$ n ocurrencias de $\ulcorner p \urcorner$ por ocurrencias respectivas de $\ulcorner p' \urcorner$; y si $\ulcorner s \urcorner$ es el resultado de reemplazar en $\ulcorner q \urcorner$ m ocurrencias de $\ulcorner r \urcorner$ por ocurrencias respectivas de $\ulcorner r' \urcorner$

Prueba similar.

A250/3 $q \supset p$ (si $\ulcorner q \urcorner$ es el resultado de reemplazar en $\ulcorner q \urcorner$ n ocurrencias de $\ulcorner p \urcorner$ por ocurrencias respectivas de $\ulcorner r \urcorner$)

A250/4 $q \supset p$ (si $\ulcorner q \urcorner$ y $\ulcorner s \urcorner$ son como en A250)

A250/5 $q \supset p$ (si $\ulcorner s \urcorner$ y $\ulcorner s' \urcorner$ son como en A250/2)

A251 $\neg p \text{ INLp}$

Prueba:

$\neg p \text{ I} \neg p$

$\text{INN} \neg p$

$\text{IN} \neg \neg p$

INLp

A252 Hp INLNp

Prueba:

$\text{Hpl} \neg \text{Np}$

INLNp

A251

A253 Lp INHNp

Prueba:

(2) HNp INLNNp

A252

INLp

NHNp INNLp

(2)

ILp

A254 $\neg p \text{ I} \neg \text{Hp}$

Prueba:

$\neg p \text{ I} \text{LNp}$

df13

INHNNp

A253

INHp

$\text{I} \neg \text{Hp}$

A255 $\neg p \text{ INLNLNp}$ Prueba: A254, A252, A252.

A256 $\neg p \text{ I } H \neg p$

Prueba:

 $H \neg p \text{ I } \neg N \neg p$ $I \neg p$ A257 $H H p \text{ I } H p$

Prueba:

 $H H p \text{ I } \neg N \neg N p$ $I \neg \neg \neg N p$ $I \neg N \neg N p$ $I \neg N p$ $I H p$ A257/2 $\neg p \text{ I } H \neg p$ Prueba: A256, A257A258 $L L p \text{ I } L p$

Prueba:

 $L L p \text{ I } \neg \neg \neg \neg p$ $I \neg \neg \neg \neg p$ $I \neg p$ $I L p$ A259 $L(p \vee q) \text{ I } L p \wedge L q$

Prueba:

 $L(p \vee q) \text{ I } \neg(p \vee q)$ $I \neg(\neg p \wedge \neg q)$ $I. \neg p \vee \neg q$ $I. L p \vee L q$ A261 $L(p \wedge q) \text{ I } L p \wedge L q$

Prueba:

 $L(p \wedge q) \text{ I } \neg \neg(p \wedge q)$ $I \neg(\neg p \vee \neg q)$ $I. \neg \neg p \wedge \neg \neg q$ $I. L p \wedge L q$ A263 $p \& q \text{ I } L p \wedge q$

Prueba:

 $p \& q \text{ I } \neg(p \supset \neg q)$ $I \neg(\neg p \vee \neg q)$ $I. \neg p \wedge \neg \neg q$ $I. L p \wedge q$ A265 $p \supset q \text{ I } \neg(L p \wedge \neg q)$

Prueba:

 $p \supset q \text{ I } \neg p \vee q$ $I. \neg \neg p \vee \neg \neg q$ $I \neg(\neg p \wedge \neg q)$ $I \neg(L p \wedge \neg q)$

A251

A260 $H(p \vee q) \text{ I } H p \vee H q$

Prueba:

 $H(p \vee q) \text{ I } \neg N(p \vee q)$ $I \neg(N p \wedge N q)$ $I. \neg N p \vee \neg N q$ $I. H p \vee H q$ A262 $H(p \wedge q) \text{ I } H p \wedge H q$

Prueba:

 $H(p \wedge q) \text{ I } \neg N(p \wedge q)$ $I. \neg N p \wedge \neg N q$ $I. H p \wedge H q$ A264 $p \vee q \text{ I } H p \vee q$

Prueba:

 $p \vee q \text{ I } \neg(N p \& N q)$ $I \neg(L N p \wedge N q)$ $I. \neg L N p \vee q$ $I. H p \vee q$ A266 $p \supset q \text{ I } \neg(p \& \neg q)$ A267 $L p \text{ I } H L p$

Prueba:

 $L p \text{ I } \neg \neg p$ $I \neg \neg \neg p$ $I H L p$

A252

A268 $HplLHp$

Prueba:

 $Hpl\neg Np$ $I\neg\neg Np$ $ILHp$ A270 $H1$ Prueba: $\neg NN0$ A271 $pl1\supset Hp$

Prueba:

 $pl1\supset.HplH1$ $\supset.H1IHp$ $A270\supset]pl1\supset Hp$ A273 $Lpl1\equiv p$

Prueba:

(2) $p\supset Lp$ $\supset H Lp$ $\supset.Lpl1$ (3) $Lpl1\supset H Lp$ $\supset Lp$ $\supset p$ A274 $pl0\equiv\neg p$ A274/3 $Npl0\equiv Hp$ A274/5 $\neg pl\neg Lp$ A275 $p\supset Hp\supset.p\rightarrow Hp$

Prueba:

(2) $\neg p\vee Hp\supset.pl0\vee.Hpl1$ (3) $pl0\supset.p\wedge HqIp$ $\supset.p\rightarrow Hq$ (4) $HqI1\supset.p\wedge HqIp$ $\supset.p\rightarrow Hq$ (5) $\delta 2\supset.p\rightarrow Hq$

A275

A276 $p\rightarrow Lp$

Prueba:

(2) $p\supset H Lp$ (3) $2\supset.]p\rightarrow H Lp$ $p\rightarrow Lp$ A277 $Hp\rightarrow p$

Prueba:

(2) $Np\rightarrow L Np$ (3) $N L Np\rightarrow N Np$ A269 $Hp\supset.pl1$

Prueba:

 $\neg Np\supset.Npl0$ $\supset.N Npl1$ A272 $pl1\equiv Hp$

Prueba: A269, A271

A267

A272

A271

A267

A274/2 $Npl1\equiv\neg p$ A274/4 $\neg pl1\equiv Np$ A274/6 $pl0\equiv.Lpl0$

A257, A272, A274

(4), (3)

(5), (2)

A267

A275

(3), A267

A276

(2)

(4) $NLNp \rightarrow p$	(3)
$Hp \rightarrow p$	(4), A252
A277/2 $Hp \supset p$ Prueba: A277, A198	A277/3 1 Prueba: A270, A277/2
A278 $\neg p \rightarrow Np$	A278/2 $\neg p \supset Np$
A279 $Hp \supset .HpIp$	
Prueba:	
(2) $Hp \supset .pI1$	A272
(3) $HHp \supset .HpI1$	A272
(4) $Hp \supset .HpI1$	(3), A257
(5) $Hp \supset .pI1 \wedge .HpI1$	(2), (4)
$\supset .HpIp$	
A280 $1IH1$	
Prueba:	
$1I1 \supset .JJ1$	A272, A257
$2 \supset .JH1I1$	A272
A281 $1I \neg 0$ Prueba: A280	A282 $0I \neg 1$ A283 $0IH0$
A284 $\neg p \supset .p \supset pI1$	
Prueba:	
$\neg p \supset .\neg pI1$ A257/2, A272	
$\supset .\neg p \vee pI1$	
$\supset .p \supset pI1$	
A284/1 $p \supset pIp \vee .p \supset pI1$ Prueba: A284, A217	
A285 $p \equiv qI . q \equiv p$	
A286 $p \equiv q \wedge (p \supset p)I . p \equiv q$	
Prueba:	
(2) $\neg p \supset .p \equiv q \wedge (p \supset p)I . p \equiv q$	A284, A250
(3) $p \wedge \neg q \supset .\neg pI0 \wedge .qI0$	
$\supset .\neg p \vee qI0$	A250
$\supset .p \equiv qI0$	
$\supset .p \equiv q \wedge (p \supset p)I . p \equiv q$	

- (4) $p \wedge q \supset \neg p \vee q \wedge (\neg q \vee p) \vdash 0 \vee q \wedge 0 \vee p$ A179, A250
 $\quad \quad \quad \vdash q \wedge p$
- (5) $p \wedge q \supset p \supset p \mid p$ A165, A217
- (6) $p \wedge q \supset \delta 4 \wedge \delta 5$ (4), (5)
 $\quad \quad \quad \supset p \equiv q \wedge (p \supset p) \vdash q \wedge p \wedge p$ A250
 $\quad \quad \quad \vdash p \wedge q$
- (7) $\delta 4 \wedge \delta 6 \supset p \equiv q \wedge (p \supset p) \vdash p \equiv q$
- (8) $p \wedge q \supset \delta 4 \wedge \delta 6$ (4), (6)
 $\quad \quad \quad \supset \delta 7$ (7)
- (9) $p \supset p \wedge q \vee p \wedge \neg q$ A178
 $\quad \quad \quad \supset p \equiv q \wedge (p \supset p) \vdash p \equiv q$ (8), (3), A201/2
A286 (2), (9)
- A287 $p \equiv q \wedge (q \supset p) \vdash p \equiv q$
- A288 $p \equiv q \wedge (p \supset p \wedge q \supset q) \vdash p \equiv q$
- A289 $p \equiv q \vdash p \wedge q \vee \neg(p \vee q)$
- Prueba:
 $p \wedge q \vee \neg(p \vee q) \vdash p \supset q \wedge q \supset p \wedge p \supset p \wedge q \supset q$
 $\quad \quad \quad \vdash p \equiv q \wedge p \supset p \wedge q \supset p$
 $\quad \quad \quad \vdash p \equiv q$ A288
- A289/2 $p \wedge q \supset p \equiv q$ A289/3 $\neg p \wedge \neg q \supset p \equiv q$
- A290 $p \vee Np$ (Prueba: A278/2) A291 $N(p \wedge Np)$ Prueba: A290
- A292 $p \supset p \wedge q \equiv q$
- Prueba:
(2) $p \supset q \supset p \wedge q$ A135
(3) $p \supset p \wedge q \supset q$
 $\quad \quad \quad p \supset \delta 3 \wedge \delta 2$ (3), (2)
- A293 $Np \wedge q \vee p \equiv p \vee q$
- Prueba:
(2) $Np \wedge q \vee p \vdash p \vee Np \wedge p \vee q$
(3) $A290 \supset \delta 2 \equiv p \vee q$ A292
 $\quad \quad \quad \sigma 2 \equiv p \vee q$ (3), (2)
- A294 $\neg p \wedge q \vee p \equiv p \vee q$ (Prueba similar)
- A295 $\neg p \vee q \wedge p \vdash p \wedge q$
- Prueba:
 $\neg p \vee q \wedge p \vdash \neg p \wedge p \vee q \wedge p$
 $\quad \quad \quad \vdash 0 \vee q \wedge p$
 $\quad \quad \quad \vdash p \wedge q$
- A296 $p \mid q \supset p \mid p \mid p \mid q$ (Prueba: A200)

A297 $plq \supset . plq \sqcap N(plq)$

Prueba:

- | | |
|---|----------------|
| (2) $plq \supset . N(plp) \sqcap . plq$ | A296, A13, A12 |
| $\supset . plp \sqcap N(plq)$ | A200, A102 |
| (3) $plq \supset . plq \supset . plq \sqcap N(plq)$ | A200, (2) |
| A297 | (3), A145 |

A298 $plq \wedge (rls) \supset . plq \sqcap . rls$

Prueba:

- | | |
|---|----------------|
| (2) $plq \supset . plq \sqcap N(plq)$ | A297 |
| (3) $rls \supset . rls \sqcap N(rls)$ | A297 |
| (4) $rls \rightarrow (plq) \vee . plq \rightarrow . rls$ | A13 |
| (5) $plq \wedge (rls) \supset . rls \rightarrow N(plq) \vee . plq \rightarrow N(rls)$ | (2), (3), A250 |
| $\supset . rls \rightarrow N(plq) \vee . rls \rightarrow N(plq)$ | A209 |
| $\supset . rls \rightarrow N(plq)$ | |
| (6) $\delta 2 \supset . 5 \supset . plq \wedge (rls) \supset . rls \rightarrow . plq$ | |
| (7) $5 \supset .] \delta 2 \supset \delta \delta 6$ | (6) |
| (8) $plq \supset \delta 7$ | (2), (7) |
| (9) $plq \wedge \sigma 5 \supset . rls \rightarrow . plq$ | (8) |
| (22) $plq \wedge (rls) \supset . rls \rightarrow . plq$ | (9) |
| (23) $plq \wedge (rls) \supset . plq \rightarrow . rls$ | similarmente |
| (24) $plq \wedge (rls) \supset . \delta 23 \wedge \delta 22$ | (23), (22) |
| $\supset . plq \sqcap . rls$ | A207 |

A299 $plq \supset . p \equiv q$

Prueba:

- | | |
|---------------------------------|----------|
| (2) $plq \supset . p \supset q$ | |
| (3) $plq \supset . q \supset p$ | |
| $plq \supset . p \equiv q$ | (2), (3) |

A300 $plq \equiv (rls) \equiv . plq \sqcap . rls$

Prueba:

- | | |
|---|-----------|
| (2) $plq \equiv (rls) \sqcap . plq \wedge (rls) \vee \neg (plq \vee . rls)$ | A289 |
| (3) $\sigma \delta 2 \supset . plq \sqcap . rls$ | A298 |
| (4) $\delta \delta 2 \supset . \neg (plq) \wedge \neg (rls)$ | |
| $\supset . plq \sqcap 0 \wedge . rls \sqcap 0$ | |
| $\supset . plq \sqcap . rls$ | |
| (5) $\delta 2 \supset \delta 3$ | (3), (4) |
| (6) $\sigma 2 \supset \delta 3$ | (2), (5) |
| A300 | (6), A299 |

A301 $p \rightarrow Nql.q \rightarrow Np$ Prueba: A209, A300

A302 $p \rightarrow ql.p \vee qlq$ Prueba: A208, A300

A303 $p \supset (p \rightarrow q) \equiv p \rightarrow q$

Prueba:

- | | |
|---|--------------|
| (2) $\neg p \supset \neg p \vee 1$ | A257/2, A272 |
| $\supset p \supset (p \rightarrow q) \vee 1$ | |
| $\supset p \supset p \vee 0$ | A277/3 |
| (3) $\neg p \supset p \vee 0$ | A274 |
| $\supset p \rightarrow q$ | |
| (4) $\neg p \supset p \supset (p \rightarrow q) \wedge p \rightarrow q$ | (2), (3) |
| $\supset A303$ | A289/2 |
| (5) $p \supset \neg p \vee 0$ | A274 |
| $\supset \neg p \vee (p \rightarrow q) \vee p \rightarrow q$ | |
| $\supset p \supset (p \rightarrow q) \vee p \rightarrow q$ | |
| $\supset A303$ | A299 |
| A303 | (4), (5) |

A304 $p \vee q \supset r \rightarrow p \vee q$

Prueba:

- | | |
|--|------|
| $p \vee q \supset r \supset p \vee q \vee r$ | A298 |
| $\supset r \supset r \rightarrow p \vee q$ | A206 |
| $\supset r \rightarrow p \vee q$ | A303 |

A305 $p \vee q \supset (r \vee s) \supset p \vee q \rightarrow r \vee s$

Prueba:

- | | |
|--|----------|
| (2) $\neg(p \vee q) \supset p \vee q \vee 0$ | |
| $\supset p \vee q \rightarrow r \vee s$ | |
| $\supset A305$ | |
| (3) $r \vee s \wedge (p \vee q) \supset p \vee q \vee r \vee s$ | A298 |
| $\supset p \vee q \rightarrow r \vee s$ | A206 |
| (4) $p \vee q \wedge \neg(r \vee s) \supset \neg(p \vee q \supset r \vee s)$ | |
| $\supset A305$ | |
| (5) $p \vee q \supset \sigma_3 \vee \sigma_4$ | |
| $\supset p \vee q \supset (r \vee s) \supset p \vee q \rightarrow r \vee s$ | (3), (4) |
| A305 | (2), (5) |

A306 $p \rightarrow q \rightarrow r \rightarrow p \rightarrow r$

Prueba:

- | | |
|---|----------|
| (2) $p \wedge q \vee p \supset p \wedge r \vee p \wedge q \wedge r$ | |
| (3) $q \wedge r \vee q \supset q \wedge p \vee q \wedge r \wedge p$ | |
| (4) $\sigma_2 \wedge \sigma_3 \supset \delta_2 \wedge \delta_3$ | (2), (3) |
| $\supset p \wedge r \vee p \wedge q$ | |
| (5) $\sigma_4 \supset p \wedge q \vee p \wedge r \vee p \wedge q$ | A166 |
| $\supset p \wedge r \vee p$ | |

(6) $p \wedge q \vdash q \wedge r \vdash p \wedge r$
 $\vdash q \wedge r \rightarrow p \wedge r$
A306

A305
(6), A305

Capítulo 9

Una extensión de *At*: el sistema infinivalente y tensorial *Aj*

Vamos ahora a exponer un sistema que constituye una extensión no conservativa de *At*, a saber: *Aj*. Es extensión no conservativa de *At* porque toda fbf de *At* lo es de *Aj* y cada una de **esas** fórmulas que es un teorema de *At* también lo es de *Aj*; pero no viceversa. (P.ej., el esquema $\lceil p \supset a \rightarrow p \rceil$ sólo contiene —aparte de la letra esquemática ‘p’— signos del vocabulario de *At*; sin embargo, siendo como es un esquema teorematizado de *Aj*, no lo es de *At*.) Además, *Aj* es una extensión recia de *At* (según la terminología del Capítulo 6) ya que cada regla de inferencia de *At* es también una regla de inferencia de *Aj*.

Aj es una aproximación a la lógica *Abp* esbozada al final del Capítulo 5 (si bien, por ser más claro el modelo que para tal lógica vino expuesto al final del Capítulo 6, convendrá que el lector, en lo que sigue, se remita a menudo a ese modelo para comprobar la validez en él de cada esquema axiomático de *Aj*). Sólo que *Abp* es un sistema semánticamente definido, al paso que *Aj* es un sistema axiomático, sintácticamente definido. Los modelos expuestos para *Abp* tanto en el Capítulo 5 como en el 6 son modelos idóneos de *Aj* pero no modelos característicos —según la terminología introducida en el cap. 6. En el capítulo 12 veremos un procedimiento para descubrir modelos característicos de sistemas como *Aj*.

Si bien *Aj* es una extensión de *At*, no voy a presentar *Aj* indicando qué hay que añadir a *At* para obtener el sistema nuevo, sino que voy a seguir otro procedimiento, que resulta más elegante porque con él se disminuye el número de signos primitivos y de esquemas axiomáticos.

Nuevas lecturas

Usaremos, además de signos ya presentados en el capítulo anterior —o más atrás incluso en este trabajo— y que, naturalmente, conservarán ahora las lecturas brindadas para ellos, otros signos, que tendrán estas lecturas:

$\lceil p \downarrow q \rceil$ se leerá: «Ni p ni q»

$\lceil Bp \rceil$ se leerá: «Es afirmable con verdad que p»; «Es verdad en todos los aspectos que p».

$\lceil Jp \rceil$ se leerá: «Es [al menos] relativamente cierto [=verdadero] que p»; «En cierto modo [es verdad que] p»; «En algunos aspectos, p».

$\lceil p = q \rceil$ se leerá: «El hecho de que p es estrictamente equivalente al de que q».

$\lceil p \Rightarrow q \rceil$ se leerá: «El hecho de que p implica estrictamente al de que q»; «El hecho de que p es, en todos los aspectos, a lo sumo tan verdadero como el de que q»

$\lceil fp \rceil$ se leerá: «Es un tanto cierto y falso a la vez que p».

$\lceil \uparrow p \rceil$ se leerá: «El hecho de que p es un sí es no falso [o irreal]» $\approx$ «Es infinitesimalmente falso que p».

$\lceil Xp \rceil$ se leerá: «Es muy cierto que p».

$\lceil Kp \rceil$ se leerá: «Es [al menos] un poco cierto que p».

Una observación incidental que conviene hacer sobre la antepenúltima lectura: en español la conjunción ‘que’ tiene varios usos; uno de ellos es similar al latín ‘quam’ (o al inglés ‘than’) para regir el segundo término de una comparación de desigualdad (inferioridad o superioridad, principalmente): ‘mayor que’, ‘menos guapa que’; otro uso es el de nominalizadora de una oración: ‘Tengo frío’ al venir precedida de ‘que’ da como resultado la “oración subordinada” ‘que tengo frío’, la cual es un sintagma nominal que puede ser sujeto o complemento en una oración principal (‘Dices que tengo frío’, ‘Que tengo frío es obvio’, ‘Si me arropo es por una razón, a saber que tengo frío’, etc.). Pues bien, ¿qué sucede cuando, por la estructura de la frase general, han de figurar dos ‘que’ sucesivos, cada uno en uno de esos dos usos? (‘Prefiero que me traigas una chamarra que que enciendas el fogón’, ‘Es más sano que comas limones

que que te inyectes vitamina C', etc.). Hay varios procedimientos para evitar ese contraste que resulta acaso cacofónico o poco perspicuo. Uno es la inserción de un 'no' expletivo entre ambos 'que's. ('Más vale que sobre que no que falte'.) (El 'no' expletivo es en el español contemporáneo menos frecuente que en otras épocas, pero está muy vivo en el habla popular. Como un testimonio de su amplio uso en otro tiempo, recuérdese esta famosa frase del Romance de los siete Infantes de Lara: 'Que más vale un caballero de los de la Flor de Lara que no ciento ni doscientos de Bureba la preciada'.) Otro procedimiento es elidir uno de los dos 'que's; otro, intercalar entre ellos un 'el': 'Prefiere que vivas tu bien que el que le llegue el sueldo a fin de mes' (o también 'a que le llegue...'; pero este uso de la preposición 'a' no es posible con construcciones comparativas con 'más' o 'menos': *'Es más importante que cumplas bien tu tarea a que queden satisfechos tus superiores').

Reglas de formación

El sistema A_j es un dúo $\langle \mathfrak{S}, \mathfrak{R} \rangle$ donde $\mathfrak{S}$ es aquel más pequeño subconjunto de $\mathfrak{S}$ que: (1º) abarque a todas las instancias de cada uno de los esquemas axiomáticos A01 a A23 expuestos más abajo; y (2º) esté cerrado con respecto a cada regla de inferencia perteneciente a $\mathfrak{R}$ —que es el conjunto de rinf1 y rinf2, expuestas más abajo—; donde $\mathfrak{S}$ viene caracterizado por las dos siguientes reglas de formación:

(1ª) 'a' $\in \mathfrak{S}$;

(2ª) Si 'p', 'q' $\in \mathfrak{S}$, entonces también son miembros de $\mathfrak{S}$ estas ristras: 'p↓q', 'Bp', 'Hp', 'plq', 'p•q'.

Definiciones

'Np' abr 'p↓p'	'p∨q' abr 'N(p↓q)'	'p∧q' abr 'Np↓Nq'
'v' abr 'Na'	'¬p' abr 'HNp'	'½' abr 'ala'
'Lp' abr 'N¬p'	'0' abr '½la∨¬(½lN½)'	'Xp' abr 'p•p'
'Kp' abr 'NXNp'	'1' abr 'N0'	'p⊃q' abr '¬p∨q'
'Sp' abr 'p∧Np'	'np' abr 'p•v'	'mp' abr 'NnNp'
'p→q' abr 'q∧plp'	'p≡q' abr 'p⊃q∧.q⊃p'	'Yp' abr 'pla∧p'
'fp' abr '¬Yp∧p'	'p↓q' abr 'p→q∧ ¬(q→p)'	'Jp' abr '¬B¬p'
'p↔q' abr 'B(plq)'	'p⊥q' abr '¬(p=q)'	'p⇒q' abr 'B(p→q)'
'↑p' abr 'YNp'	'fp' abr 'fSp'	'p&q' abr 'Lp∧q'
'p↑q' abr '↑p∧fq∨.↑q∧fq'		

Reglas de inferencia

rinf1 p, p⊃q ⊢ q

rinf2 p ⊢ Bp

Sobre la primera de esas dos reglas, el *Modus Ponens* (rinf1), nada hay que añadir a lo ya dicho al respecto en capítulos anteriores. La segunda regla, en cambio, merece unos comentarios. Una de las lecturas del funtor 'B' es 'Es afirmable con verdad que'; otra es 'Es en todos los aspectos verdad que' o 'Es desde todos los puntos de vista [objetivos] verdad que'. La introducción de este funtor tiene sentido sólo si uno entiende al mundo como

englobando aspectos diversos, de suerte que un mismo hecho o estado de cosas pueda tener más realidad en unos aspectos, menos realidad en otros. P.ej., aunque es, en todos los aspectos, verdad que Rangún es una ciudad bonita o no lo es, cabe que en algunos aspectos sea preciosa —y hasta en algunos de ellos totalmente bonita, o sea tanto que nada sea más bonito—, en otros bastante fea, en otros intermedia; en algunos incluso totalmente exenta de belleza; de suerte que no cabría afirmar con verdad ‘Rangún es una ciudad bonita’ —porque en algunos aspectos de lo real eso sería del todo falso— ni tampoco ‘Rangún no es bonita’ porque en algún aspecto sería esto último lo que fuera del todo falso (a saber: en aquellos aspectos en los cuales Rangún sea plenamente bonita).

Si el mundo es pluriaspectual, entonces tiene sentido diferenciar entre una oración ‘ p ’ y otra ‘ Bp ’ que diga que es en todos los aspectos verdadero (e.d. real) el estado de cosas consistente en que exista o suceda que p . Muchos ven al mundo como pluriaspectual y creen que, aunque siempre cabe —en virtud del tercio excluso— preguntar ‘¿Sí o no?’, a veces no cabe en absoluto (porque no es afirmable con verdad, al no ser verdad en todos los aspectos) responder ‘Sí’, pero tampoco ‘No’, salvo relativizando con respecto a tal o cual aspecto.

Una concepción que acepte tales situaciones —verdaderas o existentes tan sólo en ciertos aspectos— puede considerarse un género de relativismo. Sin embargo, no se trata del relativismo usual, pues éste último es subjetivista, al concebir la existencia de verdades relativas al sujeto («verdadero para mí, falso para tí»), y en su forma más radical rechaza verdades que no sean relativas. En cambio el relativismo alético articulable con el empleo del functor ‘ B ’ es un relativismo de cuño objetivista y no subjetivista: hay —habría según esa concepción— verdades objetivamente relativas; relativas a tales o cuales aspectos de la realidad.

Si eso es acertado, entonces se entiende el empleo del functor ‘ B ’ y la aceptación de rinf2 . De la premisa ‘ p ’ cabe inferir correctamente que es afirmable con verdad que p ; porque el **afirmar** esa premisa se hace creyendo —o al menos dando a entender— que la misma es [verídicamente] **afirmable**.

Ahora bien, la introducción de esta regla de inferencia acarreará una falla del metateorema de la deducción. Porque, aunque siempre de ‘ p ’ cabe lícitamente inferir ‘ Bp ’, no siempre será verdad ‘ $p \supset Bp$ ’. De hecho este esquema no es teorema en A_j . Por ende no es verdad en general que, si de las premisas ‘ p^1 ’, ‘ p^2 ’, ..., ‘ p^n ’, se infiere —en virtud de las reglas de inferencia de A_j — la conclusión ‘ q ’, entonces es teorema en A_j el esquema ‘ $p^1 \supset p^2 \supset \dots \supset p^n \supset q$ ’. Eso es verdad con ciertas restricciones no más; p.ej. que la inferibilidad de ‘ q ’ a partir de las premisas venga autorizada por rinf1 más los esquemas axiomáticos (o sea: sin que intervenga rinf2); o bien que, si no, cada premisa empiece por el functor ‘ B ’ (o por la ristra ‘ $\neg B$ ’). Y también es correcto el metateorema de la deducción en esta otra enunciación del mismo: si $p^1, \dots, p^n \vdash q$, entonces $Bp^1 \supset \dots \supset Bp^n \supset q$. (La corrección de estos resultados no vendrá demostrada en el presente capítulo; queda eso como ejercicio para el lector —al cual le será fácil hacerlo a partir de los procedimientos de prueba expuestos en el capítulo precedente.)

Esquemas axiomáticos

A01 $p \wedge q \supset p$

A02 $r \wedge s \vdash p, p \vdash q, q \vdash s \vdash q$

A03 $p \vdash q, r \vdash q \vdash p$

A04 $p \bullet q \vdash (Kp \bullet Kq)$

A05 $p \wedge q \supset p \bullet q$

A06 $q \wedge p \vdash p \vdash q$

A13 $p \vdash q \supset q \supset p$

A14 $p' \wedge p \vdash q \bullet r \bullet s \vdash s \bullet r \bullet p \wedge p'$

A15 $mp \rightarrow np \equiv Yp \vee \uparrow p$

A16 $Np \supset mp \rightarrow mnp$

A17 $p \& (q \rightarrow np \vee plmq) \vee p \rightarrow q$

A18 $Bp \vee B \neg B Lp$

A07 $H p \wedge H q \mid L H(p \wedge q)$ A08 $p \mid q \supset H p \vee H r \mid H(q \vee r)$ A09 $p \bullet q \rightarrow p$ A10 $p \bullet 1 \mid p$ A11 $p \mid N q \mid N p \mid q$ A12 $p \mid p \mid 1/2$ A19 $B p \supset B p \mid p$ A20 $p \rightarrow q \& B p \rightarrow B q$ A21 $q \bullet s \rightarrow (p \bullet r) \wedge (p \mid q) \supset s \rightarrow r \wedge r \rightarrow s \wedge f s \supset q \mid m p \wedge \uparrow r$ A22 $f p \wedge f q \supset p \bullet q \mid p$ A23 $q \bullet m p \mid m(p \bullet q) \vee p \uparrow q$

Comentarios sobre la base axiomática de *Aj*

Parece preferible —a fin de atenernos a los límites aconsejables para un librito de naturaleza meramente introductoria— no detenernos ya en demostraciones de esquemas teorematísicos. Con algunos rodeos previos, pueden demostrarse en *Aj* todos los esquemas teorematísicos demostrados en *At* (con la salvedad de que el metateorema de la deducción ha de aplicarse aquí con restricciones —sin que ello afecte empero a la obtención de teoremas dentro del sistema); luego, a partir de ellos —más los otros axiomas del sistema—, se prueban nuevos esquemas teorematísicos. Vale más aprovechar el poco espacio que queda antes de abordar escuetísimamente —en el capítulo siguiente— el cálculo cuantificacional para hacer unos pocos comentarios sobre la motivación y la aplicación de los funtores que forman parte del vocabulario de *Aj* y sobre alguno de los axiomas del sistema —así como también sobre algún que otro teorema. (Pero no se trata en estas pocas líneas de ofrecer un comentario detallado, sino que dejará en el silencio a la mayor parte de los esquemas axiomáticos; para una discusión punto por punto sería menester el espacio de todo un libro.) En verdad es filosóficamente más esclarecedor comentar los teoremas que los axiomas, puesto que el mérito principal de un pequeño cúmulo de esquemas oracionales que se postulan axiomáticamente es el de que con ellos más las reglas de inferencia del sistema [se aproxime uno lo más posible a la meta consistente en que] vengan demostrados **sólo todos** los teoremas que convenga demostrar (e.d. todos ellos pero sólo ellos); optar por tal base axiomática en vez de otra es, pues, algo que se hace menos en función de la plausibilidad propia de los axiomas en sí mismos que a tenor de su capacidad para producir sólo todos los teoremas que se desea probar; o sea: de cuánto permitan probar de aquellos que se desea tener como teoremas y de cuánto no permitan probar de aquello que se desea no tener como teoremas; son dos fines que hay que equilibrar y sopesar en la balanza; tomando además en consideración un principio de economía —a saber: que la base axiomática que se postule para alcanzarlos sea tan parsimoniosa como quepa; lo cual a su vez resulta de la interacción de varios factores que no siempre corren parejos: número de esquemas, longitud de los mismos, complejidad o dificultad de las definiciones en ellos empleadas. Al lector puede parecerle —y de hecho ése es el principal reproche que se ha dirigido contra este sistema— que 23 son muchos esquemas axiomáticos —y no todos ellos cortos— así como también que las definiciones son un tantico enrevesadas; aunque se reconocerá unánimemente la sencillez de las reglas de inferencia. Ahora bien, lo que se trata de calibrar es cuán costosa es la base axiomática para el resultado que con ella quepa obtener. Y, habida cuenta de ello, *Aj* es el más económico sistema de lógica existente. En efecto: hay un metateorema (todavía no publicado) que revela que *Aj* es una **extensión cuasiconservativa** de cada uno de los sistemas de lógica finivalentes; donde una teoría $\mathfrak{J}$ es una extensión **cuasiconservativa** de $\mathfrak{J}'$ si además de ser una extensión de $\mathfrak{J}'$ es tal que existe en $\mathfrak{J}$ un functor de afirmación $\$$ tal que sólo todos los teoremas de $\mathfrak{J}'$ son tales que las fórmulas $\lceil p \rceil$ de $\mathfrak{J}$ que tengan el mismo vocabulario de $\mathfrak{J}'$ son tales que, $\lceil \$p \rceil$ es un teorema de $\mathfrak{J}$. Lo cual quiere decir que —en un sentido apropiado— $\mathfrak{J}$ contiene a $\mathfrak{J}'$, y lo hace de tal manera que $\mathfrak{J}'$ es simplemente como una faceta o una vertiente de $\mathfrak{J}$, o sea: $\mathfrak{J}$ refleja un punto de vista superior que engloba al de $\mathfrak{J}'$ pero trascendiéndolo —trascendiéndolo por englobar

también a cualquier otro punto de vista limitado (limitado o ceñido a un cálculo verifuncional finivalente).

La capacidad expresiva de *Aj* es infinita. Existe en efecto otro metateorema según el cual pueden definirse en *Aj* infinitos funtores demostrablemente no equivalentes entre sí. Sin que con ello agote *Aj* la riqueza de la lengua natural, ni mucho menos, sí constituye una mejor y mayor aproximación que ningún otro sistema axiomático hasta ahora puesto en pie al ideal de representar mediante la notación simbólica lo más posible del vocabulario lógico usado en la lengua natural y dar cuenta de la corrección o incorrección de inferencias que involucren nada más que ese vocabulario. Existen otros tratamientos con aspiraciones parecidas aunque más modestas —ciertas lógicas de lo difuso que presentan una serie de puntos de disparidad respecto a *Aj*—; pero no poseen ninguna de estas características de *Aj*: constituir un sistema axiomático —recursivamente axiomatizado—; ser una extensión cuasiconservativa de cualquier sistema finivalente verifuncional; poseer —en el sentido recién apuntado— una capacidad expresiva infinita; y, por último, ser un sistema paraconsistente. Hasta donde alcanza el conocimiento del autor de estas páginas todavía no existe ningún otro sistema —excepto versiones anteriores de *Aj* (a las que se dio la misma denominación pese a haberse entre tanto modificado algo la base axiomática) o sistemas de la misma familia *A*— que posea al menos dos de esas cuatro características (salvo la combinación de la primera y la última, en un sistema que se parece mucho a *At* y que ha sido puesto en pie —de manera totalmente independiente— por la investigadora brasileña Ítala d'Ottaviano con la colaboración de Newton da Costa; mas ese sistema carece de las otras dos características).

Diferencia entre la conyunción natural ' $\wedge$ ' y la superconyunción ' $\cdot$ '

Uno de los rasgos más sorprendentes —para algunos habituados a la sencillez de la lógica clásica y de otros sistemas así— es la existencia en *Aj* de varias conyunciones; en particular, el contraste que se da entre ' $\&$ ' y ' $\wedge$ ' y ' $\cdot$ '. Surgen aquí varios problemas. Uno es el de cuán verdad sea que —a tenor de las características que poseen dentro del sistema— esos símbolos puedan venir presentados como sendas escrituras de expresiones de la lengua natural. Otro es el de si, independientemente de ello, se justifica su presencia en el sistema —quizá como signos para los que no haya ninguna lectura compacta en lengua natural. Un tercer problema es el de si, aun aceptando que, por algunos de sus rasgos, son aceptables simbolizaciones —o escrituras logográficas— de sendas expresiones copulativas de la lengua natural, son plausibles todas las características de cada uno de tales funtores —o sea: el problema de cuán plausibles sean los esquemas teorematizados en que aparezcan.

Con respecto al primer problema, cabe decir que lo más saliente de ' $\&$ ' —por oposición a ' $\wedge$ '— es que ' $p \& q$ ' será: cuando ' p ' sea del todo falso, también del todo falso; y, cuando no, tan verdadero como lo sea ' q '; o sea: ' $p \& q$ ' es una fórmula conyuntiva (copulativa) cuyo grado de verdad o falsedad depende de sólo dos factores: de que el primer conyunto sea poco o mucho verdadero, y de cuán verdadero sea el segundo si el primero no es del todo falso. Por ello en cada aspecto de lo real el grado de verdad de ' $p \& q$ ' es o bien el de ' p ' o bien el de ' q ' (el de ' p ' si éste es nulo o inexistente). Es una conyunción sensible a cuán verdadero sea el segundo conyunto pero, con respecto al primer conyunto, sensible no más a si éste es, **poco o mucho** verdadero. ¿Hay alguna locución copulativa (e.d. conyuntiva) así en la lengua natural? Aparentemente sí. En español hay seguramente varias con respecto a las cuales cabe argüir que son así. Una es el gerundio antepuesto a la "oración principal" —al menos en uno de sus usos frecuentes (puede que haya otros): 'Gustándole la astrofísica, Jacinto consagra sus estudios principalmente al álgebra'. En esa oración se enfatiza el segundo conyunto (la "oración principal"), pero la oración total sería del todo falsa si fuera enteramente falso lo dicho por la "oración gerundiva", al paso que, si no lo es, cuán verdadera sea la oración

total parece depender sólo de cuánto lo sea la principal. Por supuesto en la lengua natural cabe invertir el orden, con tal de mantener la asimetría por medio de ese formante —el gerundio—; luego cabría tener otra conjunción en *Aj*, ‘ $\wedge$ ’, definible así: ‘ $p \wedge q$ ’ abr ‘ $p \wedge Lq$ ’; pero no hace falta, y para nuestros propósitos resultaría ociosa. Otra expresión de la lengua natural que parece interpretable de la misma manera es ‘mientras que’ —especialmente cuando va correlacionada con una ocurrencia, en la oración principal, de ‘sobre todo’; cierto que puede comportar una “connotación” adversativa, mas no forzosamente: ‘Mientras que Luis se defiende en francés, [sobre todo] habla bien el inglés’. (Por cierto, en inglés la conjunción ‘while’ tiene aún más acusadamente ese carácter que la española ‘mientras que’.) En otros idiomas existen procedimientos parecidos a los del castellano para expresar algo así como ‘&’: ciertos usos del ablativo absoluto latino; o, en latín también, ciertos usos de ‘cum’, correlacionados o no con sendas ocurrencias de ‘tum’ en la oración principal (si bien habría argumentos para alegar que, a veces al menos, esa correlación ‘cum... tum’ representa más bien el ‘no sólo... sino también’, o sea algo parecido a la conjunción ‘•’ de *Aj*; son temas harto complejos); cuando el ‘cum’ tiene ese uso —si es que lo tiene— suele venir correlacionada con un ‘praesertim’ [e.d. ‘sobre todo’] en la oración principal, o alguna partícula similar. Podrían seguramente aducirse construcciones alegablemente formalizables con ‘&’ en muchos otros idiomas, indoeuropeos o no.

Pero, ¿no se trata —como alegaría el clasicista— de meras variaciones estilísticas, o simples alomorfos de ‘y’ en distribución libre o en distribución complementaria, pragmática o estilísticamente condicionada? Es difícil hallar argumentos contundentes a favor o en contra de esa hipótesis. Aquí interviene un problema general de metodología. Si tenemos una concepción del lenguaje que sea, en algún sentido, “figurativa” —que tienda a concebir al lenguaje de alguna manera afín a como lo vio el atomismo lógico de Russell y el Wittgenstein del *Tractatus*—, entonces parece preferible no relegar a la pragmática más que lo imprescindible; al menos sólo aquello para lo que no se hayan encontrado tratamientos semánticos claros. P.ej., piénsese en la oposición de muchos a reglas de inferencia como la de adición [a saber: $p \vdash p \vee q$], alegando lo improcedente en general de concluir, p.ej., que Alicia es venezolana o belga a partir del aserto de que es venezolana; pero, admitiendo lo “raro” de tal inferencia [en muchos entornos de elocución], no hay [hasta ahora] para ese problema ninguna solución satisfactoria puramente semántica, ni siquiera muy clara, al paso que sí hay una solución pragmática clara y convincente: que la conclusión no suele ser comunicacionalmente pertinente en los más contextos o entornos de elocución, y ello a tenor de cierta regla de economía comunicacional (tentativamente formulable así: No proferir un enunciado más largo que otro, en vez de éste último, si el primero vehicula menos información interesante que el segundo —interesante para el interlocutor en ese contexto).

Mas no sucede nada de eso con el caso que nos ocupa. Sí, el contexto de elocución puede hacer más interesante un conyunto que el otro; y por eso puede el locutor, al aseverar una conyunción de ambos, recalcar más el más interesante. Sin embargo, cuando es comunicacionalmente pertinente una conyunción de dos enunciados, también lo es cualquier otra conyunción entre ellos (aunque acaso menos pertinente). Y hay muchos contextos en los cuales no es seguro que la mayor pertinencia comunicacional de uno de los conyuntos sea **lo único** que hace optar por una conyunción en vez de otra. Además, aun suponiendo que así sea, es perfectamente explicable optar por ‘ $p \& q$ ’ en vez de ‘ $p \wedge q$ ’ cuando ‘ q ’ es comunicacionalmente mucho más pertinente que ‘ p ’ si ‘&’ tiene las características que le hemos atribuido.

(Es más: incluso el recurso a procedimientos prosódicos de énfasis o entonación, en casos así —en lugar del empleo del gerundio o construcciones similares— puede también explicarse mejor de manera que se reconozca una diferencia semántica y no meramente pragmática

o estilística: «**p y q**» —con énfasis aquí representado por la negrilla— puede ser un modo normal de decir lo que escribimos $\lceil p \wedge q \rceil$.)

Paso ahora a la conyunción de insistencia, o **superconyunción**, ' $\bullet$ '. Es mucho más interesante —y también más discutible— que '&', pues ésta se define con ' $\wedge$ ' y con el functor afirmativo 'L'. ' $\bullet$ ' es muy especial. Cuando uno de los dos conjuntos es o completamente verdadero o infinitamente falso (e.d. cuando es: o plenamente verdadero, o sólo infinitesimalmente verdadero, o del todo falso), entonces $\lceil p \bullet q \rceil$ tiene el mismo grado de verdad o falsedad que $\lceil p \wedge q \rceil$. Pero —según lo dice el esquema axiomático A22— cuando no sucede eso [en absoluto] —ni tampoco que sean **ambos** sólo infinitesimalmente falsos—, entonces hay diferencia; en tales casos, $\lceil p \bullet q \rceil$ es una conyunción menos verdadera que $\lceil p \wedge q \rceil$. Y es que el valor de verdad de $\lceil p \bullet q \rceil$ resulta de una interacción de los conjuntos; por expresarse esa conyunción como una insistencia (el un conjunto **más** el otro, el uno **y también** el otro, el uno **así como** el otro, el uno **y** el otro [aquí el énfasis sólo afecta a la conyunción], o sea: **no sólo** el uno **sino también** el otro) lo así expresado cuadra menos con el par de situaciones representadas, respectivamente, por $\lceil p \rceil$ y por $\lceil q \rceil$ cuando éstas distan de ser infinitamente verdaderas, aunque ninguna sea tampoco infinitamente falsa o inexistente. P.ej., si Joaquín es sólo a medias (no quiere decirse: exactamente en un 50%) aficionado a la akadiología y sólo a medias hábil en resolver problemas algebraicos, quien diga 'Joaquín es no sólo aficionado a la akadiología sino también hábil en resolver problemas algebraicos', al insistir —según lo está haciendo— en la verdad conjunta de ambos conjuntos, en la situación resultante de tal darse el uno con el otro, estará diciendo algo todavía más falso, más alejado de la verdad plena, que si meramente dijera 'y' en vez de 'no sólo... sino también'; porque, si dijera 'y', pondría, aseveraría, a cada uno de los conjuntos y los uniría por un lazo que, no recalcando su unión, sino meramente señalando su copresencia o coexistencia en la realidad, sería menos "exigente" que lo es ese lazo más unitivo, más interactivo, que se vehicula con el 'no sólo... sino también'; un lazo que para alcanzar un grado de verdad requiere más verdad de los conjuntos (salvo —cabe reiterar— si uno de ellos es infinitamente verdadero o falso). Decir «p y q» es casi como decir «p, q», yuxtaponiendo los conjuntos. Decir que no sólo p sino que además q es atribuir a esos dos hechos [cuya existencia, al decir eso, está aseverando uno] una relación más fuerte, un darse el uno **con** el otro; sólo que, eso sí, si es verdad $\lceil p \wedge q \rceil$ también lo es $\lceil p \bullet q \rceil$. De ahí el esquema axiomático A05: $\lceil p \wedge q \rceil \supset \lceil p \bullet q \rceil$. En virtud del esquema A09, $\lceil p \bullet q \rceil \rightarrow p$ cabe demostrar como teorema el esquema $\lceil p \bullet q \rceil \rightarrow \lceil p \wedge q \rceil$; pero no lo recíproco, o sea no $\lceil p \wedge q \rceil \rightarrow \lceil p \bullet q \rceil$. Así pues, si son teoremas tres de las cuatro combinaciones posibles con las fórmulas conyuntivas $\lceil p \wedge q \rceil$ y $\lceil p \bullet q \rceil$ enlazándolas por uno u otro de los dos funtores condicionales, ' $\supset$ ' y ' $\rightarrow$ '. La única que no es teorema es precisamente $\lceil p \wedge q \rceil \rightarrow \lceil p \bullet q \rceil$, porque en ella la apódosis implicativa puede ser **menos** verdadera que la prótasis (y, para que sea verdad $\lceil p \rightarrow q \rceil$ es menester que el grado de falsedad de $\lceil q \rceil$ sea a lo sumo tan grande como el de $\lceil p \rceil$). No es verdad, p.ej., que, en la medida en que Grecia es un país adelantado y próspero, es no sólo un país adelantado, sino también próspero; la apódosis, al ser una conyunción "insistencial" —una que recalca enfáticamente el darse de cada uno de los dos conjuntos **con** el otro—, es menos verdadera, si es que (según cabe suponer) cada uno de esos dos conjuntos dista hoy un tanto de ser infinitamente verdadero, aunque también dista de ser infinitamente falso.

A lo mejor es errada toda esta interpretación de esas construcciones de la lengua natural (que, desde luego, parecen existir, con unos u otros matices, en todos los idiomas). A lo mejor semánticamente 'y' y 'no sólo... sino también' son exactamente sinónimos, siendo la disparidad puramente estilística, pragmática. Eso es en todo caso lo que alegarán los classicistas; y no serán los únicos, ni mucho menos (pues el autor de este opúsculo no conoce ningún sistema axiomático fuera de la familia A que posea esa dualidad de conyunciones —aunque sí hay

teorías no axiomatizadas de conjuntos difusos, en la línea de Lofti Zadeh, que tienen una dualidad así). Lo que es menester es, en vez de sentar una tesis u otra como dogma de fe, ofrecer argumentos. Lo que precede es un razonamiento que parece bastante plausible a favor de la tesis aquí propuesta (la diferencia **semántica** entre 'y' y 'no sólo... sino también', y la representabilidad [siquiera aproximada] de esas partículas, respectivamente, por ' $\wedge$ ' y por ' $\bullet$ ', en el sistema A_j). Es dudosísimo que baste con aducir consideraciones de economía para, convincentemente, argüir a favor de la tesis opuesta —la de que la diferencia es meramente pragmática. Porque: (1º) de haber ahí diferencia semántica, se explica mucho mejor la opción en un contexto —pragmáticamente condicionado— a favor del uso de una de las dos conyunciones (sin ayuda de la diferencia semántica, la explicación pragmática será bastante complicada, quizá embrollada); (2º) los factores pragmáticos no parecen ahí los únicos pertinentes o explicativos —en todo caso está disponible, siendo clara y sencilla, la explicación semántica aquí brindada; (3º) el principio metodológico de recurso parsimonioso a soluciones pragmáticas —en la medida en que sea un principio sano y recomendable (como lo es a juicio de quien esto escribe)— abona en contra de desechar una solución semántica simplemente porque con ella **la semántica** es más complicada que si se busca exclusivamente una solución pragmática.

Consideraciones sobre la pluralidad de funtores bicondicionales

Tenemos en A_j tres funtores bicondicionales: ' $\equiv$ ', ' I ' y ' $=$ '. La diferencia es ésta: ' $p \equiv q$ ' es verdadero (y, cuando lo es, lo es sólo en la medida en que lo sean ambos miembros, o sea en la medida en que sea verdadera la [mera] conyunción ' $p \wedge q$ ', a menos que ambos conjuntos sean **totalmente** falsos, y entonces ' $p \equiv q$ ' será plenamente verdadero) sólo si, o bien ambos miembros, ' p ' y ' q ', son verdaderos (poco o mucho), o bien ambos son enteramente falsos. La verdad de ' $p \equiv q$ ' tienen como condición necesaria y suficiente que no suceda en absoluto que uno de los dos miembros sea verdadero o existente mientras el otro sea totalmente falso o inexistente. ' $p \equiv q$ ' es verdad cuando no se da en absoluto una situación en la que uno de los dos esté presente y el otro completamente ausente. ' $p \equiv q$ ' es demostrablemente equivalente al esquema ' $p \wedge q \vee \neg(p \vee q)$ ': o ambos son reales (=verdaderos) o ninguno lo es en absoluto. (Por este último disyunto sucede que si es del todo falso que p y también del todo falso que q , entonces ' $p \equiv q$ ' es completamente verdad.) Saber que, en tal aspecto de lo real, es verdad ' $p \equiv q$ ', o sea « p ssi q », es saber que en ese aspecto de la realidad o están ambos o no está, en absoluto, ninguno de ellos. Aun suponiendo que se sepa **en qué medida es verdad** ' $p \equiv q$ ', no por ello se sabe cuán próximos o alejados estén los grados de verdad de ' p ' y de ' q '. Así, p.ej., saber que Norberto está de mal humor si y sólo si le han llevado la contraria no es saber cuánto sea su mal humor en proporción a cuánto le hayan llevado la contraria (quizá Norberto es tan atrabiliario y presuntuoso que en muchos de tales casos sea diez veces más cierto que está de mal humor que no que le han llevado la contraria).

En cambio, ' I ' es un functor que requiere, para que sea verdad ' $p I q$ ', que, en aquellos aspectos (momentos, lapsos de tiempo, o lo que sea) en que esto sea verdad, ambos miembros, ' p ' y ' q ', sean igual de verdaderos (o falsos) el uno que el otro, sin que ninguno exceda al otro.

En verdad hay unos cuantos esquemas de A_j que son teorematícos y en los cuales ' I ' es el functor principal. P.ej. éstos: $A06$ (o sea ' $q \wedge p \vee p I p$ ') y su "dual" ' $q \vee p \wedge p I p$ ' (son los llamados 'principios de absorción'); la idempotencia: ' $p \vee p I p$ ' y ' $p \wedge p I p$ '; la asociatividad ' $p \wedge (q \wedge r) I p \wedge q \wedge r$ ' y similarmente para la disyunción; la conmutatividad o simetría de la conyunción y la de la disyunción, etc. También valen la conmutatividad y la asociatividad para la conyunción fuerte, ' $\bullet$ ': son en efecto teorematícos los esquemas ' $p \bullet (q \bullet r) I p \bullet q \bullet r$ ' y ' $p \bullet q I q \bullet p$ '. Además, la superconyunción es distributiva sobre la conyunción simple y también sobre la disyunción: ' $p \wedge q \bullet r I p \bullet r \wedge q \bullet r$ '

y $\lceil p \vee q \cdot r \rceil$. (Tales son algunos de los resultados que se obtienen gracias al esquema axiomático A14.) Mas no vale la idempotencia para $\cdot$: digamos que $\lceil p \cdot p \rceil$ (o sea $\lceil Xp \rceil$, leído como «Es muy cierto que p») puede ser más falso que $\lceil p \rceil$; y ello por lo ya dicho más arriba. (Según el modelo diseñado al final del Capítulo 6, cuando $\lceil p \rceil$ no sea, en absoluto, ni infinitamente verdadero ni infinitamente falso $\lceil Xp \rceil$ será **el doble** de falso que $\lceil p \rceil$.) Para muchas fórmulas, $\lceil s \rceil$, en vez de $\lceil p \rceil$, cabe demostrar en A_j $\lceil \neg(s \cdot s) \rceil$ (p.ej. para $\lceil s \rceil = \frac{1}{2}, X\frac{1}{2}, XX\frac{1}{2}, \dots, K\frac{1}{2}, KK\frac{1}{2}, \dots$).

La diferencia entre ' $\lceil$ ' y ' $=$ ' estriba en que $\lceil plq \rceil$ puede ser verdadero en unos aspectos y totalmente falso en otros, mientras que $\lceil p=q \rceil$ es verdadero en algún aspecto sólo si $\lceil plq \rceil$ es verdad en todos los aspectos; $\lceil p=q \rceil$ es o verdadero en todos los aspectos o del todo falso en todos los aspectos. Sin embargo, en virtud de rinf2 se deriva esta regla $plq \vdash p=q$; mas no es teorematizado el esquema $\lceil plq \supset p=q \rceil$.

Diversos grados de certeza de los axiomas

Vimos más atrás que en nuestro idioma «Es \$ cierto que p» (donde '\$' hace las veces de un functor monádico cualquiera) quiere de hecho decir «Es \$ verdadero que p»; 'cierto' es en tales contextos un alomorfo en distribución parcialmente complementaria de 'verdadero' —pues en tales contextos 'verdadero' resulta forzado. Mas eso no significa que, porque reconozcamos grados de verdad, vayamos a perder grados de certeza, o de seguridad, o de plausibilidad. Una cosa es cuán seguro sea que un enunciado es verdadero y otra es cuán verdadero sea (él o lo por él expresado, o el hecho por él significado —no cabe, naturalmente, en este trabajo una discusión filosófica sobre la verdad, sobre si, o en qué sentidos, se atribuye verdad a enunciados, pensamientos, “proposiciones”, hechos o lo que sea; en cualquier caso, las expresiones aquí utilizadas al respecto —y **aquí** es: a lo largo de todo este opúsculo— han de tomarse, según las preferencias filosóficas de cada uno, como meramente propedéuticas y, en caso necesario, susceptibles o menesterosas de paráfrasis). Claro que los grados de seguridad que se tengan son sendos grados de verdad del hecho de que uno está seguro de [la verdad o existencia de] aquello de lo que se trate. Pero naturalmente una cosa es el grado de verdad de «Andrés está seguro de que p» y otra el grado de verdad de «p».

Ahora bien, cabe hablar no sólo de grados de seguridad (o sea: de certeza) de algo para tal o cual persona —o grupo de personas— en particular, sino también a secas, o en general: cuán seguro (=cierto = plausible) sea tal o cual aserto. Sin embargo, también en eso hay alguna relativización, al menos implícita o contextual. Lo plausible o seguro de un aserto varía en efecto según quién lo diga, a quién se dirija y en qué circunstancias profiera el aserto.

Igual sucede con esa palabreja casi mágica (tal como se suele usar) que es la de lo 'intuitivo'. Algunos matemáticos y hasta desgraciadamente algunos filósofos profieren asertos que asignan “intuitividad” a ciertas tesis como si así se las erigiera en incontrovertibles, obvias, evidentes de suyo. Otros (como los fenomenólogos) tratan de legitimar ese empleo un tanto espúreo con una teoría cognoscitiva de la “intuición”, concebida como captación puramente inmediata de ciertas entidades o ciertos contenidos o ciertas verdades o lo que sea. Todo eso es de lo más discutible y al autor de este estudio le parece, más que dudoso, rotundamente equivocado. Vale más, si acaso, usar la palabra 'intuitivo' como significando 'muy plausible', 'verosímil', 'avalado por indicios dignos de consideración' o cosa así; o, acaso más modestamente: 'conforme con las opiniones que uno tiene o ya tenía al abordar la sistematización teórica aquí expuesta'.

Pues bien, hechas estas aclaraciones, cabe reconocer que no todos los esquemas axiomáticos de A_j poseen el mismo grado de certeza o plausibilidad (de “intuitividad”). Podemos agruparlos en dos conjuntos (pero **difusos**): 1) el de los esquemas poco disputables, aquellos que serán aceptados sin discusión por la gran mayoría de los lógicos, aquellos que menos

lugar dan a dudas suscitadas por reflexiones filosóficas u otras; 2) el de los esquemas “extraños”, “sorprendentes”, aquellos que a los lógicos clásicos —y a algunos otros también— les “sonarán a chino”, les parecerían como frases sacadas de un grimorio. Y, naturalmente, siendo difusos esos conjuntos o cúmulos, hay grados de pertenencia a cada uno de ellos y axiomas que pertenecen a ambos en alguna medida.

Cabría delimitar un cierto “núcleo duro” de axiomas que pertenecen al primer conjunto en medida “suficiente” (p.ej. aquellos que vienen abarcados por ese conjunto en una medida de al menos 50%). En esa situación están esquemas como: A01; A02 (si bien el principio de distributividad que es una de las consecuencias más directas de A02 [a saber: $\lceil r \vee s \wedge q \lceil q \wedge s \vee q \wedge r \rceil$] ha sido rechazado en algunas “lógicas cuánticas” —un problema en el que no cabe entrar aquí); A03; A06; A11, A13. Quizá también A12 (que, una vez despejada la definición, es: $\lceil p \mid p \mid a \rceil$): el clasicista rechazará acaso la introducción de esa constante ‘a’, pero admitirá que, de haber alguna constante sentencial, ‘ ψ ’, sea la que fuere, será verdad eso: $\lceil p \mid p \mid \psi \mid \psi \rceil$; porque él admitirá el esquema $\lceil p \mid p \mid q \mid q \rceil$ —sólo que obstinándose en ver a ‘l’ como una manera aberrante de escribir ‘ $\equiv$ ’, ya que él no admitirá ese distingo entre los dos funtores bicondicionales.

Una posición intermedia viene ocupada por esquemas que involucran a la superconyunción ‘ $\bullet$ ’ —algo de por sí ya chocante para muchos: la mera existencia de esa conyunción primitiva irreducible a la mera conyunción natural ‘ $\wedge$ ’— pero que no le atribuyen más que rasgos que en cualquier caso son poseídos también por la conyunción natural: esquemas como A05, A09, A10, A14. (El caso de A04 es algo más complicado porque al clasicista esos signos [‘K’ y ‘X’] le parecerán inaceptables. Sin embargo, si él se empeña en ver en la superconyunción ‘ $\bullet$ ’ un mero alógrafo de ‘ $\wedge$ ’ y también en la negación simple o natural ‘N’ un alógrafo de ‘ $\neg$ ’, así como en ‘l’ un alógrafo ‘ $\equiv$ ’, entonces él aceptará ese esquema A04; sólo que en tal caso para él decir $\lceil Kp \rceil$ [«Es [al menos] un poco cierto que p»], $\lceil Xp \rceil$ [«Es muy cierto que p»] y simplemente $\lceil p \rceil$ será, en los tres casos, decir lo mismo, con sendas variantes estilísticas contextualmente condicionadas.) También andan por ahí —por el medio, más o menos— esquemas que involucran a los funtores monádicos ‘H’ (‘completamente’, ‘del todo’), ‘L’ (‘hasta cierto punto [por lo menos]’), ‘B’ (‘es afirmable con verdad que’, ‘en todos los aspectos’): esquemas como A07, A08, A18, A19, A20: alguien de mentalidad más o menos clasicista estará inclinado a ver en tales funtores, o en sus lecturas en lengua natural, meros alógrafos, variantes estilísticas de ‘se da el caso de que’ o ‘Es verdad que’ (él dirá que en cada caso viene usada una expresión u otra por variación estilística contextualmente condicionada); mas, si así fuera, todos esos esquemas serían verdaderos (aunque también serían verdaderos otros esquemas que no son teorematizados en A_j , como $\lceil p \supset H p \rceil$, p.ej.); un clasicista a lo sumo admitirá a ‘B’ como un operador modal, o “intensional”, que signifique ‘necesariamente’ o ‘siempre’ o algo así; pero entonces rechazará (salvo restringidamente) la regla rinf_2 , alegando que de ‘Llueve’ no cabe concluir ‘Siempre llueve’, ni nada similar. (Pero un adepto de A_j podría replicar que no es en tal contexto [afirmable con] verdad que llueve sino que lo afirmable con verdad es algo así como esto: ‘Ahora aquí llueve’; sólo que por elipsis se omiten, en el contexto, esos deícticos, y quizá otros como ‘En este mundo de la experiencia cotidiana’, que han de catalizarse —en el sentido de ‘catálisis’ que usan los lingüistas, o sea: la operación inversa a la elipsis— para proferir una oración correctamente formulada.)

Llegamos así al conjunto de axiomas que —en el presente contexto— cabe reputar poco seguros, en la medida en que no sólo los mirarán con ceño los adeptos de la lógica clásica, sino que muchos cultivadores de lógicas no clásicas también sentirán recelo con respecto a ellos. Son sobre todo los esquemas que involucran ocurrencias de cuantos funtores se definen mediante la constante sentencial ‘a’: funtores como ‘Y’, ‘ $\hat{\imath}$ ’, ‘f’, ‘m’, ‘n’, ‘f’. Por ende, esquemas

como A23, A22, A21, A17, A16, A15; aunque los más aceptarán muchos esquemas teorematócos que se deducen fácil y rápidamente de uno o varios de esos esquemas axiomáticos —con ayuda de otros esquemas axiomáticos, claro. P.ej., una consecuencia pronto alcanzada a partir de A17 es el esquema teorematóco $\lceil p \rightarrow q \vee q \rightarrow p \rceil$, que el clasicista aceptará gustoso. (No así muchos no clasicistas, como los intuicionistas y los relevantistas; para los últimos en particular ese esquema es digno de rechazo.) Igualmente el esquema teorematóco —también fácilmente probado a partir de A17— $\lceil \neg p \supset p \rightarrow q \rceil$, que (aun rechazado por los relevantistas y otros) será aceptado por la mayoría —el clasicista se limitará a protestar por esa dualidad de funtores condicionales, ' $\supset$ ' versus ' $\rightarrow$ ', alegando que debe verse al segundo como un modo aberrante o pintoresco de reescribir al primero, o viceversa, pero que en cualquier caso son un solo functor [en dos variantes alográficas].

No cabe duda de que la introducción de esos matices aléticos como 'Y' ('Es un sí es no cierto que' o 'Es infinitesimalmente verdad que') no sólo acarrea complicaciones (de ahí la existencia de varios de los más largos esquemas de A_j , como A21 y A23) sino que, además, enriquece el vocabulario de la lógica con expresiones que —hasta la aparición de los sistemas de esta familia A — nadie antes había soñado en incluir en este ámbito (ni en ningún otro, salvo acaso en la lexicografía): de expresiones como 'un sí es no' ('Y'), 'un tanto' ('f'), 'viene a ser verdad que' ('m') y otros así ¿tiene que ocuparse el lógico? ¿Qué gana éste habiéndoselas con tales expresiones?

Gana bastante. P.ej., gana poder construir un cálculo cuantificacional —expuesto escuetísimamente en el capítulo siguiente— sorteando escollos indecibles que, en lógicas infinivalentes sin una constante como 'a', convierten tal construcción en un quebradero de cabeza. Gana el tener gracias a tales matices instrumentos más buidos para construir teorías axiomáticas de conjuntos. Gana el poder ofrecer un tratamiento riguroso a esas expresiones de matiz alético. Y gana también (¿por qué desdeñarla?) la elaboración de una teoría más bonita; no con la belleza de los desiertos, pero tampoco con la de las junglas, sino más bien con la de jardines frondosos pero con sus estructuras geométricas, sus simetrías, sus puntos de transición. Que eso son, para cualquier ' p ', ' np ' y ' mp ': puntos de transición, respectivamente hacia lo menos verdadero que p y hacia lo más verdadero que p ; sólo que en ciertos casos uno u otro —o ambos— coinciden con el propio p ; a saber: cuando p sea totalmente real o verdadero, entonces $p=mp$; cuando sea totalmente falso, $p=np$; cuando sea infinita mas no totalmente verdadero [o sea: cuando esté infinitamente próximo a ser completamente verdadero, mas sin alcanzarlo], entonces $np=p=mp$; lo mismo sucede cuando sea infinita mas no totalmente falso; cuando para cierto q se tenga que $nq=p$, entonces $np=p$; similarmente $mq=mmq$: el punto de transición hacia arriba de un punto de transición hacia arriba es él mismo, y similarmente con el punto de transición hacia abajo; en los demás casos $np \neq p \neq mp$. Comprendido eso, se entenderá perfectamente a qué viene cada uno de los esquemas axiomáticos aludidos, los cuales aseguran esas características precisamente.

Así pues —de aceptarse una concepción de lo verdadero articulable así— cada grado de verdad g será tal que haya al menos un grado de verdad g' que se toque con g —o sea: contiguo a g , tal, pues, que ningún otro se interponga entre ellos en la escala u orden que va de menos verdadero a más verdadero. Sin que eso sea óbice para la continuidad —en un sentido lato—: para cualesquiera dos grados de verdad, g , g' , tales que $g < ng'$, hay una infinidad de grados intermedios. (Como eso es así únicamente para grados de verdad, no se aplica a 0.) De ahí que se tenga en A_j el esquema teorematóco: $\lceil p \wedge nq \wedge p \supset p \wedge K(p \cdot q) \wedge K(p \cdot q) \wedge q \rceil$ —que es una consecuencia de A21—: si el hecho de que p es un tanto menos existente que el de que q , entonces el que sea al menos un poco existente el hecho de que no sólo p sino también q será algo cuya verdad estará entre la verdad o existencia de p y la de q ; y a su

vez habrá otro intermedio entre ello y cualquiera de los dos previamente dados; y así al infinito. Esta estructura nos da, pues, a la vez profusión infinita, riqueza ontológica, y sin embargo una agradable ordenación de los grados de verdad, con asideros o puntos de contacto que evitan el infinito aislamiento a que están, en cambio, sujetos los puntos o los números reales en un intervalo tal como suele concebirse, estando cada uno separado de cualquier otro por una infinidad de puntos intermedios. (Ciertamente que no he demostrado que una estructura como la aquí propuesta —que en términos algebraicos se denomina **atómica** [vide el capítulo 12]— sea más bella que esas otras estructuras, de Cantor o de Dedekind. ¿Será una opinión o una “intuición”?)

El principio de Heráclito

Podemos llamar ‘principio de Heráclito’ a este esquema teorematizado de A_j : $\lceil N(plp) \rceil$. Es un principio infrecuentísimo en los sistemas de lógica —aunque no son los sistemas de la familia A los únicos que lo han postulado (o que han postulado algo que puede entenderse como “traducible” de ese modo). El llamarlo como lo acabo de proponer se debe a aquella frase de Heráclito de que nadie se baña dos veces en el mismo río: por el flujo en que se halla cada ente, por la impermanencia de tantas determinaciones suyas, no es el mismo en dos momentos; pero, entonces, no es el mismo ni aun en un lapso, ya que durante tal lapso también sufre algún cambio.

Sin detenernos aquí a discutir esas ideas —que no son tan menospreciables como lo han supuesto muchos adeptos de la lógica clásica ni tan menesterosas de interpretación caritativa como lo piensan otros de tales adeptos—, centrémonos en este esquema, $\lceil N(plp) \rceil$, e.d. la tesis de que nada sucede en la medida en que sucede, o sea: que no es verdad que p suceda sólo en toda la medida en que p (dicho de otro modo: que no es verdad que p sea verdadero al menos y a lo sumo en la medida en que p), un esquema cuya afinidad con la tesis heraclítica de la autodistinción de cada ente nos autoriza a darle la denominación que le damos.

En A_j se demuestra este principio gracias a la definición de $\lceil 0 \rceil$, a saber $\lceil \frac{1}{2}la \vee \neg(\frac{1}{2}lN\frac{1}{2}) \rceil$ dada la definición de $\frac{1}{2}$ como $\lceil ala \rceil$ y dado el esquema axiomático A12 a tenor del cual se prueba para cualesquiera fórmulas $\lceil p \rceil$, $\lceil q \rceil$ que $\lceil plpl.qlq \rceil$ (una tesis harto plausible: la autoequivalencia es autoequivalente, por decirlo así, o sea: cualquier autoequivalencia equivale a cualquier otra: no puede ser ni más ni menos verdad el que plp que el que qlq , sino que el que un hecho sea tan verdadero como sí mismo es igual de verdadero que el que lo sea otro cualquiera). Supongamos que, a fin de evitar la demostración del principio de Heráclito, modificamos la definición de $\lceil 0 \rceil$ extirpando de ella el disyunto derecho y dejando no más $\lceil \frac{1}{2}la \rceil$. (Como en A_j $\lceil \neg 0 \rceil$ es teorematizado, si $\lceil 0 \rceil$ es una disyunción, $\lceil r \vee s \rceil$, se habrá probado $\lceil \neg(r \vee s) \rceil$ y de ahí rápidamente se concluirá $\lceil \neg r \wedge \neg s \rceil$ y, por ende, tanto $\lceil \neg r \rceil$ como $\lceil \neg s \rceil$; si $\lceil s \rceil$ es una fórmula supernegativa —como sucede en este caso, pues es $\lceil \neg(\frac{1}{2}lN\frac{1}{2}) \rceil$, entonces (por la **regla de apencamiento** fácilmente derivable en A_j , a saber la regla $\lceil \neg\neg p \rceil \vdash p$: lo que no es del todo falso es verdadero) hemos probado aquella fórmula (en este caso $\lceil \frac{1}{2}lN\frac{1}{2} \rceil$) cuya supernegación era el $\lceil s \rceil$ en cuestión; y de ahí se demuestra $\lceil N\frac{1}{2} \rceil$ a partir del esquema teorematizado $\lceil plp \rceil$ y de una regla de inferencia derivada.

Pues bien, una vez extirpado de la definición de $\lceil 0 \rceil$ todo el segmento $\lceil \vee \neg(\frac{1}{2}lN\frac{1}{2}) \rceil$, vamos a añadir un axioma más: definimos un functor ‘ P ’ así (proponiendo para $\lceil Pp \rceil$ la lectura: «Es más bien cierto que p »; la ‘ P ’ evoca el latín ‘potius’; en francés se diría ‘plutôt’ y en inglés ‘rather’): $\lceil Pp \rceil$ abr. $\lceil Np \rightarrow p \& p \rceil$. En virtud de tal definición, $\lceil Pp \rceil$ será: totalmente falso cuando y donde [y en los aspectos en que] $\lceil p \rceil$ sea más falso que verdadero; y exactamente igual de verdadero que $\lceil p \rceil$ en caso contrario. El esquema axiomático que añadimos entonces es éste $\lceil P(plp) \rceil$. Nos dice que cada hecho es más bien equivalente a sí mismo; que el que algo

—sea lo que fuere— suceda en la medida en que sucede es por lo menos tan verdadero como falso.

Supuesto eso, introduzcamos ahora una constante nueva, ' $\heartsuit$ ', que denotará un hecho cuyo grado de realidad equidiste entre la verdad completa y la total falsedad. Hay pocos hechos así, mas eso no nos importa. Podemos suponer que sea el salir victorioso Pirro en la batalla de Ásculo o la victoria de las tropas federales estadounidenses en Chickamauga durante la Guerra de Secesión; o bien el salir vencedor Karpov en una partida con Kasparov en la que se declararon tablas (suponiendo —como parece verosímil— que las tablas, o el empate, es, no un *tertium quid* que excluya por completo la victoria y la derrota de los contendientes, sino una situación en la que cada uno de éstos sale igual de vencedor —así como también igual de derrotado— que el otro). Añadimos el axioma ' $\heartsuit IN \heartsuit$ '. Entonces uno cualquiera de los dos siguientes esquemas, si se añade a la base axiomática así formada, restablece *Aj* en toda su integridad —o sea: permite probar el esquema ' $\frac{1}{2} IN \frac{1}{2}$ ' y, a partir de él, el principio de Heráclito—: (1º) el principio de Clavius o de abducción, a saber: ' $p \rightarrow Np \rightarrow Np$ ' (o alguna variante del mismo como: ' $Np \rightarrow p \rightarrow p$ ', ' $p \rightarrow N(p \rightarrow Np)$ ', ' $Np \rightarrow N(Np \rightarrow p)$ '); (2º) el principio implicativo de no contradicción, PINC (a saber: que lo verdadero que implica algo no implica la negación de ese algo): ' $p \rightarrow q \wedge p \rightarrow N(p \rightarrow Nq)$ ' (o algún esquema parecido a ése de entre los siguientes: ' $p \wedge (p \rightarrow q) \rightarrow N(p \rightarrow Nq)$ ', ' $p \supset p \rightarrow q \rightarrow N(p \rightarrow Nq)$ ', ' $p \rightarrow Nq \rightarrow p \supset N(p \rightarrow q)$ ', ' $q \rightarrow p \rightarrow N(Nq \rightarrow p) \vee Hp$ ').

Cada uno de esos dos esquemas —o de las variantes respectivas— es plausible de lo más. Por supuesto son teorematícos en la lógica clásica los resultados de reemplazar en ellos ' $\rightarrow$ ' por ' $\supset$ ' y ' N ' por ' $\neg$ '. Por ende, de todo lo que se ha postulado lo único que parece muy discutible desde un punto de vista clásico o similar es esa constante ' $\heartsuit$ ' con el axioma ' $\heartsuit IN \heartsuit$ '.

Un clasicista rechazará que haya sucesos que sean tan verdaderos como falsos, o en general que los haya verdaderos y falsos. Algunos no-clasicistas, aun admitiendo hechos verdaderos y falsos, rechazarán que pueda un hecho tener ambas determinaciones en la misma medida (para esos autores cada hecho o es más verdadero que falso o es más falso que verdadero; pero eso introduce una discontinuidad en la escala veritativa que acarrea consecuencias casi catastróficas). Ahora bien, notemos que para probar, no ' $\frac{1}{2} IN \frac{1}{2}$ ', pero sí el principio de Heráclito, bastaría (dados los otros esquemas que habíamos añadido al resultado de empobrecer *Aj* desmochando la definición de ' 0 ' de la manera indicada, e.d. truncándole el trozo final ' $\vee \neg(\frac{1}{2} IN \frac{1}{2})$ ') postular ' $S \heartsuit$ ', o sea que el hecho significado por ' $\heartsuit$ ' fuera [hasta cierto punto] falso pero verdadero también [en alguna medida]; porque entonces demostraríamos en el sistema así resultante: ' $\heartsuit \rightarrow N \heartsuit I(plp) \vee (N \heartsuit \rightarrow \heartsuit I.plp) \wedge N(\heartsuit \rightarrow N \heartsuit \vee N \heartsuit \rightarrow \heartsuit)$ ', de donde se infiere fácilmente ' $N(plp)$ '. (Y obsérvese que eso se demostraría aduciendo sólo esquemas de los que forman el "núcleo duro" a que aludíamos en el acápite precedente, o sea: sin venir esencialmente involucrados los funtores que, como ' f ' etc., son más discutibles o más susceptibles de ser rechazados por personas de inclinación conservadora.)

De ahí que el principio de Heráclito resulte de la hipótesis de que cierto hecho existe o es verdadero pero no totalmente, dados otros principios que no sólo son **poco** discutibles sino que en verdad casi nadie ha discutido: cierto que los relevantistas rechazan el principio de abducción, ' $p \rightarrow Np \rightarrow Np$ ', pero por consideraciones que son argumentos algo tortuosos, no constituyendo una objeción clara contra ese mismo principio; y en algunas lógicas como las de Łukasiewicz [consideradas más atrás en este trabajo] falla también ese principio, aunque los elaboradores de tales lógicas nunca han proporcionado una argumentación a favor de esa falla: parece que ésta es una lamentable consecuencia indirecta e indeseada en tales sistemas que resulta de, admitiendo la existencia de valores veritativos intermedios, reputar designado, sin embargo, sólo al valor supremo, 1, y de la asignación de valores a ' $\rightarrow$ ' (o a ' $\supset$ ', pues en

tal sistema hay un solo condicional) que viene impuesta por ese monopolio de que está en tales sistemas investido el valor supremo.

Ahora bien —se alegrará—, aunque así sea, ¿no es obvio que esa constante no puede ser una constante lógica? No, no es obvio. Porque en cualquier teoría $\mathfrak{B}$ que sea una extensión del resultado de modificar (empobrecer) Aj de la manera indicada en la cual sea teorematizado 'S♥' se demostrará el esquema $\lceil N(plp) \rceil$ pese a que, si este esquema es verdadero (más correctamente expresado: si es verdadera cada instancia del mismo), entonces es una verdad **de lógica** puesto que en cada una de tales fórmulas las únicas expresiones con ocurrencias esenciales serían 'N' y 'I', que forman parte del vocabulario de la lógica. O sea, esa teoría $\mathfrak{B}$ no sería una extensión conservativa del sistema de lógica en cuestión; habría verdades lógicas aseverables sólo en disciplinas exteriores a la lógica. Si aceptamos, en cambio, la tesis de que la lógica es el estudio de las verdades lógicas, entonces, si es verdad en general que hay algún hecho que sea verdadero y falso, como de eso se deduce que, para cada fórmula $\lceil p \rceil$, $\lceil plp \rceil$ es precisamente **un** hecho verdadero y falso, esta conclusión —siendo (dado ese supuesto) verdadera— ha de ser, en todo caso, un teorema del sistema lógico escogido.

Es más: si hay alguna teoría $\mathfrak{B}$ verdadera en la que para cierto $\lceil q \rceil$ es verdad $\lceil qINq \rceil$, entonces en esa teoría $\mathfrak{B}$ es afirmable con verdad el esquema $\lceil plpIN(plp) \rceil$, que, siendo verdadero, será una **verdad lógica**; si eso es así, ha de ser una verdad **de** [la teoría] **lógica**, o sea: un teorema del sistema lógico escogido; que es lo que sucede en Aj . Pero entonces resulta que $\lceil \heartsuit \rceil = \lceil \frac{1}{2} \rceil = \lceil ala \rceil$.

(A la objeción de que ha de haber algún paralogismo en el razonamiento precedente porque las verdades lógicas son necesarias y en cambio sería contingente la existencia de un hecho o suceso verdadero y falso —o, alternatively, **tan** verdadero **como** falso—, cabe responder que lo necesario se sigue de lo contingente; si un hecho contingente (¡supongámoslo o aceptémoslo!) entraña cierta conclusión, y si ésta es tal que, si es verdadera, lo es necesariamente, entonces si el primer hecho es verdadero (aunque sea contingentemente), la conclusión es necesariamente verdadera. Aparte de eso, en el marco de las consideraciones del presente trabajo no habemos menester de ocuparnos más de esas nociones de necesidad y contingencia, cuya discusión queda para otros trabajos de filosofía de la lógica y de lógica modal.)

Para cerrar ya este acápite —y, con él, este capítulo— cabe mencionar que un corolario del principio de Heráclito es el principio llamado de Boecio, a saber $\lceil p \rightarrow q \rightarrow N(p \rightarrow Nq) \rceil$. De él a su vez se desprende el de Aristóteles: $\lceil N(p \rightarrow Np) \rceil$. Ambos son, pues, esquemas teorematizados de Aj . No nos interesa aquí la base histórica para atribuir tales principios respectivamente a Boecio y a Aristóteles. En todo caso, existe una corriente de lógica no estudiada en el presente trabajo —la lógica **conexivista**— que entiende la implicación ' $\rightarrow$ ' precisamente de tal manera que resulten válidos el esquema de Boecio y, por lo tanto, también el de Aristóteles. (Lo que separa a esa corriente del espíritu que anima a una lógica como Aj es, entre otras cosas, que los sistemas conexivistas no son paraconsistentes, y, por ende, no pueden admitir para **ningún** p a la vez $\lceil p \rceil$ y $\lceil Np \rceil$; y que sacrifican el principio implicativo de simplificación, $\lceil p \wedge q \rightarrow p \rceil$, y el de adición, $\lceil p \rightarrow p \vee q \rceil$.) Nótese, por último, que el principio de Aristóteles es una versión reforzada del PINC. Huelga decir que, en la demostración del principio de Heráclito más arriba aludida, el papel que en ella se adjudicaba a PINC hubiera podido desempeñarlo el principio de Boecio.

Capítulo 10

El cálculo cuantificacional de primer orden

Consideraciones preliminares

El vocabulario de la lógica abarca a una palabra que hasta ahora no ha entrado en nuestro estudio: el adjetivo indefinido ‘cada’ o su equivalente ‘todo’ (‘todo’ en su acepción distributiva: ‘Todo estudiante sabe que sus profesores enseñan bien o mal’ etc.). Con ayuda de esa palabra podremos definir otro miembro más del vocabulario lógico, a saber: ‘algún’ (o ‘hay [al menos] un... tal que...’; o ‘se da algún [o al menos un]... tal que...’).

A diferencia de las expresiones hasta ahora consideradas, estas a las que se hace referencia en el párrafo precedente no afectan a las oraciones tomadas como bloques sino que se meten —como cuñas, por decirlo así— dentro de las oraciones. De ahí que lo hasta ahora estudiado sea el cálculo sentencial (del latín ‘sententia’, **oración**) o cálculo de enunciados, al paso que lo que en este capítulo vamos a presentar —muy escuetamente— es el cálculo **cuantificacional**. Así se llama porque esos indefinidos, el ‘cada’ (o ‘todo’) y el ‘algún’ (o ‘hay [al menos] un... tal que...’) son, respectivamente, **cuantificadores**: el primero es el cuantificador **universal**; el segundo el **existencial** (también llamado ‘particular’). Motejar al segundo de ‘existencial’ se debe a que una de las lecturas del mismo, el ‘hay ...’, es reemplazable por una, por lo demás igual, donde el ‘hay’ viene reemplazado por un ‘existe’. (No quiere ello decir que sea ésa la única acepción de ‘existe’. Seguramente ‘Existe Ten Xiao-ping’ o ‘No existe Alí Babá’ no pueden parafrasearse reemplazando ‘existe’ por ‘hay’; ello es indicio —no prueba— de que se da ese otro ‘existe’, uno que podríamos llamar **predicativo** o **determinativo**, mientras que el ‘existe’ que aquí nos interesa es indeterminativo: «existe un [o algún] ...».)

Un cuantificador es algo que cuantifica, o sea es una expresión tal que, en virtud de que la misma figure en una oración —en determinado lugar— lo dicho en la oración se aplica a **tantos** o a **cuantos** entes: a **todos**, si el cuantificador es universal; a sólo algun(os), si es existencial o particular. Se dan también otros cuantificadores, pero en ellos no entraremos aquí. Cuantificadores no estándar como ‘muchos’, ‘varios’, ‘la mayoría de los...’, ‘pocos’, etc.; además de que cabe concebir a los numerales como sendos cuantificadores: ‘Dos ámbitos de estudio se dan en la ciencia’, p.ej. (o ‘tres’, etc.). Ahora bien, los más autores prefieren no tomar a los números como cuantificadores sino brindar para ellos un tratamiento diverso, dentro de una teoría axiomática de conjuntos; y ése es también el parecer del autor del presente trabajo. Y aun para los cuantificadores no estándar quizá lo mejor sea definirlos dentro de una teoría axiomática de conjuntos, o de algún tratamiento de igual potencia. Con lo cual nos quedaríamos tan sólo con dos cuantificadores, los estándar: el universal y el existencial. Interdefinibles, por cierto.

Surgen muchos problemas filosóficos en torno a los cuantificadores. No entraremos en ellos. Algunos lógicos no clásicos —los intuicionistas— rechazan la interdefinibilidad de los dos cuantificadores. Lléalos a ello una tendencia idealista a concebir la verdad como dependiendo no sólo de cómo sean las cosas en la realidad sino, tanto o más, de cómo pueda descubrirla el sujeto —un descubrimiento que, desde esa perspectiva, es una semicreación subjetiva. Más que verdad en sí, verdad para mí. De ahí que los intuicionistas rechacen, p.ej., que siempre que no sea verdad que todo ente [o sea —desde su prisma—, todo ente conocido o “construido”] sea así o asá haya de ser verdad que algún ente [descubierto o “construido”] no es así o asá; porque —alegan— puede que no sea verdad ‘Todo ente es tal que p’ porque se ha demostrado en efecto que ‘Todo ente es tal que p’ es absurdo; pero sin que forzosamente por ello sea construible o conocible un ente en particular tal que sea demostrable de ese ente que no-p. En este trabajo no entraremos en esa discusión. Pero queda avisado el lector de que lo que aquí se expondrá sobre la interdefinibilidad de los cuantificadores —y muchas otras

cosas—, siendo como es una doctrina que cuenta con la aceptación de la abrumadora mayoría de los lógicos, no es empero patrimonio común o un acervo de verdades unánimemente reconocidas.

Otro problema se refiere a si el cuantificador particular ‘algún’ (o ‘al menos un’) es de veras equivalente a una partícula existencial como ‘hay’ o ‘existe’. Se dan opiniones para todos los gustos. Los más lo identificamos; pero unos cuantos lo rechazan. Alguno que otro acepta la equivalencia entre ‘algún’ y ‘Hay algún... que’, pero rechaza la equivalencia con el ‘existe’. De entre quienes aceptamos la equivalencia, los más aceptamos reglas como la que de «Tal cosa es así o asá» —donde ‘tal cosa’ hace las veces de un nombre propio— permite concluir «Algo es así o asá»; otros —los adeptos de “lógicas libres”— ponen restricciones a esa regla de inferencia. Todos esos debates tienen motivaciones filosóficas que no son de menospreciar. Pero no podemos entrar aquí en ninguna de esas controversias. Nos atendremos a la opinión mayoritaria que, en todo esto, es también la del autor del presente opúsculo.

Sea o no conveniente ese título de ‘cuantificadores’ que se da a las expresiones con que vamos a trabajar (a los indefinidos ‘cada’ y ‘algún’), el hecho es que es un uso estándar y a él nos atendremos.

Introducción de nueva terminología

Seguiremos utilizando letras esquemáticas. Sólo que una letra esquemática sentencial (como ‘p’, ‘q’, etc.) ahora hará las veces no sólo de oraciones propiamente dichas sino de **fórmulas** cualesquiera. Llamaremos ‘fórmula’ a un signo que sea por lo demás igual que una oración pero que puede que contenga algún pronombre terciopersonal (un ‘él’, ‘ella’ o ‘ello’) que no se remita a un cuantificador previo (o —en la jerga que en seguida se introducirá— que no esté **ligado** por un cuantificador previo). Así ‘El está enfermo’ o ‘Está enfermo’ (pues es nuestro idioma el pronombre puede elidirse) es una fórmula, pero no una oración. Toda oración es una fórmula, pero no viceversa. Una fórmula que no sea oración viene normalmente desambiguada en su sentido por el contexto (—‘¿Has hablado con Martín?’ —‘No’ —‘¿Por qué?’ —‘Está enfermo’.)

Lo interesante aquí es que puede haber varios pronombres terciopersonales en una fórmula. ‘El lo vio’, o sea ‘El vio a él’: en un contexto y entorno de elocución dados el primer ‘él’ puede referirse a Hilario, el segundo a Maximino, p.ej. En idiomas como el nuestro vienen utilizados varios procedimientos para dar pautas o claves de cómo cabe entender tales fórmulas: la variación de género (‘Se paseaban Luisa y Álvaro; ella lo miró con insistencia’). También vienen a veces reemplazados los ‘él’ por alomorfos en distribución parcialmente complementaria, como ‘ése’, ‘éste’, ‘áquel’, ‘el primero’, ‘el segundo’ etc. (‘Higinio era más trabajador que Lucio. El primero no se acostaba nunca antes de terminar sus ejercicios, mientras que el segundo dormía la siesta cada día’.)

A los pronombres terciopersonales —o a lo que los representa en las escrituras lógico-matemáticas— llámaselos en lógica **variables**. Una ocurrencia de una variable está **ligada** a un cuantificador cuando figura en la fórmula en que esté dicha ocurrencia una marca en virtud de la cual esa ocurrencia de la variable en cuestión se refiere, no a un ente en particular, sino a **cualquier** ente —si el cuantificador en cuestión es universal— o —si el cuantificador es existencial— a **algún** ente —pero a ninguno en particular o sea a **uno u otro** ente. Así, en ‘Todo ente es tal que, si [él] es mamífero, [él] carece de alas’ —oración falsa, desde luego— sendas ocurrencias de ‘él’ remiten al cuantificador universal, y por ende se refieren a todo ente, a cualquier ente sin restricción ninguna. En ‘Todo ente es tal que hay alguna cosa [tal que ella es] parecida a él’, tenemos: el ‘él’ remite al cuantificador universal inicial; el ‘ella’ al cuantificador existencial que sigue inmediatamente a ese cuantificador universal.

Ahora expresémonos con mayor rigor con aplicación a nuestra notación formal —en vez de a sus lecturas en la lengua natural. Introduciremos una infinidad de variables, a saber las letras 'x', 'y', 'z', 'u', 'v', más el resultado de añadir a cada una uno o más acentos o apóstrofes (pero en aras de la claridad dos acentos seguidos serán reemplazados por un '2' superescrito; tres, por un '3' superescrito, etc.).

Un cuantificador universal (respectivamente, existencial) es un signo que consiste en una ocurrencia de una variable precedida inmediatamente de ' $\forall$ ' (respectivamente, de ' $\exists$ '). Un cuantificador universal (respectivamente, existencial) como ' $\forall x$ ' (respectivamente ' $\exists x$ ') vendrá leído como 'Todo ente, x, es tal que' (respectivamente como 'Algún ente, x, es tal que') haciendo los cambios de variables que sean del caso. (Así ' $\exists z$ ' se lee 'Todo ente, z, es tal que', etc.)

A ' $\exists$ ', ' $\forall$ ', los llamaremos **prefijos cuantificacionales** —respectivamente **existencial** y **universal**. Diremos que la fórmula ' p ' que aparezca en la fórmula ' $\exists \alpha p$ ' o ' $\forall \alpha p$ ' (siendo ' α ' una variable cualquiera) es en esta última fórmula el **alcance** [de esa ocurrencia] del cuantificador dado, o sea de ' $\exists \alpha$ ' o de ' $\forall \alpha$ '. Cuando se tenga una fórmula ' p ' con ocurrencias de ' α ' y ' p ' venga precedida inmediatamente por ' $\exists \alpha$ ' o por ' $\forall \alpha$ ' pero dentro de ' p ' no aparezca ninguna ocurrencia ni de ' $\exists \alpha$ ' ni de ' $\forall \alpha$ ' diremos que cada ocurrencia de ' α ' en ' p ' está **ligada** a esa ocurrencia de ese cuantificador (o sea a o bien ' $\exists \alpha$ ' o bien ' $\forall \alpha$ '). Además en una ocurrencia de ' $\forall \alpha$ ' la única ocurrencia de ' α ' está ligada al propio cuantificador de que forma parte; y lo propio sucede con ' $\exists \alpha$ '.

Si una ocurrencia de una variable está ligada a un cuantificador en una fórmula, entonces **está ligada** [a secas]. Y, cuando una ocurrencia de una variable está ligada en una fórmula que es subfórmula de otra —o sea, que figura en ésta última—, entonces esa ocurrencia de esa variable está también ligada en la última. Si una variable tiene alguna ocurrencia ligada en una fórmula se dice que figura ligada en esa fórmula.

Una ocurrencia de una variable en una fórmula que no esté ligada en esa fórmula está **libre en** esa fórmula. Una ocurrencia de una variable en una fórmula que no esté ligada ni en esa fórmula ni en ninguna fórmula más amplia (o sea en ninguna de la cual forme parte la fórmula dada en cuestión) es una ocurrencia **libre** [a secas].

Una variable puede tener en una misma fórmula ocurrencias libres y ocurrencias ligadas. Así en ' x^2 pasa frío $\wedge \exists x^2$ (x^2 es una estufa)'¹ se dice que él o ello, x^2 , pasa frío (en un contexto dado puede esa ocurrencia de ' x^2 ' referirse p.ej. a Ezequiel) aunque hay estufas. Nótese que no podría decirse ' x^2 pasa frío $\wedge \exists x^2$ (x^2 es una estufa $\wedge x^2$ tiene x^2)' para decirse que ese ente, el que sea —Ezequiel p.ej.— pasa frío aunque [ese mismo ente] tiene alguna estufa; desde el momento en que hemos puesto el cuantificador existencial ' $\exists x^2$ ', al abrir un paréntesis inmediatamente detrás delimitamos el alcance de tal cuantificador existencial, alcance que no concluye sino con el paréntesis derecho que sea el compañero de ese paréntesis izquierdo; y cada ocurrencia de ' x^2 ' en ese alcance de ' $\exists x^2$ ' está ligada a ese cuantificador. Por lo cual esa fórmula diría antes bien esto: que Ezequiel (pongamos por caso que a él nos referimos con la primera ocurrencia —la única libre— de ' x^2 ') pasa frío pero que alguna estufa se tiene a sí misma.

Si una variable tiene una ocurrencia libre en una fórmula se dirá que figura (o aparece) libre en esa fórmula; y, si tiene en una fórmula una ocurrencia ligada, se dirá que aparece ligada en la fórmula. Puede una misma variable figurar a la vez libre y ligada en una misma fórmula, según lo hemos visto.

Variación alfabética

Llamamos **cuantificación** a una fórmula de la forma $\lceil \forall \alpha p \rceil$ donde $\lceil p \rceil$ es una fórmula y ' α ' es una variable. (En adelante prescindo de referirme a los cuantificadores existenciales —en estas definiciones—, ya que el cuantificador existencial será, en el sistema de cálculo cuantificacional aquí presentado, definido a partir del universal. Desde luego podría procederse inversamente, tomando al existencial como el único cuantificador primitivo.)

Llamamos **matriz** a una fórmula cualquiera; y, cuando está inmediatamente precedida por un cuantificador, decimos que es matriz de ese cuantificador (**matriz de** equivaldrá a **alcance de**).

Una variable puede estar en la matriz $\lceil p \rceil$ (o sea **en** el alcance) de un cuantificador sin estar ligada al mismo, porque exista otro cuantificador con la misma variable situado más a la derecha y que tenga como alcance suyo una fórmula en la cual figure esa ocurrencia de dicha variable.

Sea $\lceil q \rceil$ una fórmula en la que figura una cuantificación, siendo ' α ' una variable. Entonces, ssi $\lceil p \rceil$ difiere de $\lceil p \rceil$ por contener ocurrencias libres de una variable ' β ' (libres **en** $\lceil p \rceil$) sólo dondequiera que $\lceil p \rceil$ contiene ocurrencias libres de ' α ' (libres **en** $\lceil p \rceil$) diremos que, ssi $\lceil q \rceil$ difiere de $\lceil q \rceil$ por el uniforme reemplazo de $\lceil \forall \alpha p \rceil$ por $\lceil \forall \beta p \rceil$, $\lceil q \rceil$ es una **variante alfabética** [en sentido estrecho] de $\lceil q \rceil$. (Nótese que puede haber en $\lceil p \rceil$ otras ocurrencias de ' α ', que estén ligadas **en** $\lceil p \rceil$; éstas seguirán estando en $\lceil p \rceil$, y también ligadas **en** $\lceil p \rceil$, pues ellas no vienen afectadas por el tránsito de $\lceil p \rceil$ a $\lceil p \rceil$.) La relación entre dos variantes alfabéticas es la de **variación alfabética**. Nótese que es una relación simétrica.

Por una concatenación de variaciones alfabéticas pueden invertirse las variables de varios cuantificadores en una misma fórmula; pero no puede hacerse eso directamente. P.ej., $\lceil \forall x \forall z (x=z) \rceil$ tendrá como variante alfabética suya, según la definición dada, a esta cuantificación: $\lceil \forall u \forall z (u=z) \rceil$; en ésta está la cuantificación $\lceil \forall z (u=z) \rceil$ que tendrá como variante alfabética suya a $\lceil \forall x (u=x) \rceil$; luego $\lceil \forall u \forall x (u=x) \rceil$ será una variante alfabética de $\lceil \forall u \forall z (u=z) \rceil$. Ahora bien, $\lceil \forall u \forall x (u=x) \rceil$ tiene como variante alfabética suya propia a $\lceil \forall z \forall x (z=x) \rceil$. Pues bien, llamaremos **variante alfabética** [en sentido amplio] de una fórmula a otra tal que exista una cadena de variantes alfabéticas [en sentido estrecho] de la primera a la segunda (e.d. a otra que sea, o idéntica a la dada, o variante alfabética suya [en el sentido estrecho], o variante alfabética [en sentido estrecho] de una variante alfabética [en sentido estrecho] de la dada, o así sucesivamente).

Lo que hace que una cuantificación que tenga ciertas variables en determinados lugares sea una variante alfabética de otra (o sea —llanamente expresado— que diga lo mismo con otros signos, o más bien que sea un alomorfo en distribución libre de la cuantificación dada) es que las variables (que no son otra cosa que pronombres terciopersonales indizados, como si dijéramos ' él^1 ' o ' ella^1 ', ' él^2 ' o ' ella^2 ', ' él^3 ' o ' ella^3 ' etc.) no son signos que denoten a tal ente en particular —salvo las variables libres, y eso en tal contexto determinado y tal entorno de elocución preciso—, sino signos que se refieren indeterminadamente a cualquier cosa o alguna cosa tal que **esa cosa** sea así o asá; son **indefinidos**. (Decir ' Todo^1 es tal que ello^1 es tal que hay algo² tal que ello^2 es más grande que ello^1 ' es decir lo mismo que si se profiere el resultado de invertir uniformemente en tal mensaje los superíndices.)

Otra noción afín a la de variación alfabética es la de **versión alternativa**; sólo que ésta se aplica a variables libres. Decimos que ssi $\lceil p \rceil$ contiene ocurrencias libres de una variable ' α ' y $\lceil p \rceil$ únicamente difiere de $\lceil p \rceil$ por contener sendas ocurrencias también libres de otra variable ' β ' sólo dondequiera que $\lceil p \rceil$ contiene esas ocurrencias libres de ' α ', $\lceil p \rceil$ es una **versión alternativa** de $\lceil p \rceil$. (Nótese que, desde luego, aunque $\lceil p \rceil$ sea una versión alternativa de $\lceil p \rceil$,

no por ello otra fórmula $\lceil q \rceil$ que contenga a $\lceil p \rceil$ va forzosamente a ser una versión alternativa de $\lceil q \rceil$ aunque $\lceil q \rceil$ difiera de $\lceil q \rceil$ sólo por el uniforme reemplazo de $\lceil p \rceil$ por $\lceil p' \rceil$.)

De la pluralidad de cálculos sentenciales a la opción entre diversos cálculos cuantificacionales

Hemos visto en los capítulos precedentes del presente trabajo que hay una pluralidad —en verdad infinita— de cálculos sentenciales, todos ellos verosímiles, plausibles, apuntalados por consideraciones de mayor o menor peso, aunque en ciertos aspectos haya algunos cálculos que sean los más plausibles de todos —no siéndolo forzosamente en todos los aspectos.

Ahora bien, para cada cálculo sentencial hay varios cálculos cuantificacionales alternativos que son extensiones de ese cálculo sentencial. P.ej., hay dos modos de extender la lógica sentencial clásica en el terreno del cálculo cuantificacional: uno, el estándar; otro, el de las lógicas libres. El segundo difiere del primero en carecer de la regla de inferencia $p \vdash \exists x p$, o sea de la regla de generalización existencial ($\lceil p' \rceil$ ahí difiere de $\lceil p \rceil$ sólo por contener ocurrencias libres de 'x' únicamente en lugares en los que $\lceil p \rceil$ contenga sendas ocurrencias libres de cierta variable o constante individual). Pero también existen otras maneras de extender la lógica clásica en el terreno del cálculo cuantificacional. Una de ellas, p.ej., abandonaría la interdefinibilidad de $\exists$ y $\forall$ (abandonando el esquema, teorematizado en los cálculos cuantificacionales estándar: $\lceil \forall x p \equiv \neg \exists x \neg p \rceil$) y añadiendo acaso principios como el aristotélico de subalternación ($\lceil \forall x (p \supset q) \supset \exists x (p \wedge q) \rceil$) que no son teorematizados en ningún cálculo estándar. (Nótese que ese principio de subalternación, cuando se cercenan de él los dos cuantificadores —el universal de la prótasis y el existencial de la apódosis— se transforma en el esquema $\lceil p \supset q \supset p \wedge q \rceil$; mas como es un esquema tautológico en la lógica clásica (y en *Aj* también, claro —pues *Aj* subsume en sí toda la lógica clásica) este principio bicondicional, a saber $\lceil p \wedge q \equiv \neg (p \supset \neg q) \rceil$, resulta que el esquema dado está bicondicionalmente unido a éste otro: $\lceil p \supset q \supset \neg (p \supset \neg q) \rceil$, que es precisamente PA (Principio de Aristóteles). Ni en la lógica clásica ni en *Aj* es teorematizado ese esquema; pero en *Aj* sí es teorematizada una “versión” del mismo, aquella en la cual el condicional ' $\supset$ ' viene reemplazado por la implicación ' $\rightarrow$ ' al paso que la negación fuerte ' $\neg$ ' viene reemplazada por la negación natural 'N'.)

Ciñéndonos ya al sistema *Aj* presentado en el capítulo precedente, hay varios modos de, tomándolo como base, construir sendos cálculos cuantificacionales. P.ej., podemos, o no, hacer que sea teorematizado el esquema $\lceil \forall x p \rightarrow p \rceil$ (donde $\lceil p' \rceil$ difiere de $\lceil p \rceil$ sólo por contener ocurrencias libres de alguna variable dondequiera que $\lceil p \rceil$ las contiene de 'x'). Podríamos contentarnos con $\lceil \forall x p \supset p \rceil$. Pero entonces resultarían cosas raras. P.ej. podría ser enteramente cierto que todo ente sea así o asá sin que fuera enteramente cierto que este o aquel ente en particular sea así o asá; e.d. sería más verdadero que todos los entes, colectivamente tomados, sean así o asá que no que, distributivamente tomados, unos u otros entes sean así o asá. Por ello de $\lceil H \forall x p \rceil$ no cabría concluir $\lceil \forall x H p \rceil$. Y eso es raro, ¿verdad? ¿Cómo, de haber un ente que no sea completamente así o asá, va a ser empero completamente cierto que todos los entes sí son así o asá? ¿Qué hace que, al tomarlos todos juntos, resulte esa verdad total si hay entre ellos al menos una oveja negra que no es así o asá? (Imaginemos esto: 'Es plenamente verdad que cada miembro de esa familia es generoso; pero Recesvinto, que es miembro de esa familia, no es generoso'. ¿No parece eso **supercontradictorio**?)

Sin embargo, esas consideraciones no son tan sin vuelta de hoja que no quepa lícitamente intentar otras construcciones de cálculos cuantificacionales sobre la base de *Aj*. Habría que aquilatar las ventajas y los inconvenientes de esas construcciones, procediendo a una estimación comparativa.

Dado, empero, el volumen ya alcanzado por lo que se proponía ser un mero opúsculo iniciativo, me abstendré de desarrollar alternativas, limitándome a exponer una manera plausible

de construir un cálculo cuantificacional sobre la base de A_j , ya que, de entre los sistemas lógicos examinados en este trabajo, es A_j el que más adecuado parece para una más amplia gama de aplicaciones, conteniendo —según se vio— como subsistema suyo a cada uno de los cálculos sentenciales finivalentes. (A_j parece una aproximación buena, razonable, al ideal de constituir “la” gran lógica sentencial que subsuma —como subsistemas propios— a [las más de] cuantas lógicas sentenciales ofrezcan credenciales de plausibilidad).

El sistema A_q

Reglas de formación

- 1ª) Si $\ulcorner p \urcorner$ es una fbf de A_j , también lo es de A_q .
- 2ª) Si $\ulcorner p \urcorner$ es una fbf de A_q , también lo es $\ulcorner \forall \alpha p \urcorner$, donde ‘ α ’ es una variable cualquiera.
- 3ª) Si $\ulcorner p \urcorner$, $\ulcorner q \urcorner$ son fbfs de A_q , también lo son $\ulcorner Bp \urcorner$, $\ulcorner p \downarrow q \urcorner$, $\ulcorner plq \urcorner$, $\ulcorner Hp \urcorner$, $\ulcorner p \bullet q \urcorner$.

Esquema definicional

$\ulcorner \exists \alpha p \urcorner$ abr. $\ulcorner N \forall x (1 \bullet Np) \urcorner$

(siendo ‘ α ’ una variable cualquiera),

Reglas de inferencia

- (1º) Cada regla de inferencia de A_j es también una regla de inferencia de A_q .
- (2º) Las siguientes son reglas de A_q :
 $\text{rinf3 } p \vdash q$ (donde $\ulcorner q \urcorner$ es una variación alfabética de $\ulcorner p \urcorner$)
 $\text{rinf4 } p \vdash q$ (donde $\ulcorner q \urcorner$ es una versión alternativa de $\ulcorner p \urcorner$)
 $\text{rinf5 } p \vdash q$ (donde $\ulcorner q \urcorner$ es el resultado de prefijar a $\ulcorner p \urcorner$ un número finito de cuantificadores universales).

Esquemas axiomáticos

- (1º) Si $\ulcorner p \urcorner$ es un esquema axiomático de A_j , también lo es de A_q .
- (2º) Los siguientes esquemas son axiomáticos en A_q :

$$\text{A31 } \forall x (\exists x p \bullet q) \vdash \exists x (\forall x p \bullet q)$$

$$\text{A32 } \forall x (p \bullet q) \rightarrow \forall x p \bullet q$$

$$\text{A33 } \forall x p \wedge \exists x q \rightarrow \exists x (p \wedge q)$$

$$\text{A34 } \forall x Bp \rightarrow B \forall x p$$

$$\text{A35 } \forall x \neg p \rightarrow \neg \exists x p$$

$$\text{A36 } \forall x p \downarrow q \supset \exists x (p \downarrow q) \wedge \exists x (\forall x p \rightarrow q \rightarrow .mq \wedge p \rightarrow q)$$

En A36 $\ulcorner q \urcorner$ no ha de contener ninguna ocurrencia libre de ‘ x ’.

Comentarios

Huelga, tras las consideraciones de unas páginas más atrás, detenerse a argumentar a favor de rinf3 y de rinf4, reglas cuya corrección es ya sobradamente clara. Por lo que hace a rinf5, ésta se justifica así. Si bien en un determinado entorno de elocución proferir una fórmula con variables libres puede servir para hablar de tales o cuales entes en particular a los que —en ese contexto y en ese entorno— se refieran dichas variables, en el contexto de una teoría —o sea de una serie de asertos teóricos, doctrinales— no cabe eso, toda vez que en esos contextos no se brindan indicadores contextuales o de entorno que permitan saber que tal ocurrencia de tal variable apunta a tal individuo o ente en vez de apuntar a algún otro. Por ende, en el marco de aseveraciones puramente teóricas una fórmula enunciada en la que haya variables libres ha de entenderse como aplicándose con verdad a cualesquiera entes tomados como aquellos a los que se refieran tales variables. Así, «Si [ello¹] es un mamífero,

[ello¹] es un vertebrado» (o, más coloquialmente, 'Un mamífero es un vertebrado') es una fórmula que, aseverada en un contexto puramente teórico, quiere decir que **cualquier ente** que sea un mamífero es un vertebrado. Por ello es lícito concluir de tal aserto el de que, precisamente, todo ente es así (tal que, si es un mamífero, es un vertebrado). Con otras palabras: una fórmula teórica con variables libres —si es que es verdadera— aplícase con verdad a cualquier ente (por separado); y por consiguiente cabe concluir que lo mismo es verdad de todos los entes conjuntamente tomados. Eso es lo que hace, precisamente, esta regla, rinf5. (Por lo tanto la regla sólo es de aplicación en contextos teóricos. Fuera de ellos, no concluiríamos de que alguien diga 'Ella es sobremanera bondadosa' que todo ente es sobremanera bondadoso.)

Pasemos ahora a decir unas palabras sobre los esquemas axiomáticos. El A31 nos dice que el que todo ente, x , sea tal que no sólo hay algo, x , tal que p sino que también q es equivalente a que haya algo, x , tal que no sólo todo ente, x , es tal que q sino que también p . Sintéticamente tenemos ahí —sólo que las consecuencias se van sacando gracias a los otros esquemas axiomáticos, a la definición del cuantificador existencial, y a las reglas de inferencia del sistema— varias equivalencias fundidas en una. El esquema A32 es también un principio bastante obvio: el que todo ente, x , sea tal que no sólo p sino además q es algo a lo sumo tan verdadero como que no sólo todo ente es tal que p sino que además [tal ente en particular es tal que] q . De ahí sale el principio de aplicación: $\lceil \forall x p \rightarrow p \rceil$ donde $\lceil p \rceil$ difiere de $\lceil p \rceil$ a lo sumo por el reemplazo uniforme de las ocurrencias libres de ' x ' en ' p ' por sendas ocurrencias libres de alguna variable.

Nótese que la conyunción de A31 y A32 equivale —dados los otros esquemas axiomáticos y dadas las reglas de inferencia— a este esquema:

$\lceil \forall x p \bullet \forall x q \mid \forall x (p \bullet q) \wedge \forall x p \rightarrow p \wedge \exists x p \bullet r \mid \exists x (p \bullet r) \wedge r \mid \forall x r \rceil$, con tal de que en este esquema ' r ' no tenga ninguna ocurrencia libre de ' x '. Naturalmente el postular los dos esquemas A31 y A32 es más económico (por ende, más elegante), si bien la postulación de este otro esquema tan largo comportaría la ventaja de que cada uno de los conjuntos que lo forman es más perspicuo, más claramente verdadero a simple vista, según cabe mostrar con instancias de esos conjuntos. (Déjasele el hacerlo al lector, como ejercicio.)

El esquema A33 es un principio de adjunción cuantificacional: en la medida en que todo ente sea así o asá y haya un ente de estas o las otras características, [en esa medida al menos] hay un ente que es así o asá y es de estas o las otras características.

El esquema A34 nos dice que en la medida en que sea de cada ente afirmable con verdad que es así o asá, [en esa medida por lo menos] es afirmable con verdad que todos los entes son así o asá.

Por último A36 tiene dos partes, dos conjuntos. El primero dice que, si es menos verdad que todos los entes son así o asá que no que se cumple tal condición, entonces hay algún ente tal que su ser así o asá es menos verdadero que el cumplirse dicha condición. (Una instancia de ese conyunto izquierdo de A36 es ésta: 'Si el que todos los turcos sean asiáticos es algo menos verdadero que el hecho de que Kemal es asiático, entonces hay algún turco que es menos asiático que Kemal'.) El segundo conyunto dice que hay algo tal que, en la medida en que, en tanto en cuanto todos los entes sean así o asá, se cumple cierta condición, en esa medida por lo menos es cierto que, en tanto en cuanto la citada condición venga a cumplirse y sea verdad que ello (el algo mencionado al comienzo) es así o asá, esa condición se cumple. Lo cual equivale a este otro esquema, más largo pero más claro —siempre con la salvedad de que ' q ' no contenga ocurrencias libres de ' x '—, a saber: $\lceil q \mid m q \rightarrow \bullet \forall x p \rightarrow q \rightarrow \exists x (p \rightarrow q) \rceil$, que es un principio de **prenexación** (o sea de traslado del cuantificador de la prótasis a la oración implicativa total, pero cambiando el cuantificador, que de universal pasa a ser

existencial). En la lógica clásica ese principio de prenexación vale sin reservas; aquí vale con esa reserva: para condiciones $\lceil q \rceil$ tales que el hecho de que q sea menos existente o verdadero que el [mero] venir a ser cierto que q . ¿Por qué tal restricción? Porque, si no, tendríamos esto. Tomemos lo infinitesimalmente existente o verdadero, a . Supongamos que el que todo ente es así o asá ($\forall xp$) es sólo infinitesimalmente verdadero; por ende: $\lceil \forall xp \rightarrow a \rceil$: en tanto en cuanto todo ente es tal que p , existe o es verdad lo infinitesimalmente real; sin embargo es posible que ningún ente sea tal que el ser ese ente así o asá sea algo infinitesimalmente verdadero; puede que para cada ente así o asá haya otro también así o asá pero menos, formándose una serie decreciente cuyo límite sea a (lo infinitesimalmente real), aunque cada elemento alético de esa serie esté por encima de a —e.d. cada ente es suprainfinitesimalmente así o asá, pero para cada uno hay otro que, aun siendo también suprainfinitesimalmente así o asá, está por debajo en su ser así o asá, y no hay ningún grado superior a a tal que todo sea así o asá por lo menos en ese grado. Pero eso puede pasar únicamente porque lo infinitesimalmente real no tiene un grado que sea su tope superior y difiera de él mismo, sino que él es su tope superior ($\lceil \text{alma} \rceil$). En cambio no es posible que pase eso con $\frac{1}{2}$ en vez de a , porque $\lceil \frac{1}{2} \setminus m \frac{1}{2} \rceil$: una sucesión semejante a la mencionada pero que tienda desde arriba a $\frac{1}{2}$, tiende también a $m \frac{1}{2}$; en tal caso no será verdadera la prótasis implicativa, $\lceil \forall xp \rightarrow \frac{1}{2} \rceil$, sino que, antes bien, se tendría $\lceil \forall xplm \frac{1}{2} \rceil$.

Nótese que para el mero condicional, $\supset$, es teoremató en A_j el principio irrestricto de prenexación, o sea $\lceil \forall xp \supset q \supset \exists x(p \supset q) \rceil$ (y sus variantes derivadas, como $\lceil q \supset \exists xp \supset \exists x(q \supset p) \rceil$, siempre con la salvedad ya indicada acerca de $\lceil q \rceil$). Es que $\supset$ no es sensible como $\rightarrow$ a los grados de verdad, a que la apódosis sea al menos tan verdadera como la prótasis. Un corolario del principio implicativo irrestricto de prenexación ($\lceil \forall xp \rightarrow q \rightarrow \exists x(p \rightarrow q) \rceil$), no teoremató en A_j , es el esquema $\lceil \exists x(p \rightarrow \forall xp) \rceil$, una instancia del cual es ésta: ‘Hay un cuerpo tal que en tanto en cuanto él es grande, todos los cuerpos lo son’. Eso es del todo incompatible con la tesis de que para cada cuerpo se da otro mayor (una tesis que, aunque hoy por hoy pasada de moda y rechazada por los físicos, a muchos nos parece verosímil y atractiva, a la espera de que venga a reentronizarla un nuevo viraje en la investigación física —un cambio de paradigma—).

La índole del presente opúsculo hace desaconsejable detenernos más en comentarios sobre los esquemas axiomáticos de A_q , así como también impide que nos dediquemos a demostrar esquemas teoremató. Son tareas que quedan postergadas a un trabajo ulterior.

Capítulo 11

La lógica combinatoria

Existe un tipo de sistema lógico muy diverso del género de sistemas que hemos venido considerando hasta aquí. Son las lógicas combinatorias.

En los sistemas usuales de lógica se construye por pisos. Un primer piso o planta baja: el cálculo sentencial —al que hemos consagrado todo este trabajo, hasta este punto, salvo el Capítulo 10. Un segundo piso: el del cálculo cuantificacional [de primer orden], estudiado en el Capítulo 10. Luego viene un tercer piso: la teoría de conjuntos, que es una extensión conservativa del cálculo cuantificacional. (Alternativamente, en vez de teoría de conjuntos puede construirse un cálculo cuantificacional de orden superior a 1.)

La teoría de conjuntos añade al cálculo cuantificacional un predicado diádico, '∈' 'ε', o sea un signo tal que, al interponerse entre dos signos individuales —variables o constantes—, 'α', 'β', formándose la ristra 'α ∈ β', se tiene una **fórmula**. Si la fórmula carece de variables libres, entonces es una oración.

Una teoría de conjuntos añade ciertos axiomas o esquemas axiomáticos que involucran a ese predicado '∈'. El *desideratum* óptimo sería éste: tener un operador λ tal que yuxtapuesto a una variable 'x' p.ej., y a una fórmula, 'p', resultara un **término** —una constante individual definida—, 'λxp', siendo teorematícos estos esquemas:

- (1) $\forall x(x \in \lambda z p | q)$, donde «q» resulta de «p» sin más que reemplazar cada ocurrencia libre de 'z' por una ocurrencia libre de 'x'. (A (1) lo llamamos '**principio de abstracción**' (P.A.).)
- (2) $\exists x(x = \lambda z p)$ donde '=' es el signo de identidad (que se define p.ej. así: ' $\alpha = \beta$ ' abrevia a ' $\forall x(x \in \alpha | \beta \in x)$ '). (Este es el **principio de comprensión**, P.C.)
- (3) $\exists x(x \in \alpha | x \in \beta \supset \alpha = \beta)$. Este es el **principio de extensionalidad**, P.E.: para que dos conjuntos difieran ha de haber algo que pertenezca al uno pero no al otro, o bien a uno de ellos más que al otro.

Lo malo es que con esos tres principios se llega a una paradoja, a una fórmula que resulta demostrable pero cuya suprenegación también lo será. Con lo cual el sistema se hará delicuescente, **absolutamente** inconsistente. En efecto, por (2) habrá algo que sea $\lambda x \neg(x \in x)$, el conjunto de cosas que no se abarquen en absoluto a sí mismas. Por ende, y a tenor de un esquema teorematíco del cálculo cuantificacional (a saber: ' $\exists x p \supset \forall x q \supset \exists x(p \wedge q)$ ') se tendrá, en virtud de (1):

$$\exists x(x \in \lambda x \neg(x \in x) | \neg(x \in x) \wedge x = \lambda x \neg(x \in x)).$$

Por la definición de la identidad, '=', y las reglas del cálculo cuantificacional se tendrá entonces:

$$\lambda x \neg(x \in x) \in \lambda x \neg(x \in x) | \neg(\lambda x \neg(x \in x) \in \lambda x \neg(x \in x))$$

Y de ahí, por las reglas del cálculo sentencial, se deduce:

$$\lambda x \neg(x \in x) \in \lambda x \neg(x \in x) \wedge \neg(\lambda x \neg(x \in x) \in \lambda x \neg(x \in x)).$$

Pero ésa es una fórmula **super**contradictoria de la forma ' $p \wedge \neg p$ ': p y es **del todo** falso que p. Tal es la paradoja de Russell. Al descubrirse, resultó ser delicuescente el primer sistema axiomático de cálculo cuantificacional y teoría de conjuntos: el de Gottlob Frege.

En honor a la verdad hay que decir que, como ni Frege ni Russell diferenciaban la negación fuerte '¬' de la negación natural o simple 'N', se creyó entonces demostrado que no podía existir $\lambda x N(x \in x)$. Sin embargo no es así. Una teoría de conjuntos que se articule sobre la base de una extensión cuantificacional de At (o de Ap o de Aj) puede perfectamente, sin desmoronarse (sin hacerse delicuescente) contener el teorema:

$$\lambda x N(x \in x) \in \lambda x N(x \in x) \wedge N(\lambda x N(x \in x) \in \lambda x N(x \in x)).$$

Será una teoría contradictorial, mas no por ello forzosamente incoherente, o sea no absolutamente inconsistente.

El problema está entonces en cómo articular teorías que se aproximen lo más posible a poseer algo parecido a (1), (2), (3) —o a versiones tan poco restringidas como sea posible de tales principios— sin incurrir en incoherencia, en delicuescencia.

En el marco de la lógica clásica (en la cual ' $\rightarrow$ ' pasa simplemente a ser ' $\supset$ '; ' $\neg$ ', simplemente a ser ' $\equiv$ '; y, por supuesto, ' $\neg$ ' simplemente a ser ' $\neg$ ') lo que se suele hacer es construir teorías de conjuntos alejadas de postular uno u otro de esos tres principios. P.ej. teorías como ML de Quine que afirma algo parecido a (2) pero está alejadísima de afirmar en general (1). O teorías como ZF (el sistema de Zermelo-Fraenkel) más próximas a afirmar algo parecido a (1) pero infinitamente alejadas de afirmar (2).

Las lógicas combinatorias optan por otro camino. Desembocan en resultados que equivalen a tener siempre como teoremas los esquemas (1), (2) y (3) —o variantes de los mismos. Pero el punto de partida es otro. En vez de construirse por pisos, de entrada introducen, como primitivos, unos operadores mediante cuya **combinación** (yuxtaposición, concatenación) cabe definir los signos del cálculo sentencial, del cuantificacional y de la propia teoría de conjuntos. Para evitar las paradojas acuden a procedimientos a veces complicados de restringir el campo de aplicabilidad de reglas de inferencia o al abandono de ciertos principios del cálculo sentencial como el principio de tercio excluido en cualquiera de sus variantes.

Uno de los sistemas más bonitos de lógica combinatoria es el de Fitch. Hé aquí algo que, aun siendo diferente de la presentación del propio Fitch, se aproxima a ella y es muy elegante.

Sea la concatenación de dos signos cualesquiera, α , β , un signo. La concatenación es asociativa hacia la izquierda, viniendo tal asociatividad interrumpida por paréntesis. Sean Σ , Λ , dos signos **combinadores** para los que se postula lo siguiente (tomándose '=' como signo primitivo pero con un principio de reemplazo: ' $x=zy.p \supset p$ ', donde ' p ' es como ' p ' salvo por el reemplazo de las ocurrencias libres de ' x ' por sendas ocurrencias libres de ' z ')

$$(4) \Lambda xz = x$$

$$(5) \Sigma xzu = xu(zu)$$

Defínense entonces otros combinadores: Δ , Γ , Ω , así:

' Δ ' abr. ' $\Sigma(\Lambda\Sigma)\Lambda$ '

' Γ ' abr. ' $\Sigma(\Delta\Delta\Sigma)(\Lambda\Lambda)$ '

' Ω ' abr. ' $\Gamma\Sigma(\Sigma\Lambda\Lambda)$ '

Es interesante constatar entonces que para cualesquiera signos b , c , d : ' Δbcd ' = ' $b(cd)$ '; ' Γbcd ' = ' bdc '; ' $\Omega bc = bcc$ '.

Además, si se define ' 1 ' como ' $\Sigma\Lambda\Lambda$ ' se tendrá: ' $1b=b$ '.

Defínese entonces ' λrp ' así, donde ' r ', ' p ' son signos cualesquiera:

1º) ' $r=p$ ': entonces ' $\lambda rp=1$;

2º) ' r ' no figura en ' p ': entonces ' $\lambda rp = \Lambda p$ ';

3º) no se da ninguno de los dos casos precedentes: entonces ' p ' ha de ser un signo compuesto, ' qs ', tal que ' r ' tiene alguna ocurrencia en ' q ' o en ' s ' (o en ambos); entonces ' $\lambda rp = \Sigma\lambda rq\lambda rs$ '.

Se trata, con esas tres cláusulas (1º a 3º)), de una definición. Según como sean ' r ', ' p ', ' λrp ' abreviará o bien a 1 , o bien a ' Λp ', o bien a ' $\Sigma\lambda rq\lambda rs$ ' (siendo ' $p=q$ ', en el caso 3º, siempre y cuando se cumpla la condición de que ' $p \neq r$ ' pero ' r ' tenga ocurrencia en ' p ').

En una lógica así no hay variables. La cuantificación defínese así. Habrá un signo primitivo, 'E', que denotará a la propiedad de ser no-vacío. (No hacemos diferencia alguna entre propiedades y clases o conjuntos. Por lo menos a los presentes efectos identificamos el que algo tenga la propiedad de blancura y el que pertenezca al conjunto de cosas blancas.) Defínese entonces: $\lceil \exists r p \rceil$ abr. $\lceil E \lambda r p \rceil$. Decir que hay algo, r, tal que p será decir que es no-vacío el conjunto de entes, r, tales que p.

Y ¿cómo se evita la paradoja en un sistema así? Postulando para los signos de ese sistema que hagan las veces de signos de cálculo sentencial principios que en algunos puntos sean más débiles, evitando así la teorematidad de esquemas como, p.ej., el tercio excluido. En el sistema de Fitch no es teoremató el esquema $\lceil p \vee \neg p \rceil$ (o —si se traduce su negación como 'N' en vez de '¬'— $\lceil p \vee Np \rceil$). Conque, si bien ese sistema entraña que $\lambda x \neg (x \in x) \in \lambda x \neg (x \in x) = \neg (\lambda x \neg (x \in x) \in \lambda x \neg (x \in x))$, así y todo no se deduce ninguna antinomia de la forma $\lceil p \wedge \neg p \rceil$ (o $\lceil p \wedge Np \rceil$ si entendemos esa negación de Fitch como negación simple, 'N', en vez de fuerte, '¬').

Obsérvese que en la lógica combinatoria el signo de membría, $\in$, se entiende así: $\lceil x \in z \rceil$ abr. $\lceil zx \rceil$: z abarca a x. En un sistema de lógica combinatoria cualquier cosa puede atribuirse a cualquier cosa —a lo mejor sin verdad, pero puede. P.ej., decir $\lceil \text{no } p \rceil$ es atribuir a p la negatividad, la propiedad denotada por 'no'. Igualmente decir que una cosa x tiene una relación r es decir eso y nada más; lo que sucede es que en las aplicaciones de tales lógicas suele o puede entenderse, p.ej., que el que x guarde con z la relación r es que el tener x esa relación r sea algo que, a su vez, abarque a z. Así, el estar enamorado Marcial será la propiedad de ser un ente, z, del cual esté enamorado Marcial.

Tomando en particular a las relaciones denotadas por funtores diádicos del cálculo sentencial, tenemos esto. Sea el functor de conjunción ' $\wedge$ '. Para el adepto de una lógica combinatoria, ' $\wedge$ ' denota una relación, la relación que se da entre dos entes en la medida en que el uno y el otro son hechos verdaderos, e.d. existenciales. P.ej. si $x =$ el estar apenado Elías, mientras que $z =$ el ser engendrada Marta (o sea: el conjunto de sus engendradores), entonces $x \wedge z$ (que combinatoriamente se escribiría más bien así: ' $\wedge zx$ ') sería verdad en la medida en que lo sea que Elías está apenado y Marta es engendrada. $\wedge z$ sería que el ser engendrada Marta posea esa propiedad o relación copulativa o conjuntiva; y eso sería el conjunto de cosas o hechos con los que guarda esa relación, o sea aquellos hechos cada uno de los cuales es tal que Marta es engendrada y ese hecho también sucede.

Para hacer más uniforme el tratamiento ontológico articulable en un marco así cabe identificar a cada ente con [el hecho de] su existencia. Con lo cual todo ente es un hecho. Además, hay razones de mayor peso para postular esa identificación o reducción ontológica, al paso que las dificultades contra la misma parecen de poca monta y casi únicamente estriban en peculiaridades de expresión que se explican cómodamente como alomorfía en distribución parcialmente complementaria debida a constreñimientos pragmáticos. (Alomorfía es —según lo sabe el lector— la alternancia entre dos formas de un mismo signo, como —en ciertos contextos— 'hubiera'/'hubiese', **en distribución complementaria** es que la alternancia se da en función del contexto, como en francés 'beau'/'bel'.)

Sea ello como fuere, interesante aquí es señalar lo elegante que resulta un tratamiento como el de las lógicas combinatorias. Y, a fuer de elegante, filosóficamente plausible también (a tenor al menos del criterio epistemológico de optar por la **mejor** explicación de las cosas, o sea: la más elegante).

En varios trabajos citados en la bibliografía del presente opúsculo se ha emprendido la tarea de elaborar sistemas de lógica combinatoria que, siendo extensiones conservativas de Aq , no sean delicuescentes y, sin embargo, se aproximen asintóticamente al ideal de contener

como esquemas teorematizados los principios (1), (2) y (3), o al menos versiones matizadas de los mismos. De momento lo conseguido al respecto está en estadio de experimentación, y, si bien se han formulado esbozos de tratamiento semántico de alguna de esas construcciones, el tratamiento es condicional (bajo tales o cuales condiciones, existen modelos de tales sistemas), no habiéndose ofrecido todavía ninguna prueba de no-delicuescencia de ninguno de esos sistemas (aunque es probable que alguno de ellos por lo menos no sea delicuescente).

La línea general seguida en tales construcciones consiste en que, en vez de postularse sin restricciones (4) (o sea: $\lceil \wedge pqlp \rceil$) y (5) (o sea: $\lceil \Sigma pql.pr(qr) \rceil$), vienen restringidos ambos principios. En vez de (4) y (5), se postulan muchos casos parciales de (4) y de (5), quedando abierto el sistema (que será, pues, un conjunto de teoremas **no** recursivamente enumerable), pudiéndose añadir en todo momento más y más casos parciales de (4) y de (5), cada uno de ellos restringido de cierta manera. Trátase así de aproximarse uno lo más posible a tener (1), el P.A., así como versiones matizadas de (2) y de (3) (ninguno de los tres principios es, tal cual, teorematizado en los sistemas que estoy ahora considerando; así (3), el principio de extensionalidad, viene modificado insertándose una ocurrencia del functor 'B' que afecta a la prótasis).

Ahora bien, hay en (1) dos partes implicativas, a saber:

(1a) $\forall x(x \in \lambda zp \rightarrow q)$

(1b) $\forall x(q \rightarrow x \in \lambda zp)$.

Hay un enfoque (de Rescher y Brandom) en el cual se reconocen esas dos mitades pero cambiando en una de ellas el operador ' λ ' por otro, digamos ' μ ': habrá —digamos— un **conjunto** que abarque a todos los entes que p, pero también a otros tal vez; y una **clase** que abarque sólo a entes que p, mas posiblemente no a todos. (Tal diferencia terminológica entre conjuntos y clases sería arbitraria, desde luego.) En vez de eso, lo que hacen los sistemas combinatorios ahora considerados es reconocer en cada caso un único **conjunto** o **clase** o **cúmulo** de entes, λxp , pero, postulando (1a) sin restricciones, no aceptar en cambio (1b) más que en una serie de versiones restringidas de diversa manera (serie infinita, abierta).

Eso quiere decir que los cúmulos reconocidos en tales sistemas son cúmulos que abarcan, cada uno de ellos, a cuantas cosas cumplan la condición definitoria de miembro del cúmulo en cuestión; pero que no siempre abarcan sólo a los entes que cumplan tal condición. Desde luego podría intentarse una vía divergente: reconociendo todos los casos de (1b), restringir (1a): un cúmulo abarcaría siempre sólo a cosas que cumplan la condición definitoria, pero acaso no a todas ellas. Un cúmulo que infrinja (1a) puede llamarse una «caterva»; uno que infrinja (1b), un «corrillo». Si no vale sin restricciones P.A., o sea (1), entonces o hay catervas o hay corrillos. (Rescher parece preferir que haya de lo uno y de lo otro.) Por consideraciones filosóficas que no es del caso exponer aquí, le parece más verosímil al autor del presente trabajo que haya catervas, pero no corrillos. Lo más interesante es que el cúmulo fuerte de Russell, a saber $\lambda x \neg (xx)$, el cúmulo de entes que no se abarquen en absoluto a sí mismos, será o bien una caterva o bien un corrillo. Según el tratamiento ofrecido en los sistemas combinatorios ahora aludidos, es una caterva: uno precisamente de los entes que vienen abarcados por él es él mismo, pese a que, por ello precisamente, no cumple [en absoluto] la condición definitoria (la de no abarcarse en absoluto a sí mismo). Según el tratamiento alternativo sería un corrillo: no se abarcaría en absoluto a sí mismo, a pesar de que cumpliría entonces la condición definitoria para venir así abarcado (la de no abarcarse a sí mismo en absoluto).

Llegados a este punto, la lógica matemática desemboca en problemas claramente metafísicos y patentiza todavía más cuánto hay en esta disciplina de "especulación", de conjetura —o, si se quiere, de elucubración. No conlleva el reconocerlo una caída en irracionalismo o anarquismo metodológico a lo Feyerabend, pero sí conduce a cobrar conciencia de lo relativo

de las justificaciones racionales, toda vez que están exentos de justificación **última** o fundacional los propios patrones lógicos que sirven para calibrar cuán racional sea un pensamiento.

Capítulo 12

Modelos algebraicos

No estriba la peculiaridad de la noción de modelo algebraico (que, en este capítulo, va a venir introducida escuetamente) sino en esto: en la noción de **operación**. Dícese que, sobre un cúmulo o conjunto C dado, está definida una operación m -aria [siendo m un número natural ≥ 0], φ , ssi para cualesquiera m miembros de C , $c_1, c_2, \dots, c_m$, existe un único miembro de C , d , tal que $d = \varphi(c_1, c_2, \dots, c_m)$. Conque un **álgebra** viene definida así: un álgebra es un dúo $A = \langle C, O \rangle$ donde C es un conjunto cualquiera y O es un cúmulo de operaciones cada una de las cuales está definida sobre C .

Nótese que una operación puede ser 0-aria; en tal caso cabe tomarla como una constante, o sea un miembro fijado, “distinguido”, del cúmulo C —suele denominarse el **portador** del álgebra en cuestión, lo cual abreviaremos como $\wp A$. De momento, sin embargo, tomamos operaciones de aridad finita nada más (aunque eso vendrá rectificado al final del capítulo); o sea, el número m en cuestión ha de ser en cada caso finito.

Las álgebras más comúnmente estudiadas en manuales introductorios son las booleanas; y es que constituyen modelos característicos de la lógica clásica. (Vide al respecto el final del Capítulo 6.) Sin embargo, con vistas a la generalidad vale más considerar aquí, para empezar, las álgebras cuasiboleanas. Un **álgebra de Morgan** es un dúo $\langle C, \Delta \rangle$, donde Δ es el cúmulo $\{\star, \oplus, \odot\}$, donde $\star$ es una operación unaria, $\oplus$ y $\odot$ son operaciones binarias, y se cumplen estos postulados para cualesquiera $x, z, y, u, v, \in C$ (escribiendo, p.ej., ‘ $x \oplus z$ ’ en vez de ‘ $\oplus(x, z)$ ’ y similarmente para ‘ $\odot$ ’ y para cualesquiera otros elementos en vez de x, z):

Idempotencia: $x \oplus x = x = x \odot x$

Asociatividad: $x \oplus (y \oplus z) = (x \oplus y) \oplus z$; $x \odot (y \odot z) = (x \odot y) \odot z$

Conmutatividad: $x \odot y = y \odot x$; $x \oplus y = y \oplus x$

Distributividad: $x \oplus (y \odot z) = (x \oplus y) \odot (x \oplus z)$; $x \odot (y \oplus z) = (x \odot y) \oplus (x \odot z)$

Absorción: $x \oplus (y \odot x) = x = x \odot (y \oplus x)$

De Morgan: $\star(x \oplus y) = \star x \odot \star y$; $\star(x \odot y) = \star x \oplus \star y$

Involutividad: $\star \star x = x$

Llámbase **cuasiboleana** un álgebra $\langle C, \Delta \rangle$ donde $\Delta = \{1, \star, \oplus, \odot\}$, donde 1 es una “operación 0-aria” —o sea, una constante, un miembro fijo de C — al paso que $\langle C, \{\star, \oplus, \odot\} \rangle$ es un álgebra de Morgan y, además, para cualquier miembro de C , x , se cumple este postulado: $x \oplus 1 = 1$. (1 es el elemento máximo del álgebra en cuestión.)

En general un álgebra $\langle A, \Gamma \rangle$ **subsume** a un álgebra $\langle A, \Delta \rangle$ ssi $\Delta \subseteq \Gamma$ (Δ es un subconjunto —propio o no— de Γ). Está claro que un álgebra cuasiboleana subsume a un álgebra de Morgan; hablando laxamente cabe decir que un álgebra cuasiboleana **es** un álgebra de Morgan que cumple cierta condición, a saber la existencia de elemento máximo.

Llámbase **de Kleene** un álgebra cuasiboleana que cumple este postulado (para cualesquiera elementos x, z): $x \odot \star x \leq z \oplus \star z$ [donde en general, para dos elementos cualesquiera del portador de esa álgebra, u, v , se define ‘ $u \leq v$ ’ como ‘ $u \odot v = u$ ’].

La mejor manera de hacer corresponder un álgebra $A = \langle \wp A, \Delta \rangle$ de cierto tipo como modelo a un sistema lógico $\mathfrak{L}$ es estipular que cada valuación de ese sistema, v , será tal que para cada functor m -ádico de $\mathfrak{L}$, ϕ , hay una operación m -aria $\ast \in \Delta$, tal que $v(\phi p^1, \dots, p^m) = \ast(vp^1, \dots, vp^m)$, donde ‘ $\ulcorner p^1 \urcorner, \dots, \ulcorner p^m \urcorner$ ’ son fbfs cualesquiera de $\mathfrak{L}$, siendo $vp^1, \dots, vp^m \in \wp A$. Las demás condiciones necesarias para que el álgebra —o su portador, más exactamente— sea un modelo del sistema dado son las mismas que ya conocemos por el Capítulo 6.

Para obtener la completez de un sistema $\mathcal{L}$ suélese definir una cierta clase Γ de álgebras y entonces se postula con respecto a $\mathcal{L}$ que es una tautología de $\mathcal{L}$ cualquier fórmula $\lceil p \rceil$ tal que para cualquier valuación [admisble] v de $\mathcal{L}$ en una cualquiera de las álgebras pertenecientes a Γ , $v(p)$ es un elemento designado [del portador] de esa álgebra.

Un álgebra booleana es un álgebra de Kleene tal que para cualquier elemento x se tiene: $x \odot \star x = \star 1$.

A título de correspondencia plausible cabe señalar que la operación $\oplus$ —llamada **junción**— de un álgebra de Kleene corresponde a la disyunción natural ' $\vee$ '; que la operación $\odot$ —llamada **cruce**— corresponde a la conjunción natural ' $\wedge$ '; que la operación unaria $\star$ corresponde a la negación natural ' N '. Sólo que si se trata de un álgebra booleana, esa operación unaria $\star$ corresponderá a la negación fuerte ' $\neg$ '. Por razón de la correspondencia en cuestión se suele escribir: ' $\cap$ ' o ' $\wedge$ ' en vez de ' $\odot$ '; ' $\cup$ ' o ' $\vee$ ' en vez de ' $\oplus$ ' (aunque también a veces '+'); para la operación unaria se escribe a veces '*'. (Por otro lado, los funtores de conjunción y de disyunción escríbense en ciertas notaciones como ' $\cdot$ ' y '+' respectivamente. La notación aquí empleada pretende ser sugestiva pero, sin embargo, diferenciadora de lo que es un functor respecto de lo que es una operación algebraica.)

Por otro lado, si queremos elaborar modelos algebraicos para sistemas lógicos más flexibles y complejos —con varias conjunciones y negaciones, p.ej., o con funtores que no sean definibles en términos de conjunción o disyunción más negación—, habrá que enriquecer la noción de álgebra que queramos estudiar con nuevas operaciones. P.ej. en los sistemas lógicos explorados en los capítulos precedentes aparece una conjunción fuerte, o superconjunción ' $\bullet$ ', que no es idempotente ($\lceil p \bullet p \rceil$ no es un esquema teorematizado de esos sistemas); será menester estudiar como modelos de sistemas así algunas álgebras donde, además del cruce "normal", $\odot$, haya un "supercruce", $\otimes$, tal que no siempre se cumpla $x \otimes x = x$. En esos sistemas hay funtores como ' I ', ' H ', ' m ', que no tienen ninguna operación que les corresponda en un álgebra de Kleene. Hay que diseñar un cúmulo de postulados que permitan definir una clase de álgebras con operaciones cuyas características correspondan exactamente a las de esos funtores —álgebras que, por lo demás, sean también álgebras de Kleene, o más exactamente: subsuman a álgebras de Kleene.

Para simplificar y no seguir añadiendo notación adicional, de ahora en adelante escribimos la juncción como ' $\vee$ ', el cruce como ' $\wedge$ ' y la operación unaria de un álgebra de Kleene (que no sea de Boole) como ' N '. El contexto aclarará más que suficientemente cuándo uno de tales signos viene empleado como functor de un sistema lógico, y cuándo hace las veces de una operación algebraica. Similarmente con respecto a nuevas operaciones que introduzcamos en lo sucesivo.

Álgebras cuasitransitivas

Un **álgebra cuasitransitiva** —a.c.t. para abreviar— es un álgebra $\langle A, \Delta \rangle$ tal que $\Delta = \{1, N, H, n, \vee, \bullet, I\}$, donde 1 es una operación 0-aria, N, H , y n son operaciones unarias y $\vee, \bullet, I$ son operaciones binarias que satisfagan los 27 postulados indicados más abajo. Introduzcamos primero algunas definiciones:

$\lceil 0 \rceil$ abr $\lceil N1 \rceil$	$\lceil x \wedge y \rceil$ abr $\lceil N(Ny \vee Nx) \rceil$	$\lceil Sx \rceil$ abr $\lceil x \wedge Nx \rceil$	$\lceil x \rightarrow y \rceil$ abr $\lceil x \wedge y \rceil$
$\lceil \neg x \rceil$ abr $\lceil HNx \rceil$	$\lceil Xx \rceil$ abr $\lceil x \bullet x \rceil$	$\lceil mx \rceil$ abr $\lceil NnNx \rceil$	$\lceil Kx \rceil$ abr $\lceil NXNx \rceil$
$\lceil a \rceil$ abr $\lceil m0 \rceil$	$\lceil fx \rceil$ abr $\lceil \neg(xIa) \wedge x \rceil$	$\lceil Lx \rceil$ abr $\lceil N\neg x \rceil$	$\lceil \frac{1}{2} \rceil$ abr $\lceil 111 \rceil$
$\lceil fx \rceil$ abr $\lceil fSx \rceil$	$\lceil x \uparrow y \rceil$ abr $\lceil YNx \wedge fy \vee YNy \wedge fx \rceil$	$\lceil x \downarrow y \rceil$ abr $\lceil x \rightarrow y \wedge \neg(y \rightarrow x) \rceil$	

$x \leq y$ significa que $y = y \vee x$; $x \downarrow y$ significa que, dándose el caso de que $x \leq y$, $x \downarrow y = 0$. Sea $D = \{x \in A: \neg x = 0\}$, o sea: el cúmulo de elementos, x , de A tales que $\neg x = 0$. D es el cúmulo de los **elementos densos** de A .

Postulados (para cualesquiera $x, y, z, u, v \in A$):

(01) $y \wedge x \vee x = x$		(02) $x \downarrow y \leq x \wedge u \vee z \downarrow (y \vee z \wedge u \vee z)$	(03) $Hx \wedge Hy = LH(y \wedge x)$
(04) $z \downarrow y \leq Hx \vee Hz \downarrow H(x \vee y)$		(05) $v \downarrow y \leq v \bullet (x \wedge u) \bullet z \downarrow (u \bullet z \wedge (x \bullet z) \bullet y)$	
(06) $x \bullet 1 = x$		(07) $x \bullet y \leq y \wedge x$	(08) $x \wedge y \wedge \neg(x \bullet y) = 0$
(09) $x \downarrow x \in D$		(10) $\frac{1}{2} = N\frac{1}{2}$	(11) $x \downarrow y \leq z \downarrow y \downarrow (x \downarrow z)$
(12) $x \downarrow y \wedge \neg x \wedge y = 0$			
(13) $\neg(x \downarrow 0 \vee x) = 0$	(14) $x \downarrow y \downarrow \frac{1}{2} \vee (x \downarrow y \downarrow 0) = \frac{1}{2}$		(15) $x \downarrow Ny = Nx \downarrow y$
(16) $x \rightarrow y \vee (y \rightarrow nx) \vee (x \downarrow my) = \frac{1}{2}$		(17) $\neg(nmx \downarrow nx) \wedge x = 0$	(18) $x \bullet y \downarrow a \leq x \downarrow a \vee (y \downarrow a)$
(19) $x \rightarrow (x \bullet y) \wedge x \wedge fy = 0$		(20) $nx = x \bullet n1$	(21) $nx \downarrow mx = x \downarrow a \vee (x \downarrow Na)$
(22) $a \downarrow \frac{1}{2}$	(23) $f x \wedge f y \leq \neg(x \wedge y \downarrow (x \bullet y))$		
(24) $x \bullet y \rightarrow (u \bullet v) \wedge (u \downarrow x) \leq y \downarrow v \vee (y \rightarrow a) \vee (x \downarrow mu \wedge (n1 \downarrow v)) \wedge (y \rightarrow v)$			
(25) $x \bullet y = X(Kx \bullet Ky)$		(26) $y \bullet mx \downarrow m(x \bullet y) \vee (x \downarrow y) \in D$	(27) $x \downarrow y \in D$ sólo si $x=y$

De esos 27 postulados, los 26 primeros son ecuaciones —o conjunciones de ecuaciones—, o sea fórmulas de la forma $\ulcorner \alpha = \beta \urcorner$. Dícese que una clase de álgebras es **ecuacional** ssi los postulados que sirven para caracterizar a las álgebras de tal clase son, todos ellos, ecuaciones (o conjunciones de ecuaciones). Una clase ecuacional de álgebras llámase también una **variedad**. Las variedades poseen importantes características, de las cuales carecen otras clases de álgebras. El postulado (27) no es ninguna ecuación, pero sí es lo que se llama un **entrañamiento ecuacional**, de la forma $\ulcorner \text{Si } \alpha = \beta, \text{ entonces } \gamma = \delta \urcorner$, para ciertos $\alpha, \beta, \gamma, \delta$.

Podríamos mitigar el postulado (27) reemplazándolo por otro menos fuerte, como sigue. Definamos primero la noción de **congruencia**: una congruencia en un álgebra $\langle A, \Gamma \rangle$ es una relación diádica, Θ , tal que para cualesquiera miembros $z^1, \dots, z^m, x^1, \dots, x^m$, todos ellos $\in A$, y para cualquier operación m -aria, $f, \in \Gamma$, se cumple esto: Si $x^i \Theta z^i$ para cada i tal que $1 \leq i \leq m$, entonces $f(x^1, \dots, x^m) \Theta f(z^1, \dots, z^m)$. Podríamos entonces reemplazar (27) por este postulado: Si $x \downarrow y \in D$, entonces para cada congruencia Θ se tendrá $x \Theta z$.

Conviene notar que son álgebras cuasi transitivas todos los modelos que se examinaron para el sistema lógico $\mathcal{A}_p$ en el cap. 5 y en el cap. 6.

Terminología suplementaria

Antes de proseguir hay que introducir —o recordar— cierta terminología.

Una **función** —según lo que se estudió en el cap. 6— es algo, ϕ , tal que, dado un **argumento**, α —siendo α un miembro **cualquiera** del dominio o campo de argumentos de ϕ —, lo envía sobre un único ente, que es el **valor** o imagen de α por ϕ , $\phi(\alpha)$ (también escrito así: $\phi\alpha$, cuando no se dan dudas sobre el alcance de esa ocurrencia de ' ϕ '). Defínese una función, ϕ , con su dominio o campo de argumentos, D , y con un conjunto (contradominio), C , en el cual está incluido su campo de valores; así: $\phi: D \rightarrow C$. Una función ϕ es **inyectiva** (una **inyección**) cuando para cualesquiera x, z , si $x \neq z$, entonces $\phi x \neq \phi z$. Es **sobreyectiva** sobre (o en, o con respecto a) su contradominio si éste coincide con su campo de valores, o sea:

la función $\phi: D \rightarrow C$ es **sobreyectiva** (una **sobreyección**) con respecto a C si cada $x \in C$ es tal que hay algún $z \in D$ tal que $\phi z = x$. Es **biyectiva** (una **biyección**) la función $\phi: D \rightarrow C$ si es inyectiva y sobreyectiva. Una biyección puede escribirse así: $\phi D \leftrightarrow C$, puesto que una biyección hace corresponder a cada miembro del dominio un solo valor y a cada miembro del contradominio un solo argumento.

(Nótese que una operación unaria es una función cuyo campo de valores está incluido en el dominio de la función —su campo de argumentos—; o sea: es una función cuyo dominio está cerrado con respecto a ella.)

Llámbase **morfismo** (o también **homomorfismo**) de un álgebra $A = \langle \wp A, \Gamma \rangle$ sobre otra $B = \langle \wp B, \Delta \rangle$ (si bien en lo sucesivo, cuando no haya confusión, podemos denominar igual a un álgebra y a su portador) una función ϕ tal que para cada operación m -aria $\star \in \Gamma$ y cualesquiera $x^1, \dots, x^m \in \wp A$, hay una operación m -aria, $\star, \in \Delta$, tal que se cumple esta ecuación: $\phi(\star(x^1, \dots, x^m)) = \star(\phi x^1, \dots, \phi x^m)$. Si, además, de ser un morfismo, ϕ es una inyección, ϕ es una **inmersión** o un **monomorfismo**; si es una sobreyección, será un **epimorfismo**; si es una biyección, un **isomorfismo**. Si el campo de argumentos es idéntico al campo de valores, un **endomorfismo**. Un automorfismo es un isomorfismo endomórfico.

Llámbase **imagen** epimórfica, —o respectivamente copia isomórfica— de un álgebra A por un epimorfismo —o respectivamente isomorfismo— ϕ al álgebra cuyo portador sea el campo de valores de ϕ .

Sean A, B , dos conjuntos o cúmulos. El **producto cartesiano** de ambos $A \times B$, es un cúmulo de pares ordenados, cada uno de cuyos miembros es un par ordenado, $\langle a, b \rangle$, tal que $a \in A$ y $b \in B$. Similarmente el producto $\prod_{s \in S} Z_s$ de una familia indizada de cúmulos Z_s ($s \in S$) es un cúmulo cada uno de cuyos miembros es una secuencia ordenada cuyo i° componente pertenece al i° miembro de Z_s .

(Si S es el cúmulo de los reyes magos, puede indizarse por él a la familia de cúmulos $\{A, B, C\} = D$, siendo A el cúmulo de los vegetales, B el de los animales invertebrados, C el de los vertebrados, siendo, p.ej., $A = D_{\text{Melchor}}$, $B = D_{\text{Gaspar}}$, $C = D_{\text{Baltasar}}$. El producto de esa familia así indizada será el cúmulo o conjunto de tríos ordenados cada uno de los cuales tenga como primer componente a un vegetal, como segundo a un animal invertebrado, y como tercero a un vertebrado.)

Llamamos **proyección** a una función p cuyo campo de argumentos sea un producto de varios cúmulos siempre y cuando haya cierto número natural i tal que para cada argumento dado, x , $p x =$ el i° componente de x . Así representaremos a esa función como p_i .

Productos directos

Llámbase **especie** de álgebras a un cúmulo de álgebras con las mismas operaciones; o sea, si $\langle A, \Gamma \rangle$ y $\langle B, \Delta \rangle$ son álgebras de la misma especie, entonces $\Gamma = \Delta$. En un sentido técnico llámbase **clase de álgebras** a una especie de álgebras caracterizada por ciertos postulados: $\langle A, \Gamma \rangle$ y $\langle B, \Gamma \rangle$ pertenecen a la misma **clase** ssi hay un cúmulo de postulados definitorios de dicha clase que se cumplen tanto en A como en B .

Sean $A_1, \dots, A_n$ álgebras de una misma clase tales que, para cada i tal que $1 \leq i \leq n$, $A_i = \langle C_i, \Gamma \rangle$, donde C_i es un cúmulo de elementos y $\Gamma = \{\gamma_1, \dots, \gamma_m\}$. Entonces el **producto directo** de las mismas es un álgebra A con las mismas operaciones cuyo portador es el producto cartesiano $C_1 \times C_2 \times \dots \times C_n$; definimos, para cada z perteneciente a ese producto cartesiano, **funciones de proyección** $p_1, \dots, p_n$, donde, para cada i tal que $1 \leq i \leq n$, $p_i(z)$ es el i° componente de z , o sea: es aquel miembro de C_i que ocupa el i° lugar entre los componentes de z (recuérdese que z es un n -tuplo ordenado, o sea: una secuencia de n componentes);

para el álgebra A se definen las operaciones pertenecientes a Γ como sigue: para cada j tal que $1 \leq j \leq m$, γ_j es una operación r -aria, para cierto r tal que $0 \leq r$; sean $x_1, \dots, x_r$ miembros del producto cartesiano de los diversos cúmulos C_i (para $1 \leq i \leq n$); entonces para cada i tal que $1 \leq i \leq n$, $p_i \gamma_j(x_1, \dots, x_r) = \gamma_j(p_i x_1, \dots, p_i x_r)$. Queda así definido cuál es el elemento $\gamma_j(x_1, \dots, x_r)$. El problema que se plantea es si el producto directo de álgebras de cierta clase es también un álgebra de esa misma clase. No siempre. Pero (en virtud de un teorema de álgebra universal que no vamos a demostrar aquí (vide el libro de Birkhoff que es el segundo ítem en la Secc. 4ª de la bibliografía del presente trabajo —en adelante abreviado aquí como [Birkhoff]), p. 149), si los postulados que rigen una clase de álgebras son ecuaciones y/o entrañamientos ecuacionales (de la forma: si $p=q$, entonces $r=s$, para ciertos p, q, r, s), entonces sí se cumple esa condición: esa clase de álgebras está cerrada con respecto a la formación de productos directos. Como los postulados que determinan la clase de las aa.cc.tt. son todos de esa índole, resulta que es un a.c.t. todo producto directo de aa.cc.tt. cualesquiera (en número finito o infinito, sean o no idénticas varias, o aun todas, las álgebras que se toman para formar un producto directo dado).

Llámase **escalar** un a.c.t. en la que la relación de orden $\leq$ es conexa, o sea. para cualesquiera x, z , tiénese que: o bien $x \leq z$ o bien $z \leq x$. Eso —demostrablemente— equivale a que, para cualesquiera x, z o bien $x \rightarrow z$ sea denso o bien lo sea $z \rightarrow x$. Pero el producto directo de dos aa.cc.tt. cualesquiera, escalares o no, es un a.c.t. no escalar. En un a.c.t. no escalar hay elementos no densos diferentes de 0 (porque, si es un a.c.t. no escalar, hay en ella dos elementos, x, z , tales que $x \not\leq z, z \not\leq x$; entonces $x \rightarrow z \notin D$, pero $x \rightarrow z \neq 0$, pues, si $x \rightarrow z = 0$, resultará que, como $x \rightarrow z \vee (z \rightarrow x)$ es un elemento denso, $z \rightarrow x$ será un elemento denso y, por ende, $z \wedge x = z$, o sea: $z \leq x$, contra la hipótesis). Sean A_1 y A_2 dos aa.cc.tt. escalares cualesquiera con sendos elementos-cero 0_1 y 0_2 , y sea A su producto directo. Los elementos no densos de A serán todos aquellos elementos x tales que o bien $p_1(x) = 0_1$ o bien $p_2(x) = 0_2$.

El cálculo sentencial Ap y las álgebras libres asociadas con la clase de las aa.cc.tt.

Tomemos el cálculo sentencial Aj y suprimamos: de entre los símbolos primitivos 'B'; de entre los esquemas axiomáticos, a $A18, A19$ y $A20$; y de entre sus reglas de inferencia primitivas a $\text{rinf}2$. El resultado es el cálculo cuasitransitivo Ap . (Nótese que **este** sistema Ap , así sintácticamente definido, no es idéntico

al sistema de la misma denominación presentado en el Capítulo 5, el cual está semánticamente definido y constituye una extensión no conservativa del ahora considerado.)

Se demuestra sencillamente —pero, eso sí, armándose de paciencia— que una fórmula de Ap es un teorema ssi es válida, entendiendo por **fórmula válida de Ap** una fórmula p tal que cualquier valuación admisible v que tenga como campo de argumentos al cúmulo de fbfs de Ap y cuyo campo de valores esté incluido en el portador de un a.c.t. es tal que $v(p)$ es un elemento denso. (Así queda probado que Ap es un sistema a la vez robusto y completo.) Las condiciones que estipulamos para que sea admisible una valuación v así son que, para cualesquiera fbfs p, q , de Ap :

$$\begin{aligned} v(p \downarrow q) &= Nv(p) \wedge Nv(q); & v(a) &= m0; & v(p \downarrow q) &= v(p) \downarrow v(q); \\ v(p \bullet q) &= v(p) \bullet v(q); & v(Hp) &= Hv(p) \end{aligned}$$

(De nuevo tenemos que, en cada una de esas ecuaciones, los signos que figuran en el miembro izquierdo son símbolos del sistema de lógica —en este caso Ap ; los que figuran en el miembro derecho, aunque puedan ser equigráficamente representados con respecto a los anteriores, son, empero, operaciones del álgebra en cuestión —en este caso el a.c.t. en cuyo portador esté incluido el campo de valores de la valuación v .)

Vamos ahora a definir las álgebras libres asociadas con una determinada clase de álgebras, para ver luego cómo se vincula la existencia de tales álgebras y el que las mismas sean aa.cc.tt. con la completez de $\mathcal{A}_p$.

En primer lugar, definimos lo que es un **álgebra libre vocabular** con r generadores. Para ello definimos una especie de álgebras como un cúmulo de álgebras para las cuales se han definido las mismas operaciones. (Una clase de álgebras es, pues, un subconjunto propio de una especie de álgebras, subconjunto caracterizado por el hecho de que todas las álgebras de una clase tienen en común el que valgan en ellas ciertos postulados.) Sea Γ un cúmulo de operaciones algebraicas. (Llamamos a la especie de álgebras con tales operaciones: la especie Γ de álgebras.) Entonces definimos, para cualquier número cardinal r , finito o infinito, el álgebra vocabular libre con r generadores, $W_r\Gamma$, como sigue. Un **polinomio** es una expresión cualquiera bien formada que consta tan sólo de constantes y de símbolos que denotan a operaciones algebraicas. (Con mayor rigor se define así un polinomio: cada letra —o constante— es un polinomio; y, si γ es una operación n -aria y ' p_1 ', ..., ' p_n ', son polinomios, también es un polinomio la expresión ' $\gamma(p_1, \dots, p_n)$ '.) En $W_r\Gamma$ cada una de las r letras ' x_1 ', ..., ' x_r ' es un polinomio de rango 0; el cúmulo de esos r polinomios de rango 0 es el **alfabeto** de $W_r\Gamma$. Para cada entero positivo s se define un polinomio- Γ de rango s recursivamente como una expresión ("vocablo") de la forma $\gamma(u_1, \dots, u_n)$, donde n es (el número de) la aridad de γ (γ es una operación n -aria) y al menos un u_j ($1 \leq j \leq n$) = $p_j(x_1, \dots, x_r)$ es un polinomio- Γ de rango $s-1$, mientras que cada u_i ($1 \leq i \leq n$) es un polinomio de rango $\leq s-1$.) La igualdad en $W_r\Gamma$ se define como identidad gráfica estricta (o sea: si $p=q$, siendo p y q polinomios, es que son el mismo polinomio). Obviamente, $W_r\Gamma$ usualmente no pertenecerá a ninguna de las clases de álgebras de la especie (que tiene en común las operaciones de) Γ usualmente definidas, pues en cada una de tales clases se cumplen ciertas ecuaciones. Pero ahora veremos el interés de estas álgebras vocabulares.

Un importante teorema de álgebra universal muestra lo siguiente: cualquier inyección $\delta: X \rightarrow A$ del alfabeto X a un álgebra de la especie Γ es tal que existe una extensión unívocamente determinada de δ que es un morfismo de $W_r\Gamma$ a A . (Una **extensión** de una función f es otra función g tal que: el campo de argumentos de f está incluido en el campo de argumentos de g y para todo argumento x de f , $f(x) = g(x)$.) La prueba es por inducción matemática con respecto al rango: la función buscada, f , enviará cada polinomio p de $W_r\Gamma$ de rango s sobre un único elemento $q=f(p) \in A$ tal que: si $p = \gamma(u_1, \dots, u_n)$ (siendo n la aridad de γ), entonces $q = \gamma(f(u_1), \dots, f(u_n))$. Por definición, tal función es un morfismo.

Antes de exponer un importante corolario de tal teorema, es menester recordar un punto terminológico que figura en las introducciones a la teoría ingenua de conjuntos, a saber: una relación de **equivalencia** es una relación reflexiva, simétrica y transitiva. Lo cual significa que, si la relación R está definida sobre un cúmulo o conjunto X , entonces, para cualesquiera $x, z, u \in X$, se tendrá: (1º) xRx ; (2º) xRz sólo si también zRx ; (3º) sucede a la vez que xRz y que zRu sólo si también sucede que xRu . Si R es una relación de equivalencia definida sobre un cúmulo B , entonces la **partición** de B por R (expresada como B/R) —también llamada la **cocientización** de B por R — es el cúmulo de **clases de equivalencia** de miembros de B , o sea es una familia $\mathfrak{S}$ de cúmulos o conjuntos incluidos en B tal que, si $C \in \mathfrak{S}$ y $x \in C$, entonces también vendrá abarcado por C cualquier elemento $z \in B$ tal que xRz . (A las clases de equivalencia abarcadas por la partición de B por R suele denotárselas así: $|x|_R$ para cada $x \in B$.)

Sentado eso, procedamos a presentar el anunciado corolario, a saber: que si un álgebra $\langle A, \Gamma \rangle$ tiene r generadores, entonces hay una congruencia Θ de $W_r\Gamma$ tal que A es isomórfica con la partición $W_r\Gamma/\Theta$; de lo cual se sigue que, en tal caso, A es una imagen epimórfica de $W_r\Gamma$ (o sea: hay un epimorfismo de $W_r\Gamma$ sobre A). (Llámanse **generadores** en un álgebra $\langle A, \Gamma \rangle$

a unos elementos $x_1, \dots, x_r$ tales que todos los elementos de A son: $x_1, \dots, x_r$ y aquellos que son denotables por polinomios en los que sólo aparezcan ' x_1 ', ' $\dots$ ', ' x_r ', o sea: los elementos de A serán, además de esos r generadores, aquellos otros que resulten de aplicarles, reiteradas veces, las operaciones de Γ .)

Esa característica, que tienen las álgebras vocabulares, es compartida por otras álgebras. Un álgebra A con la característica indicada —a saber que cualquier función del cúmulo de los generadores de A al portador de un álgebra **cualquiera** B de cierta clase de álgebras, C , puede ser extendida de un modo perfectamente determinado a un morfismo de A a B — es un álgebra libre con respecto a la clase C de álgebras. (La idea subyacente es que esa álgebra A es libre con respecto a tal clase de álgebras en el sentido de que “libremente” se puede pasar de A a cualquier álgebra de la clase en cuestión, o sea: que, comoquiera que se empiece a determinar una función de A a una de tales álgebras —empezando, claro, por tomar como argumentos los generadores de A —, tal función podrá extenderse —y, lo que es más, de un modo perfectamente determinado— hasta que resulte un morfismo de A a esa otra álgebra.) Ahora vamos a definir esa noción rigurosamente: sea C una clase de álgebras $A_c = \langle S_c, \Gamma \rangle$. Un álgebra A es llamada **álgebra libre con r generadores asociada con la clase C** de álgebras ssi: 1º) las operaciones algebraicas definidas sobre A son las pertenecientes a Γ ; 2º) hay un subconjunto X de A que abarca a r miembros y que es un cúmulo de generadores de A ; y 3º) para toda función $f: X \rightarrow S_c$ hay una extensión unívocamente determinada de f que es un morfismo $\phi: A \rightarrow A_c$ (siendo A_c un álgebra cualquiera de la clase C , con tal de que S_c sea el portador de A_c). Un álgebra A que sea un álgebra libre con r generadores asociada con la clase C de álgebras y cuyo cúmulo de r generadores sea X será llamada **libremente engendrada** (por X) ssi $A \in C$.

Pruébese —es uno de los más importantes teoremas del álgebra universal— lo siguiente: sea G una clase de álgebras; entonces son isomórficas entre sí cualesquiera dos álgebras libres con el mismo número de generadores asociadas a G .

Prueba: puesto que el cúmulo X aludido es un conjunto de generadores, está claro que la extensión está perfectamente determinada (si $x_1, \dots, x_s$ son miembros de X , y γ es una operación s -aria, $f_\gamma(x_1, \dots, x_s) = \gamma(fx_1, \dots, fx_s)$; y así sucesivamente para cualesquiera polinomios de peso > 0). Si A_1 y A_2 son engendradas por sendos subconjuntos X_1 de A_1 y X_2 de A_2 , con tal de que X_1 y X_2 sean de la misma cardinalidad (tengan el mismo número de miembros), entonces cualquier biyección $g: X_1 \leftrightarrow X_2$ puede (de un modo unívocamente determinado) ser extendida a un isomorfismo $h: A_1 \leftrightarrow A_2$.

¿Cómo descubrir a un álgebra libre con r generadores asociada con cierta clase dada de álgebras? Hay dos procedimientos. He aquí el primero: sea C de nuevo una clase de álgebras $A_c = \langle S_c, \Gamma \rangle$ y sea r un número cardinal finito tal que $X = \{x_1, \dots, x_r\}$. Sea Δ el cúmulo de funciones $\delta: X \rightarrow S_c$ (para uno cualquiera de los portadores S_c) tales que, para cualquier polinomio $p(x_1, \dots, x_r) \in W_r\Gamma$, se cumple la siguiente ecuación: $\delta(p) = p(\delta x_1, \dots, \delta x_r)$. Sea Θ una congruencia definida así para cualesquiera polinomios- Γ : $p \Theta q$ significa que $\delta(p) = \delta(q)$ para todo $\delta \in \Delta$. Θ es, en efecto, una congruencia en $W_r\Gamma$. Así que $W_r\Gamma/\Theta$ (la partición de $W_r\Gamma$ por Θ) es un álgebra libre con r generadores asociada con C ; se la denota así: F_rC . El álgebra libre así descubierta puede que no sea un álgebra libremente engendrada, e.d. puede que no pertenezca a la clase C de álgebras.

Ahora bien, por el propio modo de determinación de F_rC que hemos seguido, resulta fácilmente demostrable que F_rC es (representable como —o sea: isomórfica con—) una subálgebra de productos directos de copias de las diversas álgebras A_c . ¡Expliquemos esto! Una subálgebra de un álgebra $\langle B, \Delta \rangle$ es un subconjunto de B cerrado con respecto a cada operación algebraica perteneciente a Δ . Ya sabemos qué es un producto directo. Y una copia

es una imagen isomórfica, o sea es un álgebra isomórfica con aquella álgebra de la que es copia. (Algebraicamente podemos “identificar” a dos álgebras si son isomórficas.)

Despréndese de ahí lo siguiente: si C es una clase de álgebras cerrada con respecto a la formación de subálgebras y de productos directos (o sea: si cada subálgebra de un álgebra $A \in C$ es un álgebra $A \in C$ y si cada producto directo de $A_1 \dots A_k$ es otra álgebra $A \in C$), entonces $F_r C \in C$, y, por lo tanto, $F_r C$ es libremente engendrada por el cúmulo de los r generadores. Como, según otro resultado, ya citado más arriba, de Birkhoff (a quien debe mucho toda esta exposición —vide [Birkhoff], p. 149), las clases de álgebras cuyos postulados son ecuaciones y/o entrañamientos ecuacionales están cerradas respecto a la formación de subálgebras y de productos directos, y como los postulados de las aa.cc.tt. son todos de esa índole, se tiene el resultado siguiente: el álgebra libre con r generadores asociada con la clase de las aa.cc.tt. es un a.c.t.

Otro modo de encontrar un álgebra libre con r generadores asociada con una clase dada de álgebras es el siguiente (vide teorema VI/20 de Birkhoff, [Birkhoff], p. 151): sea, de nuevo, C una clase de álgebras $A \in C = \langle S, \Gamma \rangle$. Sea Δ un cúmulo de ecuaciones que involucren sólo operaciones de Γ y variables; y sea Δ^* la clase de aquellas álgebras de la especie Γ en cada una de las cuales se cumplen todas las ecuaciones de Δ (a una clase de álgebras así se la denomina: familia de álgebras — Δ^* es, pues, una familia de álgebras). Supongamos, además, que por aplicación reiterada de tales ecuaciones pertenecientes a Δ pueden reducirse todos los elementos (todos los polinomios) de $W_r C$ a polinomios del subconjunto B de $W_r C$. Sea $A \in \Delta^*$ un álgebra con r generadores $x_1, \dots, x_r$ en la cual, para cualesquiera p_j y $p_k \in B$, si $p_j(x_1, \dots, x_r) = p_k(x_1, \dots, x_r)$, entonces $p_j = p_k$. (Un álgebra A que satisfaga esas condiciones será denominada: **álgebra polar ecuacionalmente irreducible**.) De ser verdad todo eso, A es un álgebra libre con r generadores de Δ^* .

Es interesante, a este respecto, constatar que todos los postulados, salvo uno (el (27)) de las aa.cc.tt. son ecuaciones, siendo el restante un entrañamiento ecuacional. Tomemos la familia de todas aquellas álgebras con las mismas operaciones de las aa.cc.tt. (o sea: de la misma especie) en las que se cumplan todas las ecuaciones que se cumplen en cada una de las aa.cc.tt. Esa familia de álgebras, SQ (clase de las álgebras subcuasitransitivas) contendrá álgebras que no sean aa.cc.tt. (en esas álgebras valdrán ecuaciones adicionales que infringirán el postulado (27)). Pero —según otro teorema de Birkhoff, el VI/19, ibid. p. 151— cualquier álgebra de SQ con r generadores será una imagen epimórfica de $F_r Q$, donde Q es la clase de las álgebras cuasitransitivas. La búsqueda del álgebra libre con r generadores asociada con SQ puede hacerse explotando, precisamente, el teorema VI/20 de Birkhoff. Mas lo interesante es que, como se desprende del susodicho teorema, $F_r SQ$ es un a.c.t. (libremente engendrada por r generadores —libremente con respecto a la clase SQ); en efecto: el postulado (27) de las aa.cc.tt. no sólo no queda conculcado por las condiciones involucradas en el hecho de que un álgebra sea un álgebra polar ecuacionalmente irreducible, sino que se puede demostrar que, si dos polinomios p y q no son reducibles uno a otro —ni ambos a un tercero— en virtud de las ecuaciones polinomiales que se siguen de los postulados de las aa.cc.tt., y si $x_1, \dots, x_r$ son los generadores de un a.c.t. para los cuales no se cumplan otras ecuaciones que las que tienen que cumplirse en virtud de los postulados de la clase de aa.cc.tt., entonces se tendrá que $p(x_1, \dots, x_r) \neq q(x_1, \dots, x_r)$, ya que, si no, se tendría que $p(x_1, \dots, x_r) \mid q(x_1, \dots, x_r) = \frac{1}{2}$, sin que esta ecuación se deduzca de los postulados de las aa.cc.tt. Así pues, precisamente el que un álgebra A sea un a.c.t. nos asegura que, si para sus r generadores no se cumplen, en tal álgebra, otros postulados que aquellos cuya vigencia se deduce de los postulados de las aa.cc.tt., entonces es que A es un álgebra polar ecuacionalmente irreducible perteneciente a la mayor familia de álgebras incluida en la clase de las aa.cc.tt. (y esa familia no es otra

que la clase de las álgebras subcuasitransitivas). Pero un a.c.t. así, A , será —en virtud del también mencionado teorema VI/19 de Birkhoff— una imagen epimórfica de $F_r Q$. Entonces se demuestra que A es (isomórfica con) $F_r Q$ si se demuestra también que hay un epimorfismo de A sobre $F_r Q$, el cual es idéntico al epimorfismo inverso. Sea π éste último: a cada generador x_i de $F_r Q$ le hará corresponder el generador z_i de A , donde $1 \leq i \leq r$, siendo $x_1, \dots, x_r$ los generadores de $F_r Q$ y $z_1, \dots, z_r$ los de A . Luego, aplicando el principio de sustitución polinomial (para toda operación n -aria γ , $\pi\gamma(u_1, \dots, u_n) = \gamma(\pi u_1, \dots, \pi u_n)$), resulta evidente que π es un isomorfismo, ya que las únicas ecuaciones que se cumplen en A son las que se deducen de los postulados de las aa.cc.tt. y las únicas que lo hacen en $F_r Q$ son aquellas que valen en todas y cada una de las aa.cc.tt. —o sea: justamente cuantasquiera que se deduzcan de esos postulados, pero sólo ellas.

Corolario: un a.c.t. A con r generadores es libremente engendrada (por el cúmulo de esos r generadores) ssi las únicas ecuaciones polinomiales que valen en A son aquellas que se deducen de los postulados de las aa.cc.tt.

Para que se entienda lo que va a seguir es menester introducir o recordar algunas nociones adicionales. Recordemos que, si $\mathfrak{B}$ y $\mathfrak{B}'$ son teorías, entonces $\mathfrak{B}'$ es una **extensión** de $\mathfrak{B}$ ssi: la clase de símbolos de $\mathfrak{B}$ está incluida en la de $\mathfrak{B}'$; la clase de fbfs de $\mathfrak{B}$ está incluida en la de $\mathfrak{B}'$; y la clase de teoremas de $\mathfrak{B}$ está incluida en la de $\mathfrak{B}'$. Si, además, la clase de reglas de inferencia de $\mathfrak{B}$ está incluida en la de $\mathfrak{B}'$, se dice que $\mathfrak{B}'$ es una extensión **recia** de $\mathfrak{B}$. Sea $\mathfrak{B}$ una teoría que es un sistema de lógica. Entonces, si $\mathfrak{B}'$ es una extensión recia de $\mathfrak{B}$ y no tiene ninguna regla de inferencia primitiva sobreañadida a las de $\mathfrak{B}$, se dice que $\mathfrak{B}'$ es una $\mathfrak{B}$ -teoría.

Sea $\mathfrak{B}$ una teoría en la que existe un functor diádico Ξ tal que las siguientes reglas de inferencia pueden ser derivadas en $\mathfrak{B}$ (para cualesquiera fbfs $\ulcorner p \urcorner$, $\ulcorner q \urcorner$, $\ulcorner s \urcorner$ de $\mathfrak{B}$): $p \Xi q \vdash q \Xi p$; $p \Xi q, q \Xi s \vdash p \Xi s$; $p \Xi q \vdash p^1 \Xi q^1$ (donde $\ulcorner p^1 \urcorner$ difiere de $\ulcorner q^1 \urcorner$ sólo por el reemplazo de ciertas ocurrencias de $\ulcorner p \urcorner$ en $\ulcorner p^1 \urcorner$ por sendas ocurrencias de $\ulcorner q \urcorner$ en $\ulcorner q^1 \urcorner$). El **álgebra de Tarski** definida sobre una teoría así, $\mathfrak{B}$, es un álgebra cuyos elementos son las clases de equivalencia $|p|_\Theta$ donde $\ulcorner p \urcorner$ es una fbf de $\mathfrak{B}$ y Θ es una relación de equivalencia que se establece así: $|p| \Theta |q|$ ssi $\ulcorner p \Xi q \urcorner$ es un teorema de $\mathfrak{B}$. Las operaciones en $\xi(\mathfrak{B})$ —el álgebra de Tarski definida sobre $\mathfrak{B}$ — se establecen así: para cada functor n -ádico, γ , de $\mathfrak{B}$ hay una operación n -aria, γ , de $\xi(\mathfrak{B})$ tal que, si $\ulcorner p_1 \urcorner, \dots, \ulcorner p_n \urcorner$ son fbfs de $\mathfrak{B}$, $\gamma(|p_1|, \dots, |p_n|) = |\gamma(p_1, \dots, p_n)|$. (Salta a la vista que $\xi(\mathfrak{B}) = \mathfrak{B}/\Theta$.)

Réstanos tan sólo, para coronar este acápite, mostrar que un álgebra libre con r generadores asociada con Q es (isomórfica con) $\xi(\mathfrak{B}_r)$, siendo $\mathfrak{B}_r$ una Ap -teoría que contiene r constantes sentenciales adicionales, $\ulcorner p_1 \urcorner, \dots, \ulcorner p_r \urcorner$ (como símbolos primitivos adicionales respecto de los de Ap). Naturalmente, el signo Ξ de Ap es el functor ' $\ulcorner$ '. Suponemos que $\mathfrak{B}_r$ no tiene ningún axioma adicional respecto de los de Ap .

Dadas esas definicionales resulta manifiesto que $\mathfrak{B}_r$ es isomórfica con $W_r Q$, y que $\xi(\mathfrak{B}_r)$ es isomórfica con $F_r Q$. En efecto: hemos visto que $F_r Q$ es un a.c.t. en la que hay r generadores y únicamente se cumplen aquellas ecuaciones que se deducen de los postulados de las aa.cc.tt. Como esos postulados, en su conjunto, son “traducciones” de los esquemas axiomáticos de Ap tomados junto con rinfl_1 , y viceversa, el cúmulo de postulados algebraicos entraña sólo todas aquellas identidades que son traducciones de teoremas de Ap cuyo functor central es ' $\ulcorner$ '. (En un a.c.t., A , —siendo p, q dos polinomios— $p=q$ ssi $plq = \frac{1}{2}$; y esto sucede ssi aquellas dos fbfs $\ulcorner p^1 \urcorner, \ulcorner q^1 \urcorner$ de una Ap -teoría isomórfica con A , $\mathfrak{B}$, cuyas traducciones sean los polinomios p, q de A sean tales que $\ulcorner p^1 l q^1 l \frac{1}{2} \urcorner$ sea un teorema de $\mathfrak{B}$, lo cual sucede ssi $\ulcorner p^1 l q^1 \urcorner$ es también un teorema de $\mathfrak{B}$.) Por consiguiente, como $\xi(\mathfrak{B}_r)$ y $F_r Q$ son álgebras con el mismo número de generadores y se cumplen en ambas las mismas ecuaciones (sólo todas aquellas que se

deduzcan de los postulados de las aa.cc.tt., o lo que es lo mismo: sólo todas aquellas que deban valer en $\xi(\mathfrak{A}_r)$ por definición y en virtud de los teoremas equivalenciales de Ap , sólo podrían dejar de ser isomórficas una respecto a la otra si algún entañamiento ecuacional que se cumpliera en la una no lo hiciera en la otra. Pero —además de que se demuestra también fácilmente que tal cosa no puede suceder—, como F_rQ es un álgebra polar ecuacionalmente irreducible, puede ser caracterizada sin aludir a los entañamientos ecuacionales, como F_rQ^* , siendo Q^* la mayor familia de álgebras incluida en Q (o sea: Q^* es la clase de álgebras subcuasitransitivas). Pero $\xi(\mathfrak{A}_r)$ puede ser caracterizada similarmente con respecto a la correspondiente familia de álgebras de Tarski definidas sobre Ap -teorías, cada una de las cuales es isomórfica con un a.c.t.

Álgebras transitivas y transitivoides

Vamos ahora a pasar a la presentación de las **álgebras transitivas**. Un álgebra transitiva, a.t. para abreviar, es un álgebra $\langle A, \{1, B, N, H, n, \vee, \bullet, I\} \rangle$ donde $\langle A, \{1, N, H, n, \vee, \bullet, I\} \rangle$ es un a.c.t. y B es una operación unaria sobre A tal que, para todo $x \in A$: o bien $x \vee a = x = Bx$, o bien $\neg x \neq 0 = Bx$ (el segundo disyunto significa que, mientras que $\neg x \neq 0$, $0 = Bx$). En cada a.t., pues, cada elemento Bx es tal que: o bien Bx es denso, o bien $Bx = 0$.

Todo producto directo de aa.cc.tt. (álgebras cuasitransitivas) puédese transformar en un a.t. sin más que añadir la operación B como sigue:

Sea $A = \prod_{i \in I} (A_i)$ el producto directo de un cúmulo de aa.cc.tt. indizado por I .

B vendrá definida así: $Bx = x$ si para cada función de proyección p_j (con $j \in I$) se tiene que $p_j(x) \neq 0$ (donde 0 es el elemento nulo de A_j); y $Bx = \langle 0, 0, 0, \dots \rangle$ si hay alguna función de proyección p_j tal que $p_j(x) = 0$. (En aras de la simplicidad identificamos los elementos 0 de las diversas aa.cc.tt. A_i .)

Vamos a tomar a las aa.tt. (álgebras transitivas) como modelos de la lógica sentencial transitiva, e.d. del sistema A_j . Definimos una **valuación** (admisible) v de A_j como una función cuyo campo de argumentos es el cúmulo de fbfs de A_j y cuyo campo de valores está incluido en (el conjunto portador de) un a.t. con tal de que cumpla, para cualesquiera $\ulcorner p \urcorner$, $\ulcorner q \urcorner$ que sean fbfs de A_j , las mismas condiciones que deben cumplir las valuaciones admisibles de Ap y, además, esta condición adicional: $v(Bp) = Bv(p)$.

Nótese, de nuevo, que en esa ecuación tenemos que el signo ' B ', cuando figura en el miembro ecuacional derecho, denota a una operación en un a.t., mientras que, al figurar, dentro del paréntesis, en el miembro ecuacional izquierdo representa a un símbolo del sistema lógico A_j .

Ya es asunto de rutina y que se demuestra sin más que armarse de paciencia que una fbf $\ulcorner p \urcorner$ de A_j es un teorema ssi es tal que cada valuación (admisible) de A_j , v , es tal que $v(p)$ es denso (o sea $\neg v(p) = 0$, siendo aquí 0 el elemento-cero del a.t. en la que esté incluido el campo de valuaciones de v).

Es ocioso demostrar que hay aa.tt., pues habiendo probado que hay aa.cc.tt. y mostrado cómo todo producto de aa.cc.tt. puédese transformar en un a.t., resulta obvio que hay aa.tt. (En un a.c.t. escalar o lineal se define la operación B trivialmente así: para todo x , $Bx = x$.)

Viene ahora, sin embargo, un resultado interesante y que no deja de encerrar una dificultad. Y es que el álgebra libre con r (donde $2 \leq r$) generadores asociada con la clase de las aa.tt. no es un a.t. Definimos esa álgebra libre para la clase de las aa.tt. como lo hicimos para la clase de las aa.cc.tt. Y entonces reparamos en lo siguiente: en esa álgebra libre, F_rT (siendo T la clase de las álgebras transitivas), son verdaderas sólo aquellas ecuaciones que lo son en cada a.t. Mas es obvio que hay diversas aa.tt. en las que son verdaderas diferentes ecuaciones, pese a tener los mismos generadores (o, con mayor rigor: pese a que exista una

biyección entre los generadores de la una y los de la otra, en virtud de la cual biyección podemos considerar a esos dos cúmulos de generadores como (si fueran) idénticos): sean A_1, A_2 dos aa.tt. con los mismos generadores y tales que hay un elemento x tal que $\neg x=0$ en A_1 y $\neg x \neq 0$ en A_2 ; entonces en A_1 $0 \neq Bx = x \vee a$, pero en A_2 $x \neq Bx = 0$. Si ese elemento x es miembro de $F_r T$, entonces ninguna de esas ecuaciones puede ser verdadera en $F_r T$. Lo que hay que probar —para demostrar que $F_r T$ no es un a.t.— es que hay algún elemento de $F_r T$ tal que hay dos morfismos δ_1 y δ_2 de $F_r T$ respectivamente en dos aa.tt. A_1 y A_2 tales que $\delta_1(x)$ es denso mientras que $\delta_2(x)$ no lo es. Pues bien: sean A_1, A_2 dos aa.tt. tales que hay sendos elementos $x_1 \in A_1, x_2 \in A_2$, tales que $x_1 = Bx_1$ (o sea: $x_1 \in D_1$), pero $0 = Bx_2 \neq x_2$ (o sea: $Bx_2 \notin D_2$), siendo D_1, D_2 los cúmulos de elementos densos, respectivamente, de A_1 y de A_2 ; tomemos uno cualquiera de los r generadores, $z_1, \dots, z_r$, de $F_r T$. Sabemos que **cualquier** función de $\{z_1, \dots, z_r\}$ a $\wp A_1$ o a $\wp A_2$ puede extenderse, de manera unívoca, hasta que resulte un morfismo. Tomamos, pues, una función $f: \{z_1, \dots, z_r\} \rightarrow A_1$ tal que $f(z_1) = x_1$. Similarmente hacemos con respecto a A_2 , siendo f' también una función $\{z_1, \dots, z_r\} \rightarrow A_2$ tal que $f'(z_1) = x_2$. Sean δ_1, δ_2 los dos morfismos resultantes de sendas extensiones, respectivamente, de f y f' . Resulta claro que en $F_r T$ no podemos tener: ni $Bz_1 = z_1$, pues no tenemos $\delta_2(Bz_1) = \delta_2(z_1)$; ni tampoco $Bz_1 = 0$, pues no tenemos $\delta_1(Bz_1) = \delta_1(0)$ (es palmario que —a tenor de la propia definición de las funciones δ ofrecida unas tres páginas más atrás—, para cada morfismo δ , $\delta(Bz_1) = B\delta z_1$, y que $\delta 0 = 0$, identificando los elementos-cero de las diferentes aa.tt.) Pero de ahí se sigue que $F_r T$ no es un a.t.

Llamaremos **transitivoides** a álgebras como $F_r T$, o sea: álgebras en las que valen todas las ecuaciones y todos los entrañamientos ecuacionales que valen en todas las aa.tt.

Los postulados que valen para las aa.tt.-oides (álgebras transitivoides) son, además de los que valen para la clase de las aa.cc.tt., los siguientes: $Ba=a$; $B\frac{1}{2}=\frac{1}{2}$; $Bx \wedge Bx \vee (B \neg Bx) \in D$; $B(x \rightarrow z) \leq Bx \rightarrow Bz$. (Otro cúmulo de postulados demostrablemente equivalente al anterior —de cada uno de los dos conjuntos se deducen todos los miembros del otro— es éste: Si $x \in D$, entonces $x=Bx$; $\neg Bx = B \neg Bx$; $BLx = LBx$; $\neg Bx \vee (Bx \wedge x) \in D$; $B(x \rightarrow z) \leq Bx \rightarrow Bz$.)

La prueba de que esos postulados caracterizan a las aa.tt.-oides es como sigue. Primero se prueba fácilmente que en cada a.t. valen esos postulados. Luego hay que probar que, dado un elemento cualquiera $x \in F_r T$, si en cada a.t. en la que exista x son verdaderas respecto de x ciertas ecuaciones o ciertos entrañamientos ecuacionales, tales ecuaciones y/o entrañamientos ecuacionales se deducen a partir de ese cúmulo de postulados. Pues bien, x es denso en $F_r T$ ssi es denso en toda a.t. que contenga a x (y de nuestros postulados para las aa.tt.-oides se deduce que todo elemento denso x es tal que $Bx=x \vee a$); $x=0$ en $F_r T$ ssi $x=0$ en cada a.t. que contenga a x (y, por supuesto en cada a.t. $B0=0$; y se deduce también de nuestros postulados para las aa.tt.-oides que $B0=0$); pero $0 \neq x \neq x \vee a$ en $F_r T$, ssi hay dos aa.tt. diferentes tales que en la una $Bx=0$, en la otra $Bx=x \vee a$. Lo único común entre esas dos aa.tt., además del cumplimiento de cuantasquiera ecuaciones y entrañamientos ecuacionales que se cumplen en todas las aa.tt., será la disyunción entre ambas ecuaciones (y no sólo para x sino para infinidad de elementos, como $x \vee a$, etc., se tendrán diferentes ecuaciones en sendas aa.tt., valiendo tanto en la una como en la otra la disyunción de sendas ecuaciones en cada caso). Pero, demostrablemente, en un a.t. vale una disyunción de ecuaciones ssi es densa la unión entre los equivalenciales estrictos que ligan, respectivamente, a sendos miembros de las ecuaciones en cuestión (o sea: vale una disyunción de ecuaciones $x=z$ o $y=u$ ssi es denso el elemento $x \leftrightarrow z \vee (y \leftrightarrow u)$ donde $\ulcorner v \leftrightarrow v \urcorner$ abr $\ulcorner B(v \vee v) \urcorner$). Pues bien, de los postulados que hemos brindado para las aa.tt.-oides se deduce este teorema para todo x miembro de una de tales álgebras: $x \vee a \leftrightarrow x \vee (Bx \leftrightarrow 0) \in D$. Q.E.D.

Un a.t.-oide particularmente importante es el álgebra de Tarski definida sobre una Aj -teoría, $\mathfrak{A}_r$, con r constantes sentenciales adicionales (respecto de los símbolos de Aj), pero sin ningún axioma adicional. Tal álgebra de Tarski, $\xi(\mathfrak{A}_r)$ es isomórfica con F_rT (la prueba es similar a la brindada al final del Acápito anterior respecto a las aa.cc.tt. y a una Aj -teoría con la indicada característica). Además, como se deduce del resultado que acabamos de probar, $F_rT(\approx \xi(\mathfrak{A}_r))$, donde ' $\approx$ ' expresa isomorfismo) es también el álgebra libre con r generadores asociada con la clase de las aa.tt.-oides. De ahí se deduce lo siguiente: podemos definir la validez para Aj si cambiamos, en la definición de validez arriba propuesta, uniformemente, 'álgebra transitiva' por 'álgebra transitivoide': son teoremas de Aj sólo todas las fbfs de Aj a las que cada valuación (admisible) de Aj en un a.t.-oide haga corresponder un elemento denso de esa a.t.-oide. Si hemos modelizado Aj a través de la clase de las aa.tt. en lugar de la clase de las aa.tt.-oides es porque las primeras ofrecen la ventaja de que resulta más fácil probar ciertos resultados con relación a ellas. Y también por un segundo motivo: toda Aj -teoría completa (e.d. toda Aj -teoría $\mathfrak{A}$ tal que para toda fbf de $\mathfrak{A}$, ' p ', o bien ' p ' es un teorema de $\mathfrak{A}$ o bien ' $\neg Bp$ ' es un teorema de $\mathfrak{A}$) tiene como modelo, a un a.t. Vamos ahora a probar ese importante teorema. (Definimos primero el functor ' B ' así: ' Bp ' abr ' $\neg Bp$ '.) Probamos, en primer lugar, este **Lema 1º**: una Aj -teoría $\mathfrak{A}$ es delicuescente ssi contiene tres teoremas de las formas respectivas, ' p ', ' q ', ' $B(p \wedge q)$ '.

Prueba: Si ' p ', ' q ', ' $B(p \wedge q)$ ' son teoremas, entonces por la regla de afirmabilidad, rinf2, y por la regla de adjunción, deducimos: ' $B(p \wedge q) \wedge B(p \wedge q)$ ', lo cual es una **supercontradicción** de la que se deduce ' $\perp$ ', para cualquier fbf ' r '; o sea: la teoría resulta delicuescente si contiene a la vez ' p ', ' q ', ' $B(p \wedge q)$ '. Por otro lado supongamos que una Aj -teoría $\mathfrak{A}$ es delicuescente. Entonces en $\mathfrak{A}$ se demuestra cualquier fbf como teorema. Por consiguiente, para cualesquiera fbfs ' p ', ' q ', de $\mathfrak{A}$, se tendrá que ' p ', ' q ', ' $B(p \wedge q)$ ' serán teoremas de $\mathfrak{A}$. Q.E.D.

Lema 2º: Toda Aj -teoría $\mathfrak{A}$ no delicuescente tiene una extensión recia **completa** no delicuescente (siendo completa una teoría ssi para cada fbf ' p ', o bien ' p ' es un teorema o bien ' Bp ' es un teorema).

Prueba: sea $\mathfrak{A}$ una Aj -teoría. Es patente que hay un "buen orden" alfabético de las fbfs de $\mathfrak{A}$ (un buen orden es un orden total isomórfico con el de los números naturales: o sea tal que hay un primer elemento, un segundo, etc.). Ordénense las fbfs de $\mathfrak{A}$ según ese buen orden. Formamos entonces la siguiente cadena de teorías: $\mathfrak{A}_0 = \mathfrak{A}$. Y se pasa de $\mathfrak{A}_i$ a $\mathfrak{A}_{i+1}$ como sigue: 1) si la i^a fbf de $\mathfrak{A}$ es una fórmula ' p ' tal que hay dos teoremas de $\mathfrak{A}_i$, ' q ' y ' q ', tales que $\{p, q, q\} = \{s, r, B(s \wedge r)\}$ para ciertas fbfs ' s ', ' r ' de $\mathfrak{A}$, entonces $\mathfrak{A}_{i+1} = \mathfrak{A}_i$; 2) si la i^a fbf de $\mathfrak{A}$ es un teorema de $\mathfrak{A}_i$, también entonces $\mathfrak{A}_{i+1} = \mathfrak{A}_i$; 3) si no se da ninguno de esos dos casos, entonces $\mathfrak{A}_{i+1}$ es el resultado de añadir a los axiomas de $\mathfrak{A}_i$ la i^a fbf de $\mathfrak{A}$. En cualquier caso es evidente que, si $\mathfrak{A}_i$ es no delicuescente, también es no delicuescente $\mathfrak{A}_{i+1}$, pues, por el procedimiento de construcción, que garantiza que no haya en $\mathfrak{A}_{i+1}$ ningún trío de teoremas $\{s, r, B(s \wedge r)\}$ —si no lo había en $\mathfrak{A}_i$ — y ya vimos, por el Lema 1º, que una Aj -teoría es no delicuescente ssi no tiene ningún trío de teoremas así. Sea $\mathfrak{A}_\omega$ la unión de todas las teorías $\mathfrak{A}_i$ de esa cadena. $\mathfrak{A}_\omega$ es una teoría completa no delicuescente (si la teoría inicial $\mathfrak{A}$ era no delicuescente); es completa: en efecto: si hubiera una fbf ' p ' tal que ni ' p ' ni ' Bp ' fueran teoremas de $\mathfrak{A}_\omega$, entonces —y en virtud del Lema 1º— habría habido una teoría $\mathfrak{A}_i$ de la cadena que no contuviera como teoremas ni a ' p ' ni a ' Bp ' y tal que el añadido de ' p ' a los axiomas de $\mathfrak{A}_i$ habría dado por resultado un trío de teoremas $\{r, s, B(r \wedge s)\}$. Entonces sucedería lo siguiente: el trío $\{r, s, B(r \wedge s)\}$ sería tal que una de esas tres fbfs sería ' p '. Si ' p ' = ' r ', entonces serían teoremas de $\mathfrak{A}_i$ las fbfs ' s ' y ' $B(r \wedge s)$ ', por lo cual —y como se prueba fácilmente en Aj — ' Bp ' (o sea: por hipótesis, ' Bp ') sería también un teorema de $\mathfrak{A}_i$, contrariamente a lo supuesto. Idénticamente se prueba que ' $p \vdash s$ '. Si ' p ' = ' $B(r \wedge s)$ ', entonces ' r ' y ' s ' son teoremas

de $\mathfrak{Z}_i$, de donde resultaría —en virtud de A_j — que $\ulcorner \mathbb{B}\mathbb{B}(r \wedge s) \urcorner = \ulcorner \mathbb{J}\mathbb{B}(r \wedge s) \urcorner$ es también un teorema de $\mathfrak{Z}_i$, contrariamente a lo supuesto. Por consiguiente, no puede darse la situación indicada. Con eso se prueba que $\mathfrak{Z}_\omega$ es completa. Y es no delicuescente: si fuera delicuescente contendría un trío de teoremas $\{p, q, \mathbb{B}(p \wedge q)\}$ —en virtud del Lema 1º. Ahora bien, ninguna $\mathfrak{Z}_i$, para i finito, es delicuescente ni, por lo tanto, contiene a ese trío de teoremas. Pero si ese presunto trío de teoremas no está contenido en ninguna de las teorías $\mathfrak{Z}_i$ para i finito no puede tampoco estar contenido en $\mathfrak{Z}_\omega$, pues esas teorías están ordenadas por inclusión, de modo que si, p.ej., $\ulcorner p \urcorner$ es un teorema de $\mathfrak{Z}_j$ y $\ulcorner q \urcorner$ es un teorema de $\mathfrak{Z}_k$, entonces o bien $j < k$ o bien $k < j$, y, en el primer caso, $\mathfrak{Z}_k$ contiene tanto al teorema $\ulcorner p \urcorner$ como al teorema $\ulcorner q \urcorner$, mientras que en el segundo caso es $\mathfrak{Z}_j$ la que contiene tanto al teorema $\ulcorner p \urcorner$ como al teorema $\ulcorner q \urcorner$. (Y similarmente se razona para tres o más fbfs, en lugar de sólo dos). Así pues, $\mathfrak{Z}_\omega$ es no delicuescente.

Corolario del Lema 2º.— Si de un cúmulo Γ de fbfs no se infiere una fórmula $\ulcorner q \urcorner$ según reglas de inferencia de A_j , entonces $\Gamma \cup \{\mathbb{B}q\}$ es un cúmulo no delicuescente de fbfs (o sea: un conjunto que no incluye ningún trío $\{r, s, \mathbb{B}(r \wedge s)\}$).

Prueba: basta con ordenar “alfabéticamente” las fbfs de tal modo que $\ulcorner \mathbb{B}q \urcorner$ venga antes de $\ulcorner q \urcorner$, y antes que ninguna otra fbf que no esté en Γ . Formamos la teoría $\mathfrak{Z}_0$ cuyos axiomas son los de A_j más los miembros de Γ . Entonces, por el procedimiento indicado para construir el siguiente eslabón de una cadena de teorías que sean extensiones de $\mathfrak{Z}_0$ y cuyo límite sea una teoría completa, la primera teoría que formaremos, $\mathfrak{Z}_1$, incluirá a $\ulcorner \mathbb{B}q \urcorner$ como un axioma y será no delicuescente.

Teorema: Toda A_j -teoría no delicuescente tiene una extensión recia $\mathfrak{Z}$ tal que $\xi(\mathfrak{Z})$ (el álgebra de Tarski definida sobre $\mathfrak{Z}$) es un a.t.

Prueba: Por el Lema 2º sabemos que toda A_j -teoría tiene una extensión recia no delicuescente y completa $\mathfrak{Z}$. Sea $\xi(\mathfrak{Z})$ el álgebra de Tarski definida sobre $\mathfrak{Z}$. Para cada fbf $\ulcorner p \urcorner$ de $\mathfrak{Z}$ se tendrá que o bien $\ulcorner p \urcorner$ es un teorema de $\mathfrak{Z}$, y entonces $\ulcorner \mathbb{I}\mathbb{B}p \urcorner$ es un teorema de $\mathfrak{Z}$, con lo cual se tendrá que $|a| \vee |p| = |p| = \mathbb{B}|p|$; o bien $\mathbb{B}p$ es un teorema de $\mathfrak{Z}$, con lo cual $\ulcorner \mathbb{B}p \urcorner$ es un teorema de $\mathfrak{Z}$, de donde resulta que $\mathbb{B}|p| = |0|$. Q.E.D.

Definición: un cúmulo es de cardinalidad infinita numerable ssi existe una biyección entre ese cúmulo y el de los números naturales (o sea $\{0, 1, 2, 3, 4, \dots\}$). (Nota: según lo demostró Cantor, hay cúmulos de cardinalidad infinita no numerable. Uno de ellos es el cúmulo de los números reales el cual es isomórfico con el cúmulo de los subconjuntos del conjunto de los números naturales. Cantor probó que no hay biyección alguna entre el cúmulo de los números reales y el de los naturales —más concretamente, que no hay sobreyección alguna del cúmulo de los números reales al de los naturales.)

Corolario 1º.— Toda A_j -teoría $\mathfrak{Z}$ no delicuescente tiene, con respecto a alguna valuación, un modelo que es un a.t. de cardinalidad infinita numerable.

Prueba: a partir del Teorema. Sea $\mathfrak{Z}$ una A_j -teoría. Sea $\mathfrak{Z}'$ una extensión recia de $\mathfrak{Z}$ no delicuescente tal que $\xi(\mathfrak{Z}')$ es un a.t. Sea v un morfismo de $\mathfrak{Z}$ en $\xi(\mathfrak{Z}')$ tal que para toda fbf $\ulcorner q \urcorner$ de $\mathfrak{Z}$ $v(q) = |q|$. (Si $\mathfrak{Z}'$ es una extensión **simple** de $\mathfrak{Z}$, e.d. una extensión que no contiene otros símbolos primitivos que los de $\mathfrak{Z}$, entonces v es un epimorfismo.) Ahora bien, $\ulcorner p \urcorner$ es un teorema de $\mathfrak{Z}'$ ssi $|p| \in D$ en $\xi(\mathfrak{Z}')$; como cada teorema de $\mathfrak{Z}$ es un teorema de $\mathfrak{Z}'$, si $\ulcorner p \urcorner$ es un teorema de $\mathfrak{Z}$, $v(p)$ es un elemento denso de $\xi(\mathfrak{Z}')$. Luego $\xi(\mathfrak{Z}')$ es un modelo de $\mathfrak{Z}$, siendo $\xi(\mathfrak{Z}')$ —como lo dice el Teorema recién demostrado— un a.t. Además, la cardinalidad del cúmulo de elementos de $\xi(\mathfrak{Z}')$ es infinita numerable, ya que el cúmulo de fbfs demostrablemente no equivalentes entre sí de una A_j -teoría con un número finito, o infinito numerable, de símbolos primitivos es de cardinalidad infinita numerable.

Corolario 2º.— Toda Aj -teoría no delicuescente $\mathfrak{J}$ tiene un modelo. (Prueba: obvia, a partir del corolario anterior.)

Para formular el corolario siguiente usamos el signo ' $\vdash$ ' que significa lo siguiente: $\Gamma \vdash q$ quiere decir que para toda a.t. o a.t.-oide A y toda valuación v , si, para cada $r \in \Gamma$, $v(r)$ es un elemento denso de A , entonces es que $v(q)$ es también un elemento denso de A .

Corolario 3º.— Sea Γ un cúmulo delicuescente de fbfs de una Aj -teoría y sea ' $\ulcorner q \urcorner$ ' una fbf de esa misma Aj -teoría: si $\Gamma \vdash q$, entonces $\Gamma \vdash \ulcorner q \urcorner$ (o sea: entonces hay en Aj una regla de inferencia derivada que permite inferir ' $\ulcorner q \urcorner$ ' del cúmulo de fbfs Γ).

Prueba: Si $\Gamma \vdash q$, entonces en toda a.t. tal que haya una valuación (admisible) v tal que, para cada $r \in \Gamma$, $v(r) \in D$, se tendrá también que $v(q) \in D$. Por el corolario del Lema 2º, $\Gamma \vee \{\ulcorner Bq \urcorner\}$ será no delicuescente si es que ' $\ulcorner q \urcorner$ ' no se infiere de Γ según reglas de inferencia de Aj . Sea $\mathfrak{J}$ la teoría cuyos axiomas son los de Aj más los miembros de $\Gamma \vee \{\ulcorner Bq \urcorner\}$. Esa teoría es no delicuescente y tiene un modelo algebraico A , que es un a.t. (en virtud del Corolario 1º) con respecto a alguna valuación (admisible) v , tal que $v(\ulcorner Bq \urcorner)$ es denso. Pero ese modelo y esa valuación son, respectivamente, un a.t. A y una valuación v tales que para toda fbf ' $\ulcorner r \urcorner \in \Gamma$ ', $v(r)$ es denso. Por hipótesis, pues, $v(q)$ sería denso. Pero eso es absurdo, porque entonces serían densos a la vez $v(q)$ y $v(\ulcorner Bq \urcorner)$, o sea (por ser v una valuación admisible) $v(q)$ y $Bv(q)$ (donde la operación algebraica B se define igual que su equigráfico functor lógico). Y eso sólo puede suceder en un a.t. con un solo elemento $0=1$, situación excluida de la noción misma de modelo, por definición.

Corolario 4º.— Si $\vdash q$ (o sea: si ' $\ulcorner q \urcorner$ ' es una fbf válida, e.d. tal que cualquier valuación admisible v de una Aj -teoría que contenga la fbf ' $\ulcorner q \urcorner$ ' en un a.t. es tal que $v(q)$ es denso) entonces ' $\ulcorner q \urcorner$ ' es un teorema de Aj .

Prueba: obviamente a partir del corolario precedente, pues un teorema es una fbf que se infiere a partir (hasta) de la clase vacía de premisas.

Corolario 5º (principio de compacidad).— $\Gamma \vdash p$ ssi hay un subconjunto finito de Γ , C , tal que $C \vdash p$.

Prueba: corolario 3º más el hecho de que hay una prueba de una conclusión ' $\ulcorner p \urcorner$ ' a partir de un cúmulo Γ de premisas sólo si existe algún subconjunto finito de Γ , C , tal que $C \vdash p$.

Corolario 6º.— Sea $\mathfrak{J}$ una Aj -teoría cuyo cúmulo de axiomas no lógicos (o sea: adicionales respecto de los teoremas de Aj) es Γ . Entonces hay un a.t. que es, respecto de cierta valuación, un modelo de $\mathfrak{J}$ ssi, para cada Aj -teoría $\mathfrak{J}'$ cuyo cúmulo de axiomas no lógicos es algún subconjunto finito, C , de Γ , existen un a.t. A y una valuación v tales que A es, con respecto a v , un modelo de $\mathfrak{J}'$.

Prueba: $\mathfrak{J}$ tiene un modelo respecto de una cierta valuación ssi es no delicuescente. (Corolario 1º más el hecho obvio de que, si una Aj -teoría tiene un modelo, es que no es delicuescente, por la definición misma de modelo.) Pero, como lo muestra el Lema 1º, es no delicuescente una Aj -teoría con un cúmulo Γ de axiomas no lógicos ssi **ningún** subconjunto de Γ es de la forma $\{r, s, B(r \wedge s)\}$ o sea: ningún trío así está incluido en subconjunto finito alguno, C , de Γ ; para cada teoría $\mathfrak{J}'$ cuyo cúmulo de axiomas no lógicos sea un subconjunto C de Γ el no incluir ningún trío semejante, o sea el ser no delicuescente, entraña, y es entrañado por, el tener un modelo con respecto a cierta valuación (en virtud del Lema 1º). Por consiguiente, y aplicando la transitividad del bicondicional: $\mathfrak{J}$ tiene un modelo con respecto a alguna valuación ssi lo mismo sucede a cada una de esas teorías $\mathfrak{J}'$. Q.E.D. Pues, a tenor del Corolario 1º, todos esos modelos son aa.tt.

Consideraciones finales

Para cerrar este capítulo —y, con él, este libro— cabe mencionar un par de consideraciones más. El primero es que varios de los modelos algebraicos aquí considerados son **atómicos** en este sentido: definimos una relación de **cobertura**, $\supset$, tal que $z \supset x$ ssi $x \leq z$ sin que haya ningún elemento en absoluto, u , tal que $x \leq u \leq z$; entonces diremos que un cúmulo C ordenado por una relación de orden $\leq$ es **atómico** ssi para cualesquiera dos elementos $x, z \in C$ tales que $x \leq z$ (o sea $x \leq z$ pero $x \neq z$) hay tres elementos u, v, v' , tales que se cumple al menos una de estas dos condiciones: (1^a) $z \geq u \supset v \geq x$; (2^a) $z \geq v' \supset u \geq x$; cumpliéndose ambas a menos que $z \supset x$. Gracias a ser atómicas en ese sentido, las álgebras cuasitransitivas dan una idea de la transicionalidad de las cosas (de los procesos, de los márgenes de acceso de una determinación a otra, de las franjas) según la cual para cualesquiera dos grados no contiguos [en absoluto] hay alguno intermedio que tiene su contiguo por abajo —o tope inferior— y su contiguo por arriba; conque para cada grado hay un punto de arranque y un punto de arrime (o sea: de acceso, de contacto), por lo cual el grado no está inabordable ni separado de cualquier otro por grados intermedios; sin que eso obste a la existencia de infinitos grados (de hecho cuando dos grados no se tocan [en absoluto] hay infinidad de grados entre ellos).

La segunda y última consideración es que cabe extender la modelización algebraica considerando operaciones infinitarias, e.d. no excluyendo del cúmulo de operaciones definidas sobre los portadores de las diversas álgebras de cierta clase operaciones cada una de las cuales tome, en cada caso, simultáneamente como argumentos a elementos en número infinito (quizá inenumerable, o sea: más que números naturales hay), al paso que cada operación finitaria es tal que hay un número finito m tal que esa operación en cada caso toma como argumentos sólo a m elementos haciéndoles corresponder un valor. La ventaja de introducir operaciones infinitarias es brindar así una modelización algebraica también del cálculo cuantificacional. Pero no cualquier álgebra que sirva como modelo a un cálculo sentencial va a servir como modelo a un cálculo cuantificacional que sea extensión del mismo. No todas las aa.tt-oides son modelos de Aq , p.ej. Han de cumplir, para serlo, una serie de condiciones, en las que prefiero no entrar ya, para no alargar más el presente trabajo.

Proponíase este capítulo dar una idea exacta aunque sumaria de las técnicas de la modelización algebraica. Espera el autor que con ello hayan quedado claramente perfiladas las grandes líneas de invención de sistemas lógicos como aquellos en cuya exploración nos hemos ido centrando a medida que se acercaba la culminación del presente opúsculo. Ya en posesión de tales procedimientos, tócale ahora al lector adentrarse por nuevas sendas o avanzar más por las aquí desbrozadas.

Anejo Nº 1

El recurso a una lógica infinivalente en la defensa del realismo científico

En *Word and Object*, p. 249 (citado en el apartado 7 de la Bibliografía del presente opúsculo) señala al respecto Quine:

This quandary over ideal objects, like that over infinitesimals, has its solution in the theory of limits. When one asserts that mass points behave thus and so, he can be understood as saying roughly this: that particles of given mass behave the more nearly thus and so the smaller their volumes. When one speaks of an isolated system of particles as behaving thus and so, he can be understood as saying that a system of particles behaves the more nearly thus and so the smaller the proportion of energy transferred from or to the outside world. This, broadly speaking, is how one's succinct talk of ideal objects would presumably be paraphrased when challenged.

En esto como en tantas otras cosas Quine defiende un gradualismo que, abrazado consecuentemente, lo llevaría a una filosofía de la lógica mucho menos conservadora que la que, sobre todo en los últimos lustros, ha venido defendiendo el gran profesor de Harvard. (Si bien, en honor a la verdad, hay que decir que su conservadurismo tampoco es ni mucho menos tan consecuente o cerril como lo ven algunos. Se dan antes bien en la posición de dicho filósofo titubeos al respecto, pero en sus más lúcidos momentos ve con claridad que podría ser conveniente adoptar una lógica no-clásica, en particular una lógica plurivalente de lo difuso.)

En el pasaje citado Quine casi formula la solución correcta a la dificultad de los “objetos ideales” en las teorías científicas. Equivócase empero al creer que esa solución que él roza sin llegar a desarrollar sería idéntica o similar al tratamiento de los límites por Weierstrass. En verdad ese tratamiento del análisis por medio de la teoría de límites no es ni mucho menos necesario, pues está disponible el análisis no estándar de Robinson que reconoce la existencia de infinitésimos y de ese modo rehabilita lo bien fundado de los cálculos originarios de Leibniz y Newton, con la ventaja de una claridad, elegancia y belleza incomparablemente mayores que cualesquiera teorías de límites como la de Weierstrass. Pero, sea de ello lo que fuere, si el tratamiento de los objetos ideales que está esbozando Quine va a ser —según él lo sugiere en el pasaje citado y en el contexto, sobre todo en el párrafo que sigue— una modalidad o aplicación del procedimiento de Weierstrass o de la teoría de límites, entonces habrán de modificarse [muchos de] los predicados de cualquier teoría física en que se hable [aparentemente] de “objetos ideales”. Cada predicado de éstos que sea n -ádico habrá de venir reemplazado por uno que sea $[n+1]$ -ádico, con el último lugar argumental reservado a un numeral que denote a un número real. Pues sólo así se podrá aplicar el cálculo numeral de la teoría de límites a predicados como éstos que Quine esquematiza como «behaving thus and so».

La solución articulable con una lógica de lo difuso evita ese recargamiento de los predicados, permitiendo mantener los predicados iniciales de la teoría física. Simplemente en la semántica del sistema lógico que se escoja como subyacente a la teoría sí se introducirán los grados de verdad —tantos como sean menester. La clave está en que, mientras que el mero empleo del condicional ‘si... entonces’ (que Quine considera de pasada en el lugar citado, rechazándolo con razón) para dar cuenta de las determinaciones de los gases ideales, p.ej., daría el resultado de que —por ser totalmente falso que algo sea un gas [perfectamente] ideal— dichos “gases” poseerían cualesquiera determinaciones (al parafrasearse «los gases ideales son así o asá» como «Si algo es un gas ideal, es así o asá»), en cambio con la **implicación** de una lógica plurivalente —especialmente de una infinivalente— lo que se obtiene es esto: «En la medida [al menos] en que algo es un gas ideal, es así o asá», siendo tanto más ideal un gas cuanto

más satisface las condiciones definitorias de la gaseosidad. Y similarmente para los cuerpos rígidos, para los elásticos, para los fluidos perfectos, deslizamientos sin fricción, etc.

En verdad todos los terrenos de estudio donde han revelado mejor su fecundidad y aplicabilidad las teorías de conjuntos difusos son campos donde aquello de lo que parecía hablarse eran entidades ideales: en economía, libre mercado, planificación [perfectamente] centralizada, equilibrio de oferta y demanda; en geografía, desiertos puros, o zonas perfectamente húmedas, o fértiles, etc.; en medicina, células totalmente sanas o enfermas; en biología, especies puras perfectamente definidas, cuando los propios biólogos rechazan la existencia de fronteras nítidas determinadas por rasgos morfológicos u otros: y se da por grados en verdad hasta la capacidad de acoplamiento y reproducción, que sería el supuesto “núcleo duro” en la determinación de esas fronteras: hay casos de incapacidad total para el acoplamiento, otros de capacidad total para él y casos intermedios —diversas especies próximas, como tigres y leones—, casos que sin duda han sido los más, con grados diversísimos, a lo largo de los millares de años de evolución.

Hay también una solución alternativa a la dificultad de los “objetos ideales” en las teorías científicas, a saber el **udenismo** o **medenismo** (en inglés ‘noneism’) preconizado por el filósofo australiano Richard Routley (hoy Richard Sylvan), a cuyo juicio las más teorías se ocupan de “algos” —“objetos”, “ítems”— exentos por completo de existencia, de ser, de entidad. No es que tengan, en vez de existencia real, una “ideal” —sea eso lo que fuere— ni que, en vez de existencia, tengan mero “ser”; ni, naturalmente, que posean sólo “existencia en la mente” (ya que los gases perfectos no existen en la mente, sea ésta lo que fuere: si estuvieran en ella, quizá estallarían). No, es que sencillamente tienen verdad, vigencia veritativa, ciertos enunciados acerca de tales objetos sin que para ello éstos últimos hayan de darse o estar, ni en la realidad ni fuera de ella —nada puede estar fuera de la realidad, claro. A este respecto, Routley (en *Exploring Meinong’s Jungle*, citado en el apartado 8 de la Bibliografía del presente opúsculo, p. 458, n. 3) señala que, según esa solución, igual que la geometría se ocupa, no de círculos o triángulos reales, sino de sendos objetos ideales, igualmente la física teórica no es directamente acerca de objetos físicos reales sino ‘about ideal objects to which real objects may, to some degree, approximate’. El talón de Aquiles de semejante tratamiento está en que ni nos explica en qué consista esa aproximación ni brinda nada —ningún marco ni sintáctico ni semántico— para abordar el estudio de semejante “aproximación”; y habla de grados cuando nada en el fondo del tratamiento tiene en cuenta la gradualidad de las determinaciones. Ahora bien, si no son las determinaciones mismas, o los predicados, los que se dan (o se predicán) por grados, sino las aproximaciones a tener tales determinaciones, ¿cómo se entiende eso? Si el carecer de fricción es asunto de **todo o nada**, ¿qué sentido tiene “aproximarse en tal grado a carecer de fricción”? ¿Tener fricción sólo en tal grado? No, si el **carecer de fricción** es asunto de todo o nada, no admitiendo grados, tampoco admitirá grados su complemento, el **tener fricción**, sino que igualmente será asunto de todo o nada. Luego no puede haber aproximaciones en tal o cual grado a una situación o condición que sea asunto de todo o nada, que no admita graduación alguna. (Vide infra, a propósito de una tesis de Smart, una discusión más honda de la concepción de “aproximación a la verdad”).

Pero, suponiendo que hubiera una rama (la única en ese caso —bajo ese supuesto— prácticamente útil) de la teoría física que se ocupara de esas aproximaciones, entonces estaría de más la presunta “teoría física pura” que se ocupara de los meros objetos ideales “como tales” sin condescender a posar su mirada en las aproximaciones de las cosas reales a la posesión de esas determinaciones dizque ideales. Luego habría que arrojar por la borda esa rama “pura” para quedarse uno sólo con la teoría de las aproximaciones. (Dicho entre paréntesis: consideraciones iguales se aplican, desde luego, a la geometría. En la medida en que un objeto

es de superficie circular, es así y asá. En la medida en que es elíptico, es de estas o aquellas características.)

Resulta curioso que un partidario del empleo de lógicas no clásicas como Routley presente, con esa concepción suya, quizá el único argumento no estrictamente idealista a favor de un tratamiento epistemológico de las teorías físicas que obvie el recurso a lógicas infinivalentes o de lo difuso. La discrepancia entre ese enfoque de Sylvan y el idealista de van Fraassen, p.ej., estriba en que Sylvan cree en la verdad de los asertos de una teoría científica que quepa profesar, mientras que van Fraassen, no admitiendo “objetos [completamente] inexistentes” o “ítemes [del todo] privados de ser o entidad”, pero coincidiendo con Sylvan en que la física es “acerca de ellos”, por decirlo así, concluye que la teoría física no es verdadera, sino meramente “adecuada”, o sea: pragmáticamente justificada. Si bien van Fraassen (en el libro en que expone su concepción idealista —que él denomina ‘empirismo constructivo’—, a saber: *The Scientific Image*, Clarendon Press, 1980) no insiste en lo “ideal” de muchas entidades postuladas en teorías científicas para llegar a su conclusión epistemológica de que aceptar una teoría no conlleva creer en la verdad de la misma, el género de consideraciones que desarrolla para avalar su punto de vista va claramente en esa dirección; y algunos adeptos de su enfoque se han explayado precisamente en ese problema del realismo, a saber: que, si es verdadera una teoría física aceptable, existen entonces entes que, además de ser inobservables, vendrían caracterizados por rasgos que los situarían fuera del mundo real.

Están claros, pues, los servicios que puede prestar a la defensa de una concepción como la de los “realistas científicos” (Smart y otros) el recurso a una lógica de lo difuso. Privándose del mismo, debilitan innecesariamente su causa y dejan el flanco expuesto a los reparos (¡oh cuán fundados!) del “empirismo constructivo” y concepciones afines. Por cierto, el propio J.J.C. Smart, tras haber expuesto en *Philosophy and Scientific Realism* (Routledge, 1963) su gran defensa del realismo científico, se percató del desajuste entre los objetos de las teorías científicas —según se entienden desde la lógica clásica, sin matices— y el mundo real. De ahí que, en «Science as an Approximation to Truth» (ap. *Papers Presented to the Annual Conference*, Australasian Association for Philosophy, Melbourne, 1976, p. 14) reconozca que las teorías científicas no son sobre idealizaciones (*there being no idealizations for them to be about*), sino sobre cosas reales, sólo que aproximadamente verdaderas. Sylvan (op. cit., p. 787) critica esa propuesta pragmática de Smart alegando las dificultades de articular una teoría de la aproximación a la verdad que sea extensional, según quiere Smart que lo sea. Extensional o no, lo que al autor de este opúsculo le parece claro es que —por las razones expuestas poco más atrás— una teoría de la aproximación a la verdad es incompatible con la concepción bivalentista, clásica, de que la verdad es algo que rechaza los grados. Si el predicado $\lceil \chi \rceil$ es un predicado monádico tal que cada ente x es tal que $\lceil \chi x \rceil$ es totalmente verdadero o totalmente falso, ¿qué sentido tiene decir que un ente z se aproxima a tener χ ? ¿Quizá que tiene otra propiedad, ψ , que aproxima se a χ ? Y, ¿en qué estriba esa aproximación? No —por hipótesis— en que tener ψ en grado tal conlleve tener χ en grado cual. Otra alternativa sería que el predicado tajante $\lceil \chi \rceil$ “superviniera” sobre uno difuso, $\lceil \phi \rceil$ de tal manera que $\lceil \chi x \rceil$ equivaliera a $\lceil \phi x \rceil$, donde $\lceil \phi \rceil$ sería cierto functor monádico de matiz alético. Eso no lo puede aceptar precisamente el bivalentista, pues él no puede tolerar funtores de matiz alético o veritativo: la verdad —para él— o se da del todo o no se da en absoluto. ¿Qué queda? Seguramente nada. En eso lleva razón Sylvan: no parecen viables las teorías de aproximación a la verdad; no parecen viables dentro del marco bivalentista. Y fuera de él no hacen falta. Nada de aproximación a la verdad. Basta con la verdad. Verdad en algún grado.

Podría intentarse salvar esa idea de “aproximación a la verdad” alegando que, igual que quien está en Viena se aproxima más a estar en Jerusalén que quien está en Londres, más

aún que quien está en México, aunque ninguno de los tres está en Jerusalén, similarmente un enunciado, siendo falso, puede distar menos que otro de ser verdadero. Respondo que, desde un punto de vista bivalentista, no es cierto que quien está en Viena se aproxima más a estar en Jerusalén que quien está en Londres, sino que está igualmente lejos de estar en Jerusalén. Siendo —desde ese punto de vista, el del clasicista— totalmente falso que uno u otro estén en Jerusalén, no hay menor distancia respecto del estar en Jerusalén de quien está en Viena. Lo que sí es cierto aun desde el punto de vista bivalentista —pero eso es algo totalmente distinto— es que quien está en Viena está más cerca de Jerusalén, o sea que la distancia entre Viena y Jerusalén es menor que la que hay entre Viena y Londres. Esa distancia es cuantitativa. Mas si la verdad es cuestión de todo o nada, ¿cómo va a haber entre la Verdad y tal falsedad menos “distancia” que entre la Verdad y tal otra falsedad? ¿Distancia en qué? ¿En “parecido”? Por ahí se llega a lo mismo que si se dijera que quien está en Viena dista menos de *estar en* Jerusalén. Si de veras existe alguna relación de distancia entre la Verdad y las diversas falsedades, así como entre éstas, ¿qué puede ser sino un **orden** de grado de verdad (o de grado, inverso, de falsedad)? Las falsedades menos distantes de la Verdad [total] serán menos falsas, más verdaderas.

(Toda la compleja cuestión de si cabe articular una concepción clara y coherente de la aproximación a la verdad u otra similar —como la semejanza a la verdad o verosimilitud— ha sido abordada, bajo el impulso de sugerencias de Popper, en un sentido, y luego de Putnam, en otro, por diversos autores, con resultados negativos: de hecho ni se ha logrado aclarar la noción ni se ha elaborado un tratamiento de la misma que escape a dificultades lógicas redhibitorias. Un realista científico de la escuela de Smart, Michael Devitt, quien se ha esforzado —según él mismo reconoce, sin éxito— por elaborar una noción así, aun añorando la disponibilidad de una concepción de ese género —precisamente para poder hacer, gracias a ella, frente a las dificultades del realismo científico—, llega a una constatación del fracaso. [Vide su libro *Realism and Truth* (Blackwell, 1984), pp. 113ss, p. 122 n. 2, pp. 156-7.]

Parece, pues, justificado nuestro aserto de que es mucho lo que tiene que aportar a la filosofía de la ciencia el cultivo de algunas de las lógicas estudiadas en el presente opúsculo —principalmente el cálculo *Aj*. Pero, desde luego, hay muchos otros campos de aplicación en otros terrenos, filosóficos y no filosóficos. Lo que parece improcedente es que los filósofos de la ciencia ignoren estas lógicas y barajen alternativas abiertas como si no hubiera modo de escapar al estrecho horizonte de la lógica bivalente.

Anejo Nº 2

Nota sobre la noción quineana de verdad lógica

Inspírase la definición de **verdad lógica** presentada en la Introducción del presente opúsculo en la que ofrece Quine en la Introducción a su libro (1940) *Mathematical Logic* (citado en el apartado 2 de la Bibliografía del presente libro):

A word may be said to occur essentially in a statement if replacement of the word by another is capable of turning the statement into a falsehood.

Quine añade una nota a pie de página remitiendo a una formulación más cuidadosa en su ensayo «Truth by Convention». Este ensayo, escrito en 1935, viene reproducido en *The Ways of Paradox* (citado en el apartado 7 de la Bibliografía del presente opúsculo), en las pp. 77-106. He aquí la definición (p. 80):

An expression will be said to occur **vacuously** in a given statement if its replacement therein by any and every grammatically admissible expression leaves the truth or falsehood of the statement unchanged. Thus for any statement containing some expressions vacuously there is a class of statements, describable as **vacuous variants** of the given statement which are like it in point of truth or falsehood, like it also in point of a certain skeleton of symbolic make-up, but diverse in exhibiting all grammatically possible variations upon the vacuous constituents of the given statement. An expression will be said to occur **essentially** in a statement if it occurs in all the vacuous variants of the statement, i.e. if it forms part of the aforementioned skeleton. (Note that though an expression occur non-vacuously in a statement it may fail of essential occurrence because some of its parts occur vacuously in the statement.)

Pues bien, surge una dificultad que asedia por igual —aunque de manera algo diversa— a ambas definiciones. Hela aquí. Tomemos la oración ‘En Grecia se habla el árabe o en Grecia no se habla el árabe, o Lincoln muere asesinado en 1865’. Es una instancia sustitutiva del esquema $\lceil p \vee Np \vee q \rceil$. Si reemplazamos en él ‘Lincoln’ por ‘De Gaulle’, p.ej., la oración sigue siendo verdadera —por mucho que el último disyunto sea totalmente falso. Ahora bien, si reemplazamos en el enunciado original ‘no’ por ‘totalmente’ o por ‘sí’ (un operador redundante de afirmación), el resultado sigue siendo verdadero también, pues el último disyunto hace que, aunque la fórmula resultante ya no sea una instancia sustitutiva del mencionado esquema, así y todo siga siendo verdadera.

Según la definición que vino propuesta en la Introducción del presente trabajo, no tendría ocurrencia alguna ninguna de las expresiones que figuran en la oración ahora considerada, salvo la segunda ocurrencia de ‘o’. En efecto: ni la ocurrencia de ‘no’ ni la primera de ‘o’ (puesto que reemplazando **esa** ocurrencia de ‘o’ por una de ‘y’ la fórmula total sigue siendo verdadera) ni por supuesto los enunciados “atómicos” ni ninguno de sus componentes. Nada, pues, salvo el segundo ‘o’.

Pasemos a la definición alternativa, la citada de «Truth by Convention». Aquí se exigen dos cosas: 1º) que el reemplazo afecte (uniformemente se supone) a todas las ocurrencias de la expresión vacua; 2º) que las expresiones con ocurrencias esenciales en el enunciado reaparezcan en todas las variantes vacuas así formadas. Pues bien, según esto será una variante vacua de la oración dada ésta: ‘En Grecia se habla el árabe o en Grecia [sí] se habla el árabe o Lincoln muere asesinado en 1865’. De nuevo, pues, la única expresión con una ocurrencia esencial será ‘o’ —aunque ahora habrán de ser esenciales ambas ocurrencias de ‘o’.

Por otro lado, en una oración como ‘Lincoln muere asesinado en 1865 o el mar es salado’ sólo el ‘o’ tendría ocurrencia esencial, ya sea según la definición de «Truth by Convention» ya sea según la definición menos cuidadosa ofrecida en la Introducción del presente opúsculo.

En efecto reemplazar ese 'o' por 'ni... ni' (o —si es que hay grados— para mayor seguridad de que se obtiene el efecto, por 'es totalmente falso que... y también que...') daría por resultado una oración enteramente falsa; cualquier reemplazo de otra expresión deja a la oración verdadera (más o menos). P.ej.: 'Lincoln muere asesinado en 2065 o el mar es salado', 'Lincoln muere asesinado en 1865 o el mar es dulce', etc. Por ende, como sólo la disyunción 'o' tiene ocurrencia esencial en esa oración verdadera, es ésta una verdad de lógica.

La solución mejor (de algún modo apuntada ya por Quine en su *Philosophy of Logic* —también citado en el apartado 7 de la Bibliografía, p. 53) parece la siguiente. Definimos la **ocurrencia esencial de una expresión en un esquema correcto** (o "válido") así: llamamos **correcto** a un esquema cada una de cuyas instancias sustitutivas sea una oración verdadera. Una expresión tiene una ocurrencia esencial en un esquema correcto si hay otra expresión tal que el reemplazo gramaticalmente autorizado de dicha ocurrencia de la primera expresión por una ocurrencia de la segunda genera un esquema que no es correcto (o sea un esquema alguna de cuyas instancias sustitutivas no es verdadera).

(He dicho que esa solución ya viene de algún modo apuntada por el propio Quine en un trabajo posterior. Sin embargo, no figura ahí como solución del problema aquí abordado. Además, no se formula en términos que permitan en general decir qué verdades son tema de una disciplina, sino que se enuncia tan sólo cómo definir verdades de lógica. En tercer lugar, la enunciación que brinda Quine en ese lugar presupone un cierto análisis sintáctico —al definir la verdad lógica como la verdad de una oración si persiste con cualesquiera cambios de los predicados de la oración (ibid., p. 49). Ello limita considerablemente el campo de la verdad lógica: en una lógica combinatoria no existe esa diferencia entre predicados y otros signos. En cambio una letra esquemática puede ser de cualquier categoría gramatical que se quiera: puede ser —según se indique en cada caso— una letra sentencial, o predicativa, o que haga las veces de términos singulares, o de funtores, o de preposiciones, o de adverbios etc. etc.)

Volviendo al enfoque que aquí propongo, forman —según el mismo— parte de una disciplina de estudio o investigación (de una "ciencia") aquellos esquemas en los que sólo tienen ocurrencias esenciales las expresiones que formen el vocabulario de tal disciplina. Son **verdades de una disciplina** aquellas fórmulas verdaderas que son instancias sustitutivas de esquemas de la disciplina.

Así pues, en el esquema $\lceil p \vee Np \vee q \rceil$ las únicas expresiones con ocurrencias esenciales con ' $\vee$ ' y ' N '. La oración 'En Grecia se habla el árabe o en Grecia no se habla el árabe o Lincoln muere asesinado en 1865' es una verdad de lógica porque es una instancia sustitutiva de ese esquema. La oración 'Lincoln muere asesinado en 1865 o el mar es salado' es una verdad de historia porque es una instancia sustitutiva del esquema $\lceil \text{Lincoln muere asesinado en 1865} \vee p \rceil$.

Cierto que, trivialmente, cada oración es un esquema, con un número 0 de letras esquemáticas. Así que el **esquema** $\lceil \text{Lincoln muere asesinado en 1865} \vee \text{el mar es salado} \rceil$ es un esquema que no pertenece ni a la historia ni a la talasología o ciencia del mar; aunque la oración disyuntiva que forma el esquema pertenece a ambas disciplinas por ser una instancia sustitutiva de sendos esquemas de la una y de la otra.

Quizá resulte un poco ardua esa diferencia entre la oración 'Elena es esbelta' y el esquema $\lceil \text{Elena es es esbelta} \rceil$. Pero eso es de poca monta. Hay cómo refinar y pulir la definición ofrecida para evitar tal dualidad (diciendo p.ej. que una oración pertenece a una disciplina si: o bien (1º) es instancia sustitutiva de un esquema correcto no trivial [e.d. con al menos una letra esquemática] de esa disciplina, o bien (2º) el esquema trivial idéntico a la oración es, él mismo, un esquema correcto de dicha disciplina). Y, por otro lado, poco daño hace ésta.

Bibliografía selecta comentada

1.— Introducción a la teoría de conjuntos

Siendo la teoría [ingenua] de conjuntos lo único con lo que se supone que el lector de este tratadito se halla previamente siquiera mínimamente familiarizado, conviene indicar aquí algunos textos muy accesibles de iniciación a esa disciplina.

- Javier de Lorenzo, **Iniciación a la teoría intuitiva de conjuntos**. Madrid: Tecnos, 1972.
- Lía Oubiña, **Introducción a la teoría de conjuntos**. Buenos Aires: Eudeba, 1971.
- Seymour Lipschutz, **Teoría de conjuntos y temas afines**. (Trad. J.M. Castaña & E. Robledo.) México: McGraw-Hill (Serie Shaum), 1970.

2.— Introducciones generales a la lógica

Hay muchas. Todas pecan de presentar a **una** lógica particular —la clásica— como “la” lógica, añadiendo a lo sumo algún capítulo o apéndice acerca de algunas lógicas no clásicas.

Hay varias buenas introducciones, algunas de ellas escritas en nuestro idioma. He aquí varias introducciones dignas de mención.

- Jon Barwise, **Handbook of Mathematical Logic**. Amsterdam: North-Holland, 1977. Aunque se titula ‘manual’ y se pretende dirigido a no iniciados en lógica, ello es porque va destinado a matemáticos profesionales que no hayan ahondado en su conocimiento de la lógica. Seguramente será poco atractivo para la mayoría de los lectores del presente opúsculo. Riqueza de temática, rigor de tratamiento, abundancia de información y de resultados demostrativos. Ausencia de las lógicas no clásicas. Pobreza de las pocas consideraciones filosóficas. Está concebido para hacerles ver a los matemáticos que la lógica no es de menospreciar. Abarca trabajos interesantísimos sobre teoría de conjuntos y cálculos lambda (variedades de la lógica combinatoria).
- Alonzo Church, **Introduction to Mathematical Logic**. Princeton (New Jersey): Princeton U.P., 1956. Es el mejor tratado introductorio a la lógica clásica —incluyendo el cálculo cuantificacional clásico de orden superior o teoría de tipos—, sin desdeñar del todo a ciertas lógicas no clásicas. Es el texto más riguroso. Abunda en ejercicios muy bien diseñados y en explicaciones magníficamente concebidas. No se recomienda su lectura a quienes no estén previamente familiarizados con una amplia gama de temas, puesto que —ése es su único fallo, aparte del dogmatismo clasicista— no allana mucho el camino a los todavía no duchos en la materia.
- Irving M. Copi, **Symbolic Logic**. New York: Macmillan, 1973. Pese a su sello un tanto personal, y una serie de detalles originales, es un manual bastante típico; sólo que contiene capítulos muy interesantes, p.ej. sobre la notación polaca y sobre el sistema de Nicod —una versión de la lógica sentencial clásica en la cual hay un único functor primitivo.
- Donald Kalish & Richard Montague, **Logics: Techniques of Formal Reasoning**. Es un estudio más avanzado que un mero manual introductorio. Es excelente y rigurosa su presentación de las partes más complicadas de la lógica cuantificacional clásica de primer orden y de algunas extensiones de la misma. Claro, bien escrito, elegante y diestro en combinar la preocupación por la exactitud y la profundidad con explicaciones que allanan el camino a los no previamente iniciados.
- Benson Mates, **Lógica matemática elemental**. (Trad. C. García Trevijano.) Madrid: Tecnos, 1970. Es una excelente introducción, pese a su dogmatismo clasicista. Se atiene a una concepción de lógica bastante dispar de la aquí presentada, una en la que la lógica está expuesta en un lenguaje formal especial. Contiene capítulos informativos muy útiles.
- Jesús Mosterín, **Lógica de primer orden**. Barcelona. Ariel, 1970. Un texto bien escrito y útil, aunque es más riguroso en ciertos detalles que en ciertas cuestiones de fondo.

- Willard V.O. Quine, **Mathematical Logic** (Revised Edition.) Cambridge (Mass.): Harvard U.P., 1951. (Hay trad. castellana en Editorial Revista de Occidente, Madrid.) Es un texto sin igual: desde el piso del cálculo sentencial hasta desarrollos matemáticos de envergadura articulados en una teoría de conjuntos (el sistema ML) expuesto en este libro por primera vez, todo ello con un enorme rigor y una claridad insuperable. A este libro —y en general a toda la obra lógica de Quine— debe mucho el tratamiento aquí ofrecido. No requiere ningún conocimiento previo de la materia, pero sí un estudio muy atento, ya que no está concebido como manual docente —según suele entenderse—; p.ej. carece de ejercicios y no se exploya mucho en ilustraciones.
- J. Barkley Rosser, **Logic for Mathematicians**. N.Y.: Chelsea, 1978 (2ª ed.). Es un excelente tratado de desarrollo de amplios campos de la matemática construidos desde el cimiento de la lógica clásica y de la teoría de conjuntos NF de Quine. Esta segunda edición contiene además importantes apéndices sobre temas superiores de teoría de conjuntos y sobre el análisis no estándar de Robinson (análisis utilizado en el presente opúsculo para modelizaciones del sistema de lógica propuesto en el capítulo IX).
- James A. Thomas, **Symbolic Logic**. (Columbus (Ohio): C.E. Merrill Co., 1977. Es un manual típico, con sus lados buenos y malos. Pone mucho énfasis en procedimientos de “deducción natural”, o sea en ver a la lógica no como una teoría o corpus de teoremas —que es la concepción aquí abrazada— sino como un instrumental, un arsenal de reglas de inferencia. Tiene un tratamiento bastante exhaustivo y claro del cálculo cuantificacional clásico de primer orden y muchas indicaciones valiosas e informativas sobre aplicaciones y conexiones con otras disciplinas.

3.— Introducciones a la teoría de modelos

Casi todas las introducciones a la lógica contienen un tratamiento de teoría de modelos. Los textos aquí enumerados son algunos que van más lejos en esa tarea o que la llevan a cabo con mayor éxito, sea en lo tocante a profundidad, sea por su claridad.

- C.C. Chang & H.J. Keisler, **Model Theory**. Amsterdam: North Holland, 1973. Es un texto excelente, pero —pese a su claridad— requiere una lectura sumamente atenta. No se recomienda sino a quienes ya dominen bien los rudimentos de la teoría de modelos y estén dispuestos a estudiar a fondo los temas más arduos.
- Geoffrey Hunter, **Metalogic: An Introduction to the Metatheory of Standard First Order Logic**. Berkeley & Los Angeles: University of California Press, 1971. Un excelente manualito, claro, riguroso e incluso tímidamente oteador de algún enfoque no-clásico (el relevantista). Pero en general se atiene ciegamente a los prejuicios clasicistas.
- Stephen C. Kleene, **Introducción a la metamatemática**. (Trad. M. Garrido). Madrid: Tecnos, 1974. Es uno de los tratados más descollantes, con interesantísimos desarrollos sobre aplicaciones matemáticas —p.ej. sobre teoría de la recursión. Es un libro exigente, de lectura difícil.
- Elliott Mendelson, **Introduction to Mathematical Logic** (2d. edit.). New York: Van Nostrand, 1979. Merecidamente reputado como el mejor manual de lógica matemática clásica; aunque eso es verdad sólo en ciertos aspectos, naturalmente. (P.ej., carece de la elegancia, la enorme claridad y el enfoque de construcción sistemática de una teoría lógica que caracterizan a **Mathematical Logic** de Quine; a cambio, abarca muchos más temas y suministra una enorme cantidad de información sobre diversos hallazgos acerca de tales temas.) De lectura difícil por la hondura de su tratamiento y lo complejo de muchos de los temas que toca, es, sin embargo, muy claro en sus explicaciones y en cómo facilita al lector los primeros pasos.

- Robert Rogers, ***Mathematical Logic and Formalized Theories***. Amsterdam: North-Holland, 1971. Una presentación inigualablemente clara a temas de teoría de modelos y teoría axiomática [estándar] de conjuntos. Insuperable como trabajo propedéutico.
- Joseph R. Shoenfield, ***Mathematical Logic***. Reading (Mass.): Addison-Wesley Publ. Co., 1967. Trátase de un texto clásico en varias acepciones. Su semántica es estándar en el fondo pero muy peculiar en su presentación. Es muy profundo su tratamiento de algunos problemas avanzados de teoría de modelos y teoría de conjuntos.

4.— Modelos algebraicos

Hay muy diversos tratados de temas pertenecientes al dominio del álgebra universal e incluso una revista con esa denominación (***Algebra universalis***). Menciono aquí algunos de los textos más sobresalientes y útiles.

- Raymond Balbes & Philip Dwinger, ***Distributive Lattices***. University of Missouri Press, 1974.
- Garrett Birkhoff, ***Lattice Theory***. Providence (Rhode Island): American Mathematical Society, 1940. Texto sumamente difícil, pero también el texto clásico por antonomasia, insuperado por el enorme despliegue de resultados demostrativos que ofrece.
- Paul M. Cohn, ***Universal Algebra***. Dordrecht: Reidel, 1981. Pese a su gran valor, este texto no es el que mejor introduce al tema a lectores que no sean matemáticos o que se interesen por el álgebra universal desde la perspectiva de búsqueda de modelos algebraicos para sistemas lógicos.
- J. Kuntzmann, ***Algèbre de Boole***. París: Dunod, 1968. Es un libro concebido más bien para matemáticos, pero bastante claro.
- Helena Rasiowa, ***An Algebraic Approach to Non-Classical Logics***. Amsterdam: North-Holland, 1974. Sin duda alguna, el texto cuya lectura más calurosamente cabe recomendar a quienes deseen avanzar en un estudio de temas relacionados con la teoría de modelos algebraica a partir de una perspectiva más o menos afín a la que se plasma en el presente opúsculo.

5.— Lógica combinatoria

- H.P. Barendregt, ***The Lambda Calculus***. Amsterdam: North-Holland, 1984. (2ª edic.). Es una exposición magnífica —pero para estudios avanzados— de los problemas principales del cálculo lambda y de sus relaciones con la lógica combinatoria. Es especialmente útil para comprender el género de modelos algebraicos propuestos para una gama de lógicas combinatorias. (Se complementa este libro —que se centra en tratamientos semánticos— con el de Fitch, más abajo citado, ya que en éste no aparece ningún desarrollo sistemático de un sistema **lógico** con sus axiomas y reglas de inferencia, que es lo que en cambio ofrece, magistralmente, el libro de Fitch.)
- Haskell B. Curry & Robert Feys, ***Lógica combinatoria***. (Trad. Manuel Sacristán.) Madrid: Tecnos, 1967. Es un texto de gran valor pero no muy claro —entre otras cosas por lo peculiarísimo de la técnica expositiva y de la terminología, muy alejadas de lo normal. La enorme riqueza de su temática y lo polifacético del tratamiento —a menudo demasiado prolijo— no se ven acompañados por un trabajo suficiente de construcción sistemática.
- Frederic B. Fitch, ***Elements of Combinatory Logic***. New Haven: Yale U.P., 1974. Es el mejor texto en esta disciplina de la lógica combinatoria (aunque —a diferencia del de Barendregt, más arriba citado— se ciñe a un tratamiento sintáctico): claro, riguroso, elegante, conciso, original, con montones de resultados demostrativos y de ideas interesantes. Es uno de los mejores libros de lógica matemática en este siglo.

6.— Relaciones entre la lógica y el estudio del lenguaje

Es inabarcable todo lo muchísimo que sobre este tema se ha publicado, sobre todo en años recientes. he aquí unos pocos títulos.

- Albert E. Blumberg, ***Logic: A First Course***. Es un típico manual propedéutico de lógica, pero con más énfasis de lo usual en la formalización de mensajes en lengua natural.
- Alec Fisher, ***The Logic of Real Arguments***. Cambridge U.P., 1988. Sea grande o pequeño su éxito en la búsqueda de la estructura ilativa (inferencial) de argumentos que utilizan diversos pensadores —Carlos Marx, J. Stuart Mill, Malthus etc.—, pocas iniciaciones logran como ésta poner de relieve la importancia de la lógica desde el punto de vista de la evaluación de los argumentos que de hecho se utilizan en las disputas científicas.
- Samuel Guttenplan, ***The Languages of Logic***. Oxford: Blackwell, 1986. Sería una introducción a la lógica como hay tantas si no fuera por ese énfasis —brioso y, habida cuenta de todo, más bien exitoso— en cómo encontrar correlaciones adecuadas entre ristas de signos en notación simbólica y mensajes de la lengua natural.
- James D. McCawley, ***Everything that Linguists Have Always Wanted to Know about Logic***. Oxford: Blackwell, 1981. Una excelente introducción a la lógica desde un punto de vista de intereses investigativos en lingüística y recalando aplicaciones a ese campo. No deja de contener algún punto de discusión de tratamientos desde lógicas no-clásicas, p.ej. lógicas de lo difuso. Por todo ello es muy recomendable su lectura.
- Howard Pospesil & David Marans, ***Arguments: Deductive Logic Exercises***. Un librito sin grandes pretensiones, divertido y sencillo; un poco ingenuo en cómo concibe la tarea de determinar qué sea un argumento. Pese a sus limitaciones, como despliega una amplia gama de ejemplos, resulta útil dentro de la presente rúbrica.
- Fred Sommers, ***The Logic of Natural Language***. Oxford: Clarendon, 1982. El enfoque de Sommers es muy peculiar —en cierto sentido es un camino opuesto al de la lógica matemática cuantificacional desde Frege, un cierto retorno a la silogística anterior. Por otro camino llega a algo afín en cierto modo a las lógicas combinatorias

7.— Filosofía de la lógica

Sobre este tema es difícil escoger sólo unos pocos libros, porque según qué criterios se apliquen aparecen como más destacados unos u otros títulos —siendo empero dignos de consideración esos criterios no coincidentes en sus resultados—, al paso que, en temas más técnicos, surge menor discrepancia entre los diversos criterios. Aun así —a sabiendas de cuán arbitraria puede resultar la opción—, me atrevo a ofrecer esta pequeña selección.

- Newton C.A. da Costa, ***Ensaio sobre os fundamentos da lógica***. São Paulo: Hucitec, 1980. Es una excelente colección de ensayos filosóficos sobre la lógica en una perspectiva paraconsistente.
- Alfredo Deaño, ***Las concepciones de la lógica***. Madrid: Taurus, 1980. Es un libro muy ambicioso, pues trata de trazar un cuadro exhaustivo de las concepciones posibles de qué sea la lógica, proponiendo su propio enfoque. Sin embargo, ni el acopio hecho por el autor es suficiente para acometer tal empresa ni las alternativas están siempre bien planteadas. No aborda nunca con claridad la concepción más simple de la lógica (la lógica como ontología); ni plantea con suficiente seriedad el desafío de las lógicas no-clásicas. Pese a esas y otras lagunas, y pese al carácter un tanto farragoso de gran parte de la obra, ésta es recomendable para quien quiera meditar a fondo sobre algunas de las concepciones alternativas de la lógica.

- Susan Haack, ***Philosophy of Logics***. Cambridge U.P., 1978. Es un libro ambicioso, que abarca muchos temas y aporta consideraciones de cierto interés acerca de unos cuantos de ellos. Con todo, se viene a quedar corto en la mayor parte de esas discusiones. Peca de precipitación en las conclusiones y de simplificación en los planteamientos, sin ir casi nunca lo bastante al fondo de los problemas.
- Stephen Körner (ed.), ***Philosophy of Logic***. Oxford: Blackwell, 1976. Una colección de estudios que abarca interesantes colaboraciones de autores destacados —Fitch, Geach, Wiggins, Hintikka, Dummett— sobre temas como la lógica combinatoria y su aplicación al tratamiento científico de la lengua natural, las lógicas multivalentes, los cuantificadores, etc.
- W.V. Quine, ***Philosophy of Logic***. Cambridge (Mass.): Harvard U.P., 1970. (Hay traducción castellana.) No es ni mucho menos el mejor libro de Quine —toda cuya obra tiene que ver con la filosofía de la lógica—, pero es un tratadito claro, escrito con ese espíritu penetrante y lúcido propio de su autor —si bien en esta obra se acentúan las tendencias conservadoras más que en trabajos anteriores de Quine.
- W.V. Quine, ***The Ways of Paradox and Other Essays*** (revised and enlarged edition). Cambridge (Mass.): Harvard U.P., 1976.
- W.V. Quine, ***Word and Object***. Cambridge (Mass.): The M.I.T. Press, 1960. (Hay trad. castellana de Manuel Sacristán).

8.— Lógicas no-clásicas

Voy a enumerar unos pocos de entre los libros cuya lectura es más recomendable.

- Nicola Grana, ***Logica paraconsistente***. Nápoles: Loffredo editore, 1983. Es el primer libro dedicado a este tipo de lógicas. El capítulo último está consagrado a la **lógica transitiva**, e.d. a la familia de sistemas lógicos puestos en pie por el autor del presente opúsculo.
- Susan Haack, ***Deviant Logic***. Cambridge U.P., 1974. (Hay trad. castellana de Editorial Paraninfo). Tesis doctoral de la autora, abarca muchos temas interesantes sobre aplicaciones filosóficas de unas u otras lógicas no clásicas. No aporta ningún nuevo resultado demostrativo (técnico). No aborda casi ninguno de los temas con la profundidad que sería de desear; es más: en alguna de las discusiones incurre en marcada superficialidad, esquematizando, aligerando los problemas y las dificultades, ignorando alternativas. Quizá por ello acaba quedándose en una posición conservadora, optando por la lógica clásica frente a cualquier otra. Aun así es recomendable su lectura.
- Storrs McCall (ed.), ***Polish Logic: 1920-1939***. Es una recopilación de artículos de Łukasiewicz, Jaśkowski y otros lógicos polacos. En ella se encuentran los textos donde vino expuesta y desarrollada la lógica trivalente łukasiewicziana y otros en los que por primera vez se propusieron otros sistemas lógicos no clásicos.
- Francisco Miró Quesada & Roque Carrión (eds), ***Antología de la lógica en América Latina***. Madrid: Fundación Banco Exterior, 1988. Es una colección de estudios de diversos temas de lógica matemática; pero, dado el predominio que en Latinoamérica han adquirido las lógicas no-clásicas —principalmente las paraconsistentes—, a ellas vienen consagrados muchos de los trabajos recopilados. Cabe citar entre los autores a: Newton da Costa, Ayda Arruda, C. Alchourrón & E. Bulygin, Raúl Orayen, Mario Bunge, Tomás Moro Simpson, Francisco Miró Quesada, F.G. Asenjo, H.N. Castañeda y el autor de estas líneas (cuyo trabajo se titula «La defendibilidad lógico-filosófica de teorías contradictorias»).
- Gr. C. Moisil, ***Essais sur les logiques non chrysippiennes***. Bucarest: Editions de l'Académie de la République Socialiste de Roumanie, 1972. Es una impresionante

recopilación de estudios del autor sobre lógicas multivalentes, modelizaciones de las mismas, aplicación a temas de lógica modal y muchos otros. Aunque refleja un estado de la investigación ya algo superado, es una obra tan descollante que quedará como uno de los pilares del estudio de la lógica no clásica.

- Graham Priest, Richard Routley & Jean Norman (eds), ***Paraconsistent Logic: Essays on the Inconsistent***. Munich: Philosophia Verlag, 1989. Abarca estudios de sistemas lógicos paraconsistentes —así como de algunas aplicaciones y motivaciones filosóficas de varios de ellos— por autores como Newton da Costa, Priest, Routley (hoy llamado ‘Sylvan’), Asenjo, Bunder, Batens, Brady, Francisco Miró Quesada y el autor del presente opúsculo. Es un texto insoslayable para quien quiera captar el desafío que a los clichés y prejuicios clasicistas suscita el surgimiento de lógicas paraconsistentes.
- Wolfgang Rautenberg, ***Klassische und nichtklassische Aussagenlogik***. Braunschweig/Wiesbaden: Vieweg & Sohn, 1979. Un estudio comparativo de la lógica clásica y de varios sistemas no clásicos. El capítulo III está dedicado a la lógica multivalente, con una introducción a la semántica algebraica. El V lo está a la lógica intuicionista y sistemas afines. Desgraciadamente tiene unas cuantas lagunas que son de lamentar; pero es una obra muy sólida y bien hecha.
- Nicholas Rescher, ***Many-Valued Logic***. New York: McGraw-Hill, 1969. Obra propedéutica —y que logra serlo de manera extraordinariamente destacada, como pocas—, no sacrifica en aras de serlo ni la riqueza y variedad de la temática ni el rigor del tratamiento. Ciertamente que a menudo no va suficientemente al fondo de las cosas; y que pecan a veces de cierta ligereza, precipitación o —más que nada— esquematismo las abundantes y en general interesantísimas consideraciones filosóficas con que está adobado todo el libro; sin embargo, sería éste aquel libro que recomendaría el autor de las presentes páginas como el libro de lectura obligada [por antonomasia] y sin duda aquella obra a la que más debe el cúmulo de tratamientos que se perfilan en mi propia escritura.
- J. Barkley Rosser & Atwell R. Turquette, ***Many-Valued Logics***. Amsterdam: North-Holland, 1952. Es un tratado muy riguroso pero de difícil lectura y sin grandes motivaciones filosóficas.
- Richard Routley, ***Exploring Meinong's Jungle and Beyond***. Canberra: Australian National University, 1980. Una voluminosa obra de 1035 pp., en la que el prolífico filósofo y lógico australiano —cuyo apellido, desde entonces, ha venido reemplazado por el de ‘Sylvan’— desarrolla un sistema de lógica relevante y paraconsistente y expone muchos argumentos a favor de la aplicación de ese sistema de lógica relevante y paraconsistente y expone muchos argumentos a favor de la aplicación de ese sistema en campos de teoría de los objetos, epistemología, teoría del tiempo, semántica, teoría de conjuntos etc. Su principal inspirador es Meinong. Todo ello da idea de la envergadura de la empresa. Sin embargo —y ésta es la mayor falla—, no se aprecia auténtica unidad: trátase de una recopilación de ensayos retocados y queda subyacente la pluralidad de enfoques que se han sucedido en el pensamiento del autor; con el agravante de que la mayoría de las consideraciones y los enfoques no requieren una lógica paraconsistente ni abonan a favor de la misma, sino más bien de procedimientos en buena parte conciliables con la lógica clásica. Aun así, por la enorme riqueza de su contenido es un libro que merecería que todo el mundo sacara tiempo para leerlo.
- L. A. Zadeh, K.-S. Fu, K. Tanaka, and M. Shimura, ***Fuzzy Sets and their Applications***. New York: Academic Press, 1975. Es este título sólo un botón de muestra de la enorme bibliografía hoy disponible sobre teorías de conjuntos difusos, lógicas de lo difuso y aplicaciones de tales lógicas y teorías a múltiples campos de investigación científica, dentro y fuera de la matemática.

9.— Desarrollos y aplicaciones filosóficas de los sistemas de lógica transitiva

Los sistemas de la familia **A** presentados en los últimos capítulos del presente opúsculo han venido desarrollados y debatidos en trabajos anteriores del autor de estas líneas. He aquí varios de tales trabajos.

- «Identity, Fuzziness and Noncontradiction», *Noûs* 18/2 (mayo 1984), pp. 227-59.
- «Un enfoque no-clásico de varias antinomias deónticas», *Theoria* 7-8-9 (San Sebastián: 1988), pp. 67-94.
- *Fundamentos de ontología dialéctica*. Madrid: Siglo XXI, 1987. Pp. 427. Ver esp. Anejo IV, pp. 362-98.)
- «¿Lógica combinatoria o teoría estándar de conjuntos?», *Arbor* 520 (abril 1989), pp. 33-73.
- «(Quasi)Transitive Algebras», *Multiple-Valued Logic* 13 (Los Angeles: IEEE Computer Society, 1983), pp. 129-35.
- «Fuzzy Arithmetics», *Multiple-Valued Logic* 12 (Los Angeles: IEEE Computer Society, 1982), pp. 232-34.
- «Algunos desarrollos recientes en la articulación de lógicas temporales», apud *Lenguajes naturales y lenguajes formales IV.1*, compilado por Carlos Martín Vide. Barcelona: Universitat de Barcelona, 1989, pp. 413-39.
- «Consideraciones filosóficas sobre la teoría de conjuntos», *Contextos* 11 & 12 (Universidad de León, 1988), pp. 33-62 & 7-43.
- «Tres enfoques en lógica paraconsistente», *Contextos* 3 & 4(1984), pp. 81-130 & 49-72.
- «Características técnicas y significación filosófica de un cálculo lambda libre», apud *Lógica y filosofía del lenguaje*, compilado por S. Alvarez, F. Broncano & M.A. Quintanilla. Salamanca: Universidad de Salamanca, 1986, pp. 89-114.
- «Contribución a la lógica de los comparativos», apud *Lenguajes naturales y lenguajes formales II*, compilado por Carlos Martín Vide. Barcelona: Universitat de Barcelona, 1987, pp. 335-50.
- *Rudimentos de lógica matemática*, Madrid: Servicio de Publicaciones del CSIC, 1991.
Para una crítica de las concepciones que animan a la lógica transitiva véanse estos artículos:
- Newton C.A. da Costa, «La filosofía de la lógica de Lorenzo Peña», *Arbor*, Nº 520 (abril 1989), pp. 9-32.
- Mauricio Beuchot, «Acerca de la argumentación filosófico-metafísica», *Crítica*, Nº 53 (agosto 1986), pp. 57-66.