

INTRODUCCION A LA LOGICA
Y AL ANALISIS FORMAL

© 1964, EDICIONES ARIEL, S. A.
Barcelona

Núm. registro: B. 578-1964
Depósito legal: B. 29.198-1964

Talleres Tipográficos ARIEL, S. A. - Berlín, 46-50 - Barcelona

Hay algunos importantes conceptos que son hoy de uso frecuente en numerosas ciencias positivas y que tienen en la lógica formal el lugar de su primera introducción y aclaración. Se trata de conceptos como los de sistema deductivo, algoritmo, modelo, función, estructura. Esa primera aclaración que se encuentra en la lógica es, desde luego, muy general, y los conceptos en cuestión toman en las diversas ciencias positivas que los usan connotaciones específicas. Pero una introducción formal a esos conceptos en el marco de una iniciación a la lógica es probablemente útil para toda formación científica que quiera educar también en el espíritu de la teoría. La principal motivación con que ha sido escrito este manual es la de suministrar un texto introductorio que, a diferencia de lo que muy naturalmente suele ocurrir a los libros de lógica, no presuponga en sus lectores ningún interés especial por la filosofía ni por la matemática, ni menos una educación universitaria en ellas. El lector típico tenido presente es más bien el estudiante de nuestras facultades de ciencias positivas (naturales y sociales). Esto puede dar razón del carácter ingenuo de la información y las discusiones sobre temas filosóficos y matemáticos, así como del abandono de venerables doctrinas tradicionales (por ejemplo: de la renuncia a un tratamiento sustantivo de la silogística).

Lo que aquí se pretende en sustancia es servir a la introducción del estudio de la lógica fuera de las secciones de filosofía y de matemáticas. Salvo en las Facultades de Ciencias Económicas, que cuentan con unos *Fundamentos de Filosofía* en su primer curso, no es aún nada fácil alcanzar ese deseable objetivo. Tal vez una modesta solución podría consistir por ahora en cursillos trimestrales o cuatrimestrales con alguna selección de temas como la compuesta por los siguientes capítulos: I, II,

III, IV, V, VII, XIV, XV, XVI, XVII, XVIII. (Esta reducción acarrea la supresión de derivaciones calculísticas en los caps. XIV-XVIII, que habrá que suplir con razonamientos informales de los que se dan varios ejemplos en el capítulo XV.)

El Dr. D. José Lóbez Urquía, catedrático de *Matemáticas de las operaciones financieras* en la Facultad de Ciencias Económicas de Barcelona, ha tenido la bondad, que le agradezco, de leer el texto en pruebas y sugerirme retoques de interés didáctico que he llevado a cabo en la medida en que lo permitían los límites de espacio y de contenido impuestos al manual.

M. S. L.

Barcelona, octubre de 1964,

INDICE

PRINCIPALES ERRATAS OBSERVADAS

Pág.	Línea	Dice	Debe decir
121	última	$[y/z]$	$[y/x]$
128	10 desde abajo	'E'	'E'
136	8 desde arriba	$R \leftrightarrow$	$RI \leftrightarrow$
216	12 desde abajo	por	desde
242	fórmula (DR30a)	Ey	$\exists y$
247	2 desde arriba	$R [$	$\hat{R} [$
247	última	Ryz	Ryx
261	fórmula (DR107)	$R [$	$\hat{R} [$
282	fórmula (1)	$+$	$\vdash$
143: las dos últimas fórmulas de la demostración están intercambiadas.			

Capítulo III. — El ideal del lenguaje bien hecho

15. Las paradojas de la teoría de conjuntos	38
16. Lenguajes "mal hechos"	41
17. La paradoja de Epiménides	42
18. Lenguajes "bien hechos", Cálculos formales y lenguajes formalizados. Metalenguaje	44

III, IV, V, VII, XIV, XV, XVI, XVII, XVIII. (Esta reducción acarrea la supresión de derivaciones calculísticas en los caps. XIV-XVIII, que habrá que suplir con razonamientos informales de los que se dan varios ejemplos en el capítulo XV.)

El Dr. D. José Lóbez Urquía, catedrático de *Matemáticas de las operaciones financieras* en la Facultad de Ciencias Económicas de Barcelona, ha tenido la bondad, que le agradezco, de leer el texto en pruebas y sugerirme retoques de interés didáctico.

INDICE

PARTE PRIMERA

LA LÓGICA FORMAL Y LAS CIENCIAS REALES. CATEGORÍAS LÓGICAS

Capítulo I. — *Noción de la lógica formal*

1. Cómo se caracteriza una ciencia	13
2. La falacia de abstracción	17
3. La abstracción básica de la lógica formal	17
4. Formas de pensamiento irrelevantes para la lógica formal clásica.	18
5. La forma lógica. Esquemas	19
6. Verdad formal. Esquemas finales. La lógica formal como ciencia o teoría	22
7. Sentido de las verdades formales	25
8. La lógica como arte o técnica	27

Capítulo II. — *La lógica formal en la investigación de fundamentos*

9. La cuestión de los fundamentos	30
10. La "crisis de fundamentos" de las ciencias	31
11. Aspectos materiales y formales de una crisis de fundamentos .	32
12. La estructura de las teorías	34
13. La precisión del objeto formal de una teoría	35
14. La utilidad heurística de la investigación de fundamentos . .	36

Capítulo III. — *El ideal del lenguaje bien hecho*

15. Las paradojas de la teoría de conjuntos	38
16. Lenguajes "mal hechos"	41
17. La paradoja de Epiménides	42
18. Lenguajes "bien hechos". Cálculos formales y lenguajes formalizados. Metalenguaje	44

19. La investigación de fundamentos en lógica. Sintaxis y semántica.	49
20. Los límites del ideal algorítmico	51
21. Los frutos del programa algorítmico	53

Capítulo IV. — *Las categorías lógicas*

22. Análisis lógico y categorías	55
23. Las categorías Expresión, Fórmula, Enunciado	56
24. El nombre: la categoría Constantes. Constantes lógicas	57
25. La categoría Variables	59
26. Los cuantificadores	61
27. Categorías compositivas o conjuntivas	64
28. Funciones lógicas. Abstracción funcional	66

PARTE SEGUNDA

EL SISTEMA DE LA LÓGICA ELEMENTAL

Sección primera

EL LENGUAJE DE LA LÓGICA ELEMENTAL

Capítulo V. — *La composición de enunciados.* *Lógica de enunciados*

29. Concepto de la lógica de enunciados	71
30. Símbolos elementales y primitivos de la lógica de enunciados	71
31. Fórmulas de la lógica de enunciados	78
32. Esquematización de enunciados del lenguaje común por fórmulas de la lógica de enunciados	80
33. Sobre la notación de las funciones veritativas	86

Capítulo VI. — *La estructura de los enunciados atómicos.* *Lógica de predicados*

34. Concepto de la lógica de predicados	88
35. Símbolos elementales o primitivos de la lógica de predicados	91
36. Fórmulas de la lógica de predicados de primer orden	94
37. Esquematización de enunciados del lenguaje común por fórmulas de la lógica de predicados de primer orden	95
38. La implicación en la lógica de predicados de primer orden	101
39. Sobre la constante 'I'	102

Sección segunda

EL SISTEMA DE LA LÓGICA ELEMENTAL

Capítulo VII. — *Presentación axiomática del cálculo de predicados de primer orden*

40. El sistema axiomático	103
41. Nociones de demostración o derivación, y de teorema	104
42. Axiomas y esquemas axiomáticos	105
43. Modelos isomorfos. Monomorfismo y polimorfismo	106
44. La idea de definición implícita	108
45. El sistema axiomático de Hilbert, Bernays y Ackermann (HB) para la lógica elemental	110
46. La demostración en el sistema axiomático	114

Capítulo VIII. — *La deducción a partir de premisas*

47. Las tres concepciones de la argumentación formal	119
48. Los problemas de la deducción natural	120
49. Las nociones de demostración y teorema en el cálculo de la deducción natural	122
50. Los dos géneros de operaciones básicas del cálculo de la deducción natural	123
51. Isomorfía de fórmulas y sustitución de variables	125
52. El cálculo de la deducción natural	128
53. Observaciones a las reglas de la deducción natural	130

Capítulo IX. — *Técnica de la deducción natural.* *Algunos teoremas*

54. Algunos teoremas de la lógica de enunciados	113
55. Algunos teoremas con cuantificadores	140

Capítulo X. — *Formas normales. Comparación del sistema axiomático con el cálculo de la deducción natural*

56. Noción de formal normal	149
57. Normalización de los funtores veritativos	150
58. Normalización de los cuantificadores	153
59. Justificación de las formas normales en el sistema axiomático. Comparación del sistema axiomático con el cálculo de la deducción natural	159
60. El teorema de deducción	163

PARTE TERCERA

LIMITACIONES Y ALCANCE
DEL CÁLCULO LÓGICO

Sección primera

LAS LIMITACIONES DEL CÁLCULO LÓGICO

Capítulo XI. — *Rendimiento del cálculo lógico elemental*

61. Los metateoremas sobre el rendimiento	177
62. Consistencia del cálculo lógico elemental	178
63. Completud del cálculo lógico elemental	182
64. La lógica elemental y el programa algorítmico	185

Capítulo XII. — *La lógica de predicados de orden superior
y el teorema de incompletud de Gödel*

65. La lógica de predicados de orden superior	187
66. La gödelización	189
67. El teorema de incompletud de Gödel	192
68. Sobre la significación del teorema de incompletud de Gödel para la teoría de la ciencia	197
69. El teorema de incompletud de Gödel y el programa de Hilbert	199

Capítulo XIII. — *Decidibilidad en la lógica elemental*

70. Decidibilidad de la lógica de enunciados	201
71. La técnica de las tablas veritativas	202
72. Reducción del número de funciones veritativas diádicas	204
73. Uso de las formas normales como técnicas de decisión en la lógica de enunciados	208
74. Decisión abreviada por "reducción al absurdo"	209
75. Decidibilidad de las expresiones monádicas de la lógica de pre- dicados	210

Sección segunda

EL ALCANCE ANALÍTICO DEL CÁLCULO LÓGICO

Capítulo XIV. — *Lógica de clases*

76. El álgebra de clases	216
77. La lógica general de clases	223
78. Clase nula y clase universal	226
79. El principio de abstracción y la paradoja de Russell	228

80. Algunos conceptos fundamentales de la aritmética: números car- dinales; pares ordenados	231
81. Morfología de la lógica de clases	232

Capítulo XV. — *Lógica de relaciones*

82. Las propiedades poliádicas como relaciones	234
83. Relaciones y clases	235
84. Functores de relaciones. Conversas	237
85. Productos y potencias relacionales	238
86. Dominios y campo. Descripciones relacionales	241
87. Clases de relaciones diádicas	245
88. Relaciones y clases de equivalencia	251
89. Univocidad. Isomorfía. Estructuras	253
90. Algunos conceptos matemáticos fundamentales: número cardinal; serie; función	259

PARTE CUARTA

LÓGICA FORMAL Y METODOLOGÍA

Capítulo XVI. — *La división y la definición*

91. Lógica formal y metodología	267
92. Noción metodológica de la división	268
93. Disyuntividad y exclusión mutua	268
94. El fundamento de una división	269
95. Un ejemplo de división	273
96. Noción metodológica de la definición	275
97. Definiciones sintácticas	277
98. Definiciones semánticas	277
99. Definiciones por abstracción. Definiciones creadoras	279

Capítulo XVII. — *El análisis formal de la inducción*

100. El planteamiento clásico del problema de la inducción	282
101. El esquema reductivo de Lukasiewicz	284
102. La relación de inducción	286
103. La posibilidad de una lógica inductiva según Carnap	290
104. Descripciones de estado y ámbitos semánticos	295
105. Funciones lógicas de medición	299
106. Funciones de confirmación	301
107. Idea de una lógica inductiva	303
Relación de teoremas y reglas auxiliares que se encuentran aludidos en lugares distintos de los de su introducción	307
Bibliografía	309
Registro alfabético	313

PARTE PRIMERA

LA LOGICA FORMAL Y LAS CIENCIAS REALES
CATEGORIAS LOGICAS

CAPÍTULO I

NOCIÓN DE LA LÓGICA FORMAL

Los capítulos I-III de la Primera Parte consideran algunas cuestiones de teoría de las ciencias que son interesantes para caracterizar la lógica formal. El capítulo IV ofrece unos instrumentos del análisis lógico.

1. **Cómo se caracteriza una ciencia.** — La lógica formal es una de las ciencias que estudian el conocimiento. No es la única que tiene ese objeto: desde otros puntos de vista estudian también el conocimiento otras disciplinas, como la metodología, la teoría del conocimiento o epistemología, la psicología del conocer. Es corriente que un mismo "objeto material" — que puede ser un conjunto de cosas, o un grupo de hechos o de actividades — sea estudiado a la vez por varias ciencias, sin que éstas dejen de tener por esa coincidencia temas bien distinguibles. Así, por ejemplo, un acto de los que los juristas llaman 'compraventa' es objeto de interés para el derecho, la ciencia económica, la sociología, y sin duda también para la historia, la psicología, etc. Este es un hecho obvio que no llama la atención, y el referirse a él puede parecer ocioso. Pero a pesar de su trivialidad el hecho tiene cierta significación para la teoría de la ciencia, pues permite ver que lo que caracteriza plenamente a una ciencia, una investigación o una teoría no es el objeto material estudiado — seguramente compartido con otras varias ciencias o teorías —, sino el punto de vista desde el cual se va a considerar dicho objeto.

Este punto de vista, que suele llamarse 'objeto formal', es fruto de una abstracción. Abstraer significa aislar mentalmente. Todos los datos aparentemente simples que recibimos en la consciencia resultan complejos para la posterior reflexión: nunca tenemos, por ejemplo, una pura sensación de color, sino que toda percepción de color "contiene" al mismo tiempo, y por lo menos, datos de extensión y de duración, de espacio y de tiempo. Abstraer es tomar alguno o algunos de los elementos del dato (por ejemplo, el color), prescindiendo mentalmente de los demás.

Se habla propiamente de abstracción cuando la operación de tomar

algo del dato y prescindir del resto se realiza a un nivel mental lo suficientemente elevado como para que el acto sea consciente y pueda ser voluntario. Pero la base de esa operación, la selección de noticias del mundo, es un hecho fundamental de la situación del hombre en la realidad. Ya nuestros sentidos operan con sus umbrales de percepción una selección que define la gama de los estímulos y cuyo resultado es el aspecto cotidiano común del mundo, el cual, por ejemplo, no tiene radiación ultravioleta. Por eso se ha dicho a veces, usando de un modo amplio la idea de abstracción, que ya nuestros sentidos abstraen. Lo cierto, en todo caso, es que ya por nuestros sentidos nos encontramos sometidos a la necesidad de conocer el mundo del único modo como puede reflejarlo algo limitado como es nuestra capacidad de conocer, a saber: fragmentaria y simplificadoramente, sin tener consciencia concreta total prácticamente de nada, sino generalmente sólo noticia abstracta.

Probablemente es demasiado vago decir que los sentidos abstraen. Pero no es necesario llegar al conocimiento científico para encontrarse con la abstracción en el sentido propio de una operación intelectual más o menos consciente por la cual se toma del dato complejo alguno o algunos componentes, prescindiendo de los demás. El lenguaje común, vehículo del pensamiento cotidiano, opera a determinados niveles de abstracción. Todo nombre común es un abstracto en sentido lógico, y no sólo los nombres que llama abstractos la gramática, como 'bondad', 'maldad', etc. También 'perro', 'mesa' o 'árbol' son abstractos. El perro no es ningún dato de los sentidos. La idea de perro está obtenida por abstracción, prescindiendo de muchos elementos de ciertos complejos de datos. La imposibilidad en que estamos de percibirlo todo, e incluso de pensar todo lo que llega a nuestra percepción, hace de la abstracción la operación inicial y básica del conocimiento propiamente dicho, es decir, del saber consciente y susceptible de comunicación y discusión.

Los términos 'perro', 'mesa', 'árbol' no representan directamente ningún dato de los sentidos (sino una abstracción) cuando se *usan* como nombres completos y refiriéndose a su significado, como, por ejemplo, en la expresión:

el perro es un vertebrado.

Cuando se *usan* como parte de un nombre propio, el nombre propio completo puede representar directamente un dato de los sentidos. Por ejemplo:

este perro es de Juan.

Un término puede también representarse a sí mismo como dato de los sentidos. En este caso se dice que ese término es a la vez *usado* y *mencionado*. Por ejemplo:

perro es bisílabo.

Se ha convenido en que lo correcto es escribir entonces:

'perro' es bisílabo.

No obstante, prescindiremos de las comillas simples cuando éstas deben aplicarse a expresiones que, por destacarlas, se pongan solas en una línea.

Por último, un término puede ser *mencionado* y *no usado*, como ocurre con el término 'hombre' en el ejemplo siguiente:

la palabra castellana que significa *Homo sapiens* es bisílaba.

Esta terminología ('uso' y 'mención') procede de W. V. O. Quine.

La abstracción científica se diferencia de la vulgar en que se orienta por los fines de la investigación, por el intento sistemático y crítico de reflejar lo más adecuadamente posible la realidad. Como la realidad rebasa ampliamente por su complejidad la capacidad de conocimiento del hombre, ocurre paradójicamente que la abstracción científica, para reflejar cada vez más adecuadamente la realidad, tiene que adoptar frecuentemente formas muy artificiosas desde el punto de vista del sentido común, pero a no ser la razón científica más que sentido común críticamente afinado. La capacidad vulgar de abstracción está en efecto educada en las abstracciones aprendidas con el lenguaje étnico materno. El lenguaje étnico no usa más capacidad de abstracción que la necesaria para tratar con las cosas macroscópicas que nos rodean en la vida cotidiana. (A cambio de ello, el lenguaje común es mucho más apto que el científico para recoger experiencias directas de la vida subjetiva.) La ciencia, en cambio, tiene que enfrentarse con zonas de la realidad cuyo acceso es más difícil. Por eso tiene también que construirse niveles y tipos "artificiales" de abstracción.

La abstracción — y especialmente la científica — no es, en efecto, un hecho exclusivamente pasivo, una recepción de elementos que estuvieran naturalmente separados en la realidad. Para distinguir, entre los elementos o aspectos de un dato complejo, aquellos que son interesantes o importantes de aquellos que no lo son, hay que ir guiados ya por ideas previas, por abstracciones previas: nada es importante sino desde algún punto de vista. Así, por ejemplo, el hecho de que ante el acto de compraventa el economista se sienta atraído por unos aspectos y el jurista por otros no se explica exclusivamente por una receptividad pasiva, sino por la suma de esa receptividad pasiva más la circunstancia de que uno y otro han puesto activamente unos puntos de vista selectivos que dejan pasar unos aspectos de la realidad y excluyen otros. Igual que al filtrar un líquido, en la abstracción hay siempre una intervención activa del hombre, algo así como la posición, más o menos consciente y voluntaria, de un filtro.

Esto es muy visible en el uso científico de la abstracción; pero la situación es esencialmente la misma en todo caso. Tampoco las abstracciones propias del len-

guaje étnico de cada cual son naturales, pues ningún lenguaje es natural como pueda serlo un animal o una planta. El lenguaje es un producto de la cultura, al mismo tiempo que una condición de la misma. Así también el practicar una abstracción presupone la existencia de otras abstracciones previas que guían la operación, determinando el punto de vista de ésta.

En la filosofía de la cultura y del conocimiento se presentan frecuentemente problemas como éstos, que recuerdan el de la gallina y el huevo: ¿qué fue antes, una actividad abstractiva — que presupone abstracciones previas —, o unas abstracciones “previas” — que presuponen alguna operación abstractiva —? la solución de estos problemas no se encuentra quedándose dentro del estudio del conocimiento como cosa aislada, sino que debe hallarse probablemente por una vía en la que colaboran varias disciplinas, biológicas, fisiológicas, psicológicas e históricas: el conocimiento no es toda la vida del hombre; el “primer” trato de la especie humana con la realidad exterior no ha sido seguramente teórico, sino práctico, más o menos como el de los demás animales superiores. El “conocimiento” del mundo que tienen éstos está seguramente orientado por las necesidades de su actividad biológica. Ahora bien: la actividad biológica más específica del hombre es el trabajo, que da de sí no sólo la estricta conservación o reproducción de la vida, sino, además, la producción de las condiciones y medios de la vida. La actividad precisamente intelectual de la abstracción ha ido presumiblemente precedida por esa actividad práctica, cuyo desarrollo y fijación, a través de la formación de esquemas o “estereotipos dinámicos” hereditarios en la corteza cerebral (PÁVLOV), habrá suministrado al hombre las primeras “abstracciones”, orientadoras de las sucesivas y ya conscientes.

Aunque ninguna abstracción sea natural (espontánea) ni totalmente pasiva, sino siempre fruto de receptividad y actividad, puede decirse que el pensamiento científico tiene la peculiaridad de hacer conscientemente de dicha artificialidad de la abstracción un instrumento de conocimiento. Así lo hace ya el pensamiento científico, una vez conseguida alguna apreciable cantidad de conocimiento, para establecer el punto de vista, el objeto formal o las abstracciones básicas que caracterizan plenamente a una ciencia, investigación o teoría. Su análisis del objeto a estudiar es más agudo y menos “natural” que el análisis “espontáneamente” hecho por el pensamiento común según las abstracciones de los lenguajes étnicos. (Las dos palabras entrecuilladas no deben, como se ha visto, tomarse al pie de la letra.) Así, por ejemplo, el criterio científico que permite definir a un vegetal por la propiedad de ser autótrofo (suponiendo que baste ese criterio), o el de definición de una sociedad (la feudal, por ejemplo) por su estructura, por la presencia de determinadas combinaciones de clases, instituciones, roles, etc., no coinciden exactamente con las correspondientes representaciones del lenguaje común. Los abstractos científicos ‘vegetal’ y ‘estructura social’ están contruidos bastante artificialmente, no sólo eliminando rasgos de los vegetales concretos y de las sociedades concretas, sino, además, componiendo o sintetizando los rasgos recogidos según cierto orden de importancia.

Puede observarse que los abstractos ‘vegetal’ y ‘estructura social’ son dos ejemplos de lo que hemos llamado ‘abstracciones básicas’ caracterizadoras de una ciencia o teoría: el primero caracteriza la botánica; el segundo, la teoría de la estructura social.

Una ciencia, una investigación o una teoría se caracterizan por las especiales abstracciones básicas que constituyen el punto de vista según el cual se va a considerar la realidad.

2. La falacia de abstracción. — El abstracto suele presentarse con mucho aspecto de realidad, como si fuera algo real. Ello no debe extrañar, porque la abstracción preside todas las vías racionales de penetración en la realidad. Es seguro, por ejemplo, que un historiador que haya hecho estudios de estructura de la sociedad medieval francesa conoce más profundamente lo que concretamente fue la vida de los franceses del siglo XII que quien no posea esos estudios. Pero esto no debe hacer olvidar que lo que llamamos ‘estructura social feudal de Francia’ no ha sido nunca una cosa concreta como lo fue un francés de aquella época. Un abstracto no existe como cosa independiente: razonar como si efectivamente existiera es cometer una falacia. El abstracto convalidado por la crítica científica (por el análisis lógico y la experiencia o el experimento), aunque generalmente destinado a ser superado por el progreso de la ciencia, significa rasgos importantes de la realidad concreta. Pero siempre se trata de rasgos más o menos aislados y reconstruidos.

3. La abstracción básica de la lógica formal. — Se ha indicado ya que el objeto material de la lógica formal, la cosa que estudia, es el conocimiento. También podría decirse — y efectivamente se ha dicho — que es el pensamiento. Pero si se conviene en no usar esta palabra, puede empezarse ya a hacer una distinción entre la lógica formal y otra disciplina vecina: la psicología, que estudia el conocimiento en tanto que actividad subjetiva, o sea, en tanto que pensamiento.

La lógica no se interesa por la actividad de conocer, sino por su resultado, lo que llamamos conocimiento, el cual se encuentra normalmente fijado en un lenguaje. Esto no quiere decir que los trabajos de la psicología del conocer carezcan siempre de interés para la lógica formal, ni, a la inversa, que sean siempre irrelevantes para la psicología los resultados de la lógica formal. Pero, como se ha dicho antes, el conocer humano tiene que penetrar en la realidad mediante abstracciones, y el mantener estricta y rigurosamente éstas es una condición necesaria de la claridad científica. Por eso desde el punto de vista de la lógica formal son irrelevantes los problemas referentes a la actividad de conocer y a la génesis misma, por ejemplo, de la ciencia de la lógica.

Lo que importa a la lógica es el resultado de la actividad de conocer: el conocimiento. Y un resultado interesa por su solidez, por su validez, verdad o fundamentación. Daremos un paso más hacia la abstracción básica de la lógica formal reduciendo su tema a la validez o fundamenta-

ción del conocimiento. Pero con eso no basta aún para precisar el tema de la lógica formal, pues hay otras disciplinas más que se interesan por el mismo tema: la teoría del conocimiento y la metodología, por ejemplo. Con la reducción del tema a la fundamentación del conocimiento no hemos eliminado de nuestra lista inicial más que la psicología del conocer. Los siguientes ejemplos 1a y 2a pueden ayudarnos ahora a precisar los puntos de vista diferentes según los cuales se interesan por la fundamentación del conocimiento la lógica formal y esas otras disciplinas a las cuales englobaremos a partir de ahora bajo su nombre más clásico: 'teoría del conocimiento'.

Ejemplo 1a

Todos los árboles son vegetales;
el manzano del jardín es un árbol;
luego el manzano del jardín es vegetal.

Ejemplo 2a

Lo que está en el centro de un círculo
lo es inmóvil;
la Tierra está en el centro de un círculo;
luego la Tierra es inmóvil.

Desde el punto de vista de la teoría del conocimiento, la argumentación 1a da una *conclusión* fundada. En cambio, la argumentación 2a (que es un razonamiento anticopernicano del teólogo Melancton) no da una conclusión fundada. Desde el punto de vista de la lógica formal, por el contrario, las dos conclusiones están igualmente fundamentadas, y las dos argumentaciones son igualmente válidas. Para dar un nombre a esa "igualdad", que contiene el punto de vista o abstracción básica de la lógica formal, diremos que las dos argumentaciones son *formalmente válidas*, o que las dos conclusiones están *formalmente fundamentadas*. También puede decirse que los dos conjuntos de oraciones o *enunciados*, 1a y 2a, tienen la misma *forma lógica*.

La abstracción básica de la lógica formal es la noción de forma lógica. Su punto de vista es el de la validez o fundamentación de lo formal del conocimiento.

4. Formas de pensamiento irrelevantes para la lógica formal clásica. — El atenernos al concepto de conocimiento al describir el tema de la lógica, en vez de recurrir al más amplio concepto de pensamiento, tiene, entre otras ventajas, la de recoger el tradicional desinterés de los lógicos por las formas de lenguaje que son más aptas para manifestar estados del sujeto que para indicar hechos conocidos. Formas de lenguaje no directamente indicativas de hechos conocidos, sino manifestativas de estados del sujeto, son, por ejemplo, las exclamaciones, los imperativos, los ruegos, las interrogaciones. La limitación de la lógica formal a las formas indicativas o *apofánticas* del lenguaje procede de Aristóteles (384-322 a. n. e.).

El motivo de esa exclusión es que no puede decirse, por ejemplo, que

una interrogación sea válida, verdadera o fundada como conocimiento. La fundamentación de un enunciado indicativo de conocimiento es otro enunciado indicativo de conocimiento, o bien un conjunto de tales enunciados a los que directamente refiere el primero. En el caso extremo de un enunciado empírico individual, su fundamentación es algún estado del mundo externo observable en principio por todos los hombres de modos equivalentes. En cambio, la justificación de una exclamación o interrogación, su inteligibilidad, se desprende de algún acontecimiento al cual no refiere la exclamación o interrogación directamente, y que no es tampoco accesible a todos los hombres del mismo modo.

El mismo punto de vista excluye de la lógica clásica las argumentaciones que no tienden a fundamentar una verdad, sino a persuadir a un auditorio. Estas argumentaciones, que constituyen una parte del estilo didáctico mayor de lo que suele creerse, son tradicionalmente objeto de la retórica.

Sin embargo, los métodos de la lógica moderna pueden aplicarse con fruto a elementos no apofánticos del discurso cotidiano, y entre los lógicos contemporáneos hay autores que subrayan el interés formal de los razonamientos suasorios (= retóricos).

5. La forma lógica. Esquemas. — El concepto de forma lógica — como su contrario, el de contenido o materia significativa del conocimiento — es una noción muy básica que, al igual que muchas otras nociones fundamentales, resulta imposible definir como no sea por medio de una convención, es decir, poniéndose explícitamente de acuerdo sobre una determinada aclaración y precisión de la vaga idea intuitiva que se suscitó en 3. La noción de forma lógica es una de esas abstracciones "artificiales" o científicas que tienen que sintetizarse (fabricarse, por así decirlo) de acuerdo con las intenciones de la ciencia. Esta abstracción artificial o técnica de forma lógica puede establecerse cómodamente distinguiendo, en una expresión dada, entre elementos representativos de la forma y elementos que representan el contenido empírico. Con esto, ciertamente, atendemos, más que al conocimiento, al lenguaje en que se expresa el conocimiento. Pero ésta no es una desviación importante, pues no existe real conocimiento objetivo que no esté articulado en un lenguaje: la objetividad exige la intersubjetividad, como suele decirse, esto es, la posibilidad de comunicación lo suficientemente inequívoca como para que otras personas puedan controlar lo que una considera conocimiento. Y la intersubjetividad exige a su vez la articulación del conocimiento en lenguaje.

Los anteriores ejemplos 1a y 2a — que son, según dijimos, formalmente idénticos — pueden ofrecer una buena vía de acceso a la noción de forma lógica. Cada uno de ellos es un conjunto de tres enunciados. Puesto que lo dicho en el ejemplo 1a redundaba en una conclusión verdadera, mientras que lo dicho en el ejemplo 2a culmina en una conclusión falsa, la validez formal

que 1a y 2a tienen en común no puede ser lo que dicen sus contenidos, lo que declaran sobre la realidad. Por tanto, para llegar a la forma común de 1a y 2a habrá que eliminar los contenidos. Pero ¿cómo eliminar esa materia empírica para quedarnos con la forma?

Consideremos más detenidamente los dos conjuntos de enunciados. Es claro, por de pronto, que no son meros agregados; hay entre los tres enunciados de cada conjunto una relación bastante fácil de percibir: cada tercer enunciado vale a condición de que valgan los dos anteriores. En la lengua común solemos expresar la condición mediante la conjunción 'si', y luego anteponemos a la afirmación condicionada (apódosis) algún signo, como una coma, o la palabra 'entonces', que usaremos aquí para separar más visiblemente la prótasis de la apódosis. Volvamos a escribir los ejemplos en este nuevo estilo:

1b

Si todos los árboles son vegetales;
si el manzano del jardín es un árbol;
entonces el manzano del jardín es
vegetal.

2b

Si lo que está en el centro de un
círculo es inmóvil;
si la Tierra está en el centro de un
círculo;
entonces la Tierra es inmóvil.

Seguramente podríamos mejorar nuestro nuevo estilo suprimiendo los puntos y comas entre los enunciados primero y segundo (*premisas*) de cada conjunto, y escribiendo en su lugar la conjunción 'y'. Pues es claro que la condición o prótasis está constituida en cada caso por el producto de esos dos enunciados, y no por su existencia separada, inconexa. Por otra parte, delante de la palabra 'entonces' no es necesario, en contextos como 1b y 2b, escribir un punto y coma. Si utilizamos el tercer estilo que resulta de estas observaciones para transcribir una vez más los dos ejemplos, observaremos que no hay ninguna necesidad de disponer cada conjunto de enunciados en tres líneas. Cada uno de esos conjuntos es más bien un solo enunciado compuesto (*molecular*) de enunciados simples (*atómicos*). Escribiremos consiguientemente (poniendo en mayúsculas pequeñas los términos nuevamente introducidos):

1c

Si todos los árboles son vegetales y si el manzano del jardín es un árbol, ENTONCES el manzano del jardín es vegetal.

2c

Si lo que está en el centro de un círculo es inmóvil y si la Tierra está en el centro de un círculo, ENTONCES la Tierra es inmóvil.

Escritos los ejemplos en este tercer estilo, se ve qué puede significar el eliminar, como antes decíamos, los elementos empíricamente interesantes

de esos enunciados moleculares. Se observará que ambos tienen unos cuantos elementos comunes y situados en lugares homólogos: 'SI', 'Y', ';;', 'ENTONCES'. Además, esos elementos comunes no significan nada preciso por sí mismos, a diferencia de los elementos que varían en los dos ejemplos:

'árboles': 'lo que está en el centro de un círculo';

'vegetales': 'inmóvil';

'el manzano del jardín': 'la Tierra'.

Los términos que no significan nada preciso por sí mismos, sino sólo junto con otros términos, se llaman '*sincategoremáticos*' (palabra griega que significa con-significativos, significativos con otros). Los demás se llaman '*categoremáticos*'. Los sustantivos correspondientes son '*sincategoremas*' y '*categoremas*'.

En los ejemplos 1c y 2c hay otros términos sincategoremáticos además de 'SI', 'Y', ';;', 'ENTONCES'. Pero por el momento no atenderemos a ellos.

Un lógico medieval llamado Buridán (m. hacia 1360) concibió la forma lógica como el ensamblamiento de los términos sincategoremáticos. Podemos practicar ahora esa abstracción de Buridán, eliminando de nuestros ejemplos los elementos categoremáticos, y quedándonos con un esquema compuesto por los términos sincategoremáticos y unos trazos que recuerden que entre ellos había enunciados. Obtendremos:

Esquema 1d

SI ——— Y SI ——— , ENTONCES ——— .

Esquema 2d

SI ——— Y SI ——— , ENTONCES ——— .

Como se ve, los dos esquemas son idénticos. No son propiamente dos esquemas, sino sólo uno. Consideraremos a ese esquema, llamándole 'Esquema I', como representante de la noción, cualitativa y menos controlable, de forma lógica común a los ejemplos 1a y 2a. Con esto obtenemos el principal fruto de la consideración lingüística del conocimiento.

El Esquema I representa una forma lógica común a 1a y 2a, no la única forma lógica que tienen en común dichos ejemplos. Según la finura del análisis pueden, en efecto, atribuirse a una misma expresión varias formas lógicas. El análisis que ha llevado al Esquema I es muy rudimentario.

La inteligibilidad del Esquema I requiere que no se borren los trazos por los que aludimos al contenido de los textos iniciales. Si se suprimen esos trazos (o los espacios en blanco equivalentes) no se llega a ningún esquema útil, sino a una configuración ininteligible:

SI Y SI, ENTONCES.

Hay que aclarar, por tanto, lo que quiere decir que el punto de vista de la lógica formal prescinde de todo contenido. La realidad es que prescinde de todo contenido empírico, pero no de la idea de contenido en general.

Ello se debe a que la forma — que es una abstracción — no sólo no puede darse sin un contenido concreto, sino que, además, no puede siquiera pensarse sin pensar al mismo tiempo en un contenido en general, o, dicho de otro modo, en el abstracto 'contenido'. Forma y contenido, o forma y materia, son dos conceptos que se necesitan el uno al otro: son dos "opuestos dialécticos".

6. Verdad formal. Esquemas finales. La lógica formal como ciencia o teoría. — El Esquema I no es ni verdadero ni falso. Será verdadero o falso el resultado de sustituir los trazos por enunciados, por contenidos adecuados. Eso quiere decir que el Esquema I no es *formalmente* verdadero ni falso. Los resultados de su relleno, de la *interpretación* de sus huecos, serán *materialmente* verdaderos o falsos (casos, por ejemplo, de *1a* y *2a* respectivamente).

Pero ya antes hemos indicado que el análisis de los ejemplos *1a* y *2a* que ha llevado al Esquema I es muy rudimentario, y que en *1a* y *2a* hay en realidad más términos sincategoremáticos que los antes recogidos: son el adjetivo 'todos', el artículo determinado 'EL' (o 'LA'), y la forma del verbo 'ES' (o 'SON'). En *1a* está explícito el término 'todos'. En *2a* va implícito en la expresión 'lo que'. Pues evidentemente Melanchton quiso decir que toda cosa que está en el centro de un círculo es inmóvil. La rareza de la construcción se debe a la falta de una sencilla palabra — como 'árboles' en el lugar homólogo de *1a* — para designar la clase o el conjunto de las cosas que están en el centro de un círculo. Curiosamente, esa deficiencia del lenguaje hace que la expresión sea lógicamente más exacta en *2a* que en *1a*. 'Árbol', en efecto, es un abstracto, es el nombre de una clase de cosas, y no es correcto hablar de un abstracto diciendo que *es* de tal o cual modo, en el sentido corriente en que *son* las cosas reales. En realidad, 'todos los árboles son vegetales' quiere decir propiamente:

todas las cosas concretas que tienen la propiedad árbol tienen la propiedad vegetal.

O también: 'si una cosa es árbol, entonces es vegetal'. Por el momento escribiremos simplemente, con mayor adecuación al lenguaje vivo:

Todas las cosas que son árboles son vegetales.

De este modo lógicamente correcto se expresa el primer enunciado de *2a*, gracias precisamente a que la falta de un abstracto simple, como 'árboles',

ha impuesto el uso de un nombre abstracto compuesto ('lo que está en el centro de un círculo') y ha permitido percibir más exactamente la situación lógica:

Todas las cosas que están en el centro de un círculo son inmóviles.

Por simetría de formulación con *1a* introduciremos el abstracto simple 'central', con la significación de 'tener la propiedad de estar en el centro de un círculo'. Así tenemos:

Todas las cosas que son centrales son inmóviles.

Para evitar tener que usar unas veces 'es' y otras veces 'son', usaremos esta versión definitiva de los dos enunciados, que vale también para las otras dos premisas:

Todo lo que es árbol es vegetal;
Todo lo que es central es inmóvil.

En un nuevo estilo basado en las anteriores reflexiones, nuestros ejemplos, con algunos retoques de comodidad que no requieren comentario, se convierten en:

1e

SI TODO LO QUE ES árbol ES vegetal Y SI EL manzano del jardín ES árbol, ENTONCES EL manzano del jardín ES vegetal.

2e

SI TODO LO QUE ES central ES inmóvil Y SI LA Tierra ES central, ENTONCES LA Tierra ES inmóvil.

De *1e* y *2e* se desprende el Esquema II, como antes el Esquema I de *1d* y *2d*:

Esquema II

SI TODO LO QUE ES ① ES ② Y SI EL (LA) ③ ES ①, ENTONCES EL (LA) ③ ES ②.

El esquema II es verdadero independientemente de los términos categoremáticos (ahora no se trata de enunciados) que se utilicen para rellenar los que llamaremos 'lugares de contenido': ①, ②, ③. Llamaremos 'esquema final' de un enunciado a aquel en el cual aparezcan todos los términos sincategoremáticos de dicho enunciado. El Esquema II es final; el Esquema I no lo es. Con esta terminología podemos precisar:

Un enunciado es formalmente verdadero cuando su esquema final es verdadero para cualquier interpretación de sus lugares de contenido por categoremata.

Un enunciado es formalmente falso cuando su esquema final es falso para cualquier interpretación de sus lugares de contenido por categoremata.

Un enunciado es formalmente consistente o compatible cuando no es formalmente falso, lo que quiere decir que hay al menos una interpretación que hace verdadero a su esquema final o, como suele decirse, lo satisface o cumple.

La arbitrariedad de la interpretación tiene naturalmente la restricción de que, una vez escogida una interpretación determinada para un lugar de contenido con el número n , la ocupación de cualquier otra instancia de ' n ' tiene que hacerse con el mismo categorema. Los lugares de contenido están cualificados, no son lugares indiferentes como los puntos de un plano. Esto se debe a lo que antes dijimos sobre la presencia de contenidos abstractos o indeterminados en la misma consideración formal.

Según las últimas precisiones, los enunciados 1a y 2a son formalmente equivalentes, de acuerdo con la afirmación intuitivamente hecha en 3: son ambos formalmente verdaderos, valen lo mismo; además, tienen el mismo esquema final.

Es interesante observar que el Esquema II tiene la principal propiedad de las proposiciones (o "juicios", según la terminología tradicional) y de los enunciados que las expresan: el poder ser verdaderas o falsas. Por tanto, el Esquema II es en cierto sentido un enunciado, y no sólo un esquema: es un enunciado esquemático; recordando que los esquemas representan formas, podemos también decir que es un enunciado formal. Los enunciados de este género son propios de la lógica formal.

Sabemos de él que es "siempre" verdadero, o sea, verdadero para cualquier interpretación de los lugares de contenido mediante categoremata. Este género de verdad se llama 'verdad formal'. Entre las numerosas verdades formales o enunciados formales verdaderos, hay tres célebres enunciados llamados tradicionalmente 'principios lógicos'. Son susceptibles de formulación como verdades formales en diversos sistemas lógicos. He aquí una formulación, a título de ejemplo de verdades formales (' p ' es una letra esquemática que representa un enunciado cualquiera):

si p , entonces p	("principio de identidad");
no (p y no- p)	("principio de no-contradicción");
p o no- p	("principio de tercio excluso").

La existencia de enunciados formales verdaderos permite concebir la lógica formal como una ciencia de los mismos, como el estudio sistemático de los teoremas formales. Un sistema de teoremas es una teoría. Por eso

la lógica formal puede concebirse como una teoría o ciencia, no como una mera técnica del pensamiento correcto o del conocimiento, concepción que también se ha presentado en la historia, y que discutiremos en 8.

7. Sentido de las verdades formales. — La naturaleza de ciencia o teoría (sistema de teoremas) que acabamos de ver en la lógica plantea una cuestión bastante delicada: ¿de qué tratan los teoremas de la lógica? ¿De qué trata, por ejemplo, el Esquema II? ¿Qué sentido tiene decir que es verdadero? ¿De qué es verdadero?

La respuesta más clásica en la lógica moderna se debe a L. Wittgenstein (1889-1951), y sostiene que las verdades formales no afirman nada ni se refieren a nada; son "tautologías", consisten en decir de varios modos una misma cosa. La conclusión del Esquema II, por ejemplo, no hace más que repetir de otro modo la condición o premisa compuesta. Si a algo se refiere la conclusión, es a la premisa. Y una y otra juntas no se refieren a nada más.

Es claro que las verdades formales no dicen nada concreto, no significan nada particular de ninguna cosa concreta en particular. No enriquecen nuestro conocimiento del mundo real, el cual se compone de cosas concretas.

Por otra parte, si esa respuesta agotara la cuestión no tendría sentido decir que un enunciado formal es verdadero o falso. La principal definición de verdad realmente utilizada por la ciencia es la definición de origen aristotélico, según la cual la verdad es la adecuación con la realidad. Si un enunciado formal no dijera en ningún sentido absolutamente nada sobre la realidad, entonces no tendría tampoco sentido discutir acerca de su verdad o falsedad, y, consiguientemente, tampoco tendría sentido concebir la lógica como una ciencia del conocimiento. Sin duda puede darse también ese paso y eliminar de la noción de lógica la idea de verdad. Es posible, en efecto, construir sistemas de formas sin apelar a la noción de verdad, sino sólo al respeto de ciertas reglas. También el ajedrez, por ejemplo, es un sistema de reglas, con movimientos válidos y otros incorrectos, sin que sus reglas tengan que interpretarse en términos de verdad y falsedad. Así también para la construcción misma de sistemas formales es en principio posible prescindir de la noción de verdad.

Pero con eso, como indica el ejemplo del ajedrez, los sistemas de lógica perderán definitivamente la naturaleza de lenguaje. Y cuando lo que interesa es conseguir una noción de la lógica formal como ciencia, como rama del conocimiento, en relación con las demás ramas del mismo, esa solución del problema resulta muy pobre: más que una solución es una negativa a tener en cuenta dicho problema. Desde un punto de vista filosófico, de teoría de la ciencia, es seguramente más fecundo no prescindir así del problema, sino seguir preguntándose por la significación de las verdades formales, por escasa que esa significación pueda ser.

Ahora bien: toda ciencia real y todo acto logrado de conocimiento

vulgar respetan las verdades formales. Los teoremas formales no componen una teoría de ninguna realidad concreta, pero, en cambio, toda ciencia de realidad concreta los respeta.

Ha ocurrido en la historia de la ciencia que se descubrieran verdades materiales a través de razonamientos formalmente falsos, es decir, que violaban algún teorema formal. Pero siempre que esto ocurre se trata de un error formal subjetivo del científico, no de que para demostrar la verdad material en cuestión haya que pasar por la falsedad formal. Antes al contrario, aquella verdad material no queda fundada sistemáticamente como elemento de una teoría hasta que se corrige el error formal.

Así pues, debajo de la verdad material teórica — es decir, que sea algo más que mera percepción suelta — hay siempre verdad formal. Esto permite concebir la lógica formal, el sistema de los teoremas formales, como una determinación de las leyes más generales del comportamiento de los objetos estudiados por las ciencias o teorías. Las verdades formales darían las condiciones mínimas puestas a los objetos del conocimiento en tanto que objetos del conocimiento.

Esta interpretación de las verdades formales no es, naturalmente, un teorema de la lógica formal, ni nada interno al sistema formal mismo, al sistema de los teoremas formales. Es más bien una reflexión filosófica, una consideración de filosofía de la ciencia que arroja alguna luz sobre las relaciones entre la lógica formal y las ciencias reales, y también sobre las relaciones entre la lógica formal y la realidad, a través de las ciencias reales.

La lógica formal tiene pues un carácter elemental y fundamental para las ciencias empíricas y positivas en general. Pero esto no nos autoriza a pretender que las verdades formales sean de un origen completamente distinto del de los teoremas empíricos. Sin duda el carácter muy abstracto de las verdades formales las pone a cubierto de los frecuentes cambios que experimentan las ciencias más próximas a lo concreto. Pero puesto que nunca se ha probado la existencia de una fuente del conocimiento que no sea la percepción, las noticias mediadas por los sentidos, hay que pensar que las verdades formales tienen el mismo origen remoto que cualesquiera otras, y no multiplicar arbitrariamente las fuentes del conocimiento. Si las verdades formales tuvieran otro origen que la experiencia — si fueran fruto, como se dice, de una capacidad de conocimiento *a priori* —, entonces no se vería por qué no han sido todas conocidas desde siempre por todos los hombres: no se vería por qué la lógica tiene una historia, se ha enriquecido en el curso del tiempo. La persistencia, cualitativamente mayor, de los teoremas formales respecto de la mayoría de los enunciados materiales no tiene que ver con la cuestión de su origen, sino con la de su modo de validez. Los teoremas de la lógica valen, en efecto, en virtud de ciertas afirmaciones iniciales y unas definiciones. Por eso ninguna contraprueba empírica puede falsarlos, sino sólo mostrar la inadecuación de dichas afirmaciones iniciales,

sugiriendo así el abandono de éstas. En esa capacidad que tiene la experiencia de “sugerir” cosas acerca de las proposiciones iniciales se revela el *origen* común de la lógica con cualquier otro conocimiento. Y en la incapacidad que tiene la experiencia de refutar cualquier teorema de cualquier sistema de lógica se manifiesta el *modo de validez* peculiar de ésta y distinto del empírico.

La tendencia a confundir la efectiva *validez a priori* (o sea, por definiciones) de los teoremas lógicos con un supuesto *origen a priori* de los mismos es, sin duda, más “natural” y compatible con los prejuicios tradicionales, y puede llevar a considerar la lógica como algo “sabido desde siempre” y nacido así entero y de una vez de la cabeza del hombre. Esta tendencia se encuentra entre los más grandes filósofos, incluso entre los más críticos. Así escribía, por ejemplo, KANT en 1781: la lógica “no ha podido dar [desde ARISTÓTELES] un solo paso adelante, y por tanto parece según toda verosimilitud conclusa y perfecta”. Menos de cien años después la lógica formal iniciaba una nueva fase en la cual el sistema de la lógica aristotélica queda reducido a un manojo de teoremas elementales.

8. La lógica como arte o técnica. — El hecho de que los teoremas formales se apliquen (aunque elemental y abstractamente) a cualquier campo del conocimiento puede explicar el que en la historia se haya concebido frecuentemente la lógica como una técnica para el descubrimiento de la verdad en general. Esta concepción ha ido a menudo junta con la idea, básicamente más acertada, de que la lógica es una ciencia o teoría. Así, por ejemplo, Pedro Hispano (m. 1277) escribía que la lógica formal es “el arte de las artes y la ciencia de las ciencias”, y Tomás de Aquino (1225-1274), que ha entendido en el fondo la lógica como una teoría, escribía también: “la lógica es el arte directivo del acto mismo de la razón, arte por el cual el hombre, en [dicho] acto de la razón, procede ordenadamente, fácilmente y sin error”. Del mismo modo que el cálculo matemático — que procede también de una teoría — se aplica como técnica heurística, o de descubrimiento, a las ciencias empíricas, así también el cálculo lógico se aplicaría a todas ellas.

Esto es sin duda trivialmente verdad, pues toda teoría muy amplia puede suministrar a teorías más reducidas, a las que “contenga” en algún sentido, teoremas generales utilizables como reglas (o sea, operativamente, técnicamente) en estas ciencias de campo más reducido. Pero cuando así se pasa a la concepción de las teorías como métodos heurísticos, lo que importa es su fecundidad. Y precisamente por ser aplicables a cualquier configuración u objeto, los teoremas de la lógica formal no suministrarán nunca por sí mismos información específica sobre nada en particular.

La lógica formal no es un método de descubrimiento de la verdad empírica. Puede ayudar indirectamente a descubrir y a precisar verdad empírica. Pero precisamente en la época en que dominó la idea de lógica como “arte directiva del acto de la razón”, la lógica formal fue muy estéril para

la ciencia positiva, y hasta una rémora de la misma; al menos, su teoría era mucho más pobre que los recursos formales efectivamente en poder de los científicos. Por eso su prestigio como técnica del conocimiento fue incluso perjudicial para la ciencia, como puede ejemplificar el razonamiento anti-copernicano de Melanchton (ejemplo 2a de 3).

El hecho de que la lógica clásica haya sido estéril para la ciencia se debe sobre todo a dos motivos: a que se creyó erróneamente que su utilidad principal para la ciencia tenía que estar por el lado del descubrimiento de nuevas verdades a partir de las conocidas, cuando en realidad los servicios que la lógica formal puede prestar a las ciencias se refieren más directamente al análisis, la aclaración y la ordenación de las verdades ya conocidas; y a que el sistema de la lógica clásica era muy elemental y rudimentario, hasta el punto de carecer, por ejemplo, de un tratamiento general de los enunciados de relación, lo que la incapacitaba, entre otras cosas, para recoger adecuadamente los modos de razonamiento matemáticos. En sustancia, el sistema de la lógica clásica no recogía más que los razonamientos que consisten en comparar clases de cosas, como las clases de cosas que son árboles, vegetales, manzanos. Y estos razonamientos son tan sencillos que no ya el científico empírico, sino todo niño que hable discretamente los sabe construir sin necesidad de estudiar lógica.

La historia de las relaciones entre la ciencia moderna (desde el siglo xvi) y la lógica formal empezó con un justificado desprecio de la primera por la segunda, pues las formas de argumentación puestas en práctica por los grandes investigadores de los siglos xvi-xviii eran en realidad ignoradas por los lógicos medievales y clásicos. A esa tendencia despectiva sustituyó luego, en el siglo xix, un interés por revitalizar la lógica aplicándole técnicas matemáticas. Por último, durante la segunda mitad del siglo xix y la primera del siglo xx, los progresos del pensamiento científico, y señaladamente los de la ciencia que más segura parecía, la matemática, tropezaron con amenazadoras contradicciones que volvieron a poner de manifiesto el interés del análisis lógico.

APÉNDICE AL CAPÍTULO I

A. *Textos citados.* — Pedro Hispano: *Textus summularum Petri Hispani...*, I, Proem. — Tomás de Aquino: *In Anal. post.*, I, lect. 1. — Kant: *Crítica de la Razón Pura*, 2.^a ed., p. VIII de la paginación original.

B. *Observaciones:*

a 1: sobre la abstracción: la abstracción es un objeto de estudio de la psicología puesto que es una actividad mental. Por eso algunos autores, sobre todo K. R. Popper (*La lógica de la Investigación científica*, trad. castellana de Víctor Sánchez de Zavala, Madrid 1962) no la consideran tema propio de la teoría de la ciencia. En el presente capítulo se ha hecho una básica referencia a la abstracción por dos motivos: primero, porque incluso en ese aspecto psicológico es de interés para

familiarizarse con la noción de lógica formal; segundo y principal, porque la abstracción tiene también un aspecto relevante para la teoría de la ciencia y hasta para la lógica formal. La abstracción es un concepto propiamente lógico-formal, por ejemplo, cuando se da (muy naturalmente) ese nombre al principio lógico que legitima la formación de clases a partir de propiedades de individuos.

a 6: sobre terminología: puede provocar confusiones la discrepancia entre la terminología tradicional y la contemporánea (de influencia inglesa) a propósito de enunciados. Como esa discrepancia parece ya ineliminable, es bueno conocer las dos terminologías:

Objeto	Nombre tradicional	Nombre contemporáneo
Atribución (o negación) de algo a (o de) algo	Juicio	Proposición
Enunciado lingüístico de dicha atribución o negación	Proposición	Enunciado, sentencia

CAPÍTULO II

LA LÓGICA FORMAL EN LA INVESTIGACIÓN DE FUNDAMENTOS

9. La cuestión de los fundamentos. — La lógica formal de la tradición quedó, como se ha dicho, muy pronto rebasada por los progresos de la ciencia moderna desde el siglo xvi. Y puesto que los discípulos de Aristóteles, tanto como sus herederos más indirectos, los filósofos de otras escuelas que cultivaban la lógica, concebían a ésta como un instrumento — “órganon” — del descubrimiento científico, se llegó a una situación de desprestigio de la lógica formal entre los científicos positivos.

La versión de la lógica formal aristotélica más común entre los filósofos del siglo xvii, la *Lógica* de los filósofos cartesianos Arnauld y Nicole (1662), quiere ser incluso un órgano de la prudencia y el sentido moral. “Los hombres”, dicen los autores al definir sus intereses, “no han nacido para emplear su tiempo en medir líneas, en examinar las relaciones de los ángulos, en considerar los diversos movimientos de la materia...; en cambio, están obligados a ser justos, equitativos, juiciosos”. Esta desnaturalización del tema de la lógica les lleva a dar reglas de prudencia mientras desprecian problemas formales ya conocidos de Aristóteles. Así dan “algunas reglas para dirigir bien la razón en la creencia en los acontecimientos que dependen de la fe humana”, y hasta suministran una “aplicación de la regla anterior a la creencia en los milagros”.

Esta situación persistió mientras el avance de las ciencias, especialmente de las ciencias de la naturaleza, inspiró, con sus grandes éxitos, una sólida confianza en lo que suele llamarse sus ‘fundamentos’ teóricos.

Los fundamentos de una ciencia son, por un lado, sus conceptos más generales, los cuales recogen e interpretan algunas observaciones o, más frecuentemente, contienen algunas suposiciones, en las que descansan los demás conceptos; y, por otro lado, los razonamientos con que se relacionan los conceptos para integrar con ellos un sistema de proposiciones que expliquen y justifiquen (*fundamenten*) dichos conceptos y las observaciones a que éstos se refieren.

Desde el siglo xvii hasta finales del siglo xix la ciencia empírica recorrió un camino que, tras la crítica y el abandono de los conceptos fundamentales de la física aristotélica y medieval, pareció ser de mera acumulación de conocimientos, aunque en realidad no faltaron en esa ruta violentos cambios de nociones básicas, por ejemplo y señaladamente en el nacimiento de la química a finales del siglo xviii. La acumulación de éxitos disimuló, sin embargo, el carácter crítico y revolucionario de esos momentos.

Pero a fines del siglo xix los científicos notaron una serie de dificultades en el uso teórico de algunos conceptos básicos de la ciencia, al mismo tiempo que parecía fracasar el intento de fundamentar o justificar lógicamente los conceptos elementales de la matemática, tan necesaria para toda la ciencia moderna.

10. La “crisis de fundamentos” de las ciencias. — Desde la mitad del siglo iban vacilando, o hasta caducando, algunos conceptos que habían desempeñado un papel básico en la ciencia moderna. Algunos de esos conceptos eran muy vagos y generales, y habían sido introducidos en la mentalidad científica por filósofos. Así, por ejemplo, la idea de que el único factor del descubrimiento científico es la observación, con olvido del importante trabajo que realiza la imaginación al elaborar hipótesis. Este principio procedía, más que del desarrollo de la ciencia, de la interpretación del mismo por la filosofía empirista, señaladamente por Francis Bacon (1561-1676). El principio había sido muy útil culturalmente, para liberar la cultura moderna de la sumisión medieval a las autoridades, a los antiguos autores. Pero ahora ya la riqueza de datos conocidos exigía, para interpretarlos, conceptos cada vez más complicados, abstracciones cada vez más artificiales, de las que no podía decirse que fueran simples nombres de observaciones. La palabra ‘éter’, por ejemplo, no era (en física) nombre de nada observado, sino de un supuesto imaginado — y muy complicado, por lo demás, pues al final hubo que atribuirle un incoherente manejo de propiedades — considerado necesario para explicar las observaciones.

Lo importante fue que hicieran crisis no sólo conceptos aportados por filósofos, y de los que se pudiera en rigor prescindir en el sistema de la ciencia, sino también nociones originadas en la ciencia misma. La de éter, que acabamos de recordar, puede servir de ejemplo, pues ella protagonizó un experimento que, aunque en una de sus varias repeticiones dio un resultado favorable a la noción, acabó por eliminarla de la ciencia, al establecer la imposibilidad de identificar algún efecto de la supuesta existencia de un éter.

Al mismo tiempo que quedaban sin justificación ni contenido determinados conceptos más o menos básicos, como el de éter, se producía un cambio aún más general: junto al modo de pensar mecanicista propio de la física moderna en los siglos pasados, aparecían o se generalizaban ahora conceptos como el de campo, el cual tiende a presentar los fenómenos no

como resultado de choques mecánicos entre partículas sin cualidades y en un espacio el mismo neutro, sino como manifestación de cierta cualidad o estructura del medio en el cual se nos dan esos fenómenos; un campo magnético, por ejemplo, no es un espacio sin propiedades en el que todo fenómeno se explique por choques, sino que tiene líneas de fuerza, direcciones y sentidos.

Nuevos conceptos, como el de campo, que se basaban en un modo de pensar bastante diverso del que había imperado en la ciencia hasta la fecha, eran adecuados para resolver nuevos problemas, pero su novedad alteraba, tanto como los mismos problemas sin resolver, el repertorio de las nociones básicas de la ciencia: en un primer momento, estas nociones, viejas y nuevas, se presentaban como un conjunto de ideas heterogéneas. En el capítulo siguiente se verá que esa "crisis" afectó también a la matemática, y, a través de ella, a la lógica.

Aquí bastará con hacer dos observaciones: la primera es que la "crisis de fundamentos" no acarrea, naturalmente, la imposibilidad de seguir realizando trabajo de investigación científica, observación, experimentación, generalización por hipótesis interpretativas de los hechos. En realidad, la situación más frecuente en la historia no es la de una gran claridad de la ciencia sobre sus propios fundamentos. La "crisis de fundamentos" no consiste en que nociones básicas hasta el momento seguras se hagan de repente vacilantes, sino en que en un momento dado se descubre que fundamentos tenidos antes por sólidos y claros no lo son ni lo eran. La misma matemática y la mecánica han procedido durante más de doscientos años con el instrumento del cálculo infinitesimal, basado en el oscuro fundamento de lo que se llamaba 'infinitésimo' y era una noción insostenible: la de una "magnitud infinitamente pequeña", o una "fluxión" imperceptible de una magnitud. La crítica externa de la ciencia (es decir, la verificación o falsación de sus resultados) y la práctica (la aplicación técnica de la ciencia) desempeñan, en efecto, un papel importante en el desarrollo de la investigación, y suplen siempre, en mayor o menor medida, la falta de claridad sobre los fundamentos. Pero eso no quita que la necesidad de claridad sobre los mismos se imponga como una condición del progreso ulterior en cuanto que la marcha de la ciencia la hace sentir a los científicos y a toda la cultura.

11. Aspectos materiales y formales de una crisis de fundamentos. —

La preocupación por los fundamentos de una ciencia y la reflexión sobre los mismos presuponen una considerable acumulación de conocimientos en ella, y tienen que ver con el deseo de ordenar esos conocimientos de modo que sea posible distinguir cuáles son las nociones básicas. Cuando es posible satisfacer ese deseo, y los conceptos y enunciados de una ciencia están ordenados como jerárquicamente, de modo que unos sirvan de base a otros, se dice que esa ciencia es una *teoría*.

Suponiendo una ciencia que realice en mayor o menor grado ese concepto o ideal de teoría, una "crisis de fundamentos" en ella puede presentar dos aspectos, que en la realidad están unidos, pero que aquí conviene distinguir artificialmente.

La situación crítica puede estar desencadenada por la aparición de hechos nuevos poco compatibles con las nociones anteriormente consideradas básicas. Como los hechos nuevamente conocidos no anulan, naturalmente, los hechos anteriormente conocidos, que eran explicados con las viejas nociones fundamentales, el problema de fundamentos que entonces se plantea consiste en inventar nuevos conceptos básicos que den a la vez razón de los hechos antiguamente explicados y de los hechos nuevamente conocidos. Este aspecto de la "crisis de fundamentos", al que llamaremos 'material', es el modelo típico (simplificado) del progreso cualitativo de la ciencia, es decir, del progreso en el cual algo que al principio parece una mera acumulación — el descubrimiento de hechos nuevos — resulta acarrear un cambio en la cualidad de la ciencia, en las nociones o significaciones básicas de la misma.

Pero una "crisis de fundamentos" puede también presentarse en una ciencia, llegar a la consciencia de sus cultivadores, porque se descubra, reflexionando sobre ella e independientemente de que aparezca algún espectacular hecho antes ignorado, que el edificio teórico es poco sólido: por ejemplo, que los conceptos y enunciados supuestamente básicos son oscuros y dan lugar a absurdos o contradicciones; o que sólo aparentemente fundamentan otros enunciados derivados tenidos por verdaderos, mientras que en realidad el paso de unos a otros es más sugestivo que racional, etc. Estos son los aspectos de una "crisis de fundamentos" que llamaremos 'formales'. En la práctica han solido darse unidos a los materiales, pero aquí los distinguiremos de ellos por la siguiente razón: la principal aportación de la lógica formal a las ciencias no consiste en suministrarles herramientas para descubrir hechos nuevos o para inventar nuevos conceptos básicos, sino en darles métodos para analizar sus estructuras y sus nociones y afirmaciones fundamentales, con objeto de precisar su claridad, su capacidad de fundar otras afirmaciones, etc.

Esta es la medida en la cual la ayuda de la lógica formal puede interesar al científico positivo, como suministradora de procedimientos que, con los demás métodos de la investigación básica, o investigación de fundamentos, sirven para analizar, aclarar y ordenar los resultados de la investigación fáctica. Hay autores que consideran una innecesaria artificialidad de la lógica el presentarse en la práctica muy abstractamente, como sólo apta para analizar resultados; pero el hecho es que cuando la lógica, en tiempos antiguos y medievales, se presentó como instrumento de la investigación factual, empírica, no dio de sí ningún resultado aplicable al descubrimiento. Antes al contrario, fue a veces, como indicamos, una rémora del mismo.

Las anteriores consideraciones nos presentan a la lógica formal como una rama de la investigación de fundamentos, al mismo tiempo que indican las limitaciones que tiene en ese campo: en el acto del descubrimiento científico intervienen factores psicológicos — como la imaginación —, culturales — las ideas dominantes en el ambiente del científico —, técnicos — los medios materiales a disposición del científico — y económico-sociales — el estado y la organización de los medios de producción, el modo de producción mismo —, los cuales no son, naturalmente, agotables por un análisis meramente formal. El análisis formal tiene su campo de aplicación adecuada en las teorías ya constituidas o en las acumulaciones de conocimientos ya conseguidos que se desee constituir en teorías.

A continuación consideraremos brevemente dos aplicaciones muy generales del análisis formal a teorías factuales o empíricas.

12. La estructura de las teorías. — Una ciencia que sea rica en conocimientos, acaso porque tenga tras de sí una larga historia, suele presentarse en una forma que tiene poco que ver con el modo como se ha ido adquiriendo. Los conocimientos que componen una ciencia habrán empezado a conseguirse, en la mayoría de los casos, de un modo suelto y más o menos casual. Sólo cuando se acumularon notarían sus cultivadores que esos conocimientos tenían alguna homogeneidad por su tema (comparados con otros), y que algunos de ellos eran más básicos que otros y los fundamentaban. El análisis lógico de los conceptos y los enunciados puede intervenir entonces para aclarar las relaciones de fundamentación y derivación entre unos y otros.

Por regla general, la construcción teórica de las ciencias susceptibles de ello no ha sido obra de lógicos especializados, sino de los científicos mismos en función de lógicos, y el análisis ya propio de especialistas se aplica después, cuando se sorprenden dificultades u oscuridades o rarezas en las teorías.

Ejemplo de ambas cosas puede ser la geometría euclídea. Ésta es la primera teoría, en el sentido dicho, que ha existido en nuestra cultura, y no ha sido obra de ningún lógico discípulo de Aristóteles o de las otras escuelas lógicas antiguas, sino de geómetras. Por otra parte, esa teoría ha suscitado un análisis propiamente formal, el cual se proponía saber si un determinado postulado de los que encabezaban la teoría era imprescindible o no.

En estos casos, el análisis lógico se propone unos objetivos que podrían esquematizarse así:

1.º Identificar las nociones fundamentales de la teoría analizada. A veces se les llama 'nociones primitivas', porque de ellas no se da explicación en la teoría, sino que se ponen en ella al principio de cualquier explicación. Aristóteles decía que "no todo conocimiento es demos-

trativo", pues la demostración tiene que basarse en algo, y aquello con lo cual se empieza no puede basarse en nada. Análogamente, tampoco se puede aclarar todo concepto por otro concepto; algunos deben tomarse como primitivos.

- 2.º Precisar cómo se aclaran por las nociones primitivas las demás nociones de una teoría, o, lo que es lo mismo, cómo se construyen o "constituyen" (R. Carnap) las segundas con las primeras.
- 3.º Identificar los enunciados primitivos de la teoría. Estos son los que se renuncia a fundamentar en la teoría y se toman como dados para fundamentar los demás.
- 4.º Precisar cómo se derivan en la teoría unos enunciados a partir de otros y, en última instancia, de los primitivos.

En cada una de esas investigaciones el análisis puede tener resultados críticos. Por ejemplo, a propósito del punto 1.º el análisis puede mostrar que los conceptos primitivos de la teoría estudiada son ambiguos, o insuficientes, o innecesariamente abundantes. A propósito de los puntos 3.º y 4.º puede mostrar que los enunciados primitivos, o los modos de derivación a partir de ellos, son tan laxos que posibilitan la derivación de falsedades, o, por el contrario, que no bastan para derivar algún enunciado interesante que, por otras informaciones, se sabe verdadero.

13. La precisión del objeto formal de una teoría. — Otra aplicación del análisis lógico al estudio de las teorías científicas consiste en precisar lo que en el capítulo I llamamos 'abstracciones básicas' que definen el objeto de las mismas.

Esta utilidad del análisis lógico ha sido poco estudiada. El polaco A. Grzegorzczk le ha dedicado recientemente la atención. Lo que sigue es un resumen de su método.

Una teoría suele compartir su objeto material con otra u otras. El análisis lógico de la teoría en cuestión puede proponerse, entre otras cosas, una vez dado el conjunto M , de objetos materiales individuales sobre los que versa esa teoría, precisar cuáles son las relaciones, $R_1, R_2, \dots, R_n, \dots$, que esa teoría considera entre los miembros de M ; también puede el análisis lógico precisar las funciones, $f_1, f_2, \dots, f_r, \dots$, definidas para los miembros de M , que forma y estudia la teoría considerada. El sistema formado por ese conjunto M , las relaciones y las funciones,

$$\{M, R_1, R_2, \dots, R_n, \dots, f_1, f_2, \dots, f_r, \dots\},$$

puede considerarse como una explicación de la expresión 'objeto formal de la teoría considerada', aclaración de expresiones bastante imprecisas, como 'punto de vista' o 'abstracción básica'.

14. La utilidad heurística de la investigación de fundamentos. — La investigación de fundamentos en general, y aún más en sus aspectos que hemos llamado 'formales', y que son los únicos en que la lógica resulta útil, puede parecer un trabajo bastante bizantino, que no busca más que la elegancia o, a lo sumo, la perfección de las teorías, de lo ya sabido.

Esto es en parte verdad. Pero ese trabajo de reflexión sobre lo ya sabido resulta necesario el día en que la inseguridad o la oscuridad sobre lo ya sabido ponen en peligro el ulterior progreso de la ciencia. Pues para seguir adelante es conveniente tener en lo posible suelo sólido bajo los pies, o bien un claro conocimiento de cuáles son los puntos en que el terreno cede o se interrumpe.

Por aquí se encuentra la indirecta utilidad heurística de la investigación de fundamentos en general y del análisis lógico en particular. Descubrir los puntos oscuros, vacilantes o infundados, del cuerpo del saber poseído, o descubrir huecos en él, no es, desde luego, descubrir hechos de la realidad exterior, pero sí sirve para orientar al investigador hacia las regiones de su ciencia en las que más urge su trabajo. De este modo indirecto, y no mediante la ilusoria producción de razonamientos descubridores de hechos, colabora la lógica, como instrumento de análisis, en la obtención de conocimientos positivos.

El citado ejemplo de la geometría euclídea ilustra esta utilidad heurística indirecta del análisis formal. Pues la crítica búsqueda de si todos los postulados de la geometría euclídea eran o no imprescindibles, sugirió la construcción de las geometrías no-euclidianas, geometrías en las que se altera el postulado cuya imprescindibilidad se discutía; y esas geometrías han resultado luego útiles para la interpretación de datos factuales de las ciencias de la naturaleza.

La lógica ofrece al científico positivo ese instrumental crítico analítico de dos modos:

Primero, presentándose ella misma como teoría, como sistema, dando así un modelo general (aunque muy idealizado) de lo que es una teoría científica, y suministrando al mismo tiempo un repertorio de verdades formales respetadas en toda ciencia.

Segundo, utilizando esas verdades formales como técnicas de análisis de las teorías positivas.

Las partes segunda y tercera de este manual ofrecen una presentación de la lógica como teoría. La parte cuarta se ocupa del análisis formal de algunos modos de razonar y de construir conceptos que son frecuentes en las ciencias reales. En los dos capítulos restantes de esta primera parte se consideran algunas cuestiones referentes a cómo la "crisis de fundamentos" afectó a la lógica, a través de la matemática, y a cómo la lógica se constituyó finalmente en teoría, en el sentido en que esta palabra se ha usado en el presente capítulo.

APÉNDICE AL CAPÍTULO II

A. *Textos citados.* — Arnauld et Nicole, *La logique, ou l'art de penser*, Discurso Primero; Parte 4.^a, cap. XIII; Parte 4.^a, cap. XIV. — Aristóteles, *An. Post.*, A, 3, 72b 18-24.

CAPÍTULO III

EL IDEAL DEL LENGUAJE BIEN HECHO

15. Las paradojas de la teoría de conjuntos. — Las necesidades de la matemática han contribuido considerablemente al nacimiento de la lógica formal contemporánea, orientando hacia ésta numerosos esfuerzos de los matemáticos. Es corriente distinguir dos etapas sucesivas en esas aportaciones de los matemáticos. En la primera de ellas, se tiende a enriquecer a la lógica formal con nociones y modos de análisis y razonamiento de los que carecía, y que son en cambio necesarios en las matemáticas; también se desea aplicar a la lógica los procedimientos de simbolización, de construcción artificial del lenguaje, que eran familiares en matemáticas y en otras ciencias, como la química. Esta etapa, que tenía un importante precursor en G. Leibniz (1646-1716), está representada por la obra de G. Boole (1815-1864) y E. Schröder (1841-1902).

La segunda etapa, cuyos principales representantes son G. Frege (1848-1925) y B. Russell (n. 1872), se sitúa a fines del siglo XIX y principios del XX. La tendencia que domina en ella es aparentemente opuesta a la recién descrita, pues consiste en reducir la matemática elemental (la aritmética) a la lógica. El motivo de este deseo era la creciente sensibilidad de los matemáticos a la oscuridad o vaguedad de algunas nociones fundamentales de su ciencia. La citada noción de infinitésimo, por ejemplo, que había sido básica en el cálculo infinitesimal, iba siendo abandonada. Los conceptos de los diversos tipos de "números" iban aclarándose por reducción al de número natural. La frase de Kronecker según la cual sólo los números naturales han sido creados por el Buen Dios, mientras que todos los demás son sólo útiles artificios contruidos por el hombre a partir de aquéllos, ejemplifica la tendencia de los matemáticos a reducir las nociones menos intuitivas y más complicadas de su ciencia a las más elementales, con objeto de conseguir claridad de fundamentación para las teorías de la matemática superior. Existía en general la aspiración a aritmetizar el análisis matemático, es decir, a fundar la mayor parte posible del mismo en las nociones elementales de la teoría de los números naturales.

En realidad, la tendencia característica de la segunda fase no era tan opuesta a la de la primera como puede parecer. En primer lugar, conservaba de la primera la idea de que la lógica debe servirse de un lenguaje artificial, simbólico y preciso, como el introducido en matemáticas por Viète y los algebristas a partir del siglo XVI. Además, para realizar su programa de reducción de la matemática a la lógica, esta segunda tendencia tenía también que empezar por ampliar el repertorio de las formas lógicas. De no hacerlo así, la reducción sería de todo punto imposible.

Pero el *logicismo* de Frege y Russell presentaba una interesante novedad: el deseo de fundamentos firmes, claros y seguros para la matemática, que se había orientado ya a la aritmetización del análisis, se prolongaba con ellos en una búsqueda de los fundamentos lógicos de la aritmética misma. La idea básica de Frege y Russell es que la aritmética es idéntica con alguna parte de la lógica.

El matemático G. Cantor (1845-1918) había emprendido ya ese camino con su *teoría de conjuntos*, una de cuyas tareas principales es suministrar a la aritmética el concepto de número (cardinal). Un conjunto es según Cantor una colección de objetos cualesquiera, físicos o mentales. Los objetos se llaman 'elementos' o 'miembros' del conjunto, y se dice que pertenecen a él. Un subconjunto es, en cambio, una parte de un conjunto, y se dice que está incluido en él. Dos conjuntos son equivalentes cuando puede establecerse entre sus miembros una correlación biunívoca, o, como también se dice, cuando son coordinables: cuando comparados el uno con el otro puede hacerse corresponder a cada miembro de uno de ellos un miembro, y sólo uno, del otro.

Esos conceptos iniciales de la teoría de Cantor son conceptos lógicos que no trataba la matemática tradicional: el de conjunto es el concepto lógico de extensión de un término; la extensión de un término es la colección de los objetos a los que es aplicable (de los que es predicable) ese término (mientras que la comprensión o intensión de un término es su significación, el "concepto" o contenido de la representación mental que suscita). El concepto de pertenencia de miembro a conjunto corresponde al de predicabilidad del predicado al sujeto; el de inclusión de subconjunto en conjunto corresponde a la relación lógica entre la parte y el todo; y el concepto de equivalencia de conjuntos está construido sin usar la noción de número, sino sólo la noción lógica de correlación.

Con esos conceptos lógicos Cantor define la noción de número cardinal de un conjunto: número cardinal de un conjunto es la propiedad que el conjunto tiene en común con todos los conjuntos que le son equivalentes, y con ellos sólo. Así la noción fundamental de la aritmética quedaba a su vez fundamentada en la lógica.

Uno de los hechos más importantes de la historia de la ciencia moderna ha sido que el propio Cantor, y otros autores antes y después de él, descubrieron que en esa teoría, aparentemente tan sencilla y llamada a funda-

mentar la matemática con conceptos lógicos, era posible construir contradicciones o paradojas.

La paradoja referente a los números cardinales se basa en un teorema elemental, llamado 'teorema de Cantor'. Este teorema dice que el número cardinal del conjunto potencia, CpS , de un conjunto cualquiera, S , es mayor que el número cardinal de S . (Conjunto potencia, CpS , de un conjunto cualquiera, S , es el conjunto de todos los subconjuntos de S , incluido el propio S — como subconjunto impropio — y el conjunto vacío, que es el que no tiene ningún miembro.) Cantor demuestra ese teorema por reducción al absurdo. Pero partiendo de él se puede construir la siguiente contradicción o paradoja (no introducimos, como sería necesario, una notación especial para los números cardinales de los conjuntos, sino que los nombramos por los nombres de sus conjuntos):

Sea T el conjunto de todos los conjuntos. Por el teorema de Cantor, el conjunto potencia de T , CpT , tiene un número cardinal mayor que el de T :

$$(1) \quad CpT > T.$$

Pero CpT es un conjunto (a saber: el conjunto potencia de T). Por tanto, tiene que ser un subconjunto de T , pues T es el conjunto de todos los conjuntos. Eso quiere decir que todos los miembros de CpT son miembros de T :

$$(2) \quad CpT \leq T \quad (\leq \text{ y no } < \text{ porque } CpT \text{ podría ser, en esta versión muy simplificada, un subconjunto impropio de } T).$$

(1) y (2) juntos componen la paradoja de Cantor sobre los números cardinales.

La tesis logicista consistía en afirmar que la aritmética es una parte de la lógica. Entonces las paradojas que aparecían en la fundamentación de la aritmética por la teoría de conjuntos tenían que ser, en última instancia, de naturaleza lógica. Esto exigía una revisión de la lógica misma, una investigación de los fundamentos de la lógica. Esta investigación produjo pronto una obra que en seguida fue un clásico: los *Principia Mathematica* (1910-1913) de B. Russell y A. N. Whitehead (1861-1947).

Algunos autores contrarios al logicismo estimaron que las paradojas de la teoría de conjuntos no eran lógicas, sino estrictamente matemáticas, debidas a un uso incorrecto de la noción de infinito. Así por ejemplo, la paradoja de los números cardinales se debería a que se usa un conjunto infinito — el conjunto de todos los conjuntos — como si existiera ya hecho y completo, como infinito actual, según se dice, siendo así que el infinito no puede ser nunca dado como dato actual. Pero Russell y otros autores descubrieron pronto paradojas de las que no podía decirse que implicaran la idea de infinito. Un ejemplo de ellas es la paradoja de los adjetivos heterólogos, de Grelling: llamemos 'heterólogo' al adjetivo calificativo que no puede calificarse a sí mismo. Por ejemplo, el adjetivo 'castellano' no es heterólogo, pues puede calificarse a sí mismo: en efecto, el adjetivo 'castellano' es castellano. En cambio, 'monosilábico' es heterólogo, pues 'monosilábico' no es

monosilábico. ¿Qué es el adjetivo 'heterólogo'? ¿Es heterólogo o no lo es? Si es heterólogo, no se calificará a sí mismo. Por tanto:

si 'heterólogo' es heterólogo no es heterólogo.

Pero si no es heterólogo se calificará a sí mismo. Por tanto:

si 'heterólogo' no es heterólogo es heterólogo.

Paradojas como la de Grelling mostraban que la crisis de fundamentos afectaba también a la lógica. La paradoja de Grelling puede entenderse, en efecto, como una paradoja de la predicación, de la atribución de un predicado a un sujeto, lo cual es un tema clásico de la lógica.

16. Lenguajes "mal hechos". — Las paradojas parecían indicar que el lenguaje científico — entendiendo por lenguaje el discurso, la discursividad, la posibilidad de organizar y desarrollar articuladamente las ideas — era tan imperfecto desde el punto de vista lógico como los lenguajes étnicos o históricos, como el lenguaje común. Una imperfección lo es siempre desde algún punto de vista determinado; el que aquí interesa es el de la precisión y la regularidad sin excepciones.

En los lenguajes vivos, en castellano por ejemplo (ateniéndonos por concreción al lenguaje escrito), hay frecuentemente más de un símbolo para realizar un solo oficio gramatical, con diversos matices psicológicos que no interesan desde el punto de vista lógico. Por ejemplo, la afirmación simultánea de dos enunciados puede expresarse por la conjunción 'y', pero también por una coma, o un punto y coma, o una adversativa, etc. Así ocurre con otras varias relaciones, como la del antecedente (prótasis) al consecuente (apódosis) en las oraciones condicionales, etc. Por último, muchos nombres son sinónimos, o sea, varios símbolos pueden servir para nombrar un mismo objeto.

A la inversa, un sólo símbolo puede desempeñar en el lenguaje común varios oficios. Un mismo nombre puede aplicarse por analogía a varios objetos, y hasta partículas de enlace, como las conjunciones, pueden tener más de un sentido. Así por ejemplo, la disyunción 'o' ejecuta en castellano dos oficios disyuntivos, uno excluyente y otro no-excluyente, cada uno de los cuales contaba en latín con símbolos propios ('aut' para el excluyente, 'vel', 'sive', 'seu' para el no-excluyente).

Por otra parte, incluso cuando no se presentan esas imprecisiones — que, inspirándonos en Aristóteles, podemos llamar, respectivamente, 'sinonimia' y 'homonimia' — los nombres comunes o términos abstractos de los lenguajes étnicos están generalmente poco determinados en cuanto a su extensión.

Por último, los lenguajes étnicos no tienen reglas precisas de formación y combinación de enunciados, no tienen una sintaxis exacta. Como las anteriores deficiencias, ésta se suple por el sentido común de los que hablan la lengua. La gramática suele incluso apelar explícitamente al sentido común.

Así, cuando se decía en los manuales de gramática que una oración (un enunciado) es la expresión de un pensamiento completo, se hacía una apelación al sentido común del lector, el cual debía saber precientíficamente, por su uso del lenguaje, qué cosa merece el nombre de 'pensamiento completo'. Pero las gramáticas de los lenguajes étnicos no dan por sí mismas, sin apelación al sentido común, criterios técnicos suficientes de lo que es una oración bien hecha. Sus criterios técnicos — por ejemplo, la estructura Sujeto-Predicado — no bastan para condenar pseudooraciones como la siguiente, parecida a otros ejemplos contruidos por Russell:

cenicero ojea lenidad y estroncio bebe.

Esa serie de palabras es en efecto irreprochable desde el punto de vista de aquel criterio: cada "oración" tiene su sujeto y su predicado. Para condenar como sin-sentido, como mal hecha, esa sarta de palabras, hay que apelar no a la gramática, sino al sentido común.

La filosofía conoce desde antiguo numerosas paradojas o aporías ("callejones sin salida") debidas a todas esas causas de imprecisión. La del montón de trigo, por ejemplo, se basa en la indeterminación extensional de los abstractos: si van cayendo granos de trigo uno a uno en el mismo sitio, ¿cuándo hay un montón, dado que la diferencia en un grano más o menos es irrelevante? O la aporía del calvo: si un hombre pierde los cabellos uno a uno ¿cuándo puede decirse que es calvo? Si el abstracto 'calvo' estuviera definido con precisión — por ejemplo, definiendo el conjunto de los calvos como formado por todos y sólo los hombres que no tienen ningún cabello en el exterior del cráneo —, la aporía no podría presentarse. Lo que muestra que es la indeterminación del abstracto la raíz de esta aporía.

Otras aporías o paradojas antiguas se referían al uso de conceptos ya más teóricos, como los de espacio, tiempo y movimiento. Estas aporías, que han sido muy importantes en la historia de la filosofía, interesan menos a nuestro tema.

En cambio, también en la filosofía griega, y probablemente en el seno de una escuela filosófica que cultivó profundamente la lógica — la escuela de Megara —, se formuló una paradoja muy célebre que afecta al tema mismo de la lógica: la verdad o falsedad de enunciados. La estudiaremos a continuación.

17. La paradoja de Epiménides. — La literatura llamada 'doxográfica', que recoge las tradiciones sobre dichos y opiniones de los antiguos filósofos, tradiciones aún vivas a finales de la Antigüedad, conservó esta paradoja, también llamada "del embustero", que se atribuyó al antiguo pensador cretense Epiménides, a veces incluido entre los siete sabios. Como se ha indicado, la paradoja parece proceder de la escuela de Megara y del siglo IV a. n. e. El texto más completo ha sido transmitido por el apóstol

Pablo, el cual, hablando de los cretenses, escribe: "Bien dijo uno de ellos, su propio profeta: los cretenses, siempre embusteros, bestias malas y glotonas."

La paradoja está contenida en las cinco primeras palabras de la frase atribuida a Epiménides: éste, que es cretense, dice que [todos] los cretenses mienten siempre. ¿Es verdad lo que dice Epiménides? Si es verdad, como Epiménides es cretense, tiene que mentir siempre; luego lo que dice es falso. La paradoja de Epiménides se estudia mejor en una versión simplificada que le ha dado el lógico e historiador de la lógica Jan Łukasiewicz (1878-1956). Sea el enunciado

este enunciado es falso.

Llamémosle '*p*'. Si es falso, la verdad es la negación de lo que dice. Por tanto:

Si '*p*' es falso, '*p*' es verdadero.

Si es verdadero, la verdad será lo que dice. Por tanto:

si '*p*' es verdadero, '*p*' es falso.

Esta antigua paradoja ha sido muy recordada y estudiada en los años en que se construyó la moderna lógica simbólica, porque resultaba tener un parecido de familia con algunas paradojas recién contruidas. Tomando, por ejemplo, la paradoja de Grelling antes citada, se apreciaba que la de Epiménides tiene en común con ella una cierta reflexividad: el que un adjetivo sea o no heterólogo es algo que tiene que ver con una relación del adjetivo consigo mismo; y el enunciado '*p*' es un enunciado que afirma de sí mismo la falsedad. En la gramática de los lenguajes étnicos no hay ninguna regla que prohíba construir formaciones con ese tipo de reflexividad que parece estar en la base de paradojas como la de Epiménides o la de Grelling. Por eso una de las soluciones clásicas al problema de las paradojas consiste en introducir una regla que prohíba la formación de enunciados y la construcción de nociones con esa clase de reflexividad.

Cervantes presenta en el *Quijote* un acertijo basado también en una reflexividad pero de naturaleza empírica:

"Un caudaloso río dividía los términos de un mismo señorío (y esté v. m. atento, porque el caso es de importancia y algo dificultoso:) digo pues que sobre este río estaba una puente, y al cabo della una horca, y una como casa de Audiencia, en la cual de ordinario había cuatro Iuezes que juzgaban la ley que puso el dueño del río, de la puente y del señorío, que era en esta forma: Si alguno passare por esta puente de una parte a otra, ha de jurar primero a dónde y a qué va, y si jurare verdad, déxenle passar, y si dixere mentira, muera por ello ahorcado en la horca que allí se muestra, sin remisión alguna. Sabida esta ley, y la rigurosa condición della, passaban muchos, y luego en lo que juraban se echaua

de ver que dezían verdad, y los juezes lo dexauan passar libremente. Sucedió pues que tomando juramento a vn hombre, juró y dixo, que para el juramento que hazía, que yua a morir en aquella horca que allí estaua, y no a otra cosa. Repararon los Iuezes en el juramento y dixerón: Si a este hombre le dexamos passar libremente, mintió en su juramento, y conforme a la ley deue morir, y si le ahorcamos, el juró que yua a morir en la horca, y auiendo jurado verdad, por la misma ley deue de ser libre. Pídesse a vuesa merced, señor Gobernador, qué harán los Iuezes de tal hombre...”

La importancia de las paradojas quedaba clara por el hecho de presentarse en lenguajes de gran exactitud, como son el lenguaje matemático abstracto de Cantor, o la “escritura conceptual” de Frege, un simbolismo para la lógica inspirado en el de la matemática. Pero, recíprocamente, las paradojas iban a su vez a intensificar el interés por los lenguajes exactos y artificiales de ese género, precisamente porque estos lenguajes permitían un estudio más preciso de las paradojas.

Todo esto motivaba suficientemente el deseo de introducir en la investigación de fundamentos, y señaladamente en la lógica, el uso de lenguajes artificiales.

18. Lenguajes “bien hechos”. Cálculos formales y lenguajes formalizados. Metalenguaje. — El que las paradojas aparecieran también en estudios lógicos que utilizaban el simbolismo artificial de inspiración matemática mostraba que el simbolismo no es lo esencial para conseguir un lenguaje coherente o consistente (sin contradicciones). Un buen simbolismo puede ser útil para superar algunas deficiencias lógicas del lenguaje común. Por ejemplo, atribuyendo un símbolo, y sólo uno, a cada objeto o función del campo considerado en cada caso, del universo del discurso de cada caso, se evitan la sinonimia y la homonimia. Pero no se evitan ni la vaguedad de la noción de enunciado ni las paradojas. Para evitar también éstas hace falta algo más que un conjunto de símbolos unívocos: hace falta una gramática. Todo tiene que ser artificial, no sólo el vocabulario, si se quiere tener un lenguaje “bien hecho” desde el punto de vista lógico, que no permita vaguedades en las formas ni contradicciones en el discurso, en el uso y la combinación de las formas.

Esta idea tenía ya su historia. El filósofo Condillac (1715-1780) había sostenido que una ciencia es un lenguaje bien hecho. Pero en el ideal del lenguaje bien hecho confluía además otra tradición lógico-filosófica más antigua, a la que puede llamarse ‘tradición algorítmica’. Los principales representantes de esta tradición son Ramon Lull (1235-1315) y, posteriormente, Leibniz. El ideal algorítmico aspira a reducir el razonamiento a cálculo. El cálculo de Lull (*Ars Magna*) consistía en unas combinaciones de símbolos (que representan nociones morales y teológicas) con ayuda en algunos casos de ciertas figuras geométricas superponibles y móviles.

Leibniz, que, como más moderno, prefería un cálculo aritmético (*calculus universalis*), ha expresado muy claramente la naturaleza de la concepción algorítmica del razonamiento y de la lógica: Leibniz quiere proceder en lógica “al modo como calculamos en álgebra”, porque “el único modo de enderezar nuestros razonamientos consiste en hacerlos tan sencillos como lo son los de los matemáticos, de modo que se pueda hallar el propio error a simple vista y que, cuando haya discusiones entre las personas, se pueda decir sencillamente: contemos, sin más ceremonia, a ver quién tiene razón...”

El punto de vista algorítmico es también una versión de la idea de lenguaje bien hecho, pero más ambiciosa que la considerada hasta ahora. Su principal característica consiste en lo siguiente: en un cálculo o algoritmo es posible realizar operaciones sin saber qué significan los símbolos. Algunas letras que se usan en muchas demostraciones matemáticas no significan en ellas nada concreto. Esta es la diferencia más visible entre un cálculo y un lenguaje, por bien hecho que éste esté: que para merecer el nombre de ‘lenguaje’ un sistema de símbolos tiene que ser tal que sus formaciones signifiquen algo; mientras que un cálculo no está tan directamente vinculado a significar.

Se vio en el capítulo I que los esquemas formales por los que se interesa la lógica presentan también, como los cálculos, elementos sin significación concreta. Por eso era natural que el ideal algorítmico se introdujera en lógica, acompañado y ayudado por la nueva costumbre de trabajar en esta disciplina con simbolismos parecidos a los de los algoritmos matemáticos.

Puede observarse que la introducción de la idea de cálculo en lógica hace que ésta rebase el enfoque lingüístico. Un cálculo, como se ha dicho, no es un lenguaje, pues sus formaciones no significan directamente. Un cálculo sólo es un lenguaje cuando está interpretado, atribuyéndole significaciones. Cuando no lo está, las operaciones que se realizan o pueden realizarse con sus símbolos deben compararse más con los movimientos de un juego, como el ajedrez o las damas, que con las composiciones de palabras y oraciones en un lenguaje. — Por eso hay autores que conciben una teoría general de los cálculos o sistemas formales (H. B. Curry) como idéntica con la lógica, o como fundamento de la lógica. — En este libro se conservará el enfoque lingüístico, según una concepción que se explicará más adelante.

Las dos tradiciones del lenguaje bien hecho se encuentran actualmente en la distinción entre cálculos formales y lenguajes formalizados.

Un *cálculo formal* (o sistema formal) es un sistema que consta de lo siguiente:

- 1.º Un conjunto de símbolos elementales, análogos de los términos o palabras de un lenguaje, o de las fichas de un juego de damas; ese conjunto está definido efectivamente, es decir, de tal modo que ante un símbolo

cualquiera, siempre es posible decidir si ese símbolo pertenece o no al vocabulario del cálculo. A causa de la influencia ejercida por la matemática en la lógica, muchos de esos símbolos son familiares desde la matemática de la enseñanza media. Así por ejemplo:

para indicar individuos cualesquiera suelen utilizarse en lógica las últimas minúsculas del abecedario latino, como 'u', 'v', 'x', 'y', 'w', 'z';

para indicar propiedades cualesquiera suelen usarse mayúsculas latinas, como 'P', 'Q', 'R', 'S';

para indicar conjuntos cualesquiera (clases) de individuos, suelen usarse minúsculas griegas, como 'α', 'β', 'γ'.

Menos inspirados en la matemática están otros signos que iremos introduciendo a medida que los necesitemos.

A continuación de la lista de los símbolos elementales o primitivos suele indicarse en la exposición de un cálculo el procedimiento por el cual se introducirán nuevos símbolos a partir de los primitivos. Teóricamente, esto no es necesario, pues se podría siempre operar con los símbolos primitivos. Pero en la práctica el expediente de la definición es imprescindible. Las definiciones de un cálculo no suelen dar explícitamente nociones, significaciones, sino que son meras autorizaciones para usar unos símbolos en vez de otros que la definición declara equivalentes a los primeros. Por eso las definiciones suelen expresarse mediante un símbolo de equivalencia. Aquí usaremos ' $\leftrightarrow_{df}$ ', colocado entre el nombre del símbolo o la expresión que se desea introducir (definiendum, situado a la izquierda de ' $\leftrightarrow_{df}$ ') y el del símbolo o la expresión ya conocidos a los que sustituyen (definiens, situado a la derecha de ' $\leftrightarrow_{df}$ '). He aquí un ejemplo (cfr. 15):

$$'CpT \leq T' \leftrightarrow_{df} 'CpT < T, \text{ o } CpT = T'.$$

- 2.º Un conjunto de reglas, llamadas 'de formación', que indican cómo pueden combinarse los símbolos elementales en formaciones compuestas; el análogo de estas reglas en un lenguaje es la definición de oración; en un juego de damas, las reglas sobre las posiciones que pueden ocupar las fichas. El conjunto de las reglas de formación tiene que ser tal que el concepto de enunciado, o fórmula, o esquema, sea también efectivo, de modo que se pueda decidir siempre si una determinada combinación de símbolos elementales está o no permitida en el cálculo.
- 3.º Un conjunto de reglas, llamadas 'de transformación', que indican cómo puede pasarse de una combinación de símbolos elementales a otra, de un enunciado, fórmula o esquema a otros, lo que equivale a transformar la primera; lo análogo a estas reglas en un lenguaje son los principios sintácticos de la composición de oraciones, composición que expresa el razonamiento; y en un juego de damas son las reglas sobre el movi-

miento de las fichas. Este conjunto de reglas tiene que permitir que el concepto de transformación sea también efectivo, de modo que se pueda decidir siempre si una determinada transformación está o no permitida en el cálculo.

La construcción de un cálculo se hace en esos tres pasos sin atender a nada ajeno al cálculo, a significaciones ni aplicaciones del mismo. Pero normalmente un cálculo se construye para alguna aplicación o algunas aplicaciones, por lo menos a los conceptos lógicos de enunciado, predicado, sujeto, etc. Por eso, una vez hecho el cálculo, interesa establecer de una manera general cómo se comportará en su funcionamiento y aplicación. Esto se establece estudiando el cálculo desde el punto de vista de tres propiedades que, en diversos grados, es deseable que tenga. Esas tres propiedades son la consistencia, la completud y la decidibilidad.

Un cálculo es *consistente* cuando es imposible demostrar en él una contradicción, es decir, un enunciado y su negación.

Un cálculo es *completo* cuando se pueden demostrar en él como teoremas todos los enunciados formalmente verdaderos construibles con sus símbolos. También puede decirse, desde el punto de vista de la aplicación: cuando, aplicado a los principios o axiomas de la teoría para la cual ha sido construido, produce como teoremas todas las verdades de esta teoría.

Un cálculo es *decidible* cuando es siempre posible establecer, en un número finito de pasos normados, si una determinada fórmula pertenece a su lenguaje es o no es un teorema de dicho cálculo. (Usamos pues 'decidibilidad' en el sentido de 'efectividad de la noción de teorema'.)

La consistencia es en general una propiedad necesaria (aunque a veces haya que renunciar a su demostración). En general no tiene interés un cálculo que sea inconsistente, pues al demostrar a la vez un enunciado y su negación, lo demuestra todo: todo enunciado es un teorema de ese cálculo, el cual, por tanto, no distingue entre verdades y falsedades.

La completud y la decidibilidad son menos necesarias. Son más bien concretizaciones del ideal algorítmico del lenguaje "bien hecho". Completud y decidibilidad se diferencian en lo siguiente: cuando un cálculo es completo, sus reglas permiten obtener todos los enunciados verdaderos de su universo del discurso. Pero, dado un enunciado, no se sabrá si es un teorema hasta que se consiga construir una demostración del mismo con las reglas de transformación del cálculo. Si no se encuentra esa demostración, no se podrá estar seguro de que el enunciado no es teorema, pues no estamos seguros de que la demostración no exista. En cambio, si el cálculo es decidible, cuenta con un procedimiento para averiguar, dado cualquier enunciado y aunque no se le tenga aún construido con las reglas de transformación, si ese enunciado es un teorema o no. Un cálculo que ante cualquier enunciado sea capaz de demostrar ese enunciado o su negación es pues un cálculo decidible. Por eso admitiremos que todo cálculo decidible es completo, pero no todo cálculo completo es decidible.

Las nociones de consistencia, completud y decidibilidad han quedado expuestas de un modo intuitivo. Pero se pueden dar varias formulaciones técnicas de ellas, desde los puntos de vista sintáctico y semántico que se consideran en 19.

Un *lenguaje formalizado* es un cálculo interpretado. Es lenguaje porque la interpretación le hace significativo. Es formalizado porque tiene en todo lo demás la estructura exacta efectiva del cálculo. En un lenguaje formalizado una operación se justifica, como en el cálculo, mostrando efectivamente que hay una regla que la justifica, y no, como en el lenguaje étnico, apelando al sentido común para que aprecie que la operación tiene un sentido aceptable. 'Formalizar' significa precisamente dar esa estructura exacta y efectiva. Simbolizar es en cambio sólo dar símbolos, cosa que puede hacerse también con un lenguaje no exacto. Pero la formalización es mucho más cómoda de realizar si va acompañada de simbolización.

Un principio fundamental del método con el que se construyen tanto los cálculos cuanto los lenguajes formalizados consiste en distinguir cuidadosamente entre los símbolos y formaciones del cálculo o lenguaje formalizado y el lenguaje en el cual se habla de ellos. Así pueden evitarse paradojas como la de Epiménides. La afirmación de la verdad o falsedad de formaciones del cálculo o lenguaje formalizado debe hacerse en otro lenguaje — llamado frecuentemente '*metalenguaje*' del primero, que es el *lenguaje-objeto*, o lenguaje de grado cero.

Con todos esos elementos, el ideal del lenguaje bien hecho se presentaba, señaladamente en la obra del matemático y lógico David Hilbert (1862-1943), en la forma más ambiciosa que había tenido hasta entonces en la historia de las ciencias formales: se trataba de construir la lógica y la matemática fundamental como cálculos, como sistemas de símbolos sin significación determinada, sino según las necesidades de cada caso (incluso para sillas o fieltros de cerveza, decía Hilbert), y dotado de reglas de formación y de transformación suficientes.

El principal objetivo del programa de Hilbert era conseguir demostraciones de consistencia, es decir, garantías de que en una teoría dada se evitarían las paradojas. Pero aquí nos interesa más otro aspecto del programa de Hilbert: ese programa algorítmico o *formalista* (tal es el nombre que ha recibido la escuela de Hilbert) tendía a concebir la deducción como algo mecanizable, al modo de Lull o Leibniz. La diferencia es que ahora no se pretendía mecanizar cualquier argumentación y en cualquier tema (Lull y Leibniz pretendieron hacerlo incluso con las argumentaciones morales); pero subsistía la aspiración a mecanizar la *deducción*. La exactitud del cálculo, o del lenguaje formalizado, debía liberar el pensamiento del científico, como escribe Frege, de la sumisión al lenguaje (común), y su nueva y segura eficacia evitaría al hombre el trabajo intelectual deductivo, que es sin duda en cierta medida mecánico, no creador.

19. La investigación de fundamentos en lógica. Sintaxis y semántica. — Estudiar la posibilidad y el modo de construir la lógica como lenguaje "bien hecho", como cálculo o lenguaje formalizado, es aplicar a la lógica misma la investigación de fundamentos, pues supone buscar sus nociones primitivas o elementales y averiguar el modo como éstas fundan las demás. La investigación de los fundamentos de los sistemas o lenguajes lógicos recibe frecuentemente el nombre de '*metalógica*', y comprende dos investigaciones principales: la sintaxis y la semántica.

El nombre '*metalógica*' está inspirado en '*metamatemática*', utilizado por el formalismo para nombrar la teoría fundamental (la investigación de fundamentos) de la matemática.

La *sintaxis* estudia los lenguajes desde un punto de vista algorítmico, o sea, sin interesarse más que por las relaciones entre los símbolos de un determinado lenguaje (sintaxis especial) o por las relaciones posibles entre símbolos lingüísticos cualesquiera (sintaxis general). No se interesa por las significaciones de los símbolos, por sus referencias a entidades ajenas al sistema mismo de los símbolos. Esta referencia de los símbolos y las formaciones lingüísticas a algo que no sea el sistema de símbolos mismo, es decir, el oficio significativo del lenguaje, es en cambio objeto de la *semántica* (que también puede ser especial o general).

Según esto, la presentación que antes vimos de lo que es un cálculo formal o un lenguaje formalizado era sintáctica, pues se basaba sólo en el establecimiento de reglas para usar y combinar (relacionar) símbolos. En efecto, el método comúnmente seguido en la presentación de un sistema formal es sintáctico: consiste en construir un lenguaje no interpretado — o sea, más propiamente, un cálculo — sin hacer consideración alguna sobre las posibles aplicaciones del mismo, lo cual sería ya semántico, puesto que referiría el cálculo a algo ajeno a él. El método sintáctico de la construcción garantiza que el sistema es realmente formal, que toda formación y toda operación (transformación) en él se justifican o fundamentan por alguna regla explícita, y no por la intelección del que las realiza.

La sintaxis de un cálculo o sistema es el conjunto de reglas indicado, con las nociones y notaciones necesarias para formularlas. Por ejemplo, para dar reglas sintácticas de formación de expresiones correctas a partir de símbolos elementales de un determinado cálculo, hace falta la idea sintáctica de concatenación. Con ella, una vez dados (mediante otra regla sintáctica aún más elemental) símbolos como '*x*', '*P*', una regla de formación puede ser del tenor siguiente: 'la concatenación de '*P*' con '*x*' es una expresión bien hecha'. La misma noción de expresión bien hecha, o fórmula, es sintáctica.

Notaciones importantes de la sintaxis son las llamadas '*variables sintácticas*'. Éstas son letras que se usan para referirse a expresiones del cálculo

de que se trate. También ellas sirven, por ejemplo, para enunciar reglas de formación. Así, para enunciar la regla de que en un determinado cálculo el resultado de escribir la conjunción 'y' entre dos expresiones bien hechas es a su vez una expresión bien hecha, puede decirse: 'si X, Y son expresiones bien hechas, entonces la concatenación de X, 'y', Y es una expresión bien hecha'. 'X' e 'Y' son en esa frase variables sintácticas.

Partiremos del principio de que lo lógico no es el cálculo mismo construido por la sintaxis, sino el lenguaje formalizado que resulta de interpretar ciertos cálculos con los conceptos tradicionalmente llamados 'lógicos', como son los de verdad, falsedad, enunciado, sujeto, predicado, clase, relación, etc. Los cálculos serán, según esto, la formalización sintáctica de la lógica. Esta concepción, que H. Scholz (1884-1956) sentó por motivos filosóficos y R. Carnap ha desarrollado por motivos técnicos, hace que conservemos el enfoque lingüístico en lógica.

El estudio semántico es siempre estudio de un sistema de símbolos: un lenguaje o un cálculo. Los conceptos del método semántico que más nos interesan son los siguientes:

Denotación o significación de un símbolo es el objeto representado por él. Por ejemplo, la denotación del término (símbolo nominal) 'Miguel de Cervantes Saavedra' es Miguel de Cervantes Saavedra.

Sentido de un término es el modo como ese término significa. Así por ejemplo, el término 'el autor del *Quijote*' significa Miguel de Cervantes Saavedra en un sentido distinto que el término 'el Manco de Lepanto', que también significa Miguel de Cervantes Saavedra.

Una *interpretación* es una atribución de una significación a un símbolo, o un sistema de atribuciones de significaciones a símbolos.

Un concepto importante del método semántico es el concepto de *verdad*. El concepto semántico de verdad ha sido expuesto técnicamente por A. Tarski. Su origen está en Aristóteles, y aquí puede bastarnos la versión más intuitiva dada por este filósofo en su *Metafísica*, donde se lee:

decir que lo que es no es, o que lo que no es es, es lo falso; decir que lo que es es y que lo que no es no es, es lo verdadero.

En la semántica de los cálculos tienen mucha importancia los enunciados formales, o esquemáticos, que son verdaderos para cualquier interpretación; éstos son los que en el capítulo I llamamos 'verdades formales'. Los enunciados formales que son falsos para cualquier interpretación se llaman 'contradicciones formales'. Un ejemplo sencillo es

x es P y no es P .

Por último, hay esquemas que no son ni verdaderos para toda interpretación ni falsos para toda interpretación, sino verdaderos para unas y falsos para otras. El Esquema I del capítulo I era de esta clase. Las interpreta-

ciones que hacen verdadero a un esquema se llaman '*modelos*' semánticos del mismo.

Los puntos de vista sintáctico y semántico deben distinguirse cuidadosamente. Pero es útil también usarlos uno tras otro en un mismo contexto, como haremos con frecuencia.

No se debe perder de vista que la sintaxis y la semántica (la metalógica en general) son, usando una terminología que ya conocemos, metalingüísticas respecto del cálculo que estudian. Las reglas sintácticas son enunciados que hablan de combinaciones de símbolos del cálculo. Las reglas semánticas hablan de cómo denotan ciertas combinaciones de símbolos del cálculo, etcétera. En todo caso, enunciados sintácticos y enunciados semánticos no son enunciados del lenguaje formalizado, sino sobre el lenguaje formalizado (o cálculo): pertenecen pues a metalenguajes de éste. Por tanto, para evitar que se produzcan formaciones con el tipo de reflexividad que vimos en la raíz de ciertas paradojas (llamadas, precisamente, 'metalógicas'), las reglas y, en general, los enunciados de la sintaxis y de la semántica de un lenguaje se escriben en otra notación distinta: en un metalenguaje sintáctico y un metalenguaje semántico. Esos metalenguajes, por obra sobre todo de Tarski y de Carnap, han sido formalizados a su vez en mayor o menor medida, lo que quiere decir que la sintaxis y la semántica son ellas mismas sistemas más o menos formales (para cada cálculo). Pero en este manual las consideraciones sintácticas y semánticas se harán en lenguaje castellano común, a un nivel intuitivo, aunque usando algunos expedientes metalingüísticos (señaladamente, variables sintácticas).

E. W. Beth ha llamado (1962) '*hermenéutica*' a las consideraciones semánticas elementales que se hacen intuitivamente en el lenguaje común. Análogamente podríamos llamar '*gramática*', en vez de '*sintaxis*', a las observaciones sintácticas que haremos sobre la base intuitiva del lenguaje común. Pero, aunque haremos propiamente gramática y hermenéutica en ese sentido, seguiremos hablando (laxamente) de sintaxis y semántica.

A la sintaxis y a la semántica se añade la *pragmática*, estudio de las relaciones entre el lenguaje y los que lo usan; las tres constituyen la *semiótica* (R. Carnap), que es una teoría lógica general de los lenguajes. No se hará en este manual ninguna referencia a la pragmática.

A pesar de lo dicho sobre el carácter metalingüístico de la sintaxis, existen técnicas para formalizar la sintaxis de un cálculo dentro del cálculo mismo, sin caer en paradojas. Así se puede formalizar, por ejemplo, en un cierto cálculo, la noción sintáctica de fórmula no demostrable en ese cálculo.

20. Los límites del ideal algorítmico. — La aspiración a mecanizar la inferencia se ha limitado siempre en la práctica, incluso en los casos de Lull y Leibniz, más ambiciosos en teoría, a la inferencia deductiva. Esta

se definía clásicamente como la inferencia que se funda en enunciados generales para afirmar enunciados menos generales. Los "silogismos" de los ejemplos 1a y 2a del capítulo I son deducciones. (La deducción puede definirse más generalmente como el tipo de inferencia que vale por razones puramente formales.) Pero la deducción no es el único género de inferencia usado en la ciencia. Las ciencias reales no disponen siempre de enunciados generales que les permitan inferir para cada caso otros particulares. Más bien es su tarea principal el conseguir enunciados generales a partir de enunciados menos generales o incluso singulares (de sujeto individual) referentes a hechos observados. El tipo de razonamiento por el cual se pasa de enunciados singulares, o particulares o menos generales a otros más generales se llama 'inducción', y el anterior es el modo tradicional de describirlo. Esta clase de razonamiento ha sido objeto de mucha discusión en teoría de la ciencia, y sigue siéndolo hoy. Pero es tan frecuente como el deductivo, tanto en la investigación cuanto en la exposición didáctica. He aquí un ejemplo de inferencia inductiva, tomado de un libro de química estudiado en la Universidad de Barcelona. El autor da en forma de tabla la siguiente serie de enunciados singulares:

<i>Helio a 0° C</i>		
<i>Presión en atm.</i>	<i>Vol. en litros</i>	<i>Producto $V \times P$</i>
1,0020	22,37	22,41
0,8067	27,78	22,41
0,6847	32,73	22,41
0,5387	41,61	22,41
0,3550	63,10	22,41
0,1937	115,65	22,41

(Cada una de las líneas de la tabla constituye, efectivamente, un enunciado, reducible a un esquema de la forma siguiente: 'a 0° C y una presión de P atm., una masa x de helio seco ocupa un volumen V, y el producto $V \times P$ es 22,41').

De esa serie de enunciados particulares (más otros referentes a otras presiones, otras temperaturas, otras masas y otros gases), el autor infiere el enunciado general llamado 'ley de Boyle-Mariotte' en la formulación siguiente: "Para cualquier masa de gas seco a temperatura constante, el producto del volumen por la presión ($V \times P$) es constante".

En las discusiones sobre la inducción no suele ponerse en duda que este tipo de inferencia no es válido por razones formales, como lo es la deducción: su validez depende del conocimiento de los particulares objetos estudiados, de la materia del conocimiento. Por eso mismo, la validez de los conocimientos obtenidos por inferencia inductiva no es tan permanente

como la de los que se consiguen por deducción, aunque aquéllos son, por lo general, más interesantes que éstos.

El ideal algorítmico de mecanización de la inferencia no puede, pues, aplicarse sin más a la inducción. Ni tampoco, naturalmente, a operaciones más elementales que la inducción y relacionadas, como ella, con la observación directa de los fenómenos: la descripción de éstos, su análisis concreto, su clasificación. Tales son los límites del ideal algorítmico por lo que hace a su aplicación a las ciencias reales.

Por lo que hace a la deducción, a la lógica misma, se pensó durante algún tiempo que el programa algorítmico fuera plenamente realizable, que toda la lógica formal, como teoría de la deducción, fuera reducible a cálculo, con lo que la deducción habría dejado de ser tarea interesante para el pensamiento. Pero a partir de 1930 varios autores demostraron que también esa suposición era excesiva (en los capítulos XI-XIII se consideran los principales resultados de esas demostraciones): sólo puede reducirse de un modo general a algoritmo una parte de la lógica que no llega al grado de complicación de la aritmética. Para niveles de complicación mayor, es posible reducir a algoritmos ramas o teorías más o menos amplias, pero no construir algoritmos que abarquen a todas las teorías de uno de esos niveles.

21. Los frutos del programa algorítmico. — Una empresa verdaderamente científica no es nunca estéril, aunque no alcance nada de su objetivo inicial. Así ocurre con el programa algorítmico, el cual, por lo demás, consigne algo relacionado con su objetivo, a saber, un considerable aumento de la potencia deductiva de la lógica, un enriquecimiento del arsenal de los métodos formales.

También para la aplicación a las ciencias reales ha sido fecundo el ideal algorítmico. Pues la limitación a las partes de una ciencia que se consideran conclusas y construibles deductivamente no es, como vimos, poca cosa para el progreso de la investigación, al que contribuye indirectamente.

Pero, sobre todo, al demostrar la inviabilidad de un programa de algoritmización de toda la inferencia deductiva, la lógica ha conseguido una claridad sobre los límites de lo formal que no había existido antes, como prueban las aspiraciones de Lull o de Leibniz. En el futuro no es probable que ningún filósofo vuelva a soñar con zanjar *cualquier* discusión mediante cálculos, como esperó Leibniz. Este resultado tiene pues incluso relevancia filosófica.

Otro fruto de los trabajos algorítmicos nos interesa aquí especialmente. Como se dijo en el capítulo II, la principal aportación de la lógica formal a las ciencias reales es indirecta: consiste en suministrarles los instrumentos para analizar sus propios conceptos y construcciones, y aclarar así sus fundamentos, localizar sus puntos oscuros y precisar sus necesidades. Pues bien, la construcción de cálculos y lenguajes formalizados tiene como conse-

cuencia un afinamiento de esa capacidad analítica. Ello se debe a lo siguiente.

Aunque un lenguaje formalizado es un sistema que funciona — o se usa — independientemente de la significación de sus símbolos, pues lo que funciona es el cálculo interpretado en ese lenguaje, sin embargo, la lógica construye esos lenguajes y los cálculos teniendo en cuenta posibles aplicaciones, por lo menos la aplicación a los conceptos lógicos. Del cálculo, o del lenguaje formalizado, se espera que dé de sí la forma de toda la teoría (normalmente preexistente en lenguaje común) a la que se desea aplicarlo. Una tal exigencia, aunque la mayoría de las veces no se cumpla, impone un conocimiento muy preciso del sector de lenguaje natural que se trata de formalizar, de llevar a la exactitud del cálculo. Y esto a su vez exige un análisis de la mayor finura posible.

Por lo común el análisis se limitará al sector del lenguaje común que sea relevante para la deducción, para la transformación de la teoría en lenguaje formalizado, en teoría formal. Pero en este reducido sector, el análisis tiene que ser de una finura desconocida por la lógica tradicional. Así la construcción de cálculos, aunque es una actividad sintética, o sea, una composición, facilita un apreciable progreso en el análisis de los elementos y la estructura de las teorías científicas y del lenguaje común en general. Los resultados básicos de ese progreso del análisis han renovado sustancialmente la teoría de las categorías lógicas, objeto del capítulo IV.

APÉNDICE AL CAPÍTULO III

A. *Tetxos citados*. — Pablo, 1.^a ep. a Tito, 1-2. — Cervantes, *Quijote*, Segunda parte, cap. LI. — Leibniz, "Elementa characteristicae universalis", en Couturat: *Opusculs et fragments inédits de Leibniz*, p. 43; "Projets et essais pour arriver à quelque certitude...", en Couturat, p. 176. — Aristóteles, *Met.*, I^o 7, 1011b 25-27. — J. Babor y J. Ibarz, *Química general moderna*, 4.^a ed., p. 5.

B. *Observaciones*. — a 16: la versión dada de la paradoja del montón de trigo no es la antigua, pero está contenida en ésta. — a 18: la palabra 'algoritmo' procede del nombre del algebrista árabe Mohamed ben Musa Alkarismi (siglo ix).

CAPÍTULO IV

LAS CATEGORÍAS LÓGICAS

22. Análisis lógico y categorías. — Una ciencia tiene que penetrar en la realidad que le es dada para desmenuzar, según su punto de vista, el vago dato que es esa realidad: para buscar los componentes de ésta que convenga considerar elementales para los fines de la investigación de que se trate. Esa actividad se llama análisis o resolución, y tiene su ejemplo más material en la química: el análisis químico de un cuerpo dado es la búsqueda de los "elementos" que lo componen, para hacer luego afirmaciones sobre la naturaleza y la estructura de ese cuerpo.

De acuerdo con lo que en el capítulo I dijimos sobre la abstracción, también el análisis supone ideas previas, abstracciones que lo dirijan. Pues un análisis es como una abstracción simultáneamente múltiple: al abstraer se toma de una cosa un aspecto o un componente; al analizar, se aspira a tomar por separado (o sea, abstractamente) todos los aspectos o todos los componentes de la cosa analizada. En el caso del análisis lógico del lenguaje, la abstracción orientadora es la idea de forma lógica, y lo que el análisis lógico busca son los elementos constitutivos de la forma, los puntos o nudos con que se sostiene, por así decirlo, la red o estructura que es la forma lógica.

A las clases de esos elementos se llama 'categorías'. Categorías son clases de términos o, en general, de símbolos. Si, por ejemplo, nos interesara analizar el lenguaje desde el punto de vista de las cosas significadas por sus términos, podríamos obtener una tabla de categorías, una clasificación de significatividades, como la de Aristóteles. Es ésta una tabla de diez categorías — sustancia, cualidad, cantidad, relación, acción, recepción de acción, lugar, tiempo, posición y hábito — las cuales clasifican a los términos por la clase de entidad que significan. Términos como 'mesa', 'árbol', 'Juan', que significan entidades en sí, pertenecen a la categoría de sustancia; términos como 'blanco', 'bueno', a la categoría de cualidad, etc.

Es claro que desde el punto de vista de la lógica formal esa tabla de categorías no interesa primariamente a la construcción de los cálculos

lógicos, a la sintaxis, sino, si acaso, a su interpretación, a su semántica.

El cuadro de categorías que la lógica necesita primariamente no se basa en la idea de significación, sino en la del papel gramatical (sintáctico) de los símbolos y las formaciones apofánticas del lenguaje. Las categorías lógicas son las clases de elementos que son necesarios para que exista un lenguaje apofántico; por ejemplo: las categorías 'nombre', 'enunciado' 'con-junción', etc.

Se trata de categorías gramaticales. Y efectivamente es la sintaxis lógica una especie de gramática. Su diferencia respecto de lo que corrientemente entendemos por 'gramática' consiste en que la sintaxis lógica no es la gramática de ningún lenguaje histórico, sino de unos lenguajes artificiales y contruidos a propósito para cada caso, los cuales son puramente apofánticos, están considerablemente empobrecidos respecto de los lenguajes étnicos, pero realizan con claridad lo que hemos llamado 'discursividad'.

El carácter artificialmente reducido de este tipo de lenguaje suele manifestarse en la práctica ya por el hecho de que sus símbolos no son de uso corriente en los lenguajes étnicos.

23. Las categorías Expresión, Fórmula, Enunciado. — Lo primero que puede hacerse con los símbolos primitivos de un cálculo es combinarlos. La combinación o composición de símbolos, que sintácticamente puede entenderse como la operación de ponerlos unos tras otros (concatenación), está regida en un cálculo o en un lenguaje formalizado por precisas reglas de formación.

Cualquier combinación de símbolos de un cálculo (incluso un símbolo sólo) se llama una 'expresión' ('Ausdruck' en la literatura alemana). Si la expresión ha sido correctamente construida según las reglas de formación, se dice que es una *fórmula* ('Formel' o 'zulässige Formel' en la literatura alemana, 'well-formed formula' en la terminología anglosajona [A. Church]).

Ante una expresión cualquiera de un cálculo, la única manera de decidir si es (o no) una fórmula, consiste en mostrar que está (o no está) construida según las reglas de formación de dicho cálculo.

Por ejemplo: con los símbolos corrientes que antes hemos visto a título de ilustración, los cálculos lógicos hoy utilizados suelen tener reglas de formación por las cuales una expresión como

xyz

(tres símbolos de individuo solos consecutivos) no es una fórmula, mientras que una expresión como

Px

o como

$Pxyz,$

' Px ' es la forma más corriente de simbolizar el sencillo tipo de enunciado (un símbolo predicativo seguido de otro u otros individuales) sí lo es.

que consta de un sujeto, ' x ', y un predicado, ' P ', que se le atribuye. El orden en que aparecen en esa fórmula el predicado y el individuo (sujeto) es el mismo que se encuentra en el texto griego de Aristóteles. Un enunciado como 'Sócrates es mortal' se encuentra escrito en el texto de Aristóteles en la forma: 'lo mortal se da en Sócrates'; en esquema: ' A se da en B ' ' A ὑπάρχει τῷ B ').

Cuando una fórmula es tal que tiene sentido decir de ella que es verdadera, o que es falsa, se llama 'enunciado'. Esta es la noción aristotélica de enunciado (apofántico), que se ha mantenido siempre en lógica desde entonces.

En el *Organon* aristotélico se lee: "no todo enunciado es apofántico, sino sólo aquellos en los que se da el ser verdaderos o falsos".

Una fórmula puede ser un enunciado de dos maneras: o bien porque sus lugares de contenido sean ocupados por símbolos con significaciones determinadas (categoremas); o bien porque, aun siendo una fórmula sin interpretar, resulte ser verdadera para cualquier interpretación posible (en cuyo caso es una verdad formal, o tautología), o falsa para toda interpretación posible (en cuyo caso es una falsedad formal). El Esquema II del capítulo I fue por esta razón llamado 'enunciado' (esquemático o formal), y no sólo 'esquema'.

Lo que en el capítulo I llamamos 'esquema' puede pues coincidir en la práctica con lo que aquí llamamos 'fórmula'. La diferencia entre las dos terminologías, tal como aquí las usamos, es de intención: esquema es el resultado de despojar de contenido significativo (eliminando sus categoremas) a un enunciado del lenguaje vivo; fórmula es el resultado de componer, según reglas de formación sintácticas, los símbolos de un cálculo. Usamos 'esquema' en un contexto de análisis lógico del lenguaje vivo, y 'fórmula' en un contexto de construcción o síntesis de lenguajes artificiales o cálculos.

24. El nombre: la categoría Constantes. Constantes lógicas. — La tarea básica de la parte apofántica de todo lenguaje es nombrar, denotar. La gramática de los lenguajes étnicos llama 'nombre sustantivo' a las palabras que cumplen ese oficio. Aquí diremos simplemente 'nombres'.

La distinción gramatical entre nombres propios y nombres comunes tiene importancia semántica porque lo denotable por unos y por otros se encuentra a distintos niveles: el nombre propio denota individuos, entidades concretas, mientras que el nombre común se presenta como significativo de propiedades de objetos concretos, por ejemplo, la propiedad Ser-una-mesa, o Ser-un-árbol. Un nombre común significa según esto una clase de individuos, por ejemplo, la clase de los individuos que son mesas, o la clase de los individuos que son árboles. Pero una propiedad, o una clase, no son entidades concretas: la propiedad Ser-una-mesa no existe como individuo; tampoco la clase Mesa, que es una entidad abstracta construida mentalmente.

Por tanto, si se admite que también los nombres comunes son verdaderos nombres, se está admitiendo al mismo tiempo que los abstractos pueden ser denotados igual que los individuos concretos.

Sobre este punto ha habido en la historia de la filosofía diversas opiniones. En la Edad Media esas discrepancias cristalizaron en tres modos de concebir la cuestión, los cuales siguen siendo los básicos hasta hoy. Para el *realismo*, los abstractos ('universales', en la terminología medieval) denotan entidades existentes: la propiedad Ser-mesa tiene tanta realidad como los individuos a los que llamamos 'mesas' (o más que ellos). Esta tradición filosófica se remonta a Platón (427-347 a. n. e.). El *conceptualismo* estima que el abstracto tiene una existencia intelectual, ideal: el abstracto no denota una realidad material, pero lo denotado por él tiene una consistencia o necesidad interior. El hombre no puede decidir a su voluntad sobre los abstractos: no los inventa, los descubre. El filósofo más representativo de esta tendencia fue en la Edad Media Pedro Abelardo (1079-1142). El *nominalismo* sostiene que los abstractos son meras palabras que no denotan nada, invenciones, etiquetas útiles para manejar la realidad. El principal pensador de esta escuela fue en la Edad Media Guillermo de Ockham (hacia 1300-1350), tal vez el más grande lógico entre Aristóteles y Leibniz. Diversas versiones de esas doctrinas son hoy día profesadas por importantes autores. Puede decirse que B. Russell ha tendido al realismo, A. Church al conceptualismo, y que W. V. O. Quine es nominalista.

Desde el punto de vista lógico, lo que interesa es mantenerse todo lo posible en un terreno formal, sin afirmaciones filosóficas sobre la realidad de las nociones. Por eso es conveniente retrasar todo lo posible la afirmación de que los nombres de clases y los abstractos en general denoten como los nombres propios. Esto podría entenderse, en efecto, como una afirmación filosófica muy discutible. En cambio, la semántica lógica puede admitir, sin suscitar gran discusión filosófica, que lo primariamente denotable es el individuo concreto. Salvo en los excepcionales casos de filósofos radicalmente escépticos, las diversas tradiciones filosóficas coinciden en admitir que la realidad concreta individual es la realidad en sentido propio.

Esto, sin embargo, no debe hacer creer que el concepto de individuo sea absoluto y carezca de problemas. En el uso científico, 'individuo' también es una abstracción, una construcción artificial, como indica, por ejemplo, el hecho de que un árbol, que para el botánico es un individuo, es para el físico un agregado de individuos — moléculas, átomos, partículas infraatómicas.

Las anteriores reflexiones sugieren que, aun admitiendo que sean nombres tanto los comunes cuanto los propios, es conveniente introducir en sintaxis dos clases de símbolos, unos para nombres propios, otros para nombres comunes. La diferencia semántica entre nombres propios y nombres comunes importa a la sintaxis porque va acompañada de una diferencia

en la función gramatical de unos y otros. Pero si son nombres, unos y otros denotarán algo determinado, fijo; por tanto, los símbolos usados serán *constantes* del lenguaje.

Las constantes que son nombres propios denotarán individuos; por eso se las llama 'constantes individuales'. Cumplen en fórmulas y enunciados el oficio de sujeto. Las que son nombres comunes denotarán — si denotan — clases o propiedades. En un enunciado del tipo más sencillo, el nombre de una clase o propiedad hace oficio de predicado. Por eso las constantes que son nombres comunes se llaman 'constantes predicativas'.

Por ejemplo, en el enunciado

Cervantes era poeta,

'Cervantes' es una constante individual, y su papel sintáctico es el de sujeto. 'Poeta' es el nombre — común — de la clase de los poetas: es una constante predicativa.

Los símbolos más corrientemente utilizados para constantes individuales y predicativas son:

para constantes individuales, las primeras minúsculas del abecedario latino, como 'a', 'b', 'c' ...;

para las constantes predicativas, las primeras mayúsculas del abecedario latino, como 'A', 'B', 'C' ...

Entre los elementos constantes del lenguaje tienen especial interés los que nombran relaciones lógicas, como, por ejemplo, la negación, o la conjunción. Los símbolos que denotan esas relaciones, u operaciones correspondientes, se llaman '*constantes lógicas*'.

25. La categoría Variables. — Cuando en una fórmula matemática se encuentra lo que llamamos 'una variable', entendemos que en aquel lugar de la fórmula pueden insertarse diversos valores numéricos de una determinada clase. Por ejemplo, si se trata de una fórmula que da la velocidad de un cuerpo en movimiento como función del tiempo, podremos insertar en el lugar de las variables valores numéricos de tiempos, diversos según el momento para el cual nos interese saber la velocidad; pero sólo valores numéricos de tiempo, propiamente nombres (cifras) de instantes.

La mecánica, matematizada desde comienzos de la cultura moderna y convertida en ciencia ejemplar hasta finales del siglo XIX, es precisamente la disciplina que más ha influido en el concepto de variable con el cual se opera en la práctica. Es el concepto de magnitud variable, que se encuentra en toda investigación matemática aplicada. Pero este concepto no es suficientemente claro para su uso en lógica. Cuando se dice, por ejemplo, que la temperatura de una masa dada de gas es una variable, se está usando una expresión útil, pero poco exacta: la temperatura es un número, y no puede variar, en el sentido de que no puede ser otra cosa nueva y seguir

siendo él mismo, como, por ejemplo, varía en ese sentido propio una persona. Afirmar que también el número puede variar así es profesar una dialéctica grosera que renuncia a aclarar los conceptos. En nuestro ejemplo, lo que varía no es la cifra que da en una escala termométrica la temperatura del gas en un momento dado: esa cifra se queda siempre tranquilamente en su sitio; la cifra no puede variar sin dejar de ser ella; lo que varía sin dejar de ser él mismo es el cuerpo material a cuyos estados se atribuyen esos valores numéricos, es decir, la masa de gas considerada.

Desde el punto de vista del análisis lógico del lenguaje, los símbolos llamados 'variables' no denotan o significan cosas que varían. Las variables no denotan, no son nombres. Una variable puede ser sustituida por el nombre de un objeto cualquiera de un determinado campo — así, en nuestro ejemplo de la velocidad, por el nombre de cualquier instante (ese nombre será una cifra), y en el ejemplo del gas por el nombre de cualquier temperatura (que será también una cifra). La variable reserva en un esquema un lugar para cualquiera de diversos nombres pertenecientes a un campo determinado y fijado por el tema del esquema, por su universo del discurso. Las variables son esquematizaciones de elementos del lenguaje, generalmente de nombres.

En cuanto se depura de este modo el concepto de variable como símbolo se aprecia su utilidad para la lógica. La lógica no se interesa por el contenido empírico del lenguaje, sino por su forma. Ahora bien: el contenido empírico del lenguaje está principalmente representado por los nombres. La sustitución de los nombres por variables es pues un modo natural de explicitar la forma lógica de los enunciados.

Los que antes describimos como símbolos para individuos y propiedades cualesquiera son *variables individuales* ('*u*', '*v*', '*w*', '*x*', '*y*', '*z*', etc.) y *variables predicativas* ('*P*', '*Q*', '*R*', '*S*', etc.) respectivamente.

La simbolización de las variables predicativas fue ya practicada por Aristóteles.

Otros símbolos variables frecuentes son las *variables de enunciado*: '*p*', '*q*', '*r*', '*s*', '*t*', ..., que reservan en una fórmula lugares para enunciados o fórmulas.

Un *parámetro* es un símbolo que tiene oficio de constante para cada subclase de la clase de objetos a los que es aplicable, pero es variable en el sentido de que no tiene un sustituyente único para todas las subclases de dicha clase. Por ejemplo, en la fórmula que da el alargamiento de una varilla metálica de longitud l_0 a 0° al calentarla hasta t° ,

$$(1) \quad \text{Alargamiento} = \lambda \cdot l_0 \cdot t,$$

' λ ' representa lo que se llama el 'coeficiente de dilatación lineal', el cual es distinto para cada metal. La fórmula (1) es aplicable a la clase de todas las varillas de metal. ' λ ' es una constante para cada subclase de esa clase (la subclase

de las varillas de hierro, la subclase de las varillas de plomo, etc.), pero es una variable para la clase de todas las varillas de metal. ' λ ' es un parámetro.

El papel sintáctico o gramatical de la variable en los cálculos y lenguajes formalizados es parecido al del pronombre indefinido (y demostrativo y relativo) en los lenguajes étnicos. Así por ejemplo, la fórmula

$$Px,$$

que suele leerse '*x* es *P*', o, simplemente, '*Px*' (pe equis), puede leerse, puesto que '*x*' no significa nada concreto:

cualquiera [cosa] es *P*.

La concepción pronominal de la variable se debe a W. V. O. Quine. Tiene dos ventajas: subraya el hecho de que la variable no es un nombre y permite precisar cuáles son las cosas que un cálculo admite como denotables. Este segundo motivo es el que más importa a Quine.

26. Los cuantificadores. — Es poco frecuente que en una teoría sean interesantes fórmulas como la que hemos utilizado en 25 para ilustrar el papel pronominal de la variable. Una teoría científica aspira a enunciar leyes, teoremas con una validez más o menos universal y, en todo caso, bien precisada. Le interesa poder afirmar que todos los individuos de un cierto campo son *P*, o que algunos lo son, o que ninguno lo es. O sea: afirmar que el enunciado o esquema '*Px*' vale para todo *x*, o para algún *x*, o para ningún *x*. Los enunciados teóricos son comúnmente de las siguientes formas esquemáticas (aunque con mayor complicación):

para todo *x*, *Px*, o
para algún *x*, *Px*, o
para ningún *x*, *Px*.

En la tradición lógica, la lectura de esos tres esquemas era, respectivamente:

todo *x* es *P*;
algún *x* es *P*;
ningún *x* es *P*.

Las expresiones utilizadas para cuantificar los enunciados, esquemas o fórmulas se llaman '*cuantificadores*', y desde Aristóteles suelen utilizarse dos: uno para 'todos' o 'todo', y otro para 'algunos', o 'algún'. 'Ningún' se puede expresar por la negación de 'algunos', siempre que se entienda 'algunos' precisamente en el sentido de 'al menos uno'. El cuantificador 'todos' se llama 'universal'; el cuantificador 'al menos uno' se llama 'existencial'. Otros nombres bastante corrientes son 'generalizador' y 'particularizador', respectivamente.

Hay una diferencia lógica importante entre los dos cuantificadores: el primero no hace ninguna referencia a la existencia de individuos que puedan nombrarse en el lugar de la variable cuantificada. Así, por ejemplo, aunque no existen habitantes de la Ínsula Barataria, el enunciado siguiente es verdadero:

todos los habitantes de la Ínsula Barataria son súbditos de Sancho Panza.

Ese enunciado, más detalladamente analizado según métodos que ya conocemos, puede sustituirse con ventaja por el siguiente:

- (1) todo lo que es habitante de la Ínsula Barataria es súbdito de Sancho Panza.

En cambio, el segundo cuantificador se usa de un modo que conlleva una afirmación de la existencia de al menos uno de tales individuos. Así el enunciado

algún habitante de la Ínsula Barataria es súbdito de Sancho Panza, precisado como el anterior, toma la forma

- (2) hay al menos un habitante de la Ínsula Barataria, y ese habitante es súbdito de Sancho Panza,

enunciado que es falso, pues no existe nada que sea habitante de la Ínsula Barataria.

Ese diverso alcance de los cuantificadores, sólo el segundo de los cuales tiene lo que suele llamarse 'alcance existencial', se desprende de un análisis algo más detallado que el que practicamos en el capítulo I y hemos utilizado ahora. Cuando se dice que todos los habitantes de la Ínsula Barataria son súbditos de Sancho Panza, se está enunciando una "ley" sin excepciones. El hecho de no tener ninguna excepción puede interpretarse diciendo que el primer esquema — ' x es habitante de la Ínsula Barataria' — acarrea, como condición, el segundo — ' x es súbdito de Sancho Panza' —. Con este análisis podemos poner, en lugar del enunciado (1), el enunciado

- (1') Para todo x , si x es habitante de la Ínsula Barataria, entonces x es súbdito de Sancho Panza.

Ahí no hay ninguna afirmación de existencia, sino sólo la afirmación de una relación condicional. Y ésta es verdadera o falsa independientemente de la existencia de habitantes de la Ínsula Barataria.

En cambio, cuando se dice que al menos un habitante de la Ínsula Barataria es súbdito de Sancho Panza, no se está enunciando una "ley", una regularidad universal (pues entonces se habría hecho la afirmación respecto de todos los habitantes de la Ínsula Barataria). Se está afirmando sólo un hecho empírico: el hecho casual o accidental de la coincidencia, al menos

en un individuo, de las propiedades de ambos esquemas o enunciados. Llevando el enunciado (2) al mismo nivel de análisis al que hemos llevado el enunciado (1), tendremos:

al menos un x es habitante de la Ínsula Barataria y es súbdito de Sancho Panza,

lo cual es la afirmación de un hecho concreto, por tanto, una afirmación de existencia, que puede formularse aún más explícitamente así:

- (2') hay al menos un x , tal que x es habitante de la Ínsula Barataria y x es súbdito de Sancho Panza.

Precisamente por esa afirmación de existencia que contiene se llama 'existencial' a este cuantificador.

Otro modo de ilustrar la diferencia de alcance entre los dos cuantificadores consiste en considerar el diferente efecto que tendría para uno y otro, en el caso del ejemplo, la inexistencia de habitantes de la Ínsula Barataria. Para refutar el esquema 'si x es habitante de la Ínsula Barataria, entonces x es súbdito de Sancho Panza', hace falta probar que *hay* un habitante de la Ínsula Barataria que no es súbdito de Sancho Panza. Por tanto, el hecho de que no haya habitantes de la Ínsula Barataria tiene como consecuencia la irrefutabilidad de aquel esquema universal. En cambio, ese hecho basta para refutar que algún habitante de la Ínsula Barataria es súbdito de Sancho Panza.

La simbolización más frecuente del cuantificador universal, o generalizador, consiste en unos paréntesis dentro de los cuales se escribe la variable generalizada. Así, por ejemplo, el esquema

todo x es P

se simboliza

$(x)Px$,

que es costumbre leer: 'para todo x , Px '.

El cuantificador existencial, o particularizador, suele simbolizarse anteponiendo a la variable cuantificada el símbolo ' $\exists$ ' una 'E' invertida (alusión al latín 'existit'). Así, por ejemplo, el esquema 'algún x es P ', o, más propiamente,

hay al menos un x que es P ,

se simboliza

$\exists xPx$,

que es corriente leer: 'hay [al menos] un x [tal que] Px '.

El oficio sintáctico de los cuantificadores es como el de adjetivos (como 'todos', 'algunos', o expresiones adjetivas compuestas con ellos) que determinan pronombres (es decir, variables). Así, según la paráfrasis que antes usábamos para ilustrar el oficio pronominal de la variable, la fórmula

$$(x)Px$$

puede leerse: 'todo algo es P ', en donde el adjetivo 'todo' cumple el oficio del cuantificador '()', determinando el pronombre 'algo', el cual cumple el papel de la variable ' x '.

El artículo es una especialización del adjetivo determinativo, como sugiere, por ejemplo, el hecho de que el artículo determinado castellano, 'el', 'la', 'lo', proceda del uso adjetivo del demostrativo latino 'ille', 'illa', 'illud'. Hay un tercer cuantificador interesante que puede concebirse, desde el punto de vista gramatical que hemos adoptado, como artículo determinado. Es el *descriptor*, que sirve para componer, partiendo de esquemas, una especie de nombres propios compuestos (*descripciones*). Por ejemplo: 'el autor del *Quijote*' es una descripción de Cervantes. El esquema correspondiente es

$$x \text{ es autor del } \textit{Quijote};$$

y la descripción que lo convierte en una especie de nombre se simboliza por:

$$(\iota x)(x \text{ es autor del } \textit{Quijote}),$$

o, en general,

$$(\iota x)(Px),$$

que se lee: 'el x tal que Px ', o 'el x que es P '. (Si ' P ' es 'ser autor del *Quijote*', esa descripción se leerá: 'el x tal que x es autor del *Quijote*', o 'el x que es autor del *Quijote*', o, simplemente, 'el autor del *Quijote*').

Todos los cuantificadores pueden entenderse como pertenecientes a la categoría de los operadores, que se estudiarán a continuación; pues todos ellos indican una operación practicada sobre la variable afectada: el cuantificador universal, una generalización; el existencial, una particularización; el descriptor, una singularización.

27. Categorías compositivas o conjuntivas. — Diversas "partes de la oración" desempeñan en el lenguaje común oficio de composición, señaladamente de enlace o conjunción. Las conjunciones llevan este nombre por antonomasia, pero también hay, por ejemplo, adverbios (como 'más', 'menos') que desempeñan papeles compositivos. Por ejemplo, en 'dos más cuatro', 'más' compone 'dos' con 'cuatro'. El adverbio 'no', aplicado a un

enunciado, ' p ', compone un nuevo enunciado, 'no- p '. Etc. Cualquiera que sea su clasificación en la gramática del lenguaje vivo, en la sintaxis de los lenguajes formalizados consideraremos conjunciones a todos los símbolos que desempeñen oficio de enlace, o, en general, de composición.

Distinguiremos dos grandes grupos de conjunciones en este sentido: los operadores y las conectivas.

Operadores son conjunciones que se aplican a constantes o variables individuales. Un operador es, por ejemplo, '+' en la fórmula siguiente:

$$x + y.$$

Los operadores tienen ese nombre porque representan una operación. Se trata de una operación sobre símbolos tomados como símbolos individuales, y su resultado se expresa frecuentemente por un símbolo individual. Por ejemplo:

$$x + y = z.$$

Una operación puede también concebirse como una función. Así, para atenernos al mismo ejemplo, la anterior fórmula, entendida funcionalmente, puede escribirse:

$$+ (x, y) = z,$$

ocupando '+' el lugar en el que corrientemente se pone una ' f ', o, a veces, símbolos más alusivos, como ' s ' (función suma) en el caso del ejemplo. En realidad también una función puede concebirse como una operación. Pero por la especial importancia del concepto de función, lo consideraremos aparte más adelante. Por el momento nos limitaremos a indicar que los operadores que representan funciones se llaman '*functores*' (la expresión 'característica funcional', usada por los matemáticos, no suele emplearse en lógica).

Las *conectivas* son conjunciones que se aplican a enunciados o fórmulas. Según nuestro uso de la palabra 'conjunción', las partículas que expresan, por ejemplo, un enlace condicional entre dos enunciados o fórmulas (como el par 'si..., entonces') son conjunciones, cualquiera que sea la categoría a que pertenezcan en los lenguajes étnicos. El símbolo que usaremos para expresar el enlace condicional (la conectiva condicional) es '→'. Que p es condición de q se simboliza

$$p \rightarrow q,$$

que suele leerse: 'si p , entonces q '.

Un ejemplo de conectiva muy coincidente con la noción de conjunción del lenguaje común es la que sirve para componer la afirmación simultánea de enunciados. La llamaremos por antonomasia 'conjunción', y la simbolizaremos por ' $\wedge$ '. Así la afirmación simultánea de ' p ' y ' q ' se simbolizará:

$$p \wedge q.$$

Tanto a propósito de conectivas cuanto de operadores y de cuantificadores — los cuales, como todos los adjetivos determinativos, también pueden considerarse operadores — se habla de operandos en el mismo sentido que en aritmética. En el caso de cuantificadores suele distinguirse entre operando y variable de operador. Así por ejemplo, en la fórmula

$$(x) Px$$

el operando es ' Px ' y la variable de operador es ' x '.

28. Funciones lógicas. Abstracción funcional. — La noción de función está muy relacionada con la de variable, y ha nacido también en la mecánica física. La manera corriente de hablar de funciones tiene que corregirse algo para que el concepto sea claro y útil en lógica. Así, el decir, como lectura de la expresión ' $f(a) = b$ ', que b es la función f de a , resulta incorrecto desde el punto de vista lógico. En realidad, ' b ' es el nombre del valor de la función f (y no la función f misma) para el argumento cuyo nombre es ' a '. La función, nombrada por el functor ' f ', es una operación que consiste en correlacionar dos campos de valores, llamados respectivamente de los argumentos y de los "valores", de tal modo que a cada individuo del primer campo corresponda por la función un individuo, y sólo uno (correspondencia unívoca) del segundo. (En cambio, a cada individuo del segundo pueden corresponder por la función uno o más del primero, como ocurre, por ejemplo, con la función trigonométrica seno; la correspondencia no tiene necesariamente que ser biunívoca.)

Este concepto de función se debe al matemático Lejeune Dirichlet (1805-1859), el cual no lo formula tanto en términos de operación cuanto de la correlación misma. La operación de correlatar puede entenderse ya, en efecto, como el cálculo hecho en cada caso. La función es más bien el esquema de la operación.

De esa noción de función se desprende que su uso no tiene por qué ser exclusivo del análisis matemático. La definición, en efecto, no dice que los campos de argumentos y "valores" deban ser de determinada naturaleza. Pueden ser de cualquier clase de valores, con tal de que se pueda establecer entre ellos aquella correlación unívoca.

Así precisada la idea de función, su introducción en lógica es de mucha utilidad. En la parte apofántica del lenguaje hay, en efecto, formaciones que se parecen a las funciones de las matemáticas: su valor depende de un término que, al llevar la expresión a esquema, resulta ser una variable. El resto de la expresión resulta indicar la operación o correlación que hay que practicar con la variable para obtener el valor de toda la expresión.

He aquí un ejemplo: la expresión

el autor del Quijote

puede dar el esquema

el autor de x ;

' x ' es una variable cuyo campo podemos definir como la clase de las obras literarias. Según el título de obra (con fecha de aparición) que se ponga en el lugar de ' x ', el valor de la expresión será un determinado escritor. Además: a un determinado escritor pueden corresponder varios títulos (a un "valor" varios argumentos), cosa permitida por la definición de función; pero nunca corresponderán a un mismo título varios autores (a un argumento varios "valores"), si se admite como regla sintáctica que las obras escritas en colaboración son obras de un autor compuesto, por así decirlo, o sea, de un equipo. La expresión 'el autor de' resulta ser nombre de una función, pues cumple todos los requisitos de nuestra definición del concepto. Pero es una función cuyos argumentos y "valores" no son numéricos.

Con ayuda de las funciones podemos desentendernos del sentido de las formaciones del lenguaje, para atenernos a la consideración puramente extensional de sus valores, o denotaciones. Este paso se llama '*abstracción funcional*'. Mediante ella se pueden asociar funciones a formaciones lingüísticas, y hablar de valores en vez de hablar de sentidos. Así, por ejemplo, las dos expresiones

$$\begin{aligned} x + 5 &= 2, \\ x + 7 &= 4, \end{aligned}$$

que tienen distintos sentidos o intenciones, pues son dos modos de denotar $x + 3$, poseen la misma función asociada, a saber:

$$f(x) = x + 3.$$

Una importante aplicación de esto es la asociación de funciones a las conectivas.

Función asociada a una conectiva. — Desde Aristóteles hasta la lógica moderna — y en ésta, sobre todo, por obra de Frege — el enunciado (apofántico) se define como la formación lingüística que tiene uno de dos valores: la verdad (V) o la falsedad (F). Esta atribución de *valores veritativos* a los enunciados facilita la asociación de funciones a las conectivas.

Una conectiva indica, en efecto, una composición de enunciados, y como los enunciados pueden reducirse a uno de los dos valores V o F , es posible considerar una conexión de enunciados como una función asociada con la conectiva de tal modo que cuando la conexión es verdadera, la aplicación de la función a sus argumentos, en el mismo orden en que la conectiva enlaza los enunciados, dé el valor V , y F en otro caso. Los campos de argumentos y "valores" de las funciones asociadas a conectivas están pues constituidos exclusivamente por los dos valores veritativos.

Sea la conectiva ' $\wedge$ ' (conjunción), que ya conocemos. Es una conectiva que enlaza enunciados, esquemas o fórmulas, por ejemplo:

$$p \wedge q.$$

Dada esa conexión, podemos abstraer una función $f(p, q)$, tal que cuando la fórmula ' $p \wedge q$ ' sea verdadera, $f(p, q) = V$, y cuando ' $p \wedge q$ ' sea falsa, $f(p, q) = F$. De este modo podemos prescindir del sentido de ' $\wedge$ ' (que es el de 'y' en el lenguaje común) y operar con unas tablas de valores. Puesto que p, q , no pueden asumir más que un valor cada uno en cada combinación; y puesto que todo el campo de valores que pueden asumir se limita a la clase $\{V, F\}$, tendremos los siguientes cuatro pares de valores ordenados como argumentos de f (es decir, como interpretaciones de ' (p, q) ')

	p	q
1.º	V	V
2.º	V	F
3.º	F	V
4.º	F	F

Admitiendo, de acuerdo con el uso de 'y' en el lenguaje cotidiano, que ' $p \wedge q$ ' sólo es verdadero en el caso de que ambos, ' p ' y ' q ', lo sean, podemos convenir en la siguiente definición de nuestra función:

$$\begin{aligned} f(V, V) &= V, \\ f(V, F) &= F, \\ f(F, V) &= F, \\ f(F, F) &= F, \end{aligned}$$

Con esto hemos pasado de la consideración intensional del sentido de la conectiva ' $\wedge$ ' ('y') a la consideración extensional de una *función veritativa*, f , definida convencionalmente por una tabla de valores para (pares de) argumentos.

La apelación a las nociones de verdad y falsedad no es imprescindible para la definición de funciones asociadas a conectivas. Aquí se hace por el punto de vista semántico adoptado.

APÉNDICE AL CAPÍTULO IV

A. *Texto citado*. — Aristóteles, *De Int.*, III, 17a 2-3.

B. *Observación*. — a 28: la abstracción funcional se justifica más explícita y formalmente en el capítulo XV.

PARTE SEGUNDA

EL SISTEMA DE LA LOGICA ELEMENTAL

Sección primera. — EL LENGUAJE DE LA LÓGICA ELEMENTAL

CAPÍTULO V

LA COMPOSICIÓN DE ENUNCIADOS. LÓGICA DE ENUNCIADOS

29. Concepto de la lógica de enunciados. — En el capítulo I observamos que a una formación dada del lenguaje común pueden atribuirse varias formas o esquemas (cfr. 5-6), según el detalle del análisis que se practique con ella. Así, los ejemplos *1a* y *2a* de aquel capítulo se representaron primero por el Esquema I, y luego por el Esquema II. Del Esquema I dijimos entonces que era fruto de un análisis rudimentario. Ese análisis se limitaba, en efecto, a descubrir categoremas que enlazan enunciados atómicos, pero sin entrar en el análisis de éstos. Por eso lo único que este análisis puede descubrir es la forma de los enunciados moleculares, o, más precisamente, lo molecular de esa forma. En el Esquema I del capítulo I los enunciados atómicos están indicados como categoremas, lugares de contenido (allí representados por trazos).

Pues bien: éste es el género de análisis en que se basa la lógica de enunciados. La lógica de enunciados estudia la composición, por medio de conectivas (cfr. 27), de enunciados moleculares a partir de enunciados atómicos que no se analizan.

30. Símbolos elementales o primitivos de la lógica de enunciados. — La primera fase de la construcción de un cálculo, o de un lenguaje formalizado, es el establecimiento de sus símbolos elementales o primitivos. Los de la lógica de enunciados son los siguientes, que reflejan el género de análisis en que se basará la construcción (síntesis) del cálculo:

Variables de enunciado. — Los lugares de contenido de expresiones de la lógica de enunciados están reservados para enunciados. Se usan para ocupar esos lugares las variables de enunciado que ya conocemos (cfr. 25):

$p, q, r, s, \dots$

Los demás símbolos del lenguaje son constantes lógicas.

Funciones veritativas. — La composición de enunciados moleculares a partir de enunciados se realiza en el lenguaje común mediante las partículas de composición, o conectivas. Como ya se vio (cfr. 23), para la construcción de un cálculo esas conectivas pueden sustituirse con ventaja por funciones asociadas, o funciones veritativas, que permiten un tratamiento de las conexiones como meras correlaciones de valores. Los funtores veritativos son los símbolos de esas funciones.

Funciones veritativas monádicas. — Dado un solo enunciado cualquiera, el cual no puede tomar más que uno de los dos valores V y F (sucesivamente), es posible construir las siguientes correlaciones de valores para la aplicación de una función veritativa de un solo argumento (función veritativa monádica), g_n , al mismo:

Valor del enunciado (argumento)	Valores del enunciado molecular compuesto por la aplicación de la función veritativa monádica g_n a ' p '.			
p	$g_1(p)$	$g_2(p)$	$g_3(p)$	$g_4(p)$
V	V	F	V	F
F	V	F	F	V

Esa tabla de correlaciones define cuatro funciones veritativas monádicas. g_1 y g_2 son funciones constantes: su valor es el mismo para cualquier valor del argumento. g_3 es una función que correlaciona el argumento consigo mismo. g_4 da en cambio como valor el contravalor del argumento.

Sólo la función g_4 es de uso corriente en lógica; se utiliza como función veritativa asociada a la conectiva 'no'. Se simboliza por ' $\sim$ ' antepuesto al argumento, y se lee 'no'. Así la expresión

$$\sim p$$

se lee 'no- p '. Su corriente representación por la siguiente tabla coincide bien con el uso de la conectiva 'no' en el lenguaje común

p	$\sim p$	(Lectura: ' $\sim V = F$ y ' $\sim F = V$ ').
V	F	
F	V	

La tabla puede parafrasearse del modo siguiente (en términos de conectiva, más que de función): 'cuando un enunciado es verdadero, su negación es falsa, y cuando es falso su negación es verdadera'.

' $\sim$ ' es nuestro primer functor veritativo.

Funciones veritativas diádicas. — Las funciones veritativas diádicas (de dos argumentos) se usan como asociadas a conectivas diádicas. Éstas son las que enlazan dos enunciados para componer otro molecular. Las funciones veritativas diádicas se aplican, pues, a dos valores veritativos como argumentos, y dan como valor otro valor veritativo. Hay 16 funciones veritativas diádicas (2^2 ; pues los dos valores de cada argumento, al combinarse con los dos del otro, dan 2^2 composiciones; y cada una de éstas se compone a su vez con los dos valores posibles del "valor" de la función). He aquí su tabla:

- 1: valores de los enunciados atómicos (argumentos).
- 2: valores del enunciado molecular compuesto por la aplicación de la función veritativa diádica f_n a ' p ', ' q ' (por este orden).

1		2															
p	q	f_1	f_2	f_3	f_4	f_5	f_6	f_7	f_8	f_9	f_{10}	f_{11}	f_{12}	f_{13}	f_{14}	f_{15}	f_{16}
V	V	V	V	V	V	F	V	F	V	F	V	F	V	F	F	F	F
V	F	V	V	V	F	V	V	F	F	V	F	V	F	F	F	F	F
F	V	V	V	F	V	V	F	V	V	F	F	V	F	F	V	F	F
F	F	V	F	V	V	V	F	V	F	V	V	F	F	F	F	V	F

No todas esas funciones son de uso corriente en el cálculo de enunciados. f_1 y f_{16} , funciones constantes, no se usan prácticamente nunca. Tampoco f_6 , f_7 , f_8 y f_9 se han utilizado. f_3 , f_5 , f_{11} , f_{13} , f_{14} y f_{15} se han utilizado y se utilizan algunas veces. f_2 , f_4 , f_{10} y f_{12} son de uso corriente como funciones asociadas a conectivas frecuentes en el lenguaje común. Nuestro estudio se limitará por ahora a dichas cuatro funciones diádicas, que son las normalmente usadas en la presentación de la lógica de enunciados. En el capítulo XIII se discutirá y justificará esa limitación.

f_2 se usa como asociada a la disyunción no excluyente, o sea, a ' $\vee$ ' (entre enunciados) en el sentido del 'vel' latino. Una disyunción no excluyente es verdadera con sólo que lo sea uno de sus dos miembros, y también cuando los dos componentes son verdaderos. Esa es la correspondencia que presenta la tabla de f_2 . Esta función se simboliza por ' $\vee$ ' escrito entre los dos enunciados en disyunción, y se lee 'o'. Así la expresión

$$p \vee q$$

se lee ' p o q '. La forma corriente de dibujar su tabla es

v	V	F	(Lectura: ' $V \vee V = V$, $V \vee F = V$, $F \vee V = V$, $F \vee F = F$.)
V	V	V	
F	V	F	

Los valores de la primera columna vertical empezando por la izquierda son los del primer argumento de la disyunción; los de la primera fila horizontal empezando por arriba son los del segundo argumento de la disyunción. En el ángulo superior izquierdo está el functor. El valor de la función disyunción para dos valores se encuentra en el ángulo inferior derecho, en la intersección de la fila del valor del primer argumento con la columna del valor del segundo argumento.

La función f_{11} , poco usada, es la asociada a la disyunción excluyente, al 'aut' latino (entre enunciados). Su valor es V cuando uno de los dos argumentos es verdadero y el otro falso. Y F cuando ambos son verdaderos y cuando ambos son falsos.

f_4 se usa como función asociada a la conectiva condicional, o sea, al par 'si..., entonces...'. Se simboliza por ' $\rightarrow$ ', y se lee 'si..., entonces...', o, simplemente, flecha. Así la expresión

$$p \rightarrow q$$

se lee: 'si p , entonces q ' o ' p flecha q '.

La tabla de f_4 o su dibujo más corriente

$\rightarrow$	V	F
V	V	F
F	V	V

puede parafrasearse así: 'un condicional sólo es falso cuando el antecedente es verdadero y el consecuente es falso; en los demás casos, el condicional es verdadero'.

La función $\rightarrow$ no es tan coincidente con el uso intuitivo como las funciones $\sim$ o $\vee$. La intuición suele tropezar con el caso tercero de la tabla de ' $\rightarrow$ ' ($F \rightarrow V = V$). Ello se debe a lo siguiente: en el lenguaje común puede confundirse lo que venimos llamando 'condición' con la implicación, que es la relación lógica por la cual un enunciado "contiene" ya a otro u otros, como, por ejemplo, ' a es mayor que b ' contiene a ' b es menor que a '. Y tratándose de esta relación no está, ciertamente, justificado decir que un enunciado falso implica otro verdadero. Una observación gramatical puede ser útil para distinguir entre las nociones de condicional e implicación. La implicación es una conectiva que se escribe entre nombres de enunciados o de conjuntos de enunciados, por ejemplo:

los axiomas de la geometría euclídea implican el teorema de Pitágoras.

'El teorema de Pitágoras' es un nombre del enunciado: 'la suma de los cuadrados de los catetos de un triángulo rectángulo es igual al cuadrado de la hipotenusa'. 'Los axiomas de la geometría euclídea' es un nombre de todo un conjunto de enunciados.

En cambio, el condicional ' $\rightarrow$ ' no se escribe entre nombres de enunciados sino entre enunciados, los cuales son nombres de hechos. Un condicional es, por ejemplo,

si Juan viene, Luis se va.

La relación es aquí entre hechos. En general ' $p \rightarrow q$ ' significa una determinada relación entre el hecho p (expresado por el enunciado ' p ') y el hecho q (expresado por ' q ').

Esa diferencia gramatical se debe a lo siguiente: la implicación es una relación basada en la estructura de los enunciados. Es la estructura del teorema de Pitágoras la que está implicada, la que está "contenida", en los axiomas de la geometría euclídea, o en la estructura de éstos. Por eso, dados los axiomas de la geometría euclídea, puede verse en ellos, sin más ayuda, por así decirlo, el teorema de Pitágoras, al que contienen, y que puede deducirse de ellos. En cambio, dado el hecho de que Juan viene, no puede verse sin más en él el hecho de que Luis se va. Para saber que Luis se va hay que conocer además otro hecho, a saber, que la venida de Juan es condición suficiente de la ida de Luis, o sea, la relación condicional entre los dos hechos. Y ésta es una relación empírica, factual.

Una consecuencia de esto es que no será una verdadera implicación una expresión que conste de ' $\rightarrow$ ' entre dos solos enunciados atómicos distintos y usado cada uno de ellos una sola vez; por ejemplo:

$$p \rightarrow q.$$

La razón de ello es que no conocemos la estructura de ' p ' ni la de ' q '. Pero si, además, contáramos con las definiciones:

$$\begin{aligned} 'p' &\leftrightarrow_{af} '[r \rightarrow s] \wedge [t \rightarrow r]' \\ 'q' &\leftrightarrow_{af} '[t \rightarrow s]', \end{aligned}$$

entonces podríamos decir que ' $p \rightarrow q$ ', además de un condicional, es una implicación, pues ' $[r \rightarrow s] \wedge [t \rightarrow r]$ ' "contiene" a ' $t \rightarrow s$ '.

La fórmula ' $[r \rightarrow s] \wedge [t \rightarrow r] \rightarrow [t \rightarrow s]$ ' es siempre verdadera (verdadera para cualquier interpretación posible) en razón de su estructura. Esto nos sugiere una relación entre el condicional y la implicación. Implicación es condicional siempre válido, verdadero para cualquier interpretación de las variables. Toda implicación es, por tanto, expresable mediante un condicional (mediante un condicional universalmente verdadero, que será lo que en el capítulo I llamamos una 'verdad formal'). Pero no todo condicional puede expresarse como una implicación, sino sólo aquellos

de cuyos enunciados sepamos que constituyen por su estructura condicionales válidos para cualquier interpretación.

La utilidad del condicional, que compensa su relativa oscuridad intuitiva, es que permite construir de modo extensional, por meras correlaciones de valores, una relación emparentada con la implicación, la cual es, en cambio, intensional, depende de las significaciones de los enunciados, manifiesta en sus estructuras.

La función f_{10} se usa como función asociada a la igualdad de valor, o equivalencia, que no es lo mismo que identidad de sentido. Se simboliza por ' $\leftrightarrow$ ' y se lee '...si y sólo si...' o, simplemente, 'bicondicional' o 'doble flecha'. Otra lectura corriente es '...equivale a...'. Así la expresión

$$p \leftrightarrow q$$

se lee: ' p si y sólo si q ', ' p equivale a q ', ' p doble flecha q '.

La tabla de la función veritativa $\leftrightarrow$ suele trazarse del modo siguiente, que se compadece bien con los usos intuitivos:

$\leftrightarrow$	V	F
V	V	F
F	F	V

Pero la aceptabilidad intuitiva de ' $\leftrightarrow$ ' esconde la misma "paradoja" que ' $\rightarrow$ '. En realidad, el nombre de 'doble flecha' es más adecuado a ' $\leftrightarrow$ ' que cualquier otro. Consideremos, por ejemplo, el enunciado

(1) Juan viene equivale a Luis se va,

y comparémosle con

(2) si Juan viene, Luis se va.

La diferencia entre (1) y (2) está en que el condicional no excluye que Luis se vaya por otras causas distintas de la venida de Juan, aunque éste no venga, mientras que la doble flecha de (1) excluye esa posibilidad. Por tanto, (1) puede escribirse del modo siguiente:

(1a) Luis se va si viene Juan, y sólo si viene Juan.

O, más brevemente:

(1b) Luis se va si y sólo si viene Juan.

En esta última formulación se aprecia la presencia de dos condicionales:

a) uno es muy visible, aunque con el orden de antecedente y consecuente invertido. Es el condicional: 'Si viene Juan, Luis se va'.

b) El otro está menos visible: puesto que Luis sólo se va en caso de que

venga Juan, entonces si Luis se va, es que Juan viene. Se trata del condicional: 'Si Luis se va, Juan viene'.

Esquemmatizando 'Juan viene' por ' p ' y 'Luis se va' por ' q ', el enunciado (1b), que es una traducción de la equivalencia (1), puede esquematizarse del modo siguiente:

$$[p \rightarrow q] \wedge [q \rightarrow p].$$

Esta conjunción de condicionales es lo que llamamos 'bicondicional' y simbolizamos por

$$p \leftrightarrow q.$$

La función f_{12} se usa como función veritativa asociada a la conectiva 'y'. Se simboliza, como vimos, por ' $\wedge$ ', y se lee 'y'. Así la expresión

$$p \wedge q$$

se lee ' p y q '. Se le suele trazar la siguiente tabla, que coincide con el uso corriente de 'y':

$\wedge$	V	F
V	V	F
F	F	F

El uso de paréntesis. — Además de las variables de enunciado y de las constantes lógicas (functores veritativos), necesitamos en la práctica paréntesis (usaremos corchetes) para precisar la lectura de las expresiones. Sin ellos, la expresión

$$p \rightarrow q \rightarrow r,$$

por ejemplo, es ambigua en el texto escrito, pues puede leerse como

$$p \rightarrow [q \rightarrow r],$$

o como

$$[p \rightarrow q] \rightarrow r.$$

El uso de los paréntesis en lógica es sustancialmente el mismo que tienen en matemáticas, es decir, en general, el de signos de puntuación. Para ahorrar algunos paréntesis, se conviene corrientemente en que ' $\wedge$ ' y ' $\vee$ ' ligan más fuerte que ' $\rightarrow$ ' y ' $\leftrightarrow$ ', y menos fuerte que ' $\sim$ '. Así la expresión

$$\sim p \wedge q \rightarrow r$$

debe leerse

$$[[\sim p] \wedge q] \rightarrow r.$$

Los paréntesis no son símbolos esenciales al cálculo, como lo son las variables y las constantes lógicas. Son un mero expediente de puntuación para facilitar la lectura; se puede prescindir de ellos en teoría, y se los puede sustituir en la práctica por puntos, espacios en blanco, o por pausas en el lenguaje hablado.

31. Fórmulas de la lógica de enunciados. — El segundo paso en la construcción de un cálculo consiste en definir el concepto de fórmula de dicho cálculo, lo cual se hace mediante unas(s) regla(s) de formación.

Las fórmulas de un cálculo son los elementos homólogos de las oraciones del lenguaje natural. En el lenguaje natural el concepto de oración se define de un modo intensional y apelando a significaciones. En cambio, en un cálculo la categoría Fórmula se define de un modo sintáctico, indicando simplemente, mediante la(s) regla(s) de formación, cómo está permitido componer los símbolos (concatenarlos) para que el resultado de la composición sea una fórmula, una expresión bien hecha.

Las reglas de formación tienen la forma de condicionales del metalenguaje del cálculo de que se trate. Una de ellas puede decir, por ejemplo:

Si ' p ' y ' q ' son fórmulas, entonces el resultado de concatenarlas con ' $\rightarrow$ ' en el centro, ' $p \rightarrow q$ ', es también una fórmula.

En nuestra definición no usaremos explícitamente la noción de concatenación.

Las reglas de formación, cuando se dan como varias, se apoyan unas en otras: las que permiten comprobar fórmulas más complicadas tipográficamente se apoyan en las que permiten comprobar fórmulas más sencillas. El conjunto de todas las reglas de formación en ese sentido, que da la noción general de fórmula del cálculo, constituye lo que suele llamarse una 'definición recurrente o recursiva'; el nombre puede entenderse en el sentido de que cada regla referente a fórmulas más complicadas, convertida en línea de la definición, recurre a reglas anteriores, a líneas precedentes de la misma definición, las cuales se refieren a fórmulas más sencillas.

Como es natural, la definición recursiva de la noción de fórmula de la lógica (y del cálculo) de enunciados presupone una exacta indicación de los símbolos elementales con que se cuenta para componer las fórmulas. Recapitemos pues, para empezar, esos símbolos.

Símbolos elementales o primitivos:

- 1.º Letras de enunciado, ' p ', ' q ', ' r ', ' s ', ... (con subíndices si es necesario: ' p_1 ', ' q_3 ', ...), tantas cuantas se necesiten, dentro de lo que en matemáticas se llama el infinito numerable, es decir, dentro del conjunto infinito de los números naturales.

2.º Functores veritativos: ' $\sim$ ', ' $\vee$ ', ' $\wedge$ ', ' $\rightarrow$ ', ' $\leftrightarrow$ '.

3.º Símbolos auxiliares, como ' $]$ ', ' $[$ ', que pueden ser útiles sin ser esenciales.

Hecho ese catálogo, puede pasarse a definir por recursión la noción de

Fórmulas del cálculo de enunciados:

- (I) ' p ', ' q ', ' r ', ' s ', ..., ' p_1 ', ..., ' s_n ' ..., son fórmulas (atómicas).
- (II) Si X es una fórmula, $\sim X$ también lo es.
- (III) Si X e Y son fórmulas, $X \vee Y$ también lo es.
- (IV) Si X e Y son fórmulas, $X \wedge Y$ también lo es.
- (V) Si X e Y son fórmulas, $X \rightarrow Y$ también lo es.
- (VI) Si X e Y son fórmulas, $X \leftrightarrow Y$ también lo es.
- (VII) Ninguna expresión es una fórmula del cálculo de enunciados sino en virtud de (I)-(VI).

Los símbolos ' X ', ' Y ' son variables sintácticas (cfr. 19).

La definición recursiva de la categoría Fórmula permite una decisión mecánica (algorítmica) de la cuestión de si una expresión dada es o no es una fórmula, sin necesidad de apelar al sentido de dicha expresión. La definición (entendida como el conjunto de las anteriores líneas (I)-(VII)) abraza, en efecto, todas las composiciones correctas. Así, por ejemplo, dada la expresión

$$(3) \quad [\sim p \rightarrow q] \wedge [q \rightarrow r] \rightarrow [\sim p \rightarrow r],$$

para decidir si es o no una fórmula podemos proceder del modo siguiente (sabemos que es una expresión del cálculo de enunciados porque no contiene más que símbolos del mismo):

Primer paso: La expresión es un condicional. La línea (V) de la definición nos dice que (3) será una fórmula si lo son sus dos componentes, el antecedente y el consecuente. Pasemos al consecuente, que es más corto.

Segundo paso: el consecuente es también un condicional. Por tanto, según la línea (V) de la definición, será una fórmula si lo son su antecedente, ' $\sim p$ ', y su consecuente, ' r '. El consecuente lo es por la línea (I). El antecedente es una negación. Según la línea (II) ' $\sim p$ ' será una fórmula si lo es ' p '. Y ' p ' lo es por la línea (I). El consecuente de (3), ' $\sim p \rightarrow r$ ', es pues una fórmula. Pasemos al antecedente:

Tercer paso: el antecedente de (3) es una conjunción. Según la línea (IV) será una fórmula si lo son sus dos miembros.

Cuarto paso: el segundo componente del antecedente de (3) es un condicional. Según la línea (V) de la definición, ese condicional, ' $q \rightarrow r$ ', será una fórmula si lo son ' q ' y ' r '. Estas lo son por la línea (I). Luego ' $q \rightarrow r$ ' es una fórmula.

Quinto paso: el primer componente del antecedente de (3) es un con-

dicional, ' $\sim p \rightarrow q$ '. Según la línea (V) de la definición, será una fórmula si lo son ' $\sim p$ ' y ' q '. ' q ' es una fórmula por la línea (I) de la definición. ' $\sim p$ ' será una fórmula si lo es ' p '. ' p ' lo es por la línea (I) de la definición. Luego ' $\sim p$ ' y ' $\sim p \rightarrow q$ ' son fórmulas.

Sexto paso: el antecedente de (3) es una fórmula, como resultado de los tres últimos pasos.

Séptimo paso: el antecedente y el consecuente de (3) son fórmulas. Luego (3) es una fórmula.

Cada vez que se ha dicho 'fórmula' en lo que precede se ha elidido 'del cálculo de enunciados'.

Examinemos ahora la siguiente expresión:

$$(4) \quad v p \leftrightarrow q \wedge r.$$

Se trata de un bicondicional. Según la línea (VI) de la definición, (4) será una fórmula si lo son sus dos componentes principales. El segundo componente principal es una conjunción. Es una fórmula, en virtud de las líneas (IV) y (I), aplicadas sucesivamente. El primer miembro de la equivalencia no es una fórmula por ninguna de las líneas (I)-(VI) de la definición. Consecuentemente, por virtud de la línea (VII) de la definición, ' $v p$ ' no es una fórmula del cálculo de enunciados. Por esa misma regla no lo es tampoco (4).

Tampoco en la decisión sobre (4) ha intervenido ninguna consideración acerca de su posible sentido, sino sólo la aplicación de la definición de Fórmula.

32. Esquematización de enunciados del lenguaje común por fórmulas de la lógica de enunciados. — Un cálculo es una construcción, una síntesis, como ejemplifica la definición de Fórmula. Y al ir ensamblando piezas de un lenguaje formalizado estamos preparando la máquina del cálculo. Pero, como dijimos, esa síntesis se basa en un previo análisis del discurso, análisis que nos ocupó en la Parte Primera. Ahora bien: la precisión o aclaración de categorías, que es necesaria para la construcción del cálculo, facilita a la inversa el análisis del lenguaje común en el que suele estar expresado el conocimiento positivo. Vamos a atender ahora al modo como puede prestar ese servicio analítico la lógica de enunciados.

A este propósito hay que tener ante todo en cuenta que el superficial nivel de análisis propio de la lógica de enunciados, que es un análisis, por así decirlo, "macroscópico", hace que al esquematizar en su lenguaje los enunciados del lenguaje común se pierdan numerosos rasgos de la estructura o forma de éstos. Mas, por otra parte, ese análisis "macroscópico" es también necesario cuando se analiza el discurso desde un punto de vista más detallado o "microscópico". Pues las conexiones entre enunciados — las conexiones "macroscópicas" — siguen siendo relevantes cuando se atiende a la estructura interna de los enunciados enlazados por las conectivas.

Por eso son útiles en todo caso las aclaraciones que la lógica de enunciados puede dar sobre la estructura "macroscópica" del discurso.

Un buen uso de la *conjunción*, por ejemplo, permite normalizar la forma de los enunciados del lenguaje común que parecen ser atómicos (oraciones simples), pero constan de un solo sujeto y dos verbos. La estructura molecular, y no, como aparentan, atómica, de esos enunciados, como, por ejemplo,

(5) Juan come y bebe,

se aprecia mejor, excluyendo toda ambigüedad, si se esquematiza por

$$(5a) \quad p \wedge q,$$

en donde ' p ' esquematiza 'Juan come', y ' q ' esquematiza 'Juan bebe'.

Así también es útil para la comprensión de la estructura molecular del lenguaje — en cuanto que éste ha de servir para la ciencia, para la expresión objetiva de hechos — la reducción de varias conectivas del lenguaje común, cuyo valor veritativo es el mismo, a una única función veritativa. Así por ejemplo, los enunciados

(6) Aunque Juan come, tiene hambre.

(7) Juan come, pero tiene hambre,

se esquematizan los dos por

$$(6, 7) \quad p \wedge s$$

donde ' p ' esquematiza 'Juan come' y ' s ' esquematiza 'Juan tiene hambre'. 'Aunque' y 'pero' han sido ambos esquematizados por ' $\wedge$ '. ' $\wedge$ ' es en efecto su valor, por más que 'aunque' y 'pero' signifiquen 'y' con sentidos (intenciones) ligeramente distintos. El hecho de que el valor de ambos es el de ' $\wedge$ ' puede apreciarse observando que tanto (6) cuanto (7) son verdaderos si Juan come y tiene hambre, y falsos si Juan no come o no tiene hambre. Esto es lo que indica la tabla de ' $\wedge$ '.

El siguiente enunciado permite apreciar la insuficiencia de la lógica de enunciados para recoger la estructura del discurso:

(8) Juan bebe más de lo que come.

Ese enunciado puede sin duda esquematizarse en lógica de enunciados por la fórmula

$$(8a) \quad q \wedge t,$$

esquematizando ' q ' 'Juan bebe' y ' t ' un nuevo enunciado, paráfrasis de 'más de lo que come', y que podría ser 'Juan come menos que bebe'. Pero este análisis es muy insatisfactorio, pues ' t ' esquematiza en realidad un enunciado de dos verbos: no es aún esquematización de un enunciado atómico. Si tomáramos en vez de 'Juan come menos que bebe' la traducción

'come poco', ' $q \wedge t$ ' sería finalmente un esquema satisfactorio, pues ' t ' esquematizaría, como ' q ', un enunciado atómico; pero el resultado no sería una esquematización de (8). Lo que realmente afirma (8) es una relación entre dos clases de cantidades, las cantidades que bebe Juan y las que come. En lógica de enunciados no podemos esquematizar una tal relación, sino sólo relaciones entre enunciados, nombres de hechos.

También para esquematizar la *disyunción* hay que tener presente la naturaleza molecular de enunciados que pueden parecer atómicos. Así el enunciado

(9) Juan come o bebe

debe esquematizarse por la fórmula

(9a) $p \vee q$,

esquematizando ' p ', como en los anteriores ejemplos, 'Juan come', y ' q ' 'Juan bebe'.

Análogamente,

(10) Juan es arquitecto o ingeniero

puede esquematizarse

(10a) $p_1 \vee p_2$.

Los símbolos variables que se escojan son, naturalmente, arbitrarios. Pero cuando se esquematizan enunciados en contexto, hay que tener cuidado de no esquematizar dos enunciados distintos por la misma letra. Por eso es útil acostumbrarse a introducir una letra nueva cada vez que aparece un enunciado nuevo en el texto de lenguaje natural que se está esquematizando. Esto hemos venido haciendo hasta aquí, aunque, por tratarse de enunciados sueltos y no en contexto, la precaución era innecesaria. No la seguiremos tomando.

Los ejemplos (9) y (10) corresponden a la tabla de ' $\vee$ '. Pues el que Juan coma no excluye que beba, y el que Juan sea arquitecto no excluye que sea ingeniero. Se trata de disyunciones no excluyentes, que son verdaderas siempre que lo sea alguna de sus componentes, o cuando lo son los dos.

La situación no es la misma cuando la disyunción es excluyente. En el enunciado

(11) Juan es barbudo o imberbe

es imposible lo que permite la disyunción no excluyente, a saber que sean verdaderos los dos componentes, pues 'Juan es imberbe' es la negación de 'Juan es barbudo'. De modo que la esquematización de (11) debe ser

(11a) $q \vee \sim q$

Obsérvese que 'Calixto es barbudo o Melibea es imberbe' debe esquematizarse en cambio por ' $q \vee r$ '.

La fórmula (11a) se usa frecuentemente como caso particular (instancia) del principio de *tercio excluso* (cfr. 6) en la formulación: 'un enunciado o su negación, y sólo uno de los dos, tiene que ser verdadero'. Mas, puesto que la función $\vee$ permite que ambos argumentos sean verdaderos, es claro, puede pensarse, que ' $q \vee \sim q$ ' no es una buena formulación del principio de *tercio excluso*.

Esto no significa que la afirmación (11a) sea una violación de dicho principio. Pues (11a) no afirma que ' q ' y ' $\sim q$ ' tengan que ser ambos verdaderos: sólo lo permite. (11a) sigue siendo verdadera con que lo sea uno de los dos (lo cual será el caso siempre). Consiguientemente, no hay que violar el principio de *tercio excluso* para afirmar ' $q \vee \sim q$ '.

Si se quiere una formulación más plena del principio de *tercio excluso* puede escribirse

$[q \vee \sim q] \wedge \sim [q \wedge \sim q]$.

En general, para esquematizar una disyunción excluyente de miembros no precisamente contradictorios (como lo son ' q ' y ' $\sim q$ '), por ejemplo

(12) Juan es español o sueco,

escribiremos

(12a) $[p \vee q] \wedge \sim [p \wedge q]$,

o sea: 'Juan es español o sueco, y Juan no es español y sueco'.

La esquematización de *negaciones* en el cálculo de enunciados nos impone la regla de entender toda negación, desde el punto de vista de este análisis, como negación de enunciado entero (atómico o molecular), no de partes de enunciados atómicos. Así, por ejemplo, el enunciado

(13) Juan no come

se parafraseará primero como 'no: Juan come', y se esquematizará por

(13a) $\sim p$,

si ' p ' es la esquematización de 'Juan come'.

Con ' $\sim$ ' y ' $\wedge$ ' se esquematiza la conectiva "ni" del lenguaje natural. Así el enunciado

(14) Juan no come ni bebe

se parafraseará como 'Juan no come y Juan no bebe', y se esquematizará por

(14a) $\sim p \wedge \sim q$.

Este es otro ejemplo de eliminación de homonimias del lenguaje común. La homonimia aquí suprimida es la existente entre 'ni', 'y no', 'tampoco'

(pues la misma fórmula (14a) es naturalmente el esquema de 'Juan no come; tampoco bebe').

Por último, enunciados como

- (15) Aunque Juan come, no bebe,
- (16) Juan come, pero no bebe,
- (17) Juan come y sin embargo no bebe,
- (18) Juan come sin beber,

etc., tienen todos el mismo esquema:

- (15, 16, 17, 18) $p \wedge \sim q$.

Como hemos visto, el functor ' $\rightarrow$ ' esquematiza un condicional entre nombres de hechos, no una implicación entre nombres de enunciados. Así, por ejemplo, esquematizaremos por

- (19a) $p \rightarrow q$

el enunciado

- (19) Si Juan viene, entonces Luis se va,
pero no una implicación como

- (20) si a es mayor o igual que b , entonces, si b es mayor que c , a es mayor que c ,

aunque en el texto en lenguaje común aparezca escrita como mero condicional, y no como implicación. Pues lo que (20) quiere decir es

- (20a) el enunciado ' a es mayor o igual que b ' implica ("contiene") el enunciado 'si b es mayor que c , a es mayor que c '.

(20) es una implicación: la estructura del enunciado atómico ' a es mayor o igual que b ' contiene la del enunciado 'si b es mayor que c , a es mayor que c '. En cambio, la verdad de (19) no depende de la estructura interna de sus enunciados atómicos, sino de los valores de verdad que se les atribuya, o sea, de que ocurran o no ocurran los hechos expresados por ellos.

Considerada desde el punto de vista epistemológico o metodológico — es decir, por su contenido en conocimiento y el modo como ese conocimiento ha sido adquirido — una misma afirmación puede ser unas veces un condicional y otras veces una implicación. Así, por ejemplo, el enunciado

- (21) Si Sócrates es hombre, entonces es mortal

es un condicional si la afirmación se da como afirmación de hecho, de observación. En cambio, si la afirmación se diera como resultado de una investiga-

ción fisiológica de la estructura de Sócrates, investigación que hubiera descubierto que el morir está inserto, por así decirlo, en la estructura de ese objeto Sócrates, entonces (21) sería una implicación, expresable más explícitamente por

- (21a) 'ser Sócrates' implica 'ser mortal',

o

- (21b) ' x es Sócrates' implica ' x es mortal'.

Como aplicación de las anteriores reflexiones vamos a esquematizar la frase siguiente:

- (22) Como el que está suspendido de Análisis I o no se presentó a examen de Teoría II no se puede examinar de Teoría III, si uno se puede examinar de Teoría III está aprobado de Análisis I y se presentó a examen de Teoría II.

La frase tiene su conectiva principal en la coma que sigue a 'Teoría III', la cual cierra la parte que empieza con 'Como'. Toda esta primera parte es una explicación del resto de la frase: da condiciones por las cuales es verdad el hecho indicado en la segunda mitad de la frase. Se trata, pues, de un condicional. Podemos empezar a esquematizar (22) mediante la introducción del functor ' $\rightarrow$ ' entre las dos partes de la misma:

- (22a) [El que está suspendido de Análisis I o no se presentó a examen de Teoría II no se puede examinar de Teoría III] $\rightarrow$ [si uno se puede examinar de Teoría III, está aprobado de Análisis I y se presentó a examen de Teoría II].

El antecedente de (22a) es a su vez un condicional: da en efecto una condición de no poderse examinar de Teoría III. El antecedente de este segundo condicional es estar suspendido de Análisis I o no haberse presentado a examen de Teoría II. Podemos, pues, escribir, con ligeros retoques de estilo en lenguaje común:

- (22b) [[uno está suspendido de Análisis I v uno no se presentó a examen de Teoría II] $\rightarrow$ uno no se puede examinar de Teoría III]

El consecuente de (22a) y (22b) es también un condicional, y su consecuente es una conjunción:

- (22c) [[uno está suspendido de Análisis I v uno no se presentó a examen de Teoría II] $\rightarrow$ uno no se puede examinar de Teoría III] $\rightarrow$ [uno se puede examinar de Teoría III $\rightarrow$ [uno está aprobado de Análisis I $\wedge$ uno se presentó a examen de Teoría II]].

La frase está ahora formulada como composición de enunciados atómicos mediante conectivas (que entendemos funcionalmente). Pero observamos que, por no haber corregido el texto inicial en lenguaje común, tenemos un mismo hecho expresado de dos maneras: 'estar suspendido' tiene como negación 'no estar suspendido' y no 'estar aprobado', como aún tenemos en (22c). Normalizaremos pues antes de terminar la esquematización:

(22d) [[uno está suspendido de Análisis I $\vee$ uno no se presentó a examen de Teoría II] $\rightarrow$ uno no se puede examinar de Teoría III] $\rightarrow$ [uno se puede examinar de Teoría III $\rightarrow$ [uno no está suspendido de Análisis I $\wedge$ uno se presentó a examen de Teoría II]].

Tomando ahora sucesivamente ' p ', ' q ', ' r ' como esquemas de los enunciados atómicos *no negados* en el orden en que se presentan en (22d), podemos esquematizar (22) del modo siguiente, basándonos en (22d):

(22e) [[$p \vee \sim q$] $\rightarrow \sim r$] $\rightarrow$ [$r \rightarrow [\sim p \wedge q]$].

Aplicando nuestras convenciones para el ahorro de paréntesis, tenemos finalmente:

(22f) [$p \vee \sim q \rightarrow \sim r$] $\rightarrow$ [$r \rightarrow \sim p \wedge q$].

33. Sobre la notación de las funciones veritativas. — La notación de las funciones veritativas, de acuerdo con el uso corriente en matemáticas, exigiría escribir el functor a la izquierda de los argumentos ("variables independientes"), por ejemplo, en vez de ' $p \wedge q$ ',

$\wedge(p, q)$,

que podría leerse 'conjunción de p, q ', o 'función $\wedge$ de p, q '. Esto es lo que realmente entendemos al leer ' $p \wedge q$ ', puesto que no utilizamos las conectivas mismas, sino sus funciones veritativas asociadas. Pero la notación de esos functores como si fueran conectivas tiene la ventaja de recordarnos la realidad lingüística de la cual han sido abstraídas las funciones. Lo importante es tener siempre presente que los símbolos usados son functores, no partículas conectivas, y que representan funciones de valores veritativos, no conexiones lingüísticas. Éstas están representadas sólo en segundo lugar, a través de la abstracción funcional (cfr. 23).

APÉNDICE AL CAPÍTULO V

B. Observación. — a 30 ss.: la notación aquí usada de los functores veritativos es la de H. Scholz. Hay otras también corrientes, especialmente en la literatura anglosajona. La tabla siguiente muestra las correspondencias entre las principales notaciones:

Scholz	$\sim$	$\vee$	$\wedge$	$\rightarrow$	$\leftrightarrow$	usa paréntesis
Hilbert-Ackermann	$\bar{p}$	$\vee$	$\&$	$\rightarrow$	$\sim$	usa paréntesis
Russell	$\sim$	$\vee$	.	$\supset$	$\equiv$	usa puntos
Carnap	$\sim$	$\vee$	.	$\supset$	$\equiv$	usa par. y punt.
Lukasiewicz	Np	Apq	Kpq	Cpq	Epq	ni par. ni punt.

CAPÍTULO VI

LA ESTRUCTURA DE LOS ENUNCIADOS ATÓMICOS.
LÓGICA DE PREDICADOS

34. Concepto de la lógica de predicados. — Caracterizamos antes la lógica de enunciados por el género de análisis del discurso que presupone, y vimos que ese análisis deja sin tocar los enunciados atómicos, o sea, aquellos que no contienen ninguna conectiva. Por el mismo camino puede caracterizarse la lógica llamada 'de predicados', 'de términos' o 'cuantificacional': ésta se basa en un análisis que penetra también en los enunciados atómicos, distinguiendo sus diversos elementos gramaticales.

Casi toda la lógica aristotélica, medieval y clásica es una lógica de predicados. El propio Aristóteles introdujo una primera simbolización de la estructura de los enunciados atómicos más sencillos, usando una letra variable o esquemática para indicar el lugar del predicado (nominal, generalmente), otra para el sujeto, y el verbo 'se da en', sin simbolizar, para expresar la relación entre el predicado y el sujeto. Por ejemplo

A se da en B.

La notación medieval y clásica,

S es P,

invierte, por adecuación a la lengua latina, el orden de las variables, y sustituye el verbo 'darse en', que es seguramente más apto para usos lógicos, por el verbo 'ser', cargado de connotaciones filosóficas.

La universalidad del uso del verbo 'ser' esconde una diversidad de estructuras de enunciados atómicos del tipo más sencillo. Por ejemplo, cuando se dice

(1) Juan es alto

se está generalmente queriendo decir una de dos cosas: o bien que la altura se da en Juan, como decía Aristóteles, y, consiguientemente, el adjetivo 'alto' conviene al nombre 'Juan', en cuyo caso la notación hoy co-

rriente es 'Aa', con la constante 'A' por 'alto' y la constante 'a' por 'Juan', lo que da, con variables, el esquema

(1a) Px ;

o bien se está queriendo decir que Juan pertenece a la clase de las cosas altas, en cuyo caso la notación hoy corriente es ' $a \in A$ ' ('a es un miembro de la clase A'), y, en esquema

(1b) $x \in \alpha$.

Esos son ya dos sentidos de la partícula 'es'. Cuando se dice

(2) el hombre es mortal,

no se está diciendo con 'es' lo mismo que en el caso (1). Pues 'el hombre' no es ningún símbolo de individuo, y, por tanto, la relación de Hombre con Mortal no puede ser la pertenencia de individuo a clase, $\in$, ni la relación de un individuo con una propiedad suya. No es el abstracto Hombre el que tiene la propiedad de ser mortal, sino los individuos humanos. El 'es' de (2) significa inclusión de una subclase en una clase. Es corriente simbolizar esta acepción de 'es' por ' $\subset$ '. Así (2) se simboliza ' $H \subset M$ ', o, en esquema,

(2a) $\alpha \subset \beta$, o $\alpha \subseteq \beta$

con ' $\subseteq$ ' en vez de ' $\subset$ ' para admitir la posibilidad de que α sea una subclase impropia de β . En este último caso, que se da, por ejemplo, en el enunciado

(3) todos los asturianos son de la provincia de Oviedo,

la acepción de 'es' se simboliza por ' $=$ ', para denotar la igualdad de la clase de los asturianos con la clase de los naturales de la provincia de Oviedo:

(3a) $\alpha = \beta$,

cuando la relación se contempla extensionalmente; y por ' $\leftrightarrow$ '

(3b) $(x) [Px \leftrightarrow Qx]$,

para expresar la equivalencia entre los esquemas 'x es asturiano' y 'x es natural de la provincia de Oviedo', cuando la relación se contempla como relación entre propiedades.

Cuando la equivalencia en cuestión debe entenderse como definición, escribiremos, según dijimos,

(3c) ' $Px \leftrightarrow_{\text{df}} Qx$ '.

Por último, 'es' puede significar también la identidad entre símbolos de individuos, como, por ejemplo, en el enunciado

(4) Clarín es Leopoldo Alas

La identidad entre símbolos individuales se simboliza por '=',

(4a) $x = y$,

en la notación extensional a la que también pertenece ' $\subset$ '; y por el predicado diádico ' I ', cuando interesa conservar una notación predicativa:

(4b) Ixy .

Esas simbolizaciones permiten explicitar en cada contexto la acepción de 'es' o 'son' que interesa en el universo del discurso en que uno se está moviendo. La notación que interesa en el universo del discurso de la lógica de predicados es la de (1a) y (4b). Todas las demás pueden reducirse a ella, con los correspondientes cambios de sentido, del mismo modo que, por otra parte, ella puede expresarse en esas otras. Pero en cada traducción se cambia el punto de vista (el de las clases por el de los predicados, y viceversa).

Así, por ejemplo, para expresar (2) en la notación de (1a) se escribirá:

(2b) $(x) [Hx \rightarrow Mx]$,

notación que no conserva explícita la idea de clase, las clases Hombre y Mortal, sino las propiedades correspondientes, usadas como predicados de individuos.

La fórmula (4b) permite apreciar otra diferencia importante entre la notación hoy corriente para la lógica de predicados y la tradicional. Esta segunda diferencia consiste en lo siguiente: la lógica tradicional entendía por 'predicado' un término atribuible a (o negable de) un solo sujeto cada vez, en el sentido de que, por ejemplo, el enunciado

(5) Juan, Luis y Pedro son altos

debe entenderse como

(5a) Juan es alto, y Luis es alto, y Pedro es alto,

o, en esquema,

(5b) $Px \wedge Py \wedge Pz$.

(5b) es un análisis correcto de (5), también desde el punto de vista de la notación actual. Pero hay otras nociones que no pueden analizarse como predicados de un solo sujeto, y a las que conviene, sin embargo, considerar también propiedades, y predicados, por tanto, a sus nombres; pues en todo lo demás se comportan del mismo modo. Sea, por ejemplo, el enunciado

(6) Luis es amigo de Juan.

Es claro que el esquema

$Px \wedge Py$

no refleja la estructura de (6), sino que esquematiza más bien un enunciado tan ambiguo como

Luis es amigo y Juan es amigo,

sin que se sepa de quién. Introduciendo el predicado ' A ', 'ser-amigo-de-Juan', podría, desde luego, establecerse el esquema

(6a) Ax ,

que se leería ' x es-amigo-de-Juan'. Pero ese esquema no es del todo satisfactorio, porque esconde la presencia de un segundo nombre individual en (6). En realidad, la propiedad que importa en (6) afecta al par de individuos nombrados: es una relación. Esquematizaremos (6) por

(6b) Axy ,

o, en esquema, con sólo variables,

(6c) Rxy ,

donde ' R ' (o ' A ' en (6b)) esquematiza simplemente 'ser-amigo-de'. (6c) se lee ' x tiene la relación R con y ', cuando interesa destacar la noción de relación. En otro caso se lee simplemente "erre equis y ", del mismo modo que ' Px ' se lee abreviadamente "pe equis".

Con esto se admiten como predicados, o símbolos predicativos en general, no sólo los que van seguidos por un solo símbolo de individuo (éstos son los predicados o símbolos predicativos monádicos, que, cuando son constantes, significan atributos, propiedades de un sólo individuo a la vez), sino también los que van seguidos de dos o más símbolos de individuos (predicados diádicos, triádicos, ..., n -ádicos).

35. Símbolos elementales o primitivos de la lógica de predicados. — El lenguaje de la lógica de predicados requiere por de pronto variables y constantes individuales:

$x, y, z, \dots, a, b, c, \dots$,

y predicativas:

$P, Q, R, \dots, A, B, C, \dots$,

y, entre las predicativas, tanto monádicas cuanto n -ádicas. A veces puede ser conveniente indicar de un modo explícito, en un determinado contexto, cuál es el alcance de un símbolo predicativo, si es diádico, o triádico, etc.

Cuando se necesite, se indicará mediante un subíndice, p. e.:

$$P_3xyz.$$

Como en el caso de las letras de enunciado, también en éstas podemos usar subíndices al componer el símbolo. Y entonces, en los símbolos predicativos, podremos encontrarnos con dos subíndices: uno que indica el alcance del predicado, su número de argumentos (como también diremos), y otro que es aún parte del símbolo. Convendremos en que este último se escribirá debajo del otro. (La complicación carece aquí de importancia porque nunca nos encontraremos en ese caso.)

En realidad, el uso de subíndices es la mejor solución para consideraciones teóricas. Pues, por un lado, no hay tantas minúsculas latinas (por ejemplo) cuantos individuos deben suponerse en cualquier universo del discurso científico, por elemental que sea. Y, por otro lado, el uso de subíndices simplifica la notación, al permitir no usar más que una letra para cada categoría de variables. En la práctica necesitaremos pocas letras. Por eso podemos permitirnos el uso de varias y sin subíndices (salvo para indicar el alcance de predicados).

Las variables y letras de enunciado en general son también elementos de la lógica de predicados, la cual incluye a la de enunciados para formar un solo sistema. Ello por dos razones: primera, que una mera notación de la estructura interna de los enunciados atómicos no puede constituir por sí misma un lenguaje, pues en un lenguaje los enunciados atómicos se combinan o componen para formar enunciados moleculares; segunda, que a veces puede convenir combinar enunciados atómicos analizados con otros sin analizar. Por ejemplo, si es posible demostrar tanto

$$p \rightarrow Rxy$$

cuanto

$$\sim p \rightarrow Rxy,$$

entonces se podrá afirmar como demostrado el enunciado ' Rxy ', aun sin tener analizado ' p '.

De esta reflexión se desprende que todo el lenguaje de la lógica de enunciados debe ser parte del de la lógica de predicados, y, por tanto, que todos los símbolos elementales del primero lo son del segundo. En este sentido se dice que la lógica de predicados (con la limitación que se verá en seguida) constituye el sistema de la lógica elemental (lógica de enunciados más lógica de predicados, puede decirse).

Los paréntesis tienen un papel esencial en lógica de predicados, y no sólo el de cómodo expediente de puntuación. Pues, según vimos, los paréntesis se usan como notación del cuantificador universal o generalizador. Por tanto, será conveniente distinguir a partir de ahora el uso de los parén-

tesis como notación del cuantificador: para ello usaremos paréntesis curvos, por ejemplo:

$$(x),$$

mientras que para fines de puntuación seguiremos usando los corchetes. (En el castellano que usamos como metalenguaje intuitivo, disponemos de paréntesis de un tercer aspecto tipográfico.)

Por el uso de los cuantificadores se distinguen dos niveles u órdenes en la lógica de predicados. Si los cuantificadores no pueden afectar más que a variables individuales, en las formas

$$(x), \exists x,$$

el lenguaje pertenece a la lógica de predicados de primer orden. Si los cuantificadores pueden afectar también a variables predicativas, en las formas

$$(P), \exists P,$$

las expresiones con esas cuantificaciones pertenecen a la lógica de predicados de segundo orden. Hay diferencias importantes entre las dos, y en lo que sigue, hasta el capítulo XI, consideraremos sólo la lógica de predicados de primer orden, que es la única que integra lo que hemos llamado 'sistema de la lógica elemental'.

Las variables cuantificadas en una fórmula se llaman 'ligadas'; las demás se llaman 'libres'. Así en la fórmula

$$(x)\exists yPxyz,$$

'x' e 'y' están ligadas, 'P' y 'z' están libres.

En la lógica de predicados de primer orden las únicas variables ligadas son las individuales. Esto quiere decir que ellas son las únicas verdaderas variables de ese lenguaje. Pues un símbolo que no se puede ligar (cuantificar) está de algún modo ya determinado; o sea, cumple propiamente función de constante o de parámetro. A la inversa, no tiene sentido ligar (cuantificar) una constante o un parámetro. Sería, por ejemplo, un absurdo decir 'todo Miguel de Cervantes Saavedra es autor del *Quijote*'.

De esto se desprende que ni los símbolos predicativos ni los de enunciado son verdaderas variables en la lógica de predicados de primer orden; pues no se admite su cuantificación. Unas y otras letras son más bien parámetros.

Visualizaremos el anterior repaso en la siguiente

Tabla de los símbolos elementales de la lógica de predicados

1.º Los símbolos elementales de la lógica de enunciados:

1'. Letras de enunciado: ' p ', ' q ', ' r ', ' s ', ..., ' p_1 ', ..., ' s_n ', ...

2'. Functores veritativos: ' $\sim$ ', ' $\vee$ ', ' $\wedge$ ', ' $\rightarrow$ ', ' $\leftrightarrow$ '.

- 2.º Variables y constantes individuales: ' x ', ' y ', ' z ', ..., ' x_1 ', ..., ' z_n ', ..., ' a ', ' b ', ' c ', ..., ' a_1 ', ..., ' c_n ', ...
- 3.º Letras predicativas: ' P ', ' Q ', ' R ', ..., ' P_1 ', ..., ' R_n ', ..., ' A ', ' B ', ' C ', ..., ' A_1 ', ..., ' C_n ', ...
- 4.º Cuantificadores: ' $()$ ', ' $\exists$ '.
- 5.º Signos auxiliares de puntuación: ' $]$ ', ' $[$ '.

Se observará que no recogemos al descriptor entre los símbolos elementales o primitivos de la lógica de predicados, a pesar de que alguna vez lo usaremos. El motivo de no recogerlo ahora es que las reglas de formación se complican con su presencia, y aún más las reglas de transformación. En efecto, una descripción no es propiamente un nombre, una constante, sino, como decíamos en 26 una "especie de nombre". En el uso normal de un nombre, existe un individuo nombrado por él. Éste no es siempre el caso con las descripciones. Por ejemplo, la descripción

(1x) [x es el rey de España que reinó de 1931 a 1936]

no puede cumplirse para ningún x ; es, como suele decirse, una 'descripción vacía'. Análogo problema se presenta si hay más de un x que cumple la descripción (como podría ser la opinión conjunta de un grupo de carlistas e isabelinos sobre la anterior descripción).

Pero el prescindir del descriptor no es grave, porque una descripción puede en general eliminarse mediante otros símbolos — siempre que la descripción forme parte de un enunciado —. Por ejemplo, en el enunciado

el rey de España que reinó de 1931 a 1936 era alto

la descripción 'el rey de España que reinó de 1931 a 1936' puede eliminarse reconstruyendo el enunciado del modo siguiente, que ya no tiene ninguna descripción:

Hay un x tal que x era rey de España, x reinó de 1931 a 1936, y x era alto.

36. Fórmulas de la lógica de predicados de primer orden. — En la definición siguiente, las mayúsculas negritas, como ' P ', son variables sintácticas para designar símbolos predicativos, y las negritas minúsculas, como ' x ', son variables sintácticas para símbolos de individuo. ' X ' e ' Y ' son, como de costumbre, variables sintácticas para fórmulas.

Definición de Fórmula de la lógica de predicados de primer orden:

- (I) ' p ', ' q ', ' r ', ..., ' p_1 ', ..., ' s_n ' son fórmulas (atómicas).
- (II) $P x_1 x_2 \dots x_n$ es una fórmula (atómica).
- (III) Si X es una fórmula, $\sim X$ también lo es.
- (IV) Si X e Y son fórmulas, $X \vee Y$ también lo es.
- (V) Si X e Y son fórmulas, $X \wedge Y$ también lo es.

- (VI) Si X e Y son fórmulas, $X \rightarrow Y$ también lo es.
- (VII) Si X e Y son fórmulas, $X \leftrightarrow Y$ también lo es.
- (VIII) Si X es una fórmula, $(x)X$ también lo es.
- (IX) Si X es una fórmula, $\exists x X$ también lo es.
- (X) Ninguna expresión es una fórmula de la lógica de predicados de primer orden sino en virtud de (I)-(IX).

La anterior definición permite una decisión normada de la cuestión de si una dada expresión de la lógica de predicados de primer orden es o no una fórmula de la misma. Así, por ejemplo, la expresión

(7) $(x)[Px \rightarrow Qx] \wedge (x)[Rx \rightarrow Px] \rightarrow (x)[Rx \rightarrow Qx]$,

es una fórmula, como puede establecerse por aplicación sucesiva de las líneas (VI), (VIII), (VI), (II) (dos veces), (V), (VII), (VI), (II) (dos veces), (VIII), (VI), (II) (dos veces).

En cambio la expresión

(8) $(x)[y \wedge Px]$

no es una fórmula. He aquí su análisis: por la línea (VIII), (8) es una fórmula si lo es

(8a) $y \wedge Px$.

Por la línea (IV), (8a) es una fórmula si lo son ' y ' y ' Px '. ' Px ' es una fórmula por la línea (II) de la definición. Pero ' y ' no es fórmula en virtud de ninguna línea de la definición. Entonces, por la línea (X), ' y ' no es una fórmula. Luego tampoco lo es (8a), ni tampoco, por tanto, (8).

La anterior definición excluye de la clase de las fórmulas de la lógica de predicados de primer orden muchas que pueden serlo de la de segundo orden, como, por ejemplo,

(9) $(P)(x) \exists y [Px \rightarrow Py]$,

pues ninguna de las líneas (I)-(IX) permite una cuantificación como ' (P) '. Luego, por la línea (X), una expresión con esa construcción queda excluida de la clase de las fórmulas de la lógica de predicados de primer orden.

También quedan excluidas de esa clase expresiones como

(10) PQ ,

mientras que en otros sistemas lógicos hay expresiones muy emparentadas con ésta y que son fórmulas.

37. Esquematación de enunciados del lenguaje común por fórmulas de la lógica de predicados de primer orden. — Los enunciados moleculares más sencillos esquematizados en 32 en el lenguaje de la lógica de enunciados no presentan más problema, para su esquematización por fórmu-

las de la lógica de enunciados de primer orden, que el de precisar la estructura de sus componentes atómicos. Así, por ejemplo, el enunciado

(11) Juan come y bebe,

que en 32 figuraba como (5) y tenía la esquematización

$$p \wedge q,$$

se esquematizará por

(11a) $Px \wedge Qx$.

El ejemplo (8) de 32, que numeraremos aquí

(12) Juan bebe más de lo que come,

quedaba, según vimos, mal esquematizado con los solos medios de la lógica de enunciados. El lenguaje adecuado para esquematizarlo es el de la lógica de predicados de primer orden. Para ello conviene analizarlo, con objeto de entender bien las relaciones que hay que esquematizar.

El enunciado (12) contiene el nombre de una entidad individual — 'Juan' —, pero alude además a otras entidades individuales, a saber, conjuntos de sólidos comidos por Juan y conjuntos de líquidos bebidos por Juan. De esas entidades individuales, una, Juan, está determinada unívocamente. Dicho de otro modo: 'Juan' es una constante. En cambio, las otras entidades individuales o concretas están aludidas genéricamente: el enunciado nos dice que de un par de cantidades *cualesquiera*, una de las cuales sea la del alimento sólido comido por Juan y otra la del líquido bebido por Juan, la segunda es mayor que la primera. El enunciado (12) tiene pues, aunque no se vean, lugares de argumento para cantidades de alimento sólido y para cantidades de líquido, unas y otras consumidas por Juan. Según esto, el enunciado (12) debe esquematizarse con tres símbolos individuales: una constante (para 'Juan') y dos variables individuales. Pongamos '*a*' para Juan, '*x*' e '*y*', respectivamente, para cantidades de líquido bebido por Juan y cantidades de sólido comido por Juan.

En cuanto a los símbolos predicativos, harán falta tres: uno para 'beber', otro para 'comer' y otro para 'mayor que'. Tomemos, respectivamente, '*B*', '*C*' y '>'.
 Por último, la afirmación de que la cantidad de líquido bebido por Juan es mayor que la de sólido comido por Juan no está limitada en (12) a ningún orden de magnitud preciso: vale para todo par, $\{x, y\}$, que se considere. Esto quiere decir que tendremos que cuantificar las variables del esquema con el generalizador.

Si interesa mucho una aplicación de la lógica a esta "teoría de la alimentación de Juan", aún será necesario explicitar otra variable individual más, representativa de lapsos de tiempo. 'Juan bebe más de lo que come'

se referirá en efecto, en esa "teoría de la alimentación de Juan", a un período de la vida de éste, o a toda su vida. Para dar la indicación de tiempo, que no importa en lógica pura (ni en sintaxis ni en semántica, pero sí en pragmática), habrá que considerar cuáles son los lapsos de tiempo que interesa tomar como unidades — como individuos — en el contexto particular. En una "teoría de la alimentación de Juan" sería razonable tomar como unidad de tiempo, como individuo temporal, el día, y no el segundo o el minuto. El universo del discurso de la "teoría de la alimentación de Juan" tendrá pues, suponemos, una serie de individuos que son días. Los esquematizaremos con la variable '*d*'.

El resultado del anterior análisis es la siguiente traducción de (12):

(12a) Para todo día, *d*, y para todas las cantidades, *x*, *y*, si *x* es la cantidad de líquido bebido por Juan el día *d* e *y* es la cantidad de sólido comido por Juan el día *d*, entonces *x* es mayor que *y*.

Lo cual puede esquematizarse, con los símbolos que hemos escogido, por

(12b) $(d)(x)(y)[Baxd \wedge Cayd \rightarrow x > y]$.

Esa fórmula supone un determinado universo del discurso — el de la "teoría de la alimentación de Juan" —, pues no precisa el campo del que pueden tomarse valores para las variables '*d*', '*x*', '*y*', es decir, no precisa las clases de objetos cuyos nombres pueden ocupar en la fórmula los lugares de dichas variables. Para analizar completamente (12) desde un punto de vista de lógica pura, sin presuponer el universo del discurso de la "teoría de la alimentación de Juan", habría que precisar esos campos de las variables. Esto puede hacerse introduciendo los predicados 'ser-un-día' (que simbolizaremos por '*D*'), 'ser-una-cantidad-de-líquido' (que simbolizaremos por '*L*') y 'ser-una-cantidad-de-alimento-sólido' (que simbolizaremos por '*S*'). De este ulterior análisis resulta el siguiente esquema:

(12c) $(d)(x)(y)[Dd \wedge Lx \wedge Sy \wedge Baxd \wedge Cayd \rightarrow x > y]$.

Al esquematizar enunciados que parecen atómicos y no lo son, la posibilidad de error no se da fácilmente, como en cambio ocurre en lógica de enunciados, con expresiones del tipo (11), es decir, con enunciados de un sujeto y varios verbos; sino más bien con enunciados de varios sujetos y un solo verbo, por ejemplo

(13) Juan y Luis juegan,

que puede esquematizarse por

(13a) $Px \wedge Py$,

o por

$$Ja \wedge Jb,$$

si en la esquematización no se quiere perder la alusión a nociones y sujetos determinados.

Pero no es seguro que esa esquematización sea fiel. La corrección de la esquematización depende de lo que se haya querido decir con el enunciado (13), ambiguo como suele o puede serlo todo enunciado del lenguaje común una vez desligado de su contexto. Pues si lo que se quería decir era más bien

(14) Juan y Luis juegan el uno contra el otro,

entonces la esquematización más fiel habría sido

(14a) Rxy ,

pues tal vez se quería significar la relación entre Juan y Luis disputando una partida, y no la mera coincidencia del hecho de que Luis juega y el hecho de que Juan juega. La recta comprensión del texto en lenguaje común es, naturalmente, presupuesto necesario de toda buena esquematización.

En el capítulo I vimos que un enunciado como

(15) todos los árboles son vegetales,

debe entenderse, para explicitar las variables individuales, como

(15a) todo lo que tiene la propiedad árbol tiene la propiedad vegetal.

Por nuestro conocimiento del oficio pronominal de la variable, damos en seguida con la esquematización siguiente (con 'P' por 'ser-árbol' y 'Q' por 'ser-vegetal'):

(15b) $(x)[Px \rightarrow Qx]$.

En cambio, un enunciado como

(16) algunos árboles son altos,

entendido, con pronombres (variables),

(16a) algo tiene la propiedad árbol y la propiedad alto

se esquematiza (con 'P' por 'ser-árbol' y 'Q' por 'ser-alto')

(16b) $\exists x[Px \wedge Qx]$.

La presencia de ' $\wedge$ ' en (16b), en vez de ' $\rightarrow$ ' en (15b), ilustra sobre la afirmación de existencia propia del particularizador.

Cuando el lenguaje de la lógica de predicados se utiliza para el análisis y la formalización de teorías científicas o de textos en lenguajes étnicos,

es siempre conveniente tomar símbolos constantes que representen nociones de la teoría o del texto que se está formalizando: constantes, esto es, que no son constantes lógicas, sino empíricas. También es útil no tomar las variables sobre el universo del discurso, incualificado, de la lógica pura, sino sobre campos específicos de la teoría que interese. Ya hemos usado tales variables sobre universos del discurso de teorías positivas en la discusión de (12), y hemos visto cómo se pueden eliminar para atenerse al universo del discurso de la lógica pura. Pero, además de ser conveniente, la introducción de constantes no lógicas es necesaria en el paso inicial de la formalización, o sea, al establecer la lista de símbolos primitivos (no definidos) del lenguaje a formalizar, símbolos que corresponderán a los conceptos primitivos (indefinidos) de la teoría estudiada. Para establecer la lista de esas constantes es necesario un análisis previo de la teoría en cuestión, a menos que ésta se encuentre ya relativamente formalizada, acaso sin una simbolización muy normada, por los propios científicos positivos que la han construido. En este caso, dichos científicos habrán explicitado cuáles son las nociones que quieren tomar como primitivas, y lo único que habrá que hacer en este primer paso será asegurarse de que la simbolización es coherente y cómoda.

Consideremos, por ejemplo, los axiomas de la teoría de los números naturales, o aritmética, propuestos por G. Peano (1858-1932), o sólo los cinco de sus nueve axiomas que son propiamente aritméticos. (Los otros cuatro son lógicos. También en la presentación antigua de los *Elementos* de Euclides se encuentran unas "nociones comunes", que son principios lógicos, y unos "postulados", o axiomas propiamente geométricos, que hablan de rectas, puntos, etc.).

Esos cinco axiomas (en los que 'número' quiere decir 'número natural') son:

(1.º) 1 es un número.

(2.º) El siguiente de cualquier número es un número.

(3.º) Dos números no tienen el mismo siguiente.

(4.º) 1 no es el siguiente de ningún número.

(5.º) Toda propiedad que pertenezca a 1 y al siguiente de todo número que la posea, pertenece a todo número.

Como nociones primitivas de ese conjunto axiomático podemos considerar las siguientes: 1, número (la propiedad de ser un número), siguiente (la propiedad de ser siguiente de, que es una relación — comúnmente expresada con otros términos técnicos en matemáticas) y la propiedad de ser idéntico a (la relación de identidad). Todas esas nociones tendrán que recogerse mediante constantes no-lógicas al esquematizar los enunciados (1.º)-(5.º). Tomemos, por ejemplo, los símbolos ' a ', ' N ', ' S ', ' T ' y atribuyámoslos a aquellas nociones por el mismo orden en que las hemos enume-

rado. Con esas notaciones el axioma (1.º) puede esquematizarse por

$$(1.ª) \quad Na.$$

El esquema del axioma (2.º) es

$$(2.ª) \quad (x)(y)[Nx \wedge Syx \rightarrow Ny].$$

El axioma (3.º) se refiere, naturalmente, a dos números cualesquiera, a todo par de números y a todo siguiente de ellos. Su esquema es:

$$(3.ª) \quad (x)(y)(z)(u)[Nx \wedge Ny \wedge \sim Ixy \wedge Sxz \wedge Suy \rightarrow \sim Izu].$$

El axioma (4.º) puede esquematizarse por

$$(4.ª) \quad \sim \exists x[Nx \wedge Sax].$$

Para empezar la esquematización del axioma (5.º) omitiremos por el momento la cuantificación inicial, y consideraremos la estructura de lo afirmado. Se trata de un condicional: que una propiedad pertenezca a 1 y que, si pertenece a un número, pertenezca a su siguiente, es condición de que esa propiedad pertenezca a todo número. El consecuente de (5.º) es, pues, llamando 'P' a la propiedad en cuestión:

$$(x)[Nx \rightarrow Px].$$

El antecedente de (5.º) es una conjunción de dos enunciados. El primero pone la condición de que 1 tenga la propiedad 'P':

$$Pa$$

El segundo pone la condición que hemos parafraseado ya y que puede esquematizarse así:

$$(x)(y)[Nx \wedge Syx \wedge Px \rightarrow Py].$$

El conjunto del antecedente de (5.º) es pues:

$$Pa \wedge (x)(y)[Nx \wedge Syx \wedge Px \rightarrow Py].$$

Y la totalidad del axioma, menos la cuantificación inicial, se podrá esquematizar por

$$(5.ª) \quad Pa \wedge (x)(y)[Nx \wedge Syx \wedge Px \rightarrow Py] \rightarrow (x)[Nx \rightarrow Px].$$

Veamos ahora la cuantificación inicial aún no recogida. Ésta es una cuantificación de 'P', pues (5.º) dice que (5.ª) vale de toda propiedad P. Consiguientemente, el esquema completo de (5.º), en la versión, no muy afinada técnicamente, que hemos elegido, es

$$(5.ºb) \quad (P)[Pa \wedge (x)(y)[Nx \wedge Syx \wedge Px \rightarrow Py] \rightarrow (x)[Nx \rightarrow Px]].$$

Por la cuantificación de 'P', (5.ºb) no es una fórmula de la lógica de predicados de primer orden, sino de la de segundo orden. Lo mismo ocurre entonces con el axioma (5.º). Un enunciado como (5.º) no puede, pues, esquematizarse suficientemente en el lenguaje de la lógica de predicados de primer orden. No pertenece a ella.

En relación con lo que antes se vio sobre la suposición de determinados universos del discurso en la aplicación de la lógica a la esquematización y formalización de teorías positivas, debe observarse ahora lo siguiente: en la esquematización de los axiomas de Peano hemos precisado siempre explícitamente que los objetos o individuos considerados son números. Esto quiere decir que aquellos esquemas se basan en el universo del discurso de la lógica pura, en el cual los individuos son meras entidades sin cualificar, y hay por tanto que expresar explícitamente su cualificación relevante (aquí la de número natural, 'N'). Dando, en cambio, por supuesto el universo del discurso de la aritmética, no es necesario usar la constante predicativa 'N' más que cuando esté explícitamente usada ya en el texto en lenguaje común. Con esta simplificación notacional — o sea, presupuesto el universo del discurso de la aritmética — los axiomas (1.º)-(5.º) podrían esquematizarse así:

$$(1.ºb) \quad Na.$$

$$(2.ºb) \quad (x)(y)[Syx \rightarrow Ny].$$

$$(3.ºb) \quad (x)(y)(z)(u)[\sim Ixy \wedge Sxz \wedge \rightarrow \sim Izu].$$

$$(4.ºb) \quad \sim \exists x Sax.$$

$$(5.ºc) \quad (P)[Pa \wedge (x)(y)[Syx \wedge Px \rightarrow Py] \rightarrow (x)Px].$$

38. La implicación en la lógica de predicados de primer orden. — Con la notación de la lógica de predicados de primer orden pueden escribirse simples condicionales de la misma naturaleza lógica, igualmente diversos de la implicación, que los usados en lógica de enunciados. Por ejemplo, la fórmula

$$(17) \quad Pa \rightarrow Qa$$

es un mero condicional, pues no sabemos nada de la razón por la cual existe una relación entre que *a* sea *P* y que *a* sea *Q*. Esta relación se nos presenta como mero hecho empírico.

Mas supongamos que dicha relación vale para todo individuo, es decir que

$$(18) \quad (x)[Px \rightarrow Qx].$$

Entonces sabemos además, por lo menos, que la relación entre el ser *P* y el ser *Q* es universal. Se recordará que en la discusión de 30 sobre el condicional y la implicación hallamos que implicación es condicional generalizado, universalmente verdadero. Por eso se ve comúnmente en fórmulas

como (18) un género de implicación, a la que suele llamarse 'implicación formal'.

La interpretación en que se basa ese nombre responde a la reflexión siguiente: puesto que para todo individuo el ser P hace que sea Q , debe haber una relación estructural entre la propiedad P y la propiedad Q , y esa relación estructural se expresará mediante una implicación entre el predicado ' P ' y el predicado ' Q '. Una tal relación estructural mediaría, por ejemplo, entre

$$'Px' \leftrightarrow_{\text{df}} 'Rx \wedge Sx'$$

y

$$'Qx' \leftrightarrow_{\text{df}} 'Rx',$$

en cuyo caso (18) sería la implicación

$$(19) \quad (x)[Rx \wedge Sx \rightarrow Rx].$$

39. Sobre la constante ' I '. — Hemos utilizado la constante ' I ' para simbolizar la relación de identidad. Para el mismo fin habríamos podido usar el símbolo ' $=$ '. Pero uno y otro símbolos suponen una especificación del universo del discurso puramente lógico. Pues en éste no hay dos individuos idénticos. Ningún individuo es, como tal individuo (y así se los considera en el universo del discurso de la lógica), idéntico con ningún otro. En realidad, sólo puede hablarse de identidad entre individuos desde el punto de vista de alguna teoría, es decir, dentro de un determinado universo del discurso. Leibniz estableció un criterio para dar un sentido a la idea de identidad entre individuos. El expediente de Leibniz es éste: dos individuos son idénticos, en un universo del discurso dado, cuando los dos tienen por igual todas las propiedades que se consideran en ese universo del discurso ("principio de identidad de los indiscernibles").

Por ese carácter ya "aplicado" de la constante ' I ', ésta no suele incluirse en el lenguaje de la lógica de predicados pura, sino en una ampliación del mismo a la que es corriente llamar 'lógica de predicados (de primer orden) con identidad'.

APÉNDICE AL CAPÍTULO VI

B. Observaciones. — a 34: no es del todo verdad que Aristóteles introdujera símbolos para variables predicativas y símbolos para variables individuales. Las entidades individuales no están representadas en la silogística aristotélica, excepto en ciertas demostraciones poco frecuentes (ékthesis). Aristóteles simboliza sólo predicados, como 'árbol', 'vegetal', etc.

a 37: la idea de noción primitiva es, naturalmente, relativa. La noción de número, por ejemplo, que Peano toma como primitiva, aparece como definida en el sistema de Russell (cfr. caps. XIV y XV).

Sección segunda. — CALCULOS LÓGICOS ELEMENTALES

CAPÍTULO VII

PRESENTACIÓN AXIOMÁTICA DEL CALCULO DE PREDICADOS DE PRIMER ORDEN

40. El sistema axiomático. — El sistema axiomático es, desde los tiempos de la geometría griega, la forma típica de presentarse el cálculo o lenguaje formalizado. Lo característico del sistema axiomático como realización de la idea de cálculo consiste en disponer de un conjunto de enunciados o fórmulas que se admiten sin demostración y a partir de los cuales se obtienen todas las demás afirmaciones de la teoría, las cuales se llaman 'teoremas'. Las fórmulas aceptadas sin demostración se llaman 'axiomas' o 'postulados'. Aquí usaremos la primera palabra.

El conjunto de los axiomas, más la definición de enunciado o fórmula del sistema (definición que precede al enunciado de los axiomas) y el conjunto de las reglas para la obtención de teoremas a partir de los axiomas (reglas de transformación) constituyen la *base primitiva* del sistema.

El nombre de 'reglas de transformación' para las aludidas está justificado porque las operaciones mediante las cuales se obtienen teoremas a partir de los axiomas consisten en transformaciones de éstos, como sustituciones de unas variables por otras, composición de axiomas para formar otras fórmulas, etc.

Suele distinguirse entre sistemas axiomáticos formalizados y no-formalizados. La diferencia principal entre unos y otros consiste en que los formalizados, que son los que realizan claramente la idea de cálculo, presentan explícitamente todas las reglas de transformación, mientras que los otros no lo hacen. En un sistema axiomático formalizado — y sólo a tales sistemas nos referiremos — el conjunto de los axiomas y el de las reglas de transformación son ambos efectivos, aunque en un sentido un tanto laxo de esta palabra.

Como es natural, las reglas de transformación (igual que las de forma-

ción de fórmulas o enunciados) son metalingüísticas respecto de las fórmulas del sistema, puesto que son afirmaciones acerca de lo que puede hacerse con fórmulas del sistema. También un enunciado que diga que tal o cual fórmula es un axioma será metalingüístico respecto del lenguaje al que pertenezca dicho axioma.

Por último, del mismo modo que los teoremas se obtienen de los axiomas, así también las nociones (los predicados) no-primitivos, no contenidos en los axiomas, se obtienen en el sistema a partir de las nociones primitivas, contenidas en los axiomas. (Lo mismo vale de constantes individuales, cuando éstas son necesarias.) El modo de hacerlo se especifica mediante reglas de definición.

41. Nociones de demostración o derivación, y de teorema. — Como el sistema axiomático se constituye, según el ideal del lenguaje “bien hecho”, para establecer con precisión la fundamentación de los teoremas de una teoría en sus axiomas, la noción de derivación o demostración, que es la noción del modo formal de fundamentar, tiene una gran importancia en el estudio de los sistemas axiomáticos.

Una *demostración* en un sistema axiomático es una sucesión de fórmulas, cada una de las cuales está fundamentada, lo que quiere decir que cada una de ellas es: o bien un axioma, o bien una fórmula obtenida inmediatamente de un axioma por la aplicación de una regla de transformación, o bien una fórmula obtenida de otra u otras de los dos géneros anteriores mediante una aplicación de las reglas de transformación.

Un *teorema* es, en sentido estricto, una fórmula de cualquiera de las dos clases últimamente citadas. Y, en sentido amplio, es cualquier fórmula fundamentada del sistema. En este sentido amplio, también los axiomas son teoremas. (La decisión por la cual se ponen tales o cuales axiomas se considera fundamentación de éstos. Es una manera de decir que los axiomas no tienen fundamento en el sistema.)

La anterior noción de demostración es efectiva. Dada una sucesión cualquiera de fórmulas de un cálculo, es posible indicar de cada una de ellas si es un axioma o no lo es (comparándola con la lista de los axiomas); y, caso de no serlo, es posible decidir de la fórmula en cuestión si ha sido o no obtenida a partir de alguna o algunas de las anteriores mediante la aplicación de alguna o algunas de las reglas de transformación, de cuya lista también se dispone.

En cambio, la noción de teorema no es, en general, efectiva: dada una fórmula de un cálculo, no siempre es posible averiguar si se fundamenta o no en los axiomas mediante las reglas de transformación. Es posible que no demos con la demostración que la fundamenta, aunque en realidad tenga demostración. Y es posible que, no teniendo la fórmula demostración, no resulte tampoco demostrable su negación. Exigir que la clase de los teoremas de una teoría sea efectiva, es lo mismo que exigir que la teoría

sea decidible (cfr. 18), pues equivale a pedir que exista un procedimiento efectivamente utilizable (o sea, que no requiera infinitos pasos) para establecer si una fórmula cualquiera es válida o no en el sistema. Mas no todos los cálculos son decidibles.

Las cuestiones de la consistencia, la completud y la decidibilidad de un cálculo se plantean primariamente a propósito de los sistemas axiomáticos. Así, un sistema axiomático es consistente cuando con sus reglas de transformación es imposible demostrar a partir de sus axiomas un teorema y su negación. Un sistema axiomático es completo cuando sus axiomas y sus reglas de transformación bastan para demostrar como teoremas todas las verdades formales referentes a sus nociones primitivas (y, consiguientemente, también todas las verdades relativas a las nociones constituidas por medio de las reglas de definición).

La presentación del cálculo como sistema axiomático facilita la introducción de otra noción de completud, además (no en vez) de la ya conocida: un sistema es completo en este segundo sentido cuando al añadir a sus axiomas cualquier fórmula de su mismo lenguaje, pero que no sea demostrable a partir de ellos, el nuevo conjunto de fórmulas así obtenido permite demostrar una fórmula y su negación. Esta noción de completud es en ciertas circunstancias más fuerte que la otra.

Por último, un sistema axiomático es decidible cuando para toda fórmula de su lenguaje puede averiguarse, mediante un procedimiento normado y en un número finito de pasos, si la fórmula es un teorema del sistema o no lo es.

La presentación del cálculo como sistema axiomático plantea otra cuestión emparentada con las de consistencia, completud y decidibilidad: es la cuestión de la *independencia* de los axiomas. Un axioma de un conjunto axiomático es independiente de los demás cuando no es demostrable como teorema a partir de ellos.

La independencia no es una propiedad tan importante como la consistencia, y parece responder más bien a un ideal de economía lógica y elegancia. Pero investigaciones sobre la independencia de axiomas han sido alguna vez muy fecundas en la historia de la ciencia, y son siempre de interés para averiguar la estructura de las teorías, sus supuestos realmente necesarios.

42. Axiomas y esquemas axiomáticos. — A veces, especialmente en lenguajes formalizados y simbolizados, puede ser incómodo enunciar un conjunto de axiomas usando las letras del cálculo, como sería, por otra parte, natural. La incomodidad se debe a lo siguiente: una fórmula que se quiere usar como axioma, por ejemplo, el “principio de identidad”,

$$p \rightarrow p,$$

puede ser frecuentemente necesaria, para las demostraciones, con otras letras, por ejemplo,

$$r \rightarrow r.$$

El método axiomático no permite usar ' $p \rightarrow p$ ' como si fuera ' $r \rightarrow r$ '. Por tanto, si el axioma está escrito con la letra ' p ' y hay que usarlo con otras, es necesario adoptar una regla de transformación que permita sustituir una letra de enunciado por cualquier otra. A veces no interesa una regla así, y entonces se recurre a formular los axiomas en el metalenguaje, con variables sintácticas, por ejemplo:

$$X \rightarrow X.$$

Pero ' $X \rightarrow X$ ' no es propiamente un axioma del cálculo si, como en el ejemplo, el lenguaje del cálculo en cuestión es el de la lógica de enunciados. Se dice entonces que ' $X \rightarrow X$ ' es un esquema axiomático, y la realidad es que en ese sistema hay, en vez de un axioma ' $p \rightarrow p$ ', una infinidad de axiomas: uno por cada letra de enunciado. La operación de escribir ' $r \rightarrow r$ ', no es en efecto fruto de una sustitución de ' X ' por ' r ' — pues ' X ' no pertenece al lenguaje del sistema, sino al metalenguaje —, sino lo que llamaremos una 'instanciación de X en el sistema'. Como las reglas (de formación y de transformación) son siempre metalingüísticas, todas sus aplicaciones son instanciaciones.

Por este motivo, al afirmar que los conjuntos de los axiomas y de las reglas tienen que ser efectivos, añadimos que a veces lo son en un sentido laxo de esa palabra: no siempre en el sentido de que sea posible ponerlos en una lista finita, sino sólo, en general, en el sentido de que tiene que ser posible: 1.º enumerar unas expresiones (axiomas y reglas) que a veces serán metalingüísticas; 2.º averiguar de toda formación u operación dada en el lenguaje del cálculo si es o no es una instanciación de aquéllas.

43. Modelos isomorfos. Monomorfismo y polimorfismo. — Según una terminología que ya hemos introducido (cfr. 19), aunque no es de uso universal, un modelo de una fórmula (o de un conjunto de fórmulas) es una interpretación que la (o lo) satisface o cumple, es decir que la (o lo) hace verdadera (o verdadero). La noción de modelo es aplicable a cualquier conjunto de axiomas, puesto que éstos son fórmulas.

Recordaremos también que una interpretación de una fórmula o conjunto de fórmulas (cfr. 19) es un sistema unívoco de atribuciones de entidades ajenas al cálculo de que se trate a las variables de una fórmula, de un conjunto de fórmulas o del cálculo en general. No todos los símbolos de una fórmula son interpretables, sino sólo aquellos que se toman como variables. Así, por ejemplo, en una fórmula como

$$Px \rightarrow Qa$$

no consideramos interpretable más que ' x '; ' $\rightarrow$ ' es una constante lógica, y las constantes lógicas no se interpretan por referencia a significaciones ajenas al cálculo, sino que se definen en el cálculo mismo, aunque, en ge-

neral, lo que en el cálculo se hace es precisar el uso de esas constantes, mediante la definición de 'fórmula' y mediante las reglas de transformación. (Cfr. 44, más abajo, sobre definición implícita). — ' a ' es una constante no-lógica, empírica: no necesita interpretación porque siempre está, por así decirlo, interpretada, es el nombre de una significación. — En cuanto a ' P ' y ' Q ', hemos visto que, en lógica de predicados de primer orden, son propiamente parámetros (cfr. 35). Esos parámetros están fijados en cada caso una vez determinado el campo de individuos elegido para la interpretación de las únicas variables auténticas, las individuales. Por ejemplo, elegido el campo de los individuos humanos, las letras predicativas de una fórmula representan en general nombres de propiedades que pueden tener individuos humanos. La interpretación de las letras predicativas no es pues necesaria. Pero, por el punto de vista semántico intuitivo que hemos adoptado, hablaremos también de interpretación de las letras predicativas.

Esta no es la única manera de definir lo que son interpretaciones, pero sí la que nos será más útil en los siguientes ejemplos.

Cada interpretación de una fórmula asocia a ella un universo del discurso, un campo o una clase de individuos, a la que nos referiremos mediante la mayúscula griega Ω , con subíndices, y una clase de propiedades a la que nos referiremos mediante la mayúscula griega Π , también con subíndices. Por ejemplo: una interpretación, $\mathfrak{I}_1$, de la fórmula ' $Px \rightarrow Qx$ ' puede consistir en interpretar ' x ' por 'mi dedo índice derecho', ' P ' por 'ser parte de mi mano derecha', ' Q ' por 'ser parte de mi brazo derecho'. La interpretación $\mathfrak{I}_1$ está definida sobre Ω_1 y un Π_1 :

Ω_1 = mis dedos derechos;

Π_1 = las propiedades ser parte de mi mano derecha y ser parte de mi brazo derecho.

La interpretación $\mathfrak{I}_1$ es modelo de ' $Px \rightarrow Qx$ '.

Puede ser conveniente indicar para una interpretación el Ω y el Π sobre los cuales está definida. Lo haremos con notaciones (metalingüísticas) como la siguiente:

$\mathfrak{I}_1 (\Omega_1, \Pi_1)$.

Tras esto fijaremos el significado de la expresión 'modelos isomorfos'. Dos modelos, $\mathfrak{I}_1$ e $\mathfrak{I}_2$, de una fórmula o conjunto de fórmulas, X , son isomorfos cuando entre sus respectivos campos de individuos, Ω_1 y Ω_2 , puede establecerse una correlación biunívoca tal que, para cualquier objeto, a_1 , de Ω_1 , que interprete según $\mathfrak{I}_1$ la variable ' x ' de la(s) fórmula(s) X , ocurre que el individuo b_1 de Ω_2 , que es el correlatado con a_1 , interpreta, según $\mathfrak{I}_2$, dicha variable ' x ' de X , y a la inversa. Por ejemplo: la interpre-

tación $\mathfrak{S}_2 (\Omega_2, \Pi_2)$, que interprete 'x' por 'mi dedo índice izquierdo', 'P' por 'ser parte de mi mano izquierda' y 'Q' por 'ser parte de mi brazo izquierdo', con

$\Omega_2 =$ mis dedos izquierdos;

$\Pi_2 =$ las propiedades ser parte de mi mano izquierda y ser parte de mi mano derecha,

y que es también modelo de ' $Px \rightarrow Qx$ ', es además isomorfa con $\mathfrak{S}_1 (\Omega_1, \Pi_1)$. La isomorfía es aquí muy fácil de comprobar: a cada dedo de mi mano izquierda corresponde biunívocamente un dedo de mi mano derecha. Y sus nombres se comportan como antes se dijo en la interpretación de ' $Px \rightarrow Qx$ '.

El uso del posesivo 'mi' en el anterior ejemplo tiene como finalidad conseguir que las entidades que interpretan a 'x' sean "de verdad" individuos, cosa que no serían si se tratara de objetos como dedo índice izquierdo en general.

La noción de isomorfía será estudiada más en general en el capítulo XV.

Se dice que un conjunto de axiomas es monomórfico cuando todos sus modelos son isomorfos. En otro caso se dice que el conjunto de axiomas es polimórfico. (Esta terminología procede de R. Carnap. Pero la noción de interpretación está aquí dada muy poco detalladamente.)

Se observará que $\{\Omega_1, \Pi_1\}$ puede fundar no sólo la interpretación $\mathfrak{S}_1$, sino también otras, $\mathfrak{S}_1^*$, $\mathfrak{S}_1^{**}$, etc. $\mathfrak{S}_1^*$, por ejemplo, podría interpretar 'x' por mi dedo anular derecho'. $\mathfrak{S}_1$ es pues miembro de una clase o sistema de interpretaciones. Lo mismo vale de $\mathfrak{S}_2$.

44. La idea de definición implícita.—Desde principios del siglo XIX se han entendido frecuentemente los sistemas axiomáticos monomórficos como una especie de definiciones disfrazadas de las nociones primitivas que contienen. Estas nociones primitivas, que no son definidas explícitamente (puesto que se toman como iniciales), estarían definidas de un modo implícito por los axiomas. Según esto, los cinco axiomas aritméticos de Peano, por ejemplo, serían una definición implícita de la noción de número natural, del conjunto de los números naturales (en el supuesto de que dicho conjunto de axiomas fuera monomórfico).

Esta idea tiene la siguiente justificación: como todos los modelos de un sistema monomórfico son isomorfos, las relaciones entre individuos del campo de un modelo — digamos de los individuos de Ω_1 — son las mismas que las relaciones entre los individuos del campo de cualquier otro modelo, Ω_2 . Más precisamente: las relaciones entre determinados individuos de Ω_1 , $a_1, \dots, a_n$, son formalmente las mismas que median entre los n individuos de Ω_2 , $b_1, \dots, b_n$, correspondientes a $a_1, \dots, a_n$ de Ω_1 . Esto queda claro porque los nombres de unos y otros objetos, tomados en el mismo

orden, satisfacen las mismas fórmulas del sistema, las mismas fórmulas axiomáticas por de pronto. Así, por ejemplo, si a_1, a_2, a_3 de Ω_1 se corresponden respectivamente con b_1, b_2, b_3 de Ω_2 , entonces si vale

$$Pa_1 \rightarrow Qa_2 \wedge Ra_3,$$

vale también

$$Pb_1 \rightarrow Qb_2 \wedge Rb_3.$$

Entonces puede pensarse que los individuos de Ω_1 son, en cierto sentido a precisar, los mismos de Ω_2 y que los de cualquier otro modelo isomorfo con ellos — que son todos los modelos del sistema axiomático monomórfico. Todos esos objetos se comportan, en efecto, del mismo modo, entran en las mismas relaciones formales. De aquí nace la idea de que, puesto que el sistema determina el comportamiento de cualquier clase de objetos que lo satisfaga, define a esos objetos de algún modo.

Pero la idea de definición implícita tiene una importante limitación: la noción de un objeto o clase de objetos puede ser algo más que el comportamiento de esos objetos respecto de ciertas relaciones formales, o que la aptitud de sus nombres para figurar en ciertas fórmulas. Las relaciones que en el sistema axiomático describen el comportamiento de esos objetos son, en efecto, formales: en el sistema no nos interesamos por toda la cualidad de la relación, sino sólo por su estructura. Así, en nuestro anterior ejemplo (' $Px \rightarrow Qx$ ' con los modelos isomorfos $\mathfrak{S}_1, \mathfrak{S}_2$) las relaciones entre (1.º) ser parte de la mano izquierda y ser parte del brazo izquierdo y entre (2.º) ser parte de la mano derecha y ser parte del brazo derecho, se identifican formalmente, sin atender más que a la relación ser-parte-de-parte. Dicho de otro modo: aquella fórmula ' $Px \rightarrow Qx$ ' "define implícitamente" mis dedos sólo en su comportamiento como partes de partes, e ignora otras propiedades — otros comportamientos — de dichos dedos que pueden ser interesantes desde otros puntos de vista, con otras abstracciones básicas. En resolución: lo que de verdad define implícitamente un conjunto monomórfico de axiomas es la abstracción básica de su teoría (cfr. 12). No cosas, ni comportamientos totales de cosas, ni tampoco los conceptos del lenguaje común, sin duda mucho más vagos, pero también más ricos.

Así la idea de definición implícita nos conduce a considerar la abstracción básica de una teoría, lo que a ella importa tomar de los objetos del conocimiento que se propone exponer de un modo exacto. Tratándose de la teoría que aquí nos ocupa, la lógica de predicados de primer orden, habrá que pedir de los axiomas de la misma que "definan implícitamente" determinados símbolos, como los funtores veritativos y los cuantificadores, como si aún no hubiéramos hecho ninguna reflexión metalingüística (señaladamente, semántica) sobre ellos. Esto se hace sentando unas cuantas afirmaciones básicas — los axiomas — y unas posibilidades de operar — las reglas de transformación — que contienen esos símbolos y determinan su uso

o comportamiento, es decir, las relaciones en que pueden entrar entre ellos y con otros símbolos.

Existen unos cuantos sistemas axiomáticos más o menos clásicos para la lógica de predicados de primer orden. El que se expone a continuación es el más utilizado. Se debe a D. Hilbert, P. Bernays y W. Ackermann (HB).

45. El sistema axiomático HB para la lógica de predicados de primer orden.

I. Axiomas:

- A1: $p \vee p \rightarrow p$.
- A2: $p \rightarrow p \vee q$.
- A3: $p \vee q \rightarrow q \vee p$.
- A4: $[p \rightarrow q] \rightarrow [s \vee p \rightarrow s \vee q]$.
- A5: $(x)Px \rightarrow Py$.
- A6: $Py \rightarrow \exists xPx$.

II. Reglas de transformación:

$R\alpha_1$: una variable de enunciado (letra de enunciado) puede sustituirse en cualquier fórmula por una fórmula, con las siguientes condiciones:

- α_1' : que la sustitución se haga en todos los lugares en que aparece la variable (letra) de enunciado sustituida;
- α_1'' : que el resultado de la sustitución sea una fórmula;
- α_1''' : que la fórmula sustituyente no contenga variables individuales libres que vayan a quedar ligadas en la fórmula resultante de la sustitución.

$R\alpha_2$: una variable individual libre puede sustituirse por otra, con las siguientes condiciones:

- α_2' : que la sustitución se haga en todos los lugares en que aparezca la variable individual sustituida;
- α_2'' : que la variable sustituyente no quede ligada en la fórmula resultante de la sustitución.

$R\alpha_3$: una variable (letra) predicativa n -ádica puede sustituirse por una fórmula predicativa cualquiera, aunque no sea atómica, con tal de que la sustituyente tenga al menos n variables libres. Y ello con las siguientes condiciones:

- α_3' : que el resultado sea una fórmula;
- α_3'' : que en la fórmula resultante no aparezcan nuevas ligaduras.

$R\beta$: si han sido ya demostradas dos fórmulas, $X \rightarrow Y$ y X , puede escribirse como demostrada Y . Visualizaremos esta regla con el dibujo esquemático:

$$\frac{X \rightarrow Y \quad X}{Y}$$

Esta regla suele llamarse 'de separación', o 'modus ponendo ponens'.

$R\gamma_1$: si vale ya una fórmula $X \rightarrow Y$, tal que ' x ' está libre en Y y no se presenta en X , puede escribirse como válida la fórmula $X \rightarrow (x)Y$:

$$\frac{X \rightarrow Y}{X \rightarrow (x)Y} \quad \text{con: } \gamma_1': 'x' \text{ está libre en } Y; \\ \gamma_1'': 'x' \text{ no se da en } X.$$

$R\gamma_2$: si vale ya una fórmula $Y \rightarrow X$, tal que ' x ' está libre en Y y no se presenta en X , entonces puede escribirse como válida la fórmula

$$\frac{\exists xY \rightarrow X}{\exists xY \rightarrow X} \quad \text{con: } \gamma_2': 'x' \text{ está libre en } Y; \\ \gamma_2'': 'x' \text{ no se da en } X.$$

$R\delta$: una variable individual ligada puede sustituirse por cualquier otra, con las siguientes condiciones:

- δ' : que la sustitución se realice en el operando y en el operador que afectaba a la variable sustituida;
- δ'' : que el resultado de la sustitución sea una fórmula.

III. *Regla de definición*: Las fórmulas representadas por los miembros izquierdo y derecho de cada una de las definiciones siguientes, (D1)-(D3), pueden sustituirse la una por la otra en el interior de cualquier fórmula:

- (D1) ' $X \wedge Y$ ' $\leftrightarrow_{df}$ ' $\sim [\sim X \vee \sim Y]$ '.
- (D2) ' $X \rightarrow Y$ ' $\leftrightarrow_{df}$ ' $\sim X \vee Y$ '.
- (D3) ' $X \leftrightarrow Y$ ' $\leftrightarrow_{df}$ ' $[X \rightarrow Y] \wedge [Y \rightarrow X]$ '.

He aquí algunos ejemplos de lo que permiten las reglas $R\alpha$ — $R\delta$.

Regla $R\alpha_1$. Dada una fórmula, por ejemplo,

$$(1) \quad p \wedge q \rightarrow p,$$

$R\alpha_1$ permite hacer la siguiente sustitución de ' p ' por ' $p \rightarrow r$ ' (anotaremos por mera comodidad: ' $p/p \rightarrow r$ ')

$$(1a) \quad [p \rightarrow r] \wedge q \rightarrow [p \rightarrow r];$$

y también, por ejemplo, el paso de (1) a

$$(1b) \quad s \wedge r \rightarrow s,$$

con p/s y q/r ;
y el paso de (1) a

$$(1c) \quad (x)Px \wedge \exists xRx \rightarrow (x)Px,$$

con $p/(x)Px$, $q/\exists xRx$.

En cambio, $R\alpha_1$ no permite el paso de (1) a

$$(1d) \quad (x)Px \wedge \exists xRx \rightarrow p,$$

el cual no respeta la condición α_1' .

Regla $R\alpha_2$. Dada, por ejemplo, la fórmula

$$(2) \quad (x)[Px \wedge [Qy \vee Ry]] \rightarrow (x)Sx,$$

es correcto por $R\alpha_2$ el paso a

$$(2a) \quad (x)[Px \wedge [Qz \vee Rz]] \rightarrow (x)Sx,$$

con y/z ,

pero no está autorizado el paso de (2) a

$$(2b) \quad (x)[Px \wedge [Qz \vee Ry]] \rightarrow (x)Sx,$$

que vulnera la condición α_2' , ni tampoco el paso de (2) a

$$(2c) \quad (x)[Px \wedge [Qx \vee Rx]] \rightarrow (x)Sx,$$

que no cumple la condición α_2'' .

La *regla $R\alpha_3$* permite, por ejemplo, el paso de

$$(3) \quad (y)(x)Pxy \rightarrow Pyy \vee Pxx$$

a

$$(3a) \quad (y)(x)[Qxy \wedge Rty] \rightarrow [Qyy \wedge Rty] \vee [Qxx \wedge Rtx],$$

con $Pxy/Quv \wedge Rtv$, o $P_2/[Q_2 \wedge R_2]_3$.

En el uso de esta regla se empieza por establecer una correspondencia entre las variables individuales de la fórmula sustituyente y las que siguen a la letra predicativa que se quiere sustituir por aquélla. En el ejemplo, se ha establecido la correspondencia

$$\begin{aligned} u &: x, \\ v &: y. \end{aligned}$$

Al sustituir la letra predicativa por la fórmula, se usan las variables individuales de la fórmula inicial.

La *regla $R\gamma_1$* autoriza, por ejemplo, el siguiente paso de

$$(4) \quad (y)Py \rightarrow Qx$$

a

$$(5) \quad (y)Py \rightarrow (x)Qx.$$

Pero no autoriza el paso de

$$(6) \quad (y)Pyx \rightarrow Pxx$$

a

$$(7) \quad (y)Pyx \rightarrow (x)Pxx,$$

que no cumple la condición γ_1'' .

El motivo de esta prohibición — o sea, de la prevención γ_1'' — puede verse mediante una interpretación que sea modelo de (6). Sea, por ejemplo, $\mathfrak{I}_1$, con

Ω_1 : los números naturales;
 Π_1 : las relaciones diádicas entre números naturales;
 P : $>$.

Con esa interpretación (6) significa:

$$(6a) \quad \text{si todo número natural es mayor que } x, \text{ entonces el número natural } v \text{ es mayor que } x.$$

Mientras que, para la misma interpretación, (7) significa:

$$(7a) \quad \text{si todo número natural es mayor que } x, \text{ entonces el número natural } v \text{ es mayor que todo número natural.}$$

La *regla $R\gamma_1$* se llama a veces 'regla de generalización posterior'.

La *regla $R\gamma_2$* permite pasos como el de

$$(8) \quad Px \rightarrow (y)Qy$$

a

$$(9) \quad \exists xPx \rightarrow (y)Qy;$$

pero no autoriza un paso como el de

$$(10) \quad Px \rightarrow Qyx$$

a

$$(11) \quad \exists xPx \rightarrow Qyx,$$

el cual no cumple la condición γ_2'' . El motivo de esta prohibición puede hacerse intuitivo mediante una interpretación, sobre un Ω de objetos cualesquiera, para la cual (10) signifique

$$(10a) \quad \text{si } x \text{ es un centro, } y \text{ es el círculo cuyo centro es } x.$$

Con esa misma interpretación, (11) cobra otro significado cuyo valor veritativo no será siempre el mismo (no será formalmente el mismo) que el de (10a):

(11a) si hay un centro, y es el círculo cuyo centro es x .

La regla R_{γ_2} se llama a veces 'regla de particularización anterior'.

El sentido de las reglas R_{γ_1} y R_{γ_2} puede resumirse intuitivamente así: si una fórmula que no contiene una determinada variable, ' x ', es antecedente condicional de otra que la tiene libre, entonces es que el condicional no depende de ningún valor particular de ' x ', y, por tanto, la fórmula condicionada (consecuente) lo está para todo valor de ' x ' (Regla R_{γ_1}). Cuando una fórmula que contiene libre una determinada variable ' x ' es condición de otra fórmula que no la contiene, entonces es que para que valga el condicional basta con que el antecedente (la condición) sea válido para algún valor de ' x ' (regla R_{γ_2}).

Por último, la regla R_8 permite pasos como el de

(12) $(x)\exists yPxy,$

por sustitución y/z , a

(12a) $(x)\exists zPxz;$

pero no a

(13) $(x)\exists yPxz,$

ni a

(14) $(x)\exists zPxy,$

los cuales no cumplen la condición δ' .

Interpretando ' P ' por I , la relación de identidad, (12) y (12a) significan ambas

(12b) toda cosa es idéntica con alguna cosa (a saber, consigo misma),

mientras que (13) significa

(13a) toda cosa es idéntica con z ,

y (14) significa

(14a) toda cosa es idéntica con y .

46. La demostración en el sistema axiomático. — El establecimiento de un teorema en una teoría axiomatizada es, en principio, una construcción mecánica de la fórmula correspondiente mediante la aplicación de las reglas. Por eso suele evitarse para nombrar ese proceso la palabra 'demostración',

sustituyéndola por 'derivación'. 'Demostración' quedaría así reservada a operaciones más complejas, y conservaría sus connotaciones psicológicas. Pero aquí no vamos a utilizar esa terminología, pues, primero, la palabra 'derivación' es ya de uso fijo en matemáticas, y, segundo, la demostración en el cálculo o teoría axiomatizada es el esquema lógico representativo de la operación metódica llamada 'demostración' o 'deducción'. Por eso no parece inadecuado usarla en ambos casos.

La palabra que, en cambio, debería excluirse al hablar de cálculos es 'prueba', pues este término se aplica en castellano propiamente a la fundamentación de un enunciado por vía empírica, no formal.

En la exposición de un sistema formal suele darse, tras el conjunto de axiomas y las reglas de transformación, la demostración sistemática de teoremas y reglas auxiliares. El trabajo real no procede siempre así. Lo corriente es, a la inversa, que uno se encuentre con una fórmula, que supone válida por alguna razón (por ejemplo, por haberla tomado de una teoría sin formalizar, pero sólidamente establecida y probablemente formalizable con los axiomas de que se dispone; o por experiencia, etc.), y que tenga que construir su demostración en el sistema.

Veremos a continuación, a título de ejemplo, un modo de proceder para conseguir una demostración de la fórmula (15), inspirada en lo que en lógica tradicional se llamaba 'silogismo en Barbara':

(15) $(x)[[Px \rightarrow Qx] \rightarrow [[Rx \rightarrow Px] \rightarrow [Rx \rightarrow Qx]]]$.

La suposición de que esa fórmula debe ser teorema de la lógica de predicados de primer orden nos viene de la tradición silogística.

La dificultad principal de su demostración consiste en que (15) tiene una generalización no posterior, sino que afecta a toda ella. Si, dejando para después ese problema, atendemos sólo al operando de (15), podemos obtener una rápida demostración del mismo notando su parecido con el axioma A4. Partiendo de éste, empezamos la demostración, numerando las líneas (L) para poder aludir a ellas:

L1.	$[p \rightarrow q] \rightarrow [s \vee p \rightarrow s \vee q]$	A4.
L2.	$[p \rightarrow q] \rightarrow [\sim t \vee p \rightarrow \sim t \vee q]$	$R\alpha_1 : s/\sim t; L1.$
L3.	$[p \rightarrow q] \rightarrow [[t \rightarrow p] \rightarrow [t \rightarrow q]]$	(D2): L2.
L4.	$[Px \rightarrow Qx] \rightarrow [[Rx \rightarrow Px] \rightarrow [Rx \rightarrow Qx]]$	$R\alpha_1 : p/Px; q/Qx;$ $t/Rx; L3.$

Con esto tenemos en la línea L4 el operando de (15) como teorema. Falta resolver el problema de su cuantificación. No podemos generalizarlo sin más, pues las reglas del cálculo (R_{γ_1}) sólo nos permiten generalizar el consecuente de un condicional, y aún eso a condición de que el antecedente no contenga la variable que deseamos generalizar. Para poder seguir adelante, lo que necesitamos es poder afirmar como teorema una fórmula

condicional, cuyo antecedente carezca de 'x' y cuyo consecuente sea nuestro teorema L4.

El problema podría resolverse si hubiera alguna regla que permitiera pasar de una fórmula cualquiera ya demostrada, X , a un condicional, $Y \rightarrow X$, cuyo antecedente fuera una fórmula cualquiera — de acuerdo con el principio que tradicionalmente se enunciaba diciendo: 'lo verdadero se sigue de cualquier cosa' —. De existir esa regla, en la línea L5 de nuestra demostración podríamos escribir un condicional con nuestro teorema L4 como consecuente; y luego, en la línea L6, podríamos generalizarlo por $R\gamma_1$. Por eso, antes de seguir con la demostración de (15), vamos a hacernos con una regla que nos permita dicha operación, es decir, pasar de una fórmula X , ya demostrada, ya teorema, a una fórmula $Y \rightarrow X$, con Y cualquiera. Las reglas así demostradas, mediante reflexiones metalógicas, después de las de la base primitiva, se llaman 'reglas auxiliares'. Llamaremos a ésta que buscamos ' $RA\gamma_3$ '. He aquí su demostración, que es propiamente una justificación metalingüística de operaciones a realizar en el lenguaje objeto:

La.	X	teorema ya demostrado, tomado como premisa.
Lb.	$X \rightarrow X \vee W$	nombre metalingüístico del axioma A2.
Lc.	$X \vee W \rightarrow W \vee X$	nombre metalingüístico del axioma A3.
Ld.	$X \vee W$	$R\beta$ aplicada a La y Lb.
Le.	$W \vee X$	$R\beta$ aplicada a Lc y Ld.
Lf.	$\sim Y \vee X$	$R\alpha_1 : W/\sim Y$; Le.
Lg.	$Y \rightarrow X$	(D2); Lf.

Esa argumentación metalingüística nos justifica la nueva regla $RA\gamma_3$:

$$\frac{X}{Y \rightarrow X}$$

(Se observará que, mientras para los lugares de variable hemos usado nombres metalingüísticos de las mismas — variables sintácticas — en cambio hemos escrito los funtores, y no nombres metalingüísticos de ellos. Nos atendremos siempre a esta práctica intuitiva.)

Aplicando $RA\gamma_3$ para formar la línea L5 de nuestra demostración de (15), podemos tomar el teorema L4 como consecuente de un condicional cuyo antecedente sea una fórmula cualquiera. Buscaremos como antecedente una fórmula que no contenga 'x', con objeto de poder luego generalizar por $R\gamma_1$ el consecuente (L4).

Pero con esto no queda del todo decidida la elección del antecedente de nuestra futura línea L5. Está claro que no debe tener 'x', para poder aplicarle $R\gamma_1$. Pero el resultado no será aún (15), sino (15) como consecuente de una fórmula que habrá que eliminar, para tener sólo (15). La situación será como sigue; tendremos a L4 generalizado como consecuente de una

fórmula 's', que aún hemos de elegir:

$$s \rightarrow (x)[[Px \rightarrow Qx] \rightarrow [[Rx \rightarrow Px] \rightarrow [Rx \rightarrow Qx]]],$$

y de ella queremos obtener el consecuente sólo. Esto no puede hacerse más que aplicando $R\beta$. Pero, para poder aplicar $R\beta$, 's' tiene que poder ser afirmada, tiene que ser a su vez un teorema o un axioma. Lo más cómodo será que tomemos un axioma para el lugar de 'Y' en el esquema de $RA\gamma_3$. (El lugar de 'X' en dicho esquema lo ocupa (15).) Decidido esto, sigamos con la demostración de (15):

- L5. $[p \vee p \rightarrow p] \rightarrow [[Px \rightarrow Qx] \rightarrow [[Rx \rightarrow Px] \rightarrow [Rx \rightarrow Qx]]]$
 $RA\gamma_3$ aplicada a la línea L4.
- L6. $[p \vee p \rightarrow p] \rightarrow (x)[[Px \rightarrow Qx] \rightarrow [[Rx \rightarrow Px] \rightarrow [Rx \rightarrow Qx]]]$
 $R\gamma_1$; L5.
- L7. $p \vee p \rightarrow p$
 A1.
- L8. $(x)[[Px \rightarrow Qx] \rightarrow [[Rx \rightarrow Px] \rightarrow [Rx \rightarrow Qx]]]$
 $R\beta$, aplicada a L6, L7.

En la presentación de un sistema axiomático, el camino que hemos seguido en la demostración de (15), que es un camino analítico, de búsqueda — "hacia arriba", por así decirlo — de las fórmulas y reglas necesarias para la demostración, se sustituye por un camino sintético, "descendente" desde la base primitiva del sistema hasta el teorema. En el caso de nuestro ejemplo, se habría empezado por establecer $RA\gamma_3$, y luego se habría desarrollado la demostración de (15).

Pero también en el trabajo práctico, y no sólo en el de exposición, es fecundo poder proceder constructivamente, o "hacia abajo": por ejemplo, para averiguar consecuencias de un conjunto de fórmulas o enunciados cualesquiera. Este proceder constructivo, sintético, es precisamente lo que da al método algorítmico (y al axiomático como método algorítmico) un rasgo experimental del cual careció la lógica clásica.

Ese rasgo sintético y como experimental es el esencial al método algorítmico, y está presente incluso cuando se procede, como muchas veces en la práctica, de un modo analítico. Por ejemplo: al buscar una demostración de (15) procedimos primero sintéticamente, "hacia abajo", construyendo el operando de (15). (A eso había precedido un elemental momento "hacia arriba", búsqueda del axioma más emparentado con la fórmula.) Luego buscamos de nuevo analíticamente cómo debía ser una pieza que aún nos faltaba para la demostración. Y al descubrir cómo debía ser esa pieza, la sintetizamos a partir de la base primitiva del sistema. Por último, utilizamos las dos piezas sintetizadas, la regla auxiliar y el operando de (15), para construir (15) mismo.

APÉNDICE AL CAPÍTULO VII

B. *Observación.* — a 24: el sistema axiomático HB se ofrece en esta presentación en los *Grundzüge der theoretischen Logik*, de Hilbert y W. Ackermann, sobre todo en las ediciones segunda (1937) y tercera (1949), debidas a Ackermann. Pero en la 4.^a edición (que es la traducida al castellano), Ackermann lo ha alterado sustancialmente. (Cfr. D. Hilbert y W. Ackermann, *Elementos de lógica teórica*, trad. de V. Sánchez de Zavala, Madrid, 1962.)

CAPÍTULO VIII

LA DEDUCCIÓN A PARTIR DE PREMISAS

47. Las tres concepciones de la argumentación formal. — Aristóteles llamó 'analíticos' a sus principales escritos de lógica formal, y la lógica tradicional se calificaba a sí misma de 'disciplina resolutoria', que quiere decir lo mismo que 'analítica'. Estos adjetivos suponen una concepción de la argumentación formal como una operación por la cual, dada una fórmula o dado un enunciado, se busca cuáles son los elementos y los fundamentos de la misma o del mismo. Hemos visto que un trabajo analítico es siempre presupuesto de cualquier otra actividad en lógica formal. Pero también hemos considerado ya otro tipo de operación lógica, que se presenta característicamente en el método axiomático.

La comprensión de la lógica formal como una teoría del sistema de las verdades formales — concepción que se realiza en la presentación axiomática — se basa, en efecto, en un punto de vista contrapuesto al analítico o resolutorio, aunque complementario del mismo, a saber: en un punto de vista sintético, o constructivo. La presentación del sistema axiomático no consiste en analizar las formas para aislar sus elementos y fundamentos, sino en componer o sintetizar, a partir de unos fundamentos dados (el léxico del sistema y sus axiomas) y mediante transformaciones determinadas, nuevas fórmulas o nuevos enunciados (teoremas). Este carácter sintético de la argumentación da a la lógica, como se indicó, unas ciertas posibilidades de experimentación.

Pero ya en la construcción de las ciencias positivas como teorías, incluso como teorías axiomáticas, la situación se complica. Es claro que un conjunto de fórmulas axiomáticas de la lógica no basta para deducir de él las verdades de toda ciencia positiva. Incluso en el caso de que basten las relaciones lógicas contenidas en aquellos axiomas, habrá que añadir a éstos las nociones básicas de la teoría positiva que se esté estudiando. Pero lo normal es que esto no sea aún suficiente, y que haya que añadir algunas relaciones nuevas, de tipo no-lógico, acaso empírico. Por ejemplo: en el conjunto axiomático de Peano para la aritmética (cfr. 37) hay unos cuantos axiomas

que son los de la lógica, y luego cinco que introducen nuevas relaciones y propiedades, como ser-el-siguiente-de-, ser-un-número, ser-propiedad-de-un número, ser-idéntico-con-. Las cinco afirmaciones específicas de la aritmética se añaden al conjunto de los axiomas de la lógica, para obtener del conjunto axiomático así formado las verdades de la aritmética.

Mas también los axiomas o proposiciones primitivas de cualquier otra teoría pueden añadirse de ese modo a los axiomas de la lógica. En cada aplicación concreta, éstos tienen, consiguientemente, una significación u otra, porque los axiomas añadidos, entendidos como "definiciones implícitas", delimitan los campos individuales y predicativos para los que se afirma su validez. Pero entonces el nuevo sistema axiomático se compondrá ya de afirmaciones sobre determinadas realidades, y no de fórmulas sobre cualquier campo individual, como se entiende que son las afirmaciones de la lógica (puesto que las verdades lógicas son las fórmulas válidas para cualquier interpretación posible). Los enunciados materiales añadidos como axiomas a los axiomas lógicos pueden ser, desde el punto de vista de la teoría del conocimiento, meras hipótesis, no verdades universalmente válidas.

Pues bien: una deducción que parte de enunciados de esa naturaleza, que pueden siempre volver a discutirse y de hecho se discuten frecuentemente, es una deducción a partir de supuestos materiales, o premisas. Se trata de un tipo de deducción muy frecuente en la ciencia.

Ese tipo de deducción aún puede recogerse de un modo natural, como acabamos de ver con el ejemplo del sistema de Peano, en la concepción axiomática, mediante el añadido de axiomas materiales a los formales. Pero el tipo de deducción más corriente en la vida cotidiana y en la práctica de la ciencia es todavía un poco distinto. En las decisiones que solemos tomar cotidianamente van implícitas deducciones a partir de premisas que son exclusivamente materiales, de hecho, y sólo útiles para cada caso particular; en esas deducciones la lógica aparece sólo en forma de reglas de inferencia, no de axiomas. La estructura de este género de argumentación, estudiada por Jaskowski y, sobre todo, G. Gentzen (1909-1945) y W. V. O. Quine, se llama, con terminología de Quine, 'deducción natural'.

43. Los problemas de la deducción natural. — Estudiar la estructura de la deducción natural consiste en construirla como algoritmo lógico, previamente a cualquier aplicación. Ahora bien: la deducción natural es una deducción a partir de premisas de hecho, que no tienen por qué ser universalmente válidas, y que son, además, cambiantes, distintas para cada caso. Esto plantea dos problemas cuando se quiere construir la lógica elemental como un sistema de deducción natural.

Primer problema: ¿cómo se pueden conseguir teoremas lógicos, verdades universalmente válidas, partiendo de premisas que no lo son?

Segundo problema: ¿cómo puede realizar la deducción natural la idea

de cálculo o algoritmo, si carece de un elemento de los cálculos, a saber, el conjunto de fórmulas primitivas, o axiomas?

El primer problema se resuelve cuando se consigue legitimizar una operación que consiste en eliminar las premisas de las que se ha partido. Si eso se consigue, entonces se tendrá fórmulas sin premisas, es decir, válidas en cualquier caso. He aquí un ejemplo: supongamos que partiendo de la premisa ' p ' se ha demostrado ' q ':

L1 (p) q .

Si luego se consigue demostrar también ' q ' partiendo de la premisa ' $\sim p$ ':

L2 ($\sim p$) q .

entonces se tiene una situación en la cual ' q ' vale tanto si vale ' p ' cuanto si no vale ' p ' (es decir, si vale ' $\sim p$ '). De acuerdo con el principio del tercio excluso, esta situación autoriza a afirmar ' q ' sin ninguna premisa, esto es, a eliminar las premisas de ' q ':

L3 (0) q .

El segundo problema se resuelve tomando como reglas de transformación (de transformación de premisas cualesquiera) las relaciones expuestas como fórmulas primitivas en los sistemas axiomáticos. Así, por ejemplo, un modo de convertir el sistema axiomático HB en un sistema de deducción natural sería prescindir de los axiomas (situados a la izquierda en la tabla siguiente) y tomar en su lugar las reglas de transformación adecuadas (indicadas a la derecha). (En la notación metalingüística usada para las reglas, una expresión como ' $X \leftrightarrow [x/y]Y$ ' debe entenderse así: ' X ' es la fórmula que resulta de Y al sustituir en Y la variable ' y ', en todos los lugares en que aparece, por la variable ' x '.)

TABLA DE ILUSTRACIÓN DEL EJEMPLO

A1.	$p \vee p \rightarrow p$	R1	$\frac{X \vee X}{X}$
A2.	$p \rightarrow p \vee q$	R2	$\frac{X}{X \vee Y}$
A3.	$p \vee q \rightarrow q \vee p$	R3	$\frac{X \vee Y}{Y \vee X}$
A4.	$[p \rightarrow q] \rightarrow [s \vee p \rightarrow s \vee q]$	R4	$\frac{X \rightarrow Y}{W \vee X \rightarrow W \vee Y}$
A5.	$(x)Px \rightarrow Py$	R5	$\frac{(x)X}{Y, \text{ con } Y \leftrightarrow [y/x]X}$
A6.	$Py \rightarrow \exists xPx$	R6	$\frac{Y}{\exists xX, \text{ con } Y \leftrightarrow [y/z]X}$

Como reglas $R7$, $R8$, etc., figurarían en el conjunto de reglas de ese sistema de deducción natural las demás reglas, $R\alpha_1$, etc., del sistema axiomático, o aquellas que aún hicieran falta. Pues esta transposición de sistema axiomático a cálculo de deducción natural sería muy poco económica, duplicaría reglas; sólo nos vale aquí a título de ejemplo. Pero para completarlo en esta condición de mera ilustración, puede verse que el axioma $A2$, por ejemplo, da de sí en el sistema axiomático (con la ayuda de la regla $R\beta$ del mismo) lo que da $R2$ en el cálculo de deducción natural. En el sistema axiomático, lo único que puede hacerse con $A2$ es pasar, en una demostración, de una línea

$L_n \quad p \rightarrow p \vee q \quad A2$

a una línea

$L_n + 2 \quad p \vee q,$

siempre que se cuente con una línea

$L_n + 1 \quad p$

que permita la aplicación de la regla $R\beta$ del sistema axiomático HB. Esto mismo puede conseguirse con la regla $R2$ del ejemplo de deducción natural. Según ella, si en una línea se tiene

$L_m \quad p,$

puede ya escribirse en la línea siguiente

$L_m + 1 \quad p \vee q.$

49. Las nociones de demostración y teorema en el cálculo de la deducción natural. — Las nociones de demostración y teorema se definieron (cfr. 41) haciendo intervenir la noción de axioma. Como el cálculo de la deducción natural es un cálculo sin axiomas, esas nociones deben formularse de nuevo para él, del modo siguiente:

Una demostración en el cálculo de la deducción natural es una sucesión de fórmulas, cada una de las cuales es una premisa o procede de una o varias premisas mediante un número finito de aplicaciones de las reglas del cálculo.

Un teorema del cálculo de la deducción natural es una fórmula que aparece sin premisas en una demostración (porque las premisas han sido eliminadas).

El ejemplo $L1$ - $L3$ de 48 es una demostración (aún vagamente, con reglas no formuladas explícitamente), y la fórmula $L3$ de ese mismo 48 se presenta como un teorema. En cambio, ' $p \vee q$ ' de la línea $L_m + 1$ de 48 no se presenta como un teorema sino porque no hemos indicado las pre-

misas por las cuales habíamos aceptado ' p ' en la línea L_m . Supongamos que ' p ' valía en esa línea bajo la premisa ' r '. Entonces una indicación correcta debería registrar ese hecho, escribiendo, por ejemplo:

$L_m (r) \quad p,$

lo cual puede leerse: 'bajo la premisa ' r ', vale ' p '. — De L_m la regla $R2$ permite pasar a la línea

$L_m + 1 (r) \quad p \vee q,$

pero también bajo la condición de que valga la premisa ' r ', "arrastrándola" o, como diremos, 'bajo ' r '. ' $p \vee q$ ' no es pues un teorema en la línea $L_m + 1$.

Esta indicación de las premisas de que depende el resultado de la aplicación de una regla es imprescindible. Pues una premisa no tiene por qué tomarse, al modo de los axiomas, como válida. Por tanto, lo que se obtenga de ella mediante la aplicación de las reglas del cálculo no será válido en general, sino sólo válido bajo aquella premisa, según las reglas del cálculo.

Esta es la primera de las características del cálculo de la deducción natural comparado con el sistema axiomático. A continuación estudiaremos sus demás peculiaridades.

50. Los dos géneros de operaciones básicas del cálculo de la deducción natural. — Como algoritmo de la lógica de predicados de primer orden, el cálculo de la deducción natural tiene que operar con las constantes lógicas de ese lenguaje, que son los funtores veritativos y los cuantificadores. Las dos operaciones básicas que pueden hacerse con cada una de ellas son introducirla y eliminarla en o de las fórmulas. El ejemplo de las líneas L_m - $L_m + 1$ es un caso de introducción del functor ' $\vee$ '. Está claro, en cambio, que de

$L_r \quad p \vee q$

no debe ser lícito pasar a

$L_r + 1 \quad p,$

pues ' $p \vee q$ ' puede valer con sólo que valga ' q ', aunque ' p ' sea falsa. Hechos como éste tienen que tomarse en consideración al establecer las reglas del cálculo.

Pero las dificultades principales se encuentran en la introducción de ' $()$ ' y la eliminación de ' $\exists$ '.

La introducción de ' $()$ '. — No es correcto afirmar que bajo la premisa

$L_s \quad Puy$

vale la fórmula

$Ls + 1 (s) \quad (x)Pxy,$

pues bajo la premisa

u es menor que y ,

que puede ser un modelo de la fórmula Ls , no vale la interpretación correspondiente de la fórmula $Ls + 1$:

toda cosa es menor que y .

Si a pesar de eso ese paso es en algún caso lícito, ello será porque ' u ' tenga la peculiar condición de haber sido tomada en Ls sin ninguna restricción, de tal modo que lo que vale para ella respecto de ' y ' valga también para todas las cosas. Este paso, por tanto, sólo será lícito si ' u ' se mantiene para siempre así en el resto de la demostración: como la especial variable que, por estar con ' y ' en la relación P , hace que toda cosa esté con ' y ' en la relación P . Para recordarlo, procederemos a anotar la variable ' u ' en pasos de este género, con objeto de no someterla ya nunca (en la misma demostración) a operaciones que puedan alterar ese su especial papel, al ponerla en alguna otra relación especial. Pero lo que hay que recordar no es sólo que ' u ' es muy especial, sino, además, que su universalidad se refiere precisamente a una relación con ' y '. El uso de ' u ' con esa universalidad depende de ' y '. Tal vez en relación con otra cosa no fuera ' u ' universal de ese modo. Dicho de otra manera: la generalización de ' u ' depende de ' y ', o, en general, de las demás variables libres de la fórmula generalizada.

Según esta reflexión, anotaremos el paso anterior del modo siguiente:

$Ls \quad Puy$
 $u [y] Ls + 1 \quad (x)Pxy.$

También se desprende de la consideración anterior que una variable no puede ser anotada dos veces en una misma demostración, pues la segunda anotación haría ambigua la razón por la cual se ha generalizado la variable. Anotada ésta una vez, su significación está fijada. Anotarla otra vez sería cambiarle la significación.

Por análogas razones debe prohibirse el anotar una variable ' x ' en dependencia de otra, ' y ',

$x [y],$

y luego, en la misma demostración, ' y ' en dependencia de ' x ',

$y [x],$

pues eso anularía toda dependencia, y destruiría así la significación fijada para ' x ', al mismo tiempo que impediría fijar la de ' y '.

La eliminación de ' $\exists$ ' plantea un problema parecido. En efecto, de la línea

$Lt \quad \exists xPxu,$

no debe poderse pasar sin más precauciones a la línea

$Lt + 1 (t) \quad Pyu,$

pues del enunciado

hay algún astro que es satélite de la Tierra,

que puede ser un modelo de la fórmula Lt , no puede pasarse a

Marte es satélite de la Tierra,

enunciado que puede obtenerse con la misma interpretación (sobre un Ω = los astros) aplicada a la fórmula $Lt + 1$.

Si el paso es válido, es que la ' y ' de $Lt + 1$ tiene una característica muy especial: se toma, en efecto, como nombre del algo que es satélite de la Tierra, el algo que está en la relación P con el objeto nombrado ' u '.

Como en el caso de la introducción de ' $()$ ', habrá, pues, que anotar la variable así especificada o fijada, y someterla a las mismas prohibiciones de doble anotación y de anotación en dependencia recíproca. El paso se indicará así:

$Lt \quad \exists xPxu$
 $y [u] Lt + 1 (t) \quad Pyu.$

Las situaciones de los pasos de introducción de ' $()$ ' y de eliminación de ' $\exists$ ' son corrientes en matemáticas. La primera se da, por ejemplo, cuando, después de haber demostrado un determinado enunciado para una variable cualquiera, sin poner a ésta ninguna restricción, se afirma que el enunciado es válido para cualquier individuo del campo de la variable. (Esta comparación es de H. Hermes.) — También la situación de la eliminación de ' $\exists$ ' es frecuente. En las siguientes hipotéticas palabras de una argumentación de geometría se da, por ejemplo, ese paso, y la consiguiente anotación o fijación de ' y ': "Sea el triángulo ABC. Sabemos que hay un punto en el cual se cortan las bisectrices, u , v , z , de sus tres ángulos (premisa: $\exists xPxuvz$). Sea y ese punto (paso a: $Pyuvz$)", etc. Es claro que, después de eso, el geómetra no podrá ya llamar ' y ' a ningún otro punto en la misma demostración: ' y ' está fijado, anotado, según la terminología que hemos adoptado, en dependencia de las tres bisectrices, u , v , z .

51. Isomorfía de fórmulas y sustitución de variables. — Al introducir o eliminar ' $()$ ' o ' $\exists$ ' hay que asegurarse de que se conserva la estructura del operando. La fórmula a la que se aplica la operación y la fórmula resultante

deben tener una parte, el operando de la fórmula cuantificada, isomorfa. Esto ha quedado garantizado en los ejemplos anteriores por el requisito — que no dimos explícitamente — de que se pueda pasar de una a otra versión del operando por sustitución de variables. Así, por ejemplo, la conservación de estructura en el paso de 'Puy' a '(x)Pxy' quedaba garantizada por el hecho de que de 'Pxy' puede pasarse a 'Puy' por sustitución de 'x' por 'u', lo que anotaremos

$$Puy \leftrightarrow [u/x] Pxy,$$

o, más en general, llamando a 'Pxy' X y a 'Puy' Y,

$$Y \leftrightarrow [u/x] X.$$

Obsérvese que ese requisito no impide el paso de

$$\text{La} \quad Pxy \quad (= X)$$

a

$$\text{La} + 1 \text{ (a)} \quad Pyy \quad (= Y),$$

pues vale

$$Y \leftrightarrow [y/x] X.$$

Un paso así no presenta peligro alguno en una operación como la eliminación de '(')' o la introducción de 'E', que son operadores sin dificultades:

$$\text{Lb} \quad (x)Pxy \quad (= (x)X)$$

$$\text{Lb} + 1 \text{ (b)} \quad Pyy \quad (= Y)$$

Pues si *toda* cosa está con *y* en la relación *P*, también *y* mismo estará en esa relación con *y*.

Pero, en cambio, esa posibilidad debe excluirse en la introducción de '(')' o la eliminación de 'E', pasos, como vimos, injustificados en general, y sólo justificables mediante el respeto de ciertas medidas de precaución. En una eliminación de 'E' debe excluirse aquella posibilidad. Por ejemplo:

$$\text{Ld} \quad \exists xPxy \quad (= \exists xX)$$

$$\text{Ld} + 1 \text{ (d)} \quad Pyy \quad (= Y)$$

es un paso que cumple la prevención estructural

$$Y \leftrightarrow [y/x] X,$$

pero es incorrecto, como ilustra, según un procedimiento que ya nos es familiar, una interpretación adecuada aplicada a ambas fórmulas:

para Ld: hay algo que es mayor que *y*;

para Ld + 1: *y* es mayor que *y*.

Con objeto de evitar estos resultados indeseables en los "pasos críticos" (H. Hermes), que son la introducción de '(')' y la eliminación de 'E', habrá que exigir, además de la equivalencia estructural

$$Y \leftrightarrow [y/x] X,$$

(y además de la anotación de la variable fijada, y de las prohibiciones de doble anotación y de anotación en dependencia recíproca), la segunda equivalencia estructural:

$$X \leftrightarrow [x/y] Y.$$

Con esta nueva exigencia era imposible el desastroso paso de Ld a Ld + 1. Pues aunque vale en ese caso

$$Y \leftrightarrow [y/x] X,$$

no se respeta en cambio

$$X \leftrightarrow [x/y] Y,$$

ya que

$$[x/y] Y \leftrightarrow 'Pxx',$$

mientras que

$$X \leftrightarrow 'Pxy'.$$

No siempre es necesario hacer verdaderas sustituciones. La necesidad de practicar sustituciones se presenta cuando, en una demostración larga, hay que anotar varias variables. Entonces conviene utilizar varias, practicando sustituciones. Mas cuando la demostración no requiere anotaciones, o cuando, aun requiriendo algunas, los pasos son pocos, puede ser inútil sustituir realmente. Un ejemplo es la siguiente demostración:

$$\text{L1 (1).} \quad (x)Pxy \quad (= (x)X)$$

$$\text{L2 (1).} \quad Pxy \quad (= Y)$$

$$\text{L3 (0).} \quad (x)Pxy \rightarrow Pxy$$

Otro, el siguiente fragmento de demostración:

$$\text{L1 (1).} \quad \exists xPxy \quad (= \exists xX)$$

$$x [y] \text{ L2 (1).} \quad Pxy \quad (= Y).$$

En los dos casos se respetan los requisitos de sustituibilidad que hemos establecido. En el primer caso, porque vale

$$Y \leftrightarrow [x/x] X;$$

y en el segundo porque vale

$$[Y \leftrightarrow [x/x] X] \wedge [X \leftrightarrow [x/x] Y].$$

Las variables de la última línea de una demostración. — Todas esas reglas de la teoría (metalingüística) de la sustituibilidad en el cálculo de la deducción natural no bastan aún para evitar completamente el peligro de inconsistencia. Considérese, por ejemplo, la siguiente breve serie de fórmulas que se diera como una demostración:

L1 (Py)	Py	(premisa)
y L2 (Py)	(x)Px	(introducción de '()').

Esa sucesión de fórmulas cumple todas las prevenciones puestas al paso: 'y' está anotada (sin depender de ninguna otra variable libre porque no la hay en L2: se recordará que en lógica de predicados de primer orden las letras predicativas, como no son ligables, no son verdaderas variables, sino parámetros). Vale además:

$$[Py \leftrightarrow [y/x] Px] \wedge [Px \leftrightarrow [x/y] Py].$$

Por último, 'y' no está anotada más de una vez, ni anotada, tampoco, en dependencia recíproca con ninguna otra variable.

Sin embargo, la conclusión de la presunta demostración no está justificada, pues afirma que todo es *P* basándose en que *y* es *P*.

Para evitar que fórmulas sin fundamentar, como la de ese ejemplo, puedan darse como conclusiones de demostraciones, estableceremos la siguiente última condición sobre el manejo de las variables: en la última línea de una demostración no puede aparecer libre, ni en la fórmula misma ni en sus premisas, una variable anotada. En el ejemplo anterior, 'y' aparece libre en la premisa 'Py' de '(x)Px'. En el ejemplo siguiente, la variable anotada aparece libre en la fórmula misma de la última línea:

L1 (∃xPx).	∃xPx	(premisa)
y L2 (∃xPx).	Py	(eliminación de '∃').

52. El cálculo de la deducción natural. — En el sistema de reglas del cálculo de la deducción natural que se da a continuación, según una presentación de H. Hermes (con muy pocas variaciones), la 'I' significa 'introducción', la 'E' 'eliminación', 'RP' abrevia 'Regla de introducción de premisas'. — Cuando debajo del trazo que separa las premisas del resultado hay dos fórmulas separadas por un trazo vertical, es que la regla permite pasar de las premisas a cualquiera de las dos fórmulas. Las letras entre paréntesis son premisas que se van arrastrando. — En 53 se comentarán intuitivamente estas reglas.

RP:	Lm (X)	X
RI ~ :	Lm (Z)	X → Y
	Lm + r (U)	X → ~ Y
	Lm + r + 1 (Z, U)	~ X

RE ~ :	Lm (Z)	X
	Lm + r (U)	~ X
	Lm + r + 1 (Z, U)	Y
RI v :	Lm (Z)	X
	Lm + 1 (Z)	X v Y Y v X
RE v :	Lm (Z)	X v Y
	Lm + r (U)	X → W
	Lm + r + n (T)	Y → W
	Lm + r + n + 1 (Z, U, T)	W
RI ∧ :	Lm (Z)	X
	Lm + r (U)	Y
	Lm + r + 1 (Z, U)	X ∧ Y Y ∧ X
RE ∧ :	Lm (Z)	X ∧ Y
	Lm + 1 (Z)	X Y
RI → :	Lm (X, Z)	Y
	Lm + 1 (Z)	X → Y
RE → :	Lm (Z)	X → Y
	Lm + r (U)	X
	Lm + r + 1 (Z, U)	Y
RI ↔ :	Lm (Z)	X → Y
	Lm + r (U)	Y → X
	Lm + r + 1 (Z, U)	X ↔ Y Y ↔ X
RE ↔ :	Lm (Z)	X ↔ Y
	Lm + 1 (Z)	X → Y Y → X
RI () :	Lm (Z)	Y
	Lm + 1 (Z)	(x)X

- con: a) $[Y \leftrightarrow [y/x] X] \wedge [X \leftrightarrow [x/y] Y]$;
 b) anotación de 'y' en dependencia de las demás variables libres de '(x)X';
 c) $Lm + 1$ sólo puede ser última línea de una demostración si 'y' no está libre en Z ni en '(x)X'.

RE () :	Lm (Z)	(x)X
	Lm + 1 (Z)	Y

con: $Y \leftrightarrow [y/x] X$.

$$\text{RI } \exists : \quad \begin{array}{c} \text{Lm (Z)} \\ \text{Lm + 1 (Z)} \end{array} \quad \frac{Y}{\exists x X}$$

con: $Y \leftrightarrow [y/x] X$.

$$\text{RE } \exists : \quad \begin{array}{c} \text{Lm (Z)} \\ \text{Lm + 1 (Z)} \end{array} \quad \frac{\exists x X}{Y}$$

- con: a) $[Y \leftrightarrow [y/x] X] \wedge [X \leftrightarrow [x/y] Y]$;
 b) anotación de 'y' en dependencia de las variables libres en ' $\exists x X$ ';
 c) $\text{Lm} + 1$ sólo puede ser última línea de una demostración si 'y' no está libre en Z ni en Y.

Control final de una demostración. — Una sucesión de fórmulas sólo es una demostración si todas ellas se deben a aplicaciones de las anteriores reglas y si, además, se han respetado las prohibiciones de doble anotación y de anotación en dependencia recíproca.

La demostración está constituida sólo por la sucesión de fórmulas. La numeración de las líneas y la indicación de las premisas son sólo expedientes gráficos para facilitar el control. Por eso mismo es corriente indicar las premisas ya por un mero asterisco (Quine), ya por el simple número de la línea en que fueron introducidas (Hermes). Este último expediente usaremos aquí. A la derecha de las líneas de la demostración es también útil indicar las reglas utilizadas en cada paso, y las líneas a las cuales se han aplicado esas reglas.

53. Observaciones a las reglas de la deducción natural. — La regla RP, regla de premisas, o de introducción de premisas, es esencial al cálculo de la deducción natural, que no tiene axiomas de los que partir. Esta regla autoriza a escribir en cualquier línea de una demostración una fórmula cualquiera, tomándola al mismo tiempo como premisa. Por ejemplo:

$$\text{L1 (1)} \quad p \quad \text{RP}$$

Si se aplica a esa línea L1 la regla $\text{RI} \rightarrow$, que es la que permite eliminar premisas, se tiene:

$$\text{L2 (0)} \quad p \rightarrow p \quad \text{RI} \rightarrow ; \text{L1.}$$

' $p \rightarrow p$ ' es una versión del principio de identidad. Vale sin ninguna premisa (con cero premisa), lo cual indicaremos convencionalmente por '(0). La indicación '(1)', de la línea L1, según una convención de 53, indica no que '1' — que es el número de la línea — sea premisa de ' p ', cosa sin sentido, sino que la fórmula de la línea en que se encuentra — o sea, ' p ' — tiene como premisa la fórmula de la línea L1 (en este caso, también ' p ').

La regla $\text{RI} \sim$, introducción de la negación, introduce a ésta como falsedad de una fórmula que tiene consecuentes contradictorios. Es una versión del principio de no-contradicción.

La regla $\text{RE} \sim$ elimina la negación según el dicho tradicional 'a falso sequitur quælibet', 'de lo falso se sigue cualquier cosa'. La regla autoriza a escribir cualquier fórmula como resultado de una contradicción. Esta regla es también una versión del principio de no-contradicción.

La regla $\text{RI} \vee$ puede entenderse como una definición (de uso) de ' $\vee$ '.

La regla $\text{RE} \vee$ elimina ' $\vee$ ' al modo de lo que tradicionalmente se llamaba 'dilema', un esquema de razonamiento del tipo siguiente:

no hay más que dos posibilidades, X o Y;
 si X, entonces W;
 si Y, entonces W.

Luego: W.

Las reglas $\text{RI} \wedge$ y $\text{RE} \wedge$ pueden también entenderse como definiciones (de uso) de ' $\wedge$ '.

La regla $\text{RI} \rightarrow$ se basa en la idea de que una premisa, X, de una fórmula, Y, puede concebirse como antecedente de un condicional cuyo consecuente es Y. Esta regla puede también entenderse como regla de eliminación de premisas.

La regla $\text{RE} \rightarrow$ es el equivalente del "modus ponendo ponens" tradicional, o de la regla $\text{R}\beta$ del sistema axiomático HB.

Las reglas $\text{RI} \leftrightarrow$ y $\text{RE} \leftrightarrow$ pueden entenderse como definiciones (de uso) de ' $\leftrightarrow$ ' por ' $\rightarrow$ '.

Las reglas referentes a los cuantificadores han sido ya discutidas intuitivamente en 51.

Todas las reglas pueden entenderse como definiciones de uso, y no sólo aquellas a propósito de las cuales se ha invitado explícitamente a hacerlo. Las reglas "definen" las constantes por el procedimiento de fijar su uso.

El principio de tercio excluso. — Dos de los tres "principios lógicos" tradicionales (cfr. 6) — el de identidad y el de no-contradicción — están recogidos, como se ha visto, en las reglas del cálculo. No así el de tercio excluso. Éste debe incluirse explícitamente cuando hace falta, que no es siempre. Y puesto que hay una escuela lógica — la llamada 'intuicionista' o 'constructivista' — que no utiliza el principio de tercio excluso como instrumento para sintetizar fórmulas en demostraciones cualesquiera, es interesante observar que el cálculo de la deducción natural sirve también para dicha escuela.

Quando, en cambio, se le quiere utilizar — y así haremos aquí —, se añade a las reglas anteriores una última:

RTE : Lm (0) $X \vee \sim X$,

cuyo contenido intuitivo es: en cualquier línea de una demostración puede escribirse sin premisa alguna (esto es, como teorema) una fórmula de la forma $X \vee \sim X$.

Esta situación tiene la siguiente consecuencia respecto de las relaciones entre lógica sin tercio excluso y lógica con él: los sistemas lógicos sin tercio excluso no son arbitrarios, pues respetan todas las demás reglas. Son, simplemente, sistemas que carecen de todos aquellos teoremas clásicos para cuya demostración es necesario usar la regla RTE.

APÉNDICE AL CAPÍTULO VIII

B. *Observación.* — a 53: a pesar de la simplicidad de su principio, el cálculo de la deducción natural ha sido de elaboración un tanto accidentada. La primera versión de Quine no era consistente. La presentación aquí leída, fruto de la modificación del cálculo por Quine para asegurar su consistencia, es poco simétrica, cosa que contrasta con el carácter simétrico de sus dos operaciones básicas. Hasenjäger ha dado una presentación más simétrica que resulta, en cambio, algo menos cómoda en su manejo, y también menos intuitiva.

CAPÍTULO IX

TÉCNICA DE LA DEDUCCIÓN NATURAL. ALGUNOS TEOREMAS

54. *Algunos teoremas de la lógica de enunciados.* — Los siguientes teoremas (TE1)-(TE15) no contienen cuantificadores, y pertenecen así a la parte de la lógica de predicados de primer orden que llamamos 'lógica de enunciados'.

(TE1) $p \leftrightarrow \sim \sim p$.

Demostración

(Número de la línea y premisas)	(Fórmula afirmada)	(Reglas utilizadas, con indicación de las líneas a las cuales se han aplicado)
L1 (1)	p	RP.
L2 (2)	$\sim p$	RP.
L3 (1, 2)	p	RE $\sim$; L1, L2.
L4 (1)	$\sim p \rightarrow p$	RI $\rightarrow$; L3.
L5 (0)	$\sim p \rightarrow \sim p$	RI $\rightarrow$; L2.
L6 (1)	$\sim \sim p$	RI $\sim$; L4, L5.
L7 (0)	$p \rightarrow \sim \sim p$	RI $\rightarrow$; L6.
L8 (8)	$\sim \sim p$	RP.
L9 (2, 8)	p	RE $\sim$; L2, L8.
L10 (8)	$\sim p \rightarrow p$	RI $\rightarrow$; L9.
L11 (0)	$p \rightarrow p$	RI $\rightarrow$; L1.
L12 (0)	$p \vee \sim p$	RTE.
L13 (8)	p	RE $\vee$; L10, L11, L12.
L14 (0)	$\sim \sim p \rightarrow p$	RI $\rightarrow$; L13.
L15 (0)	$p \leftrightarrow \sim \sim p$	RI $\leftrightarrow$; L7, L14.

(TE1) suele llamarse 'ley de doble negación'. Para su demostración hace falta el principio de tercio excluso, utilizado en la línea L12. Puede

servirnos como ejemplo del modo de demostrar equivalencias: primero se demuestra uno de los dos condicionales contenidos en la equivalencia, luego el otro, y finalmente se compone la equivalencia mediante la regla RI $\leftrightarrow$.

Cada condicional, por su parte, se demuestra empezando por tomar como premisa el antecedente: L1, L8.

Cada equivalencia demostrada puede luego utilizarse según una regla auxiliar que autorice a sustituir en cualquier demostración una fórmula por el equivalente demostrado de la misma. Esa "regla auxiliar número 1" tendrá el siguiente aspecto:

RA1:	Ln (Z)	X
	Lr (U)	$X \leftrightarrow Y$
	Lr + 1 (Z, U)	$\frac{Y}{X}$

La justificación de esta regla auxiliar es que siempre podríamos intercalar en la nueva demostración la demostración de la equivalencia utilizada. He aquí un ejemplo en el que (TE1) se usa según esa regla auxiliar:

(TE2) $[\sim p \rightarrow q] \rightarrow [\sim q \rightarrow p]$

Demostración

L1 (1)	$\sim p \rightarrow q$	RP.
L2 (2)	$\sim p$	RP.
L3 (3)	$\sim q$	RP.
L4 (2, 3)	$\sim p \wedge \sim q$	RI $\wedge$; L2, L3.
L5 (2, 3)	$\sim q$	RE $\wedge$; L4.
L6 (3)	$\sim p \rightarrow \sim q$	RI $\rightarrow$; L5.
L7 (1, 3)	$\sim \sim p$	RI $\sim$; L1, L6.
L8 (1, 3)	p	(TE1) aplicado según la regla RA1 a L7 (omitiendo una línea ' $p \leftrightarrow \sim \sim p$ ', que suponemos registrada en una tabla de teoremas).

L9 (1)	$\sim q \rightarrow p$	RI $\rightarrow$; L8.
L10 (0)	$[\sim p \rightarrow q] \rightarrow [\sim q \rightarrow p]$	RI $\leftrightarrow$; L9.

Esta demostración da pie para admitir como regla auxiliar otro expediente cómodo: afirmar una fórmula con más premisas de las necesarias, con objeto de conseguir luego un condicional conveniente, mediante RI $\rightarrow$. La regla podría tener el siguiente aspecto:

RA2:	Ln (Z)	X
	Ln + m (Z, U)	$\frac{X}{X}$

En la línea L3 de la anterior demostración se tenía a ' $\sim q$ ' sin más premisa que ' $\sim q$ ' misma. Mediante la regla RI $\wedge$, en la línea L5 se llega a tener a ' $\sim q$ ' bajo las premisas ' $\sim q$ ' y ' $\sim p$ '. Esto permite construir un útil condicional en L6. — Usando la regla auxiliar RA2 se habría podido abreviar en un paso la demostración de (TE1), del modo siguiente:

L1 (1)	$\sim p$	RP.
L2 (0)	$\sim p \rightarrow \sim p$	RI $\rightarrow$; L1.
L3 (1, 3)	p	RP, RA2, con la premisa superflua ' $\sim p$ '.
L4 (3)	$\sim p \rightarrow p$	RI $\rightarrow$; L3.
L5 (3)	$\sim \sim p$	RI $\sim$; L2, L4.
L6 (0)	$p \rightarrow \sim \sim p$	RI $\rightarrow$; L5.

Este resultado no se había alcanzado antes hasta la línea L7.

El teorema (TE2) es una de las llamadas 'leyes de la contraposición'. Valen otras tres sin cuantificadores:

(TE3)	$[p \rightarrow q] \rightarrow [\sim q \rightarrow \sim p].$
(TE4)	$[p \rightarrow \sim q] \rightarrow [q \rightarrow \sim p].$
(TE5)	$[\sim p \rightarrow \sim q] \rightarrow [q \rightarrow p].$

(TE5) se demuestra como (TE2), usando la ley de doble negación: como ésta a su vez se demuestra usando el principio de tercio excluso, eso quiere decir que (TE2) y (TE5) suponen dicho principio. El cual, en cambio, no es necesario para la demostración de (TE3) y (TE4).

(TE6) $p \vee q \leftrightarrow q \vee p$

Demostración

L1 (1)	$p \vee q$	RP.
L2 (2)	p	RP.
L3 (3)	q	RP.
L4 (2)	$q \vee p$	RI $\vee$; L2.
L5 (3)	$q \vee p$	RI $\vee$; L3.
L6 (0)	$p \rightarrow q \vee p$	RI $\rightarrow$; L4.
L7 (0)	$q \rightarrow q \vee p$	RI $\rightarrow$; L5.
L8 (1)	$q \vee p$	RE $\vee$; L1, L6, L7.
L9 (0)	$p \vee q \rightarrow q \vee p$	RI $\rightarrow$; L8.

La otra mitad de la demostración, o sea; la demostración de ' $q \vee p \rightarrow p \vee q$ ', procede del mismo modo. La demostración termina con una aplicación de la regla RI $\leftrightarrow$.

El teorema (TE6) se llama 'ley conmutativa de la disyunción'.

Igualmente vale la ley conmutativa de la conjunción:

$$(TE7) \quad p \wedge q \leftrightarrow q \wedge p$$

Demostración

L1 (1)	$p \wedge q$	RP.
L2 (1)	q	RE $\wedge$; L1.
L3 (1)	p	RE $\wedge$; L1.
L4 (1)	$q \wedge p$	RI $\wedge$; L2, L3.
L5 (0)	$p \wedge q \rightarrow q \wedge p$	RI $\rightarrow$; L4.

La otra mitad de la demostración, hasta ' $q \wedge p \rightarrow p \wedge q$ ', procede del mismo modo. La demostración termina con una aplicación de la regla $R \leftrightarrow$.

$$(TE8) \quad p \wedge [q \wedge r] \leftrightarrow [p \wedge q] \wedge r$$

Demostración

L1 (1)	$p \wedge [q \wedge r]$	RP.
L2 (1)	p	RE $\wedge$; L1.
L3 (1)	$q \wedge r$	RE $\wedge$; L1.
L4 (1)	q	RE $\wedge$; L3.
L5 (1)	r	RE $\wedge$; L3.
L6 (1)	$p \wedge q$	RI $\wedge$; L2, L4.
L7 (1)	$[p \wedge q] \wedge r$	RI $\wedge$; L6, L5.
L8 (0)	$p \wedge [q \wedge r] \rightarrow [p \wedge q] \wedge r$	RI $\rightarrow$; L7.

La segunda mitad de la demostración, hasta ' $[p \wedge q] \wedge r \rightarrow p \wedge [q \wedge r]$ ', procede del mismo modo. La demostración termina con una aplicación de la regla $RI \leftrightarrow$.

(TE8) se llama 'ley asociativa de la conjunción'. Igualmente vale la ley asociativa de la disyunción:

$$(TE9) \quad p \vee [q \vee r] \leftrightarrow [p \vee q] \vee r.$$

Como consecuencia de (TE8) y (TE9), los paréntesis resultan inútiles en fórmulas que no contengan más functor que ' $\wedge$ ' y en fórmulas que no contengan más functor que ' $\vee$ '.

$$(TE10) \quad p \vee [q \wedge r] \leftrightarrow [p \vee q] \wedge [p \vee r].$$

Demostración

L1 (1)	$p \vee [q \wedge r]$	RP.
L2 (2)	p	RP.
L3 (2)	$p \vee q$	RI $\vee$; L2.
L4 (2)	$p \vee r$	RI $\vee$; L2.
L5 (2)	$[p \vee q] \wedge [p \vee r]$	RI $\wedge$; L3, L4.
L6 (0)	$p \rightarrow [p \vee q] \wedge [p \vee r]$	RI $\rightarrow$; L5.
L7 (7)	$q \wedge r$	RP.
L8 (7)	q	RE $\wedge$; L7.
L9 (7)	r	RE $\wedge$; L8.
L10 (7)	$p \vee q$	RI $\vee$; L8.
L11 (7)	$p \vee r$	RI $\vee$; L9.
L12 (7)	$[p \vee q] \wedge [p \vee r]$	RI $\wedge$; L10, L11.
L13 (0)	$q \wedge r \rightarrow [p \vee q] \wedge [p \vee r]$	RI $\rightarrow$; L12.
L14 (1)	$[p \vee q] \wedge [p \vee r]$	RE $\vee$; L1, L6, L13.
L15 (0)	$p \vee [q \wedge r] \rightarrow [p \vee q] \wedge [p \vee r]$	RI $\rightarrow$; L14.
L16 (16)	$[p \vee q] \wedge [p \vee r]$	RP.
L17 (16)	$p \vee q$	RE $\wedge$; L16.
L18 (16)	$p \vee r$	RE $\wedge$; L16.
L19 (19)	q	RP.
L20 (20)	r	RP.
L21 (19, 20)	$q \wedge r$	RI $\wedge$; L19, L20.
L22 (19, 20)	$p \vee [q \wedge r]$	RI $\vee$; L21.
L23 (2)	$p \vee [q \wedge r]$	RI $\vee$; L2.
L24 (20)	$q \rightarrow p \vee [q \wedge r]$	RI $\rightarrow$; L22.
L25 (0)	$p \rightarrow p \vee [q \wedge r]$	RI $\rightarrow$; L23.
L26 (16, 20)	$p \vee [q \wedge r]$	RE $\vee$; L17, L24, L25.
L27 (16)	$r \rightarrow p \vee [q \wedge r]$	RI $\rightarrow$; L26.
L28 (16)	$p \vee [q \wedge r]$	RE $\vee$; L18, L25, L27.
L29 (0)	$[p \vee q] \wedge [p \vee r] \rightarrow p \vee [q \wedge r]$	RI $\rightarrow$; L28.
L30 (0)	$p \vee [q \wedge r] \leftrightarrow [p \vee q] \wedge [p \vee r]$	RI $\leftrightarrow$; L15, L29.

Este ejemplo muestra la técnica general adecuada para operar partiendo de una disyunción: hay que introducir como premisas todos los componentes de la disyunción; luego hay que conseguir que la fórmula buscada se encuentre como consecuente de condicionales cuyos antecedentes sean aquellos componentes de la disyunción inicial (por ejemplo, L6, L13). Entonces se puede aplicar la regla RE $\vee$, y conseguir la fórmula buscada como consecuente de la disyunción de partida.

(TE10) es la ley distributiva de la disyunción por la conjunción. A diferencia de lo que ocurre en álgebra — donde sólo hay una ley distributiva: del producto por la adición, pero no de la adición por el producto —, para

' $\wedge$ ' y ' $\vee$ ' hay dos leyes distributivas. Pues vale también la distribución de la conjunción por la disyunción:

$$(TE11) \quad p \wedge [q \vee r] \leftrightarrow [p \wedge q] \vee [p \wedge r]$$

Demostración

L1 (1)	$p \wedge [q \vee r]$	RP.
L2 (1)	p	RE $\wedge$; L1.
L3 (1)	$q \vee r$	RE $\wedge$; L1.
L4 (4)	q	RP.
L5 (5)	r	RP.
L6 (1, 4)	$p \wedge q$	RI $\wedge$; L2, L4.
L7 (1, 4)	$[p \wedge q] \vee [p \wedge r]$	RI $\vee$; L6.
L8 (1)	$q \rightarrow [p \wedge q] \vee [p \wedge r]$	RI $\rightarrow$; L7.
L9 (1, 5)	$p \wedge r$	RI $\wedge$; L2, L5.
L10 (1, 5)	$[p \wedge q] \vee [p \wedge r]$	RI $\vee$; L9.
L11 (1)	$r \rightarrow [p \wedge q] \vee [p \wedge r]$	RI $\rightarrow$; L10.
L12 (1)	$[p \wedge q] \vee [p \wedge r]$	RE $\vee$; L3, L8, L11.
L13 (0)	$p \wedge [q \vee r] \rightarrow [p \wedge q] \vee [p \wedge r]$	RI $\rightarrow$; L12.
L14 (14)	$[p \wedge q] \vee [p \wedge r]$	RP.
L15 (15)	$p \wedge q$	RP.
L16 (16)	$p \wedge r$	RP.
L17 (15)	p	RE $\wedge$; L15.
L18 (15)	q	RE $\wedge$; L15.
L19 (15)	$q \vee r$	RI $\vee$; L18.
L20 (15)	$p \wedge [q \vee r]$	RI $\wedge$; L17, L19.
L21 (0)	$p \wedge q \rightarrow p \wedge [q \vee r]$	RI $\rightarrow$; L20.
L22 (16)	p	RE $\wedge$; L16.
L23 (16)	r	RE $\wedge$; L16.
L24 (16)	$q \vee r$	RI $\vee$; L23.
L25 (16)	$p \wedge [q \vee r]$	RI $\wedge$; L22, L24.
L26 (0)	$p \wedge r \rightarrow p \wedge [q \vee r]$	RI $\rightarrow$; L25.
L27 (14)	$p \wedge [q \vee r]$	RE $\vee$; L14, L21, L26.
L28 (0)	$[p \wedge q] \vee [p \wedge r] \rightarrow p \wedge [q \vee r]$	RI $\rightarrow$; L27.
L29 (0)	$p \wedge [q \vee r] \leftrightarrow [p \wedge q] \vee [p \wedge r]$	RI $\leftrightarrow$; L13, L28.

Los dos últimos ejemplos muestran repetidamente, además de los procedimientos generales para operar con conjunciones y disyunciones, la necesidad de disponer de una misma fórmula varias veces, pero con premisas distintas. (Por ejemplo, en la última demostración, las líneas L19 y L24.)

(TE10) y (TE11) fundan dos modos de "sacar factor común" en lógica de enunciados.

$$(TE12) \quad \sim [p \wedge q] \leftrightarrow \sim p \vee \sim q.$$

Demostración

L1 (1)	$\sim [p \wedge q]$	RP.
L2 (2)	p	RP.
L3 (3)	q	RP.
L4 (2, 3)	$p \wedge q$	RI $\wedge$; L2, L3.
L5 (1, 2, 3)	$\sim p$	RE $\sim$; L1, L4.
L6 (1, 3)	$p \rightarrow \sim p$	RI $\rightarrow$; L5.
L7 (0)	$p \rightarrow p$	RI $\rightarrow$; L2.
L8 (1, 3)	$\sim p$	RI $\sim$; L6, L7.
L9 (1, 3)	$\sim p \vee \sim q$	RI $\vee$; L8.
L10 (1)	$q \rightarrow \sim p \vee \sim q$	RI $\rightarrow$; L9.
L11 (11)	$\sim q$	RP.
L12 (11)	$\sim p \vee \sim q$	RI $\vee$; L11.
L13 (0)	$\sim q \rightarrow \sim p \vee \sim q$	RI $\rightarrow$; L12.
L14 (0)	$q \vee \sim q$	RTE.
L15 (1)	$\sim p \vee \sim q$	RE $\vee$; L10, L13, L14.
L16 (0)	$\sim [p \wedge q] \rightarrow \sim p \vee \sim q$	RI $\rightarrow$; L15.
L17 (17)	$\sim p \vee \sim q$	RP.
L18 (18)	$\sim p$	RP.
L19 (19)	$\sim q$	RP.
L20 (20)	$p \wedge q$	RP.
L21 (20)	p	RE $\wedge$; L20.
L22 (18, 20)	$\sim p$	RI $\sim$; L18, L21.
L23 (18)	$p \wedge q \rightarrow \sim p$	RI $\rightarrow$; L22.
L24 (0)	$p \wedge q \rightarrow p$	RI $\rightarrow$; L21.
L25 (18)	$\sim [p \wedge q]$	RI $\sim$; L23, L24.
L26 (20)	q	RE $\wedge$; L20.
L27 (19, 20)	$\sim q$	RE $\sim$; L19, L26.
L28 (19)	$p \wedge q \rightarrow \sim q$	RI $\rightarrow$; L27.
L29 (0)	$p \wedge q \rightarrow q$	RI $\rightarrow$; L26.
L30 (19)	$\sim [p \wedge q]$	RI $\sim$; L28, L29.
L31 (0)	$\sim p \rightarrow \sim [p \wedge q]$	RI $\rightarrow$; L25.
L32 (0)	$\sim q \rightarrow \sim [p \wedge q]$	RI $\rightarrow$; L30.
L33 (17)	$\sim [p \wedge q]$	RE $\vee$; L17, L31, L32.
L34 (0)	$\sim p \vee \sim q \rightarrow \sim [p \wedge q]$	RI $\rightarrow$; L33.
L35 (0)	$\sim [p \wedge q] \leftrightarrow \sim p \vee \sim q$	RI $\leftrightarrow$; L16, L34.

(TE12) es una de las llamadas 'leyes de De Morgan' (De Morgan, 1806-1871). Otra es

$$(TE13) \quad \sim [p \vee q] \leftrightarrow \sim p \wedge \sim q.$$

Cuando se trata de demostrar una fórmula que no contiene más functor que ' $\rightarrow$ ', el procedimiento se reduce a introducir como premisas todos los

antecedentes, y luego eliminarlos de modo oportuno. He aquí un ejemplo, forma del 'principio del silogismo' en lógica de enunciados, que es una ley de transitividad de ' $\rightarrow$ ':

$$(TE14) \quad [p \rightarrow q] \rightarrow [[q \rightarrow r] \rightarrow [p \rightarrow r]]$$

Demostración

L1 (1)	$p \rightarrow q$	RP.
L2 (2)	$q \rightarrow r$	RP.
L3 (3)	p	RP.
L4 (1, 3)	q	RE $\rightarrow$; L1, L3.
L5 (1, 2, 3)	r	RE $\rightarrow$; L2, L4.
L6 (1, 2)	$p \rightarrow r$	RI $\rightarrow$; L5.
L7 (1)	$[q \rightarrow r] \rightarrow [p \rightarrow r]$	RI $\rightarrow$; L6.
L8 (0)	$[p \rightarrow q] \rightarrow [[q \rightarrow r] \rightarrow [p \rightarrow r]]$	RI $\rightarrow$; L7.

Análogamente vale

$$(TE15) \quad [p \leftrightarrow q] \wedge [q \leftrightarrow r] \rightarrow [p \leftrightarrow r],$$

que es una ley de transitividad de ' $\leftrightarrow$ '.

55. Algunos teoremas con cuantificadores.—Las dos fórmulas siguientes son los esquemas de dos "modos silogísticos" tradicionales. Su consideración aquí ilustra dos cosas: primero, que toda la silogística (es decir, el sistema de los "modos silogísticos") es un breve conjunto de teoremas de la lógica de predicados de primer orden; segundo, que no todos los modos silogísticos son correctos si se formulan sin precauciones en dicha lógica.

$$(TP16) \quad (x)[Px \rightarrow \sim Qx] \wedge (x)[Rx \rightarrow Px] \rightarrow (x)[Rx \rightarrow \sim Qx]$$

Demostración

L1 (1)	$(x)[Px \rightarrow \sim Qx] \wedge (x)[Rx \rightarrow Px]$	RP.
L2 (1)	$(x)[Px \rightarrow \sim Qx]$	RE $\wedge$; L1.
L3 (1)	$(x)[Rx \rightarrow Px]$	RE $\wedge$; L1.
L4 (1)	$Px \rightarrow \sim Qx$	RE (); L2.
L5 (1)	$Rx \rightarrow Px$	RE (); L3.
L6 (6)	Rx	RP.
L7 (1, 6)	Px	RE $\rightarrow$; L5, L6.
L8 (1, 6)	$\sim Qx$	RE $\rightarrow$; L4, L7.
L9 (1)	$Rx \rightarrow \sim Qx$	RI $\rightarrow$; L8.
x L10 (1)	$(x)[Rx \rightarrow \sim Qx]$	RI (); L9.
L11 (0)	$(x)[Px \rightarrow \sim Qx] \wedge (x)[Rx \rightarrow Px] \rightarrow (x)[Rx \rightarrow \sim Qx]$	RI $\rightarrow$; L10.

(TP16) es el modo "Celarent", de la primera figura silogística ("Ningún M es P; todo S es M; luego ningún S es P").

La fórmula

$$(x)[Rx \rightarrow Qx] \wedge (x)[Rx \rightarrow Px] \rightarrow \exists x[Px \wedge Qx],$$

que corresponde al modo silogístico "Darapti" de la tercera figura, ("Todo M es P; todo M es S; luego algún S es P"), no es válida en esa formulación, sino sólo si se añade a la premisa la afirmación de la existencia de algún x que sea R (o sea, si se le añade: ' $\exists xRx$ '). La razón de esto es que mientras la premisa sólo afirma relaciones lógicas, la conclusión afirma existencia (cfr. 26). Pero la existencia no puede inferirse de meras relaciones lógicas. Así, por ejemplo, de la premisa

todos los centauros tienen cabeza humana y todos los centauros son cuadrúpedos

no se sigue la conclusión

algunos cuadrúpedos tienen cabeza humana

más que si hay centauros.

La formulación correcta de Darapti es pues:

$$(TP17) \quad (x)[Rx \rightarrow Qx] \wedge (x)[Rx \rightarrow Px] \wedge \exists xRx \rightarrow \exists x[Px \wedge Qx]$$

Demostración

L1 (1)	$(x)[Rx \rightarrow Qx] \wedge (x)[Rx \rightarrow Px] \wedge \exists xRx$	RP.
L2 (1)	$(x)[Rx \rightarrow Qx]$	RE $\wedge$; L1.
L3 (1)	$(x)[Rx \rightarrow Px]$	RE $\wedge$; L1.
L4 (1)	$\exists xRx$	RE $\wedge$; L1.
L5 (1)	$Rx \rightarrow Qx$	RE (); L2.
L6 (1)	$Rx \rightarrow Px$	RE (); L3.
x L7 (1)	Rx	RE $\exists$; L4.
L8 (1)	Qx	RE $\rightarrow$; L5, L7.
L9 (1)	Px	RE $\rightarrow$; L6, L7.
L10 (1)	$Px \wedge Qx$	RI $\wedge$; L8, L9.
L11 (1)	$\exists x[Px \wedge Qx]$	RI $\exists$; L10.
L12 (0)	$(x)[Rx \rightarrow Qx] \wedge (x)[Rx \rightarrow Px] \wedge \exists xRx \rightarrow \exists x[Px \wedge Qx]$	RI $\rightarrow$; L11.

(TP18) $(x)[Py \rightarrow Px] \leftrightarrow [Py \rightarrow (x)Px]$.

	L1 (1)	$(x)[Py \rightarrow Px]$	RP.
	L2 (1)	$Py \rightarrow Pz$	RE () ; L1.
	L3 (3)	Py	RP.
	L4 (1, 3)	Pz	RE $\rightarrow$; L2, L3.
z	L5 (1, 3)	$(x)Px$	RI () ; L4.
	L6 (1)	$Py \rightarrow (x)Px$	RI $\rightarrow$; L5.
	L7 (0)	$(x)[Py \rightarrow Px] \rightarrow [Py \rightarrow (x)Px]$	RI $\rightarrow$; L6.
	L8 (8)	$Py \rightarrow (x)Px$	RP.
	L9 (3, 8)	$(x)Px$	RE $\rightarrow$; L3, L8.
	L10 (3, 8)	Px	RE () ; L9.
	L11 (8)	$Py \rightarrow Px$	RI $\rightarrow$; L10.
$x[y]$	L12 (8)	$(x)[Py \rightarrow Px]$	RI () ; L11.
	L13 (0)	$[Py \rightarrow (x)Px] \rightarrow (x)[Py \rightarrow Px]$	RI $\rightarrow$; L12.
	L14 (0)	$(x)[Py \rightarrow Px] \leftrightarrow [Py \rightarrow (x)Px]$	RI $\leftrightarrow$; L7, L13.

(TP19) $(x)[Px \rightarrow Qz] \leftrightarrow [\exists x Px \rightarrow Qz]$ *Demostración*

	L1 (1)	$(x)[Px \rightarrow Qz]$	RP
	L2 (1)	$Pu \rightarrow Qz$	RE () ; L1
	L3 (3)	$\exists x Px$	RP
u	L4 (3)	Pu	RE $\exists$; L3
	L5 (1, 3)	Qz	RE $\rightarrow$; L2, L4
	L6 (1)	$\exists x Px \rightarrow Qz$	RI $\rightarrow$; L5
	L7 (0)	$(x)[Px \rightarrow Qz] \rightarrow [\exists x Px \rightarrow Qz]$	RI $\rightarrow$; L6
	L8 (8)	$\exists x Px \rightarrow Qz$	RP
	L9 (9)	Px	RP
	L10 (9)	$\exists x Px$	RI $\exists$; L9
	L11 (8, 9)	Qz	RE $\rightarrow$; L8, L10
	L12 (8)	$Px \rightarrow Qz$	RI $\rightarrow$; L11
$x[z]$	L13 (8)	$(x)[Px \rightarrow Qz]$	RI () ; L12
	L14 (0)	$[\exists x Px \rightarrow Qz] \rightarrow (x)[Px \rightarrow Qz]$	RI $\rightarrow$; L13
	L15 (0)	$(x)[Px \rightarrow Qz] \leftrightarrow [\exists x Px \rightarrow Qz]$	RI $\leftrightarrow$; L7, L14

También valen:

(TP20) $\exists x[Px \rightarrow Qz] \leftrightarrow [(x)Px \rightarrow Qz]$ (TP21) $\exists x[Qz \rightarrow Px] \leftrightarrow [Qz \rightarrow \exists x Px]$

(TP20) y (TP21) se demuestran análogamente a (TP18) y (TP20), con los que están emparentados. Son todos ellos teoremas útiles para conse-

guir que los cuantificadores afecten a fórmulas enteras, en vez de a partes de fórmula enlazadas por funtores veritativos.

(TP22) $(x)Px \wedge (x)Qx \leftrightarrow (x)[Px \wedge Qx]$ *Demostración*

	L1 (1)	$(x)Px \wedge (x)Qx$	RP
	L2 (1)	$(x)Px$	RE $\wedge$; L1
	L3 (1)	$(x)Qx$	RE $\wedge$; L1
	L4 (1)	Px	RE () ; L2
	L5 (1)	Qx	RE () ; L3
	L6 (1)	$Px \wedge Qx$	RI $\wedge$; L4, L5
x	L7 (1)	$(x)[Px \wedge Qx]$	RI () ; L6
	L8 (0)	$(x)Px \wedge (x)Qx \rightarrow (x)[Px \wedge Qx]$	RI $\rightarrow$; L7
	L9 (9)	$(x)[Px \wedge Qx]$	RP
	L10 (9)	$Pz \wedge Qz$	RE () ; L9
	L11 (9)	Pz	RE $\wedge$; L10
z	L12 (9)	$(x)Px$	RI () ; L11
	L13 (9)	$Pu \wedge Qu$	RE () ; L9
	L14 (9)	Qu	RE $\wedge$; L13
u	L15 (9)	$(x)Qx$	RI () ; L14
	L16 (9)	$(x)Px \wedge (x)Qx$	RI $\wedge$; L12, L15
	L17 (0)	$(x)Px \wedge (x)Qx \leftrightarrow (x)[Px \wedge Qx]$	RI $\rightarrow$; L16
	L18 (0)	$(x)[Px \wedge Qx] \rightarrow (x)Px \wedge (x)Qx$	RI $\leftrightarrow$; L8, L17

En L13 hay que aplicar por segunda vez RE () a L9, con nuevas variables, para evitar una doble anotación de 'z'.

(TP23) $\exists x Px \vee \exists x Qx \leftrightarrow \exists x [Px \vee Qx]$ *Demostración*

	L1 (1)	$\exists x Px \vee \exists x Qx$	RP
	L2 (2)	$\exists x Px$	RP
	L3 (3)	$\exists x Qx$	RP
u	L4 (2)	Pu	RE $\exists$; L2
	L5 (2)	$Pu \vee Qu$	RI $\vee$; L4
	L6 (2)	$\exists x [Px \vee Qx]$	RI $\exists$; L5
	L7 (0)	$\exists x Px \rightarrow \exists x [Px \vee Qx]$	RI $\rightarrow$; L7
z	L8 (3)	Qz	RE $\exists$; L3
	L9 (3)	$Pz \vee Qz$	RI $\vee$; L8
	L10 (3)	$\exists x [Px \vee Qx]$	RI $\exists$; L9
	L11 (0)	$\exists x Qx \rightarrow \exists x [Px \vee Qx]$	RI $\rightarrow$; L10
	L12 (1)	$\exists x [Px \vee Qx]$	RE $\vee$; L1, L7, L11

L13 (0)	$\exists xPx \vee \exists xQx \rightarrow \exists x[Px \vee Qx]$	RI $\rightarrow$; L12
L14 (14)	$\exists x[Px \vee Qx]$	RP
x L15 (14)	$Px \vee Qx$	RE $\exists$; L14
L16 (16)	Px	RP
L17 (16)	$\exists xPx$	RI $\exists$; L16
L18 (16)	$\exists xPx \vee \exists xQx$	RI $\vee$; L17
L19 (19)	Qx	RP
L20 (19)	$\exists xQx$	RI $\exists$; L19
L21 (19)	$\exists xPx \vee \exists xQx$	RI $\vee$; L20
L22 (0)	$Px \rightarrow \exists xPx \vee \exists xQx$	RI $\rightarrow$; L18
L23 (0)	$Qx \rightarrow \exists xPx \vee \exists xQx$	RI $\rightarrow$; L21
L24 (14)	$\exists xPx \vee \exists xQx$	RE $\vee$; L15, L22, L23
L25 (0)	$\exists x[Px \vee Qx] \rightarrow \exists xPx \vee \exists xQx$	RI $\rightarrow$; L24
L26 (0)	$\exists xPx \vee \exists xQx \leftrightarrow \exists x[Px \vee Qx]$	RI $\leftrightarrow$; L13, L26

(TP22) y (TP23) son leyes de distribución de cuantificadores y sirven, como (TP18)-(TP21), para sacar cuantificadores del interior de fórmulas.

Es interesante observar que (TP22) y (TP23) ponen a ' $()$ ' y ' $\exists$ ' en relación con la conjunción y la disyunción, respectivamente. El generalizador puede, en efecto, entenderse como una conjunción total o universal. Sea, por ejemplo, un Ω finito, de tres elementos a, b, c . Dada una propiedad cualquiera, P , afirmar sobre ese universo del discurso que todo individuo de él es P ,

$$(x)Px,$$

es lo mismo que afirmar:

$$Pa \wedge Pb \wedge Pc.$$

Sobre ese mismo universo del discurso, afirmar que alguna cosa del mismo es P ,

$$\exists xPx,$$

es lo mismo que afirmar:

$$Pa \vee Pb \vee Pc.$$

El particularizador puede, pues, por su parte, entenderse como una disyunción total o universal, es decir, que abarca a todos los individuos del Ω dado.

$$(TP24) \quad (x)Px \leftrightarrow \sim \exists x \sim Px$$

Demostración

L1 (1)	$(x)Px$	RP
L2 (1)	Py	RE $()$; L1
L3 (3)	$\exists x \sim Px$	RP
y L4 (3)	$\sim Py$	RE $\exists$; L3
L5 (1, 3)	Py	RE $\sim$; L2, L4
L6 (1)	$\exists x \sim Px \rightarrow Py$	RI $\rightarrow$; L5
L7 (0)	$\exists x \sim Px \rightarrow \sim Py$	RI $\rightarrow$; L4
L8 (1)	$\sim \exists x \sim Px$	RI $\sim$; L6, L7
L9 (0)	$(x)Px \rightarrow \sim \exists x \sim Px$	RI $\rightarrow$; L8
L10 (10)	$\sim \exists x \sim Px$	RP
L11 (11)	$\sim Px$	RP
L12 (11)	$\exists x \sim Px$	RI $\exists$; L11
L13 (10, 11)	Px	RE $\sim$; L10, L12
L14 (10)	$\sim Px \rightarrow Px$	RI $\rightarrow$; L13
L15 (0)	$\sim Px \rightarrow \sim Px$	RI $\rightarrow$; L11
L16 (10)	$\sim \sim Px$	RI $\sim$; L14, L15
L17 (10)	Px	RA1, con (TE1); L16
x L18 (10)	$(x)Px$	RI $()$; L17
L19 (0)	$\sim \exists x \sim Px \rightarrow (x)Px$	RI $\rightarrow$; L18
L20 (0)	$(x)Px \leftrightarrow \sim \exists x \sim Px$	RI $\leftrightarrow$; L9, L19

Análogamente se demuestran:

$$(TP25) \quad (x) \sim Px \leftrightarrow \sim \exists xPx.$$

$$(TP26) \quad \sim (x)Px \leftrightarrow \exists x \sim Px.$$

$$(TP27) \quad \sim (x) \sim Px \leftrightarrow \exists xPx.$$

(TP24)-(TP27) muestran que es posible prescindir de uno u otro de los cuantificadores, pues basta uno (con la ayuda de la negación) para expresar ambas cuantificaciones.

De modo parecido, las leyes de De Morgan (TE12)-(TE13) muestran que es posible prescindir de ' $\wedge$ ' o de ' $\vee$ ', pues basta uno de esos dos funtores (con la ayuda de la negación) para expresar ambas funciones.

Los ejemplos siguientes se refieren al orden de ' $()$ ' y ' $\exists$ ' cuando aparecen en una misma fórmula.

$$(TP28) \quad \exists y(x)Pxy \rightarrow (x)\exists yPxy$$

Demostración

L1 (1)	$\exists y(x)Pxy$	RP.
u L2 (1)	$(x)Pxu$	RE $\exists$; L1.
L3 (1)	Pzu	RE $()$; L2.
L4 (1)	$\exists yPzy$	RI $\exists$; L3.
z L5 (1)	$(x)\exists yPxy$	RI $()$; L4.
L6 (0)	$\exists y(x)Pxy \rightarrow (x)\exists yPxy$	RI $\rightarrow$; L5.

No vale, en cambio, la implicación inversa,

$$(x)\exists yPxy \rightarrow \exists y(x)Pxy,$$

como puede verse por el intento de demostración:

	L1 (1)	$(x)\exists yPxy$	RP.
	L2 (1)	$\exists yPzy$	RE (); L1.
$u [z]$	L3 (1)	Pzu	RE $\exists$; L2.
$z [u]$	L4 (1)	$(x)Pxu$	RI (); L3.
	L5 (1)	$\exists y(x)Pxy$	RI $\exists$; L4.
	L6 (0)	$(x)\exists yPxy \rightarrow \exists y(x)Pxy$	RI $\rightarrow$; L5.

Esa sucesión de fórmulas no es una demostración, porque hay anotación de variables en dependencia recíproca (L3, L4). Interpretando 'P' por 'ser-menor-que-', con $\Omega =$ los números naturales, (TP28) sería:

si hay un número y tal que todo número x es menor que y , entonces para todo número x hay un número y tal que x es menor que y ;

o, dicho más condensadamente:

si hay un número mayor que todos los números, entonces todo número tiene uno mayor que él.

En cambio, con esa misma interpretación, la fórmula indemostrable dice:

si para todo número x hay otro, y , mayor que él, entonces hay un número, y , tal que, para todo número x , x es menor que y ;

o, en estilo más condensado:

si para todo número hay otro mayor que él, entonces hay un número mayor que todos los números.

Vamos a ver ahora una fórmula en la cual resulta, en cambio, lícito cambiar el orden de '()' y ' $\exists$ ' en el mismo sentido que en el anterior ejemplo hemos visto era ilícito:

$$(TP29) \quad (x)\exists y[Px \wedge Qyx] \rightarrow \exists y(x)[Px \vee Qyx]$$

Demostración

	L1 (1)	$(x)\exists y[Px \wedge Qyx]$	RP.
	L2 (1)	$\exists y[Px \wedge Qyx]$	RE (); L1.
$z [x]$	L3 (1)	$Px \wedge Qzx$	RE $\exists$; L2.
	L4 (1)	Px	RE $\wedge$; L3.
	L5 (1)	$Px \vee Qtx$	RI $\vee$; L4.
$x [t]$	L6 (1)	$(x)[Px \vee Qtx]$	RI (); L5.
	L7 (1)	$\exists y(x)[Px \vee Qyx]$	RI $\exists$; L6.
	L8 (0)	$(x)\exists y[Px \wedge Qyx] \rightarrow \exists y(x)[Px \vee Qyx]$	RI $\rightarrow$; L7.

Podría pensarse que (TP29) vale, a diferencia de la anterior fórmula, porque el consecuente, que tiene sólo disyunción, es más débil (afirma menos) que el antecedente, el cual lleva conjunción. Pero esta debilitación no basta siempre. Así, por ejemplo, no es válida la siguiente fórmula, a pesar de presentar la misma debilitación en el consecuente:

$$(x)\exists y[Pxy \wedge Qxy] \rightarrow \exists y(x)[Pxy \vee Qxy].$$

Que no es válida puede verse por el siguiente intento de demostración:

	L1 (1)	$(x)\exists y[Pxy \wedge Qxy]$	RP.
	L2 (1)	$\exists y[Pxy \wedge Qxy]$	RE (); L1.
$z [x]$	L3 (1)	$Pxz \wedge Qxz$	RE $\exists$; L2.
	L4 (1)	Pxz	RE $\wedge$; L3.
	L5 (1)	$Pxz \vee Qxz$	RI $\vee$; L4.
$x [z]$	L6 (1)	$(x)[Pxz \vee Qxz]$	RI (); L5.
	L7 (1)	$\exists y(x)[Pxy \vee Qxy]$	RI $\exists$; L6.
	L8 (0)	$(x)\exists y[Pxy \wedge Qxy] \rightarrow \exists y(x)[Pxy \vee Qxy]$	RI $\rightarrow$; L7.

Esa sucesión de fórmulas no es una demostración, pues hay anotación de variables en dependencia recíproca (L3, L6).

Consiguientemente, la demostrabilidad de (TP29) no se debe sólo a la debilitación de ' $\wedge$ ' en ' $\vee$ ', sino a alguna otra causa. Comparando las líneas L4 de ambas demostraciones, la correcta y la incorrecta, puede verse que esa otra causa es el hecho de contar, en la demostración de (TP29), con una fórmula predicativa monádica, lo que permite practicar un cambio de variable al debilitar ' $\wedge$ ' en ' $\vee$ ' en las líneas L5.

Este ejemplo nos indica, que, aun dentro del cálculo de predicados de primer orden, hay una diferencia importante entre su parte monádica y su parte poliádica. Esa diferencia nos ocupará más adelante. Por ahora bastará con observar que la relación anteriormente indicada entre '()' y ' $\wedge$ ', por una parte, y ' $\exists$ ' y ' $\vee$ ' por otra, en cuanto se refiere a operaciones sobre individuos, vale sólo, en la forma que vimos, para expresiones monádicas.

La expresión monádica

$$(x)Px$$

puede entenderse como una conjunción de afirmaciones individuales:

$$Px_1 \wedge Px_2 \wedge \dots$$

Pero la expresión

$$(x)(y)Pxy$$

no puede entenderse de ese modo. En este caso, si se la quiere interpretar como conjunción de expresiones sin cuantificar, las entidades tomadas como

individuos (los símbolos tomados como sujetos) no serán propiamente individuos, sino pares ordenados de ellos:

$$Px_1y_1 \wedge Px_1y_2 \wedge Px_1x_1 \wedge Py_1y_1 \wedge Px_2y_2 \wedge \dots$$

Tomando el mismo $\Omega = \{a, b, c\}$ de antes, ' $(x)(y) Pxy$ ' vale tanto como la conjunción:

$$Paa \wedge Pab \wedge Pac \wedge Pba \wedge Pca \wedge Pbb \wedge Pbc \wedge Pcb \wedge Pcc.$$

CAPÍTULO X

FORMAS NORMALES. COMPARACIÓN DEL SISTEMA AXIOMÁTICO CON EL CÁLCULO DE LA DEDUCCIÓN NATURAL

56. *Noción de forma normal.*—Varios de los teoremas (TE) y (TP) del capítulo anterior son equivalencias, y nos autorizan a utilizar una fórmula en vez de otra de la que es equivalente, de acuerdo con una regla auxiliar del tipo de las definiciones (una regla como RA1, que usaremos abreviadamente en lo que sigue, como, por lo demás, hemos hecho ya alguna vez). Por ejemplo: el teorema (TE12) nos permite usar ' $\sim p \vee \sim q$ ' en vez de ' $\sim [p \wedge q]$ '. Basándonos en este teorema, podríamos convenir en una reducción del número de símbolos elementales o primitivos del lenguaje de la lógica de enunciados, suprimiendo ' $\wedge$ '. Teniendo en cuenta (TE12) y la ley de doble negación (TE1), vale en efecto:

$$p \wedge q \leftrightarrow \sim [\sim p \vee \sim q].$$

Por convenciones como ésta puede normarse de un modo económico — es decir, con ahorro de símbolos primitivos — el lenguaje del cálculo considerado. Establecer un sistema de tales convenciones es adoptar una *forma normal* para las expresiones del cálculo.

El uso de formas normales es sobre todo útil para la teoría lógica, para los razonamientos metalógicos sobre un cálculo, porque disminuye el número de símbolos, fórmulas y reglas que estudiar. En cambio, no suele ser conveniente (salvo por una aplicación: cfr. cap. XIII, 73) para el manejo de los cálculos lógicos desde el punto de vista de su aplicación. Pues, por regla general, las formas normales, precisamente por la economía de léxico con que están escritas, son más largas y engorrosas que las fórmulas escritas con un léxico más abundante. Así, por ejemplo, ' $p \leftrightarrow q$ ' es más cómoda de manejar que su posible forma normal (cfr. 57) ' $[\sim p \vee q] \wedge [\sim q \vee p]$ '.

Las convenciones que dan como resultado formas normales en la lógica de predicados afectan, como es natural, tanto a los funtores veritativos cuanto a los cuantificadores. Empezaremos por considerar aquéllos.

57. Normalización de los funtores veritativos.— En las formas normales de uso más frecuente en lógica se prescinde de los funtores ' $\leftrightarrow$ ' y ' $\rightarrow$ '. De ' $\leftrightarrow$ ' puede prescindirse por el teorema

$$(TE30) \quad [p \leftrightarrow q] \leftrightarrow [p \rightarrow q] \wedge [q \rightarrow p].$$

Demostración

L1 (1)	$p \leftrightarrow q$	RP.
L2 (1)	$p \rightarrow q$	RE $\leftrightarrow$; L1.
L3 (1)	$q \rightarrow p$	RE $\leftrightarrow$; L1.
L4 (1)	$[p \rightarrow q] \wedge [q \rightarrow p]$	RI $\wedge$; L2, L3.
L5 (0)	$[p \leftrightarrow q] \rightarrow [p \rightarrow q] \wedge [q \rightarrow p]$	RI $\rightarrow$; L4.
L6 (6)	$[p \rightarrow q] \wedge [q \rightarrow p]$	RP.
L7 (6)	$p \rightarrow q$	RE $\wedge$; L6.
L8 (6)	$q \rightarrow p$	RE $\wedge$; L6.
L9 (6)	$p \leftrightarrow q$	RI $\leftrightarrow$; L7, L8.
L10 (0)	$[p \rightarrow q] \wedge [q \rightarrow p] \rightarrow [p \leftrightarrow q]$	RI $\rightarrow$; L9.
L11 (0)	$[p \leftrightarrow q] \leftrightarrow [p \rightarrow q] \wedge [q \rightarrow p].$	RI $\leftrightarrow$; L5, L10.

De ' $\rightarrow$ ' puede a su vez prescindirse por el teorema

$$(TE31) \quad [p \rightarrow q] \leftrightarrow \sim p \vee q$$

Demostración

L1 (1)	$p \rightarrow q$	RP.
L2 (2)	p	RP.
L3 (1, 2)	q	RE $\rightarrow$; L1, L2.
L4 (1, 2)	$\sim p \vee q$	RI $\vee$; L3.
L5 (1)	$p \rightarrow \sim p \vee q$	RI $\rightarrow$; L4.
L6 (6)	$\sim p$	RP.
L7 (6)	$\sim p \vee q$	RI $\vee$; L6.
L8 (0)	$\sim p \rightarrow \sim p \vee q$	RI $\rightarrow$; L7.
L9 (0)	$p \vee \sim p$	RTE.
L10 (1)	$\sim p \vee q$	RE $\vee$; L5, L8, L9.
L11 (0)	$[p \rightarrow q] \rightarrow \sim p \vee q$	RI $\rightarrow$; L10.
L12 (12)	$\sim p \vee q$	RP.
L13 (2, 6)	q	RE $\sim$; L2, L6.
L14 (6)	$p \rightarrow q$	RI $\rightarrow$; L13.
L15 (0)	$\sim p \rightarrow [p \rightarrow q]$	RI $\rightarrow$; L14.
L16 (16)	q	RP.
L17 (2, 16)	q	RA2; L16.
L18 (16)	$p \rightarrow q$	RI $\rightarrow$; L17.
L19 (0)	$q \rightarrow [p \rightarrow q]$	RI $\rightarrow$; L18.
L20 (12)	$p \rightarrow q$	RE $\vee$; L12, L15, L19.
L21 (0)	$\sim p \vee q \rightarrow [p \rightarrow q]$	RI $\rightarrow$; L20.
L22 (0)	$[p \rightarrow q] \leftrightarrow \sim p \vee q$	RI $\leftrightarrow$; L11, L21.

Los teoremas (TE30) y (TE31) nos autorizan a normalizar la escritura de fórmulas del cálculo de la deducción natural (designado en adelante por 'CDN') reduciendo los funtores veritativos a ' $\sim$ ', ' $\vee$ ', ' $\wedge$ '. (Obsérvese que lo mismo permiten en el sistema axiomático HB las definiciones (D2) y (D3).)

Otra convención corriente para el establecimiento de formas normales es la reducción a uno del número de negaciones que afectan a una fórmula. Basándose en la ley de doble negación, (TE1), se puede escribir, en vez de ' $\sim \sim p$ ', ' p ', en vez de ' $\sim \sim \sim p$ ', ' $\sim p$ ', en vez de ' $\sim \sim \sim \sim p$ ', ' p ', etc.

En tercer lugar, las leyes de De Morgan, (TE12) y (TE13), permiten conseguir que la negación afecte sólo a fórmulas atómicas. Por ejemplo: si se tiene ' $\sim [p \wedge q]$ ', puede escribirse en vez de ella, gracias a (TE12), ' $\sim p \vee \sim q$ ' (como veíamos antes), fórmula en la cual ' $\sim$ ' afecta sólo a fórmulas atómicas.

Por último, las leyes de conmutatividad, asociatividad y distributividad de ' $\wedge$ ' y ' $\vee$ ', (TE6)-(TE11), permiten transformar toda fórmula ya normada en cuanto a los funtores—es decir: sin más funtores que ' $\sim$ ', ' $\vee$ ' y ' $\wedge$ ', y sin fórmulas negadas que no sean atómicas, ni más de un signo de negación por fórmula negada—, para conseguir cualquiera de las dos formas normales de uso corriente en lógica de enunciados: la forma normal conjuntiva y la forma normal disyuntiva.

(En las anteriores reflexiones, palabras como 'permite', 'autoriza', encubren un uso no calculístico, sino propiamente metacalculístico, metalógico, de la operación de la regla RA1. En rigor habría que decir, por ejemplo, a propósito del abandono de ' $\sim [p \wedge q]$ ' por ' $\sim p \vee \sim q$ ': 'el teorema (TE12) garantiza que si hay una demostración de ' $\sim [p \wedge q]$ ', entonces hay también una demostración de ' $\sim p \vee \sim q$ ', y a la inversa', expresión cuyo carácter metalógico se aprecia directamente. Así pues, nuestro modo de expresarnos en este contexto es propio de una metalógica no formal, sino intuitiva (cfr. 19).)

Una fórmula está en *forma normal conjuntiva* cuando es una conjunción no negada de disyunciones no negadas de fórmulas atómicas negadas o no negadas. Cualquier fórmula de la lógica de enunciados puede llevarse a forma normal conjuntiva utilizando los medios antes dichos: (TE30)-(TE31), (TE1), (TE12)-(TE13), (TE6)-(TE11). Veamos como ejemplo el paso de la fórmula ' $p \vee q \rightarrow r$ ' a forma normal conjuntiva:

Fórmula inicial: $p \vee q \rightarrow r$

$$\text{Paso 1.º:} \quad \sim [p \vee q] \vee r \quad (\text{TE31})$$

$$\text{Paso 2.º:} \quad [\sim p \wedge \sim q] \vee r \quad (\text{TE13})$$

$$\text{Paso 3.º:} \quad [\sim p \vee r] \wedge [\sim q \vee r] \quad (\text{TE10}).$$

El siguiente ejemplo volverá a interesarnos en el capítulo XIII (72). Se trata de llevar a forma normal conjuntiva la fórmula

$$[p \rightarrow q] \rightarrow [[q \rightarrow r] \rightarrow [p \rightarrow r]]$$

- Paso 1.º: $\sim [p \rightarrow q] \vee [\sim [q \rightarrow r] \vee [p \rightarrow r]]$ (TE31)
 Paso 2.º: $\sim [p \rightarrow q] \vee \sim [q \rightarrow r] \vee [p \rightarrow r]$ (TE9)
 Paso 3.º: $\sim [\sim p \vee q] \vee \sim [\sim q \vee r] \vee [\sim p \vee r]$ (TE31)
 Paso 4.º: $[p \wedge \sim q] \vee [q \wedge \sim r] \vee [\sim p \vee r]$ (TE13)
 Paso 5.º: $[\sim p \vee r] \vee [p \wedge \sim q] \vee [q \wedge \sim r]$ (TE6)
 Paso 6.º: $[[\sim p \vee r] \vee p \vee q] \wedge [[\sim p \vee r] \vee p \vee \sim r] \wedge [[\sim p \vee r] \vee \sim q \vee q] \wedge [[\sim p \vee r] \vee \sim q \vee \sim r]$ (TE10)
 Paso 7.º: $[\sim p \vee r \vee p \vee q] \wedge [\sim p \vee r \vee p \vee \sim r] \wedge [\sim p \vee r \vee \sim q \vee q] \wedge [\sim p \vee r \vee \sim q \vee \sim r]$ (TE9)

Corrientemente no será necesaria la prolijidad de este ejemplo: los pasos 1.º, 5.º y 6.º podían saltarse, pasando directamente de la fórmula inicial al paso 2.º, y del paso 4.º al paso 7.º.

Cuando una fórmula está en forma normal conjuntiva y, además, presenta en cada disyunción todas las fórmulas atómicas componentes (negadas o no negadas), se dice que está en *forma normal conjuntiva de Schröder*. Esta forma es fácil de conseguir: si se encuentra una disyunción en la que falte alguna componente atómica de la fórmula inicial, por ejemplo 's', se añade a esa disyunción, también en disyunción, la fórmula 's $\wedge$ $\sim$ s'. Esta fórmula es formalmente falsa y, por tanto, no altera el valor de la disyunción a la cual se añade en disyunción: si era verdadera, seguirá siéndolo, y si era falsa seguirá siéndolo. Luego se distribuye 's $\wedge$ $\sim$ s' por la disyunción a la que se ha añadido. — La anterior forma conjuntiva normal de 'p $\vee$ q $\rightarrow$ r',

$$[\sim p \vee r] \wedge [\sim q \vee r],$$

no es una forma normal conjuntiva de Schröder para aquella fórmula inicial; en su primera disyunción falta 'q' o ' $\sim$ q'; en su segunda falta 'p' o ' $\sim$ p'. Y tanto 'p' cuanto 'q' son componentes atómicas de la fórmula inicial. He aquí el paso de 'p $\vee$ q $\rightarrow$ r' a la forma normal conjuntiva de Schröder mediante el expediente indicado:

Paso 1.º — Paso 3.º como antes. El Paso 3.º daba:

- Paso 3.º: $[\sim p \vee r] \wedge [\sim q \vee r]$
 Paso 4.º: $[\sim p \vee r \vee [q \wedge \sim q]] \wedge [\sim q \vee r \vee [p \wedge \sim p]]$
 Paso 5.º: $[[\sim p \vee r \vee q] \wedge [\sim p \vee r \vee \sim q]] \wedge [[\sim q \vee r \vee p] \wedge [\sim q \vee r \vee \sim p]]$ (TE10)
 Paso 6.º: $[\sim p \vee r \vee q] \wedge [\sim p \vee r \vee \sim q] \wedge [\sim q \vee r \vee p] \wedge [\sim q \vee r \vee \sim p]$ (TE8)

La fórmula del paso 6.º está en forma normal conjuntiva de Schröder.

Una fórmula está en *forma normal disyuntiva* cuando es una disyunción no negada de conjunciones no negadas de fórmulas atómicas negadas o no negadas. Una forma normal disyuntiva de una fórmula puede conseguirse redistribuyendo una forma normal conjuntiva correspondiente a

la misma fórmula. Por ejemplo, dada la fórmula 'p $\vee$ q $\rightarrow$ r', se la puede llevar a una forma normal conjuntiva, como hemos hecho antes, y luego volver a distribuir:

el Paso 3.º daba: $[\sim p \vee r] \wedge [\sim q \vee r]$

- Paso 4.º: $[[\sim p \vee r] \wedge \sim q] \vee [[\sim p \vee r] \wedge r]$ (TE11)
 Paso 5.º: $[[\sim p \wedge \sim q] \vee [r \wedge \sim q]] \vee [[\sim p \wedge r] \vee [r \wedge r]]$ (TE11)
 Paso 6.º: $[\sim p \wedge \sim q] \vee [r \wedge \sim q] \vee [\sim p \wedge r] \vee [r \wedge r]$ (TE9)

La fórmula del Paso 6.º está en forma normal disyuntiva (salvo en que habitualmente se pone en vez de una fórmula 'r $\wedge$ r' su equivalente 'r').

58. Normalización de cuantificadores. — Los teoremas (TP24)-(TP27), sugieren una normalización de la escritura de las fórmulas con cuantificadores: no usar más que un cuantificador. Pero la normalización más práctica y más frecuentemente utilizada consiste en conservar los dos cuantificadores, aunque sirviéndose al mismo tiempo de esos teoremas (TP24)-(TP27) y de los teoremas (TP18)-(TP23) para conseguir dos cosas:

a) que los cuantificadores de una fórmula estén todos delante de la misma (con lo que se quiere decir: que a la izquierda de un cuantificador no haya ningún símbolo de la fórmula, o haya otro cuantificador de la fórmula), formando un prefijo cuantificacional;

b) que ninguno de los cuantificadores esté negado.

De una fórmula con esas características se dice que está en *forma normal prenexa*.

Ejemplo. — Sea la fórmula.

$$\sim (x) \exists y \sim [Px \wedge Qxy \rightarrow \exists z Rxyz],$$

que no está en forma normal prenexa, pues tiene un cuantificador — 'Ez' — en el seno de la fórmula y otro — '(x)' — negado. Los teoremas (TP18)-(TP27) sugieren que una forma normal prenexa de esa fórmula es

$$\exists x (y) \exists z [Px \wedge Qxy \rightarrow Rxyz].$$

La sugerencia se confirma con la siguiente demostración de la equivalencia

$$\sim (x) \exists y \sim [Px \wedge Qxy \rightarrow \exists z Rxyz] \leftrightarrow \exists x (y) \exists z [Px \wedge Qxy \rightarrow Rxyz].$$

Demostración

	L1 (1)	$\sim (x) \exists y \sim [Px \wedge Qxy \rightarrow \exists z Rxyz]$	RP.
	L2 (1)	$\exists x \sim \exists y \sim [Px \wedge Qxy \rightarrow \exists z Rxyz]$	(TP26); L1.
	L3 (1)	$\exists x (y) [Px \wedge Qxy \rightarrow \exists z Rxyz]$	(TP24); L2.
u	L4 (1)	$(y) [Pu \wedge Quy \rightarrow \exists z Ruyz]$	RE $\exists$; L3.
	L5 (1)	$Pu \wedge Quv \rightarrow \exists z Ruvz$	RE (); L4.
	L6 (1)	$\exists z [Pu \wedge Quv \rightarrow Ruvz]$	(TP21); L5.

$v[u]$	L7 (1)	$(y) \exists z [Pu \wedge Quy \rightarrow Ruyz]$	RI (); L6.
	L8 (1)	$\exists x (y) \exists z [Px \wedge Qxy \rightarrow Rxyz]$	RI $\exists$; L7.
	L9 (0)	$\sim (x) \exists y \sim [Px \wedge Qxy \rightarrow \exists z Rxyz] \rightarrow$ $\rightarrow \exists x (y) \exists z [Px \wedge Qxy \rightarrow Rxyz]$	RI $\rightarrow$; L8.
	L10 (10)	$\exists x (y) \exists z [Px \wedge Qxy \rightarrow Rxyz]$	RP.
x	L11 (10)	$(y) \exists z [Px \wedge Qxy \rightarrow Rxyz]$	RE $\exists$; L10.
	L12 (10)	$\exists z [Px \wedge Qxy \rightarrow Rxyz]$	RE (); L11.
	L13 (10)	$Px \wedge Qxy \rightarrow \exists z Rxyz$	(TP21); L12
$y[x]$	L14 (10)	$(y) [Px \wedge Qxy \rightarrow \exists z Rxyz]$	RI (); L13.
	L15 (10)	$\sim \exists y \sim [Px \wedge Qxy \rightarrow \exists z Rxyz]$	(TP24); L14.
	L16 (10)	$\exists x \sim \exists y \sim [Px \wedge Qxy \rightarrow \exists z Rxyz]$	RI $\exists$; L15.
	L17 (10)	$\sim (x) \exists y \sim [Px \wedge Qxy \rightarrow \exists z Rxyz]$	(TP26); L16.
	L18 (0)	$\exists x (y) \exists z [Px \wedge Qxy \rightarrow Rxyz] \rightarrow$ $\rightarrow \sim (x) \exists y \sim [Px \wedge Qxy \rightarrow \exists z Rxyz]$	RI $\rightarrow$; L17.
	L19 (0)	$\sim (x) \exists y \sim [Px \wedge Qxy \rightarrow \exists z Rxyz] \leftrightarrow$ $\leftrightarrow \exists x (y) \exists z [Px \wedge Qxy \rightarrow Rxyz]$	RI $\leftrightarrow$; L9, L18.

La normalización respecto de los funtores veritativos debe, naturalmente, realizarse también. Ella nos lleva por último, de la fórmula

$$\sim (x) \exists y \sim [Px \wedge Qxy \rightarrow \exists z Rxyz],$$

pasando, como hemos visto, por

$$\exists x (y) \exists z [Px \wedge Qxy \rightarrow Rxyz],$$

a la forma normal (disyuntiva y prenexa)

$$\exists x (y) \exists z [\sim Px \vee \sim Qxy \vee Rxyz].$$

59. Justificación de las formas normales en HB. Comparación de HB con CDN.—Todos los teoremas de CDN en que se basa la reducción a formas normales son también fórmulas válidas de HB. Por ejemplo, el teorema (TE6) se obtiene en HB directamente del axioma A2. Los teoremas (TE30) y (TE31) están suplidos en HB por las definiciones (D3) y (D2).

Más interesante que repasar uno por uno esos teoremas de CDN para comprobar su validez en HB será aprovechar esta ocasión para llevar a cabo una comparación general de ambos sistemas, como fruto de la cual obtendremos la convicción de que uno y otro “rinden” lo mismo. Esa convicción es deseable porque los dos sistemas nos serán útiles en lo que sigue: HB para consideraciones teóricas (metalógicas), que son mucho más cómodas de hacer, con pocos medios metalógicos (sintácticos y semánticos), sobre un sistema axiomático que sobre un sistema de deducción natural, a causa de que el primero tiene menos reglas de transformación cuyo funcionamiento (más sencillo, por otra parte) haya que estudiar; también nos será útil HB para toda aplicación de la lógica a la axiomatización de teorías. Pero CDN es, por su parte, más adecuado para el trabajo práctico

analítico (en lógica pura o en aplicación a otras teorías), esto es, cuando se trata de ilustrar formalmente puntos de tal o cual teoría positiva: pues CDN no exige remontarse a lejanos principios, sino que permite empezar el análisis en el lugar o fragmento mismo que interese, con las premisas más inmediatas, y de un modo “natural”. Ahora bien: sólo si queda claro que ambos sistemas “rinden” lo mismo, que se “equivalen”, podemos tener la seguridad de que lo que las consideraciones teóricas (metalógicas) indiquen para uno vale también para el otro.

La igualdad de rendimiento de HB y CDN no debe hacer olvidar la diferencia esencial entre ambos sistemas: CDN es un cálculo para demostrar bajo premisas; HB puede servir para demostrar a partir de premisas, pero entonces las trata como axiomas o, en general, teoremas (incluyendo, como hemos hecho, a los axiomas entre los teoremas). Esto tiene una consecuencia importante: será imposible encontrar en HB nada equivalente a la regla de eliminación de premisa (RI $\rightarrow$) de CDN. Pues esta regla permite precisamente el tratamiento de la premisa como tal premisa: como algo que hay que eliminar para que lo demostrado valga universalmente, sea un teorema de la lógica. Y eso es lo que no se puede hacer en un sistema que, como HB, trata la premisa como si fuera un teorema. Por eso, para comprobar que, de todos modos, HB rinde lo mismo que CDN, nos será necesaria, a propósito de la eliminación de premisas, una reflexión metalógica más detallada que para las demás reglas (cfr. 60).

Para persuadirnos de la igualdad de rendimiento entre HB y CDN procederemos así: primero comprobaremos que todo teorema de HB es también un teorema de CDN. Luego, que todo teorema de CDN es también un teorema de HB.

Todo teorema de HB es un teorema de CDN.

a) Los axiomas de HB, base de toda otra operación y todo otro teorema en HB, son teoremas de CDN. Las demostraciones son muy sencillas. La del más largo, A4, es:

L1 (1)	$p \rightarrow q$	RP.
L2 (2)	$s \vee p$	RP.
L3 (3)	s	RP.
L4 (3)	$s \vee q$	RI $\vee$; L3.
L5 (0)	$s \rightarrow s \vee q$	RI $\rightarrow$; L4.
L6 (6)	p	RP.
L7 (1, 6)	q	RE $\rightarrow$; L1, L6.
L8 (1, 6)	$s \vee q$	RI $\vee$; L7.
L9 (1)	$p \rightarrow s \vee q$	RI $\rightarrow$; L8.
L10 (1, 2)	$s \vee q$	RE $\vee$; L2, L5, L9.
L11 (1)	$s \vee p \rightarrow s \vee q$	RI $\rightarrow$; L10.
L12 (0)	$[p \rightarrow q] \rightarrow [s \vee p \rightarrow s \vee q]$	RI $\rightarrow$; L11.

Las demostraciones de los demás axiomas de HB en CDN son aún más cortas.

b) En cuanto a las reglas de HB, una, $R\beta$, se encuentra en la misma forma en CDN (es la regla $RE \rightarrow$). Las reglas de sustitución ($R\alpha_1$, $R\alpha_2$, $R\alpha_3$, $R\delta$) de HB no permiten, en última instancia, sino escribir los axiomas de HB con otras letras. Puesto que los axiomas de HB son teoremas de CDN, lo son con cualesquiera letras. Por ejemplo, $A4$ podía haberse demostrado con ' r ' en vez de ' p '. Consiguientemente, también los resultados alcanzados con esas reglas en HB se consiguen en CDN.

Dos observaciones para resolver posibles dudas que puede haber suscitado esa argumentación expeditivamente intuitiva:

Primera. — Hemos dicho que las reglas de sustitución de HB no permiten, en última instancia, más que escribir los axiomas de HB con otras letras. Esta es una forma laxa de hablar. Los axiomas de HB — si son axiomas, y no esquemas axiomáticos (cfr. 42) — son las fórmulas que son, con las letras que tienen. Habría pues que haber dicho: las reglas de sustitución de HB no permiten sino obtener de los axiomas teoremas que también se obtienen en CDN.

Segunda. — Hemos dicho que eso es así 'en última instancia'. En "primera instancia" puede parecer que no es así: pues en HB puede aplicarse una regla de sustitución a cualquier fórmula ya escrita en la misma demostración. Pero es que en HB cualquier fórmula ya escrita en una demostración se trata como si fuera un teorema, y se reduce, por tanto, a los axiomas utilizados para afirmarla (o a la premisa utilizada como axioma para afirmarla). Por eso, dicho sea de paso, se puede demostrar directamente que las operaciones $R\alpha_1$, $R\alpha_2$, $R\alpha_3$ y $R\delta$ son realizables también en CDN:

I) $R\alpha_1$ lo es porque si ' $p \rightarrow q$ ', por ejemplo, es un teorema (está escrita con 0 premisa) en una demostración en CDN, también se conseguirá en una demostración en CDN ' $r \rightarrow s$ ', por ejemplo, lo que equivale a una aplicación de $R\alpha_1$ en HB, sobre ' $p \rightarrow q$ ', con p/r , q/s .

II) $R\alpha_2$ lo es porque es correcto en CDN un paso como el siguiente, que equivale a una aplicación en HB de $R\alpha_2$ con x/y :

	L1 (0)	Px	Teorema
x	L2 (0)	$(x) Px$	RI (); L1.
	L3 (0)	Py	RE (); L2.

('x' no está libre en L3, porque ' Px ' es un teorema si es que, como suponemos, es una fórmula escrita en una demostración de HB).

III) $R\alpha_3$, por la misma razón que $R\alpha_1$.

IV) $R\delta$, por la misma razón que $R\alpha_2$, como muestra la demostración

	L1 (0)	$(x) Px$	Teorema
	L2 (0)	Py	RE (); L1.
y	L3 (0)	$(y) Py$	RI (); L2.

Las reglas de HB sobre la cuantificación dan de sí resultados también alcanzables en CDN.

La regla $R\gamma_1$ de HB permite pasar, por ejemplo, de un teorema

$$p \rightarrow Qx$$

en el que ' p ' no contiene a ' x ', a un teorema

$$p \rightarrow (x) Qx.$$

Este paso se justifica en CDN por la siguiente demostración:

	L1 (0)	$p \rightarrow Qx$	Teorema.
	L2 (2)	p	RP.
	L3 (2)	Qx	RE $\rightarrow$; L1, L2.
x	L4 (2)	$(x) Qx$	RI (); L3.
	L5 (0)	$p \rightarrow (x) Qx$	RI $\rightarrow$; L4.

('x' no está libre en L5).

La regla $R\gamma_2$ de HB permite pasar, por ejemplo, de un teorema

$$Qx \rightarrow p,$$

en el que ' p ' no contiene a ' x ', a un teorema

$$\exists x Qx \rightarrow p.$$

Este paso se justifica en CDN por la siguiente demostración:

	L1 (0)	$Qx \rightarrow p$	Teorema.
	L2 (2)	$\exists x Qx$	RP.
x	L3 (2)	Qx	RE $\exists$; L2.
	L4 (2)	p	RE $\rightarrow$; L1, L3.
	L5 (0)	$\exists x Qx \rightarrow p$	RI $\rightarrow$; L4.

('x' no está libre en L5).

Las anteriores consideraciones muestran que todo teorema de HB es también un teorema de CDN: pues todos los axiomas de HB son teoremas de CDN, y todo lo que puede obtenerse en HB de esos axiomas mediante la aplicación de las reglas de HB puede también obtenerse en CDN.

Todo teorema de CDN es un teorema de HB

Como queda dicho, la relación que aquí nos interesa comprobar entre HB y CDN no quedará establecida hasta el apartado siguiente, 60, pues ya hemos visto que en el sistema HB es imposible que exista una operación equivalente a la de $RI \rightarrow$ en CDN.

Dejando esa cuestión a un lado hasta 60, la parte veritativo-funcional

de CDN no presenta más problema que la justificación de las reglas sobre ' $\sim$ ' y ' $\wedge$ '. La dificultad es, en realidad, debida a la notación de los axiomas y de las reglas. Pero por lo que hace a las reglas de CDN sobre los demás funtores veritativos, su justificación en HB no tiene ni esa dificultad notacional. Por ejemplo, la regla RI $\vee$ de CDN, que permite pasar de una fórmula ' p ' a una fórmula ' $p \vee q$ ', se justifica en HB por la demostración a partir de ' p ':

L1	p	Premisa.
L2	$p \rightarrow p \vee q$	A2
L3	$p \vee q$	R β ; L1, L2.

Atendamos pues a las reglas de CDN sobre ' $\sim$ ' y sobre ' $\wedge$ '.

RI $\sim$ de CDN permite afirmar, por ejemplo, ' $\sim p$ ' bajo las premisas ' $p \rightarrow q$ ' y ' $p \rightarrow \sim q$ '. La siguiente demostración a partir de esas premisas en HB muestra que la operación es también lícita en este sistema:

L1	$p \rightarrow q$	Premisa 1. ^a
L2	$p \rightarrow \sim q$	Premisa 2. ^a
L3	$[p \rightarrow q] \rightarrow [s \vee p \rightarrow s \vee q]$	A4.
L4	$s \vee p \rightarrow s \vee q$	R β ; L1, L3.
L5	$\sim p \vee \sim q \rightarrow \sim p \vee q$	R α_1 : $s/\sim p$, $p/\sim q$; L4.
L6	$[p \rightarrow \sim q] \rightarrow \sim p \vee q$	(D2); L5.
L7	$\sim p \vee q$	R β ; L2, L6.
L8	$\sim p \vee \sim p$	R α_1 : $q/\sim p$; L7.
L9	$p \vee p \rightarrow p$	A1.
L10	$\sim p \vee \sim p \rightarrow \sim p$	R α_1 : $p/\sim p$; L9.
L11	$\sim p$	R β ; L8, L10.

La regla RE $\sim$ de CDN permite, por ejemplo, pasar de las premisas ' p ' y ' $\sim p$ ' a una afirmación cualquiera, ' q '. La justificación de esta operación en HB es trivial (pues en realidad la regla R α_1 de HB da más que RE $\sim$ de CDN sobre la negación). R α_1 permite, en efecto, pasar a afirmar ' q ' con sólo la premisa ' p ', sin necesidad de usar la premisa ' $\sim p$ ' además:

L1	p	Premisa
L2	q	R α_1 : p/q

Esta trivialidad de la operación correspondiente a RE $\sim$ en HB se debe a que la afirmación de una fórmula atómica es ya una inconsistencia. En CDN la afirmación de una fórmula atómica es imposible, pues una fórmula atómica sólo puede ponerse como premisa, no afirmarse.

El que R α_1 de HB de sobre la negación más que RE $\sim$ de CDN no debe preocupar: ya hemos visto antes que, en conjunto, CDN rinde todo lo que rinde HB (que todo teorema de HB es teorema de CDN), o sea, que HB no es más potente que CDN.

Las reglas RI $\wedge$ y RE $\wedge$ de CDN son de justificación técnicamente más laboriosa en HB (aunque sólo por motivos de notación). Para hacer más cómoda la justificación de RI $\wedge$ en HB empezaremos por procurarnos en HB una regla auxiliar (transitividad de ' $\rightarrow$ ') a la que llamaremos 'RA β_1 ':

$$\text{RA}\beta_1: \frac{p \rightarrow q, q \rightarrow r}{p \rightarrow r}$$

La justificaremos mediante la siguiente demostración (que es, en realidad, una argumentación metalógica, como toda justificación de regla auxiliar, en la que en vez de expresiones del lenguaje objeto, como ' p ' o ' q ', habría que usar variables sintácticas, como ' X ' o ' Y ', y en vez de funtores del lenguaje objeto, como ' $\vee$ ', nombres suyos, como ' v '. Pero esta libertad, que suele tomarse en casos como éste, no tiene ninguna mala consecuencia porque no se producen formaciones reflexivas (cfr. 19)):

L1	$p \rightarrow q$	Premisa 1. ^a
L2	$q \rightarrow r$	Premisa 2. ^a
L3	$[p \rightarrow q] \rightarrow [s \vee p \rightarrow s \vee q]$	A4.
L4	$[p \rightarrow q] \rightarrow [\sim s \vee p \rightarrow \sim s \vee q]$	R α_1 : $s/\sim s$; L3.
L5	$[p \rightarrow q] \rightarrow [[s \rightarrow p] \rightarrow [s \rightarrow q]]$	(D2); L4.
L6	$[q \rightarrow r] \rightarrow [[p \rightarrow q] \rightarrow [p \rightarrow r]]$	R α_1 : p/q , q/r , s/p ; L5.
L7	$[p \rightarrow q] \rightarrow [p \rightarrow r]$	R β ; L2, L6.
L8	$p \rightarrow r$	R β ; L1, L7.

Contando con RA β_1 podemos justificar cómodamente RI $\wedge$ en HB. RI $\wedge$ permite en CDN, por ejemplo, pasar de las premisas ' p ' y ' q ' a la fórmula ' $p \wedge q$ '. La operación es lícita en HB, como muestra la siguiente demostración (RA β_1 se aplica en la línea L12):

L1	p	Premisa 1. ^a
L2	q	Premisa 2. ^a
L3	$[p \rightarrow q] \rightarrow [s \vee p \rightarrow s \vee q]$	A4.
L4	$[\sim p \vee q] \rightarrow [s \vee p \rightarrow s \vee q]$	(D2); L3.
L5	$p \rightarrow p \vee q$	A2.
L6	$p \vee q$	R β ; L1, L5.
L7	$\sim p \vee q$	R α_1 : $p/\sim p$; L6.
L8	$s \vee p \rightarrow s \vee q$	R β ; L4, L7.
L9	$s \vee p \rightarrow s \vee s$	R α_1 ; q/s .
L10	$p \vee p \rightarrow p$	A1.
L11	$s \vee s \rightarrow s$	R α_1 ; p/s ; L10.
L12	$s \vee p \rightarrow s$	RA β_1 ; L9, L11.
L13	$\sim [s \vee p] \vee s$	(D2); L12.
L14	$p \vee q \rightarrow q \vee p$	A3.
L15	$\sim [s \vee p] \vee s \rightarrow s \vee \sim [s \vee p]$	R α_1 : $p/\sim [s \vee p]$, q/s ; L14.

L16	$s \vee \sim [s \vee p]$	R β ; L13, L15.
L17	$\sim q \vee \sim [\sim q \vee \sim p]$	R α_1 : $s/\sim q$, $p/\sim p$; L16.
L18	$q \rightarrow \sim [\sim q \vee \sim p]$	(D2); L17.
L19	$\sim [\sim q \vee \sim p]$	R β ; L2, L18.
L20	$q \wedge p$	(D1); L19.

La justificación de la operación RIA en HB nos suministra una nueva regla auxiliar para HB, a la que llamaremos RAA₁:

$$\text{RAA}_1 \quad \frac{p, q}{p \wedge q}$$

Esta regla, junto con la definición (D3), justifica, por otra parte, cómodamente la operación RI $\leftrightarrow$ de CDN en HB:

L1	$p \rightarrow q$	Premisa 1. ^a
L2	$q \rightarrow p$	Premisa 2. ^a
L3	$[p \rightarrow q] \wedge [q \rightarrow p]$	RAA ₁ ; L1, L2.
L4	$p \leftrightarrow q$	(D3); L3.

Llamaremos RA $\leftrightarrow_1$ a la regla auxiliar así justificada para HB:

$$\text{RA } \leftrightarrow_1 \quad \frac{p \rightarrow q, q \rightarrow p}{p \leftrightarrow q}$$

Por lo que hace a RE $\wedge$ de CDN, facilitaremos su justificación en HB procurándonos antes un teorema auxiliar, la ley de doble negación, ' $p \leftrightarrow \sim \sim p$ ', que ya conocemos por CDN, donde es el teorema (TE1), pero que no hemos demostrado aún en HB (y mientras no tengamos comprobada la igualdad de rendimiento de CDN y HB no podemos tomar para HB teoremas demostrados sólo en CDN). La siguiente demostración de ' $p \leftrightarrow \sim \sim p$ ' en HB utiliza las reglas auxiliares recién justificadas:

L1	$p \rightarrow p \vee q$	A2.
L2	$p \rightarrow p \vee p$	R α_1 ; q/p ; L1.
L3	$p \vee p \rightarrow p$	A1.
L4	$p \rightarrow p$	RA β_1 ; L2, L3.
L5	$\sim p \vee p$	(D2); L4.
L6	$p \vee q \rightarrow q \vee p$	A3.
L7	$\sim p \vee p \rightarrow p \vee \sim p$	R α_1 : $p/\sim p$, q/p ; L6.
L8	$p \vee \sim p$	R β ; L5, L7.
L9	$\sim p \vee \sim \sim p$	R α_1 : $p/\sim p$; L8.
L10	$p \rightarrow \sim \sim p$	(D2); L9.
L11	$\sim p \rightarrow \sim \sim \sim p$	R α_1 : $p/\sim p$; L10.
L12	$[p \rightarrow q] \rightarrow [s \vee p \rightarrow s \vee q]$	A4.

L13	$[\sim p \rightarrow \sim \sim \sim p] \rightarrow [s \vee \sim p \rightarrow s \vee \sim \sim \sim p]$	R α_1 : $p/\sim p$, $q/\sim \sim \sim p$; L12.
L14	$s \vee \sim p \rightarrow s \vee \sim \sim \sim p$	R β ; L11, L13.
L15	$p \vee \sim p \rightarrow p \vee \sim \sim \sim p$	R α_1 : s/p ; L14.
L16	$p \vee \sim \sim \sim p$	R β ; L8, L15.
L17	$p \vee \sim \sim \sim p \rightarrow \sim \sim \sim p \vee p$	R α_1 : $q/\sim \sim \sim p$; L6.
L18	$\sim \sim \sim p \vee p$	R β ; L16, L17.
L19	$\sim \sim p \rightarrow p$	(D2); L18.
L20	$p \leftrightarrow \sim \sim p$	RA $\leftrightarrow_1$; L10, L19.

Con el auxilio de la ley de doble negación puede justificarse cómodamente RE $\wedge$ de CDN en HB. Se trata de justificar la afirmación ' p ', por ejemplo, partiendo de ' $p \wedge q$ ':

L1	$p \wedge q$	Premisa.
L2	$\sim [\sim p \vee \sim q]$	(D1); L1.
L3	$p \rightarrow p \vee q$	A2.
L4	$\sim [\sim p \vee \sim q] \rightarrow \sim [\sim p \vee \sim q] \vee p$	R α_1 : $p/\sim [\sim p \vee \sim q]$; q/p ; L3
L5	$\sim [\sim p \vee \sim q] \vee p$	R β ; L2, L4.
L6	$\sim p \vee \sim q \rightarrow p$	(D2); L5.
L7	$\sim p \rightarrow \sim p \vee \sim q$	R α_1 : $p/\sim p$, $q/\sim q$; L3.
L8	$\sim p \rightarrow p$	RA β_1 ; L7, L6.
L9	$\sim \sim p \vee p$	(D2); L8.
L10	$p \vee p$	Ley doble negación.
L11	$p \vee p \rightarrow p$	A1.
L12	p	R β ; L10, L11.

La justificación de la operación RE $\wedge$ en HB nos suministra otra regla auxiliar, a la que llamaremos RA $\wedge_2$:

$$\text{RA } \wedge_2 \quad \frac{p \wedge q}{p | q}$$

Con ella, naturalmente, puede explicitarse la justificación de RE $\leftrightarrow$ de CDN en HB (con la ayuda de (D3) de HB), como antes la de RI $\leftrightarrow$. Pero, como dijimos, se trata de operaciones que no presentan ninguna dificultad ni siquiera notacional.

Por lo que hace a las reglas de CDN sobre cuantificadores, las dos reglas no críticas, RE () y RI $\exists$, son de justificación inmediata por los axiomas A5 y A6 respectivamente.

Las reglas RI () y RE $\exists$ de CDN parecen a primera vista más fuertes que las reglas cuantificacionales de HB. Pero éste no es el caso, como se aprecia si se atiende a todas las prevenciones que acompañan al uso de esas reglas en las demostraciones en CDN. Para hacer más sencilla y rápida la comparación, debe recordarse que toda operación en HB da lugar a

teoremas, a fórmulas afirmadas; por eso lo que tenemos que comparar con HB no son usos cualesquiera de RI () y RE $\exists$ en CDN, sino usos que den lugar a teoremas sin más que una aplicación de RI $\rightarrow$ (si la hay). Hay que comparar con las operaciones posibles en HB lo que llamaremos usos, pasos o resultados 'definitivos' de RI () y RE $\exists$ en CDN. Así precisado el objeto de la comparación, se aprecia que RI () y RE $\exists$ no rinden más de lo que puede conseguirse en HB. Veámoslo a propósito de RE $\exists$.

La aplicación de RE $\exists$ no da en CDN un resultado definitivo mas que si la variable anotada no queda libre ni ha sido anotada previamente en la misma demostración. La primera condición excluye que pueda considerarse definitivo en CDN el paso a 'Py', por ejemplo, a partir de ' $\exists x Px$ ', pues 'y' quedaría libre en la fórmula 'Py':

Ln (n)	$\exists x Px$	RP.
y Ln + 1 (n)	Py	RE $\exists$; Ln.

Para evitar que 'y' quede libre en el resultado definitivo habría que ligarla. Pero la segunda condición excluye que se pueda ligar a 'y' por medio del generalizador '()', pues entonces, aunque efectivamente 'y' dejaría de estar libre, se produciría una doble anotación de dicha variable:

Ln (n)	$\exists x Px$	RP.
y Ln + 1 (n)	Py	RE $\exists$; Ln
y Ln + 2 (n)	(y) Py	RI () ; Ln + 1.

Por tanto, el único uso de RE $\exists$ que puede ser definitivo es el que vuelve a ligar la variable anotada mediante el particularizador 'E':

Ln (n)	$\exists x Px$	RP.
y Ln + 1 (n)	Py	RE $\exists$; Ln.
Ln + 2 (n)	$\exists y Py$	RI $\exists$; Ln + 1.

En este caso el resultado es efectivamente definitivo. Pues para conseguir un teorema no hay ya que aplicar ninguna regla de cuantificación, sino RI $\rightarrow$. Pero este resultado se obtiene en HB, a partir de la misma premisa, en un solo paso, por R δ :

L1	$\exists x Px$	Premisa
L2	$\exists y Py$	R δ : x/y; L1.

Una reflexión análoga (para mostrar que el único uso definitivo de RI () es aquel que procede de una aplicación de RE ()) justifica en HB los usos definitivos de RI () en CDN.

En el caso de RI () puede darse una demostración directa muy breve en HB:

L1	Py	Premisa
L2	$p \rightarrow p \vee q$	A2.
L3	$p \vee q \rightarrow q \vee p$	A3.
L4	$p \rightarrow q \vee p$	RA β ₁ ; L2, L3.
L5	$Py \rightarrow q \vee Py$	R α ₁ : p/Py; L4.
L6	$q \vee Py$	R β ; L1, L5.
L7	$\sim s \vee Py$	R α ₁ : q/ $\sim s$; L6.
L8	$s \rightarrow Py$	(D2); L7.
L9	$s \rightarrow Px$	R α ₂ : y/x; L8.
L10	$[p \rightarrow p \vee q] \rightarrow Px$	R α ₁ : s/p $\rightarrow p \vee q$; L9.
L11	$[p \rightarrow p \vee q] \rightarrow (x)Px$	R γ ₁ ; L10.
L12	(x)Px	R β ; L2, L11.

No nos queda pendiente más que la operación de la regla RI $\rightarrow$, para poder admitir que todo teorema de CDN es también un teorema de HB.

60. El teorema de deducción. — La presencia de la regla RI $\rightarrow$ en CDN representa la diferencia esencial entre este cálculo y HB. En HB una premisa se manipula exactamente igual que un axioma o un teorema del sistema. Es la persona que usa HB la que sabe que tal o cual fórmula que ha introducido es una mera premisa. HB mismo la trata como si fuera un teorema, "sin distinguir", por así decirlo, entre éstos (incluyendo en ellos los axiomas) y las premisas. En cambio, CDN "distingue" entre unos y otras, pues tiene un trato especial para las segundas: precisamente la aplicación de la regla RI $\rightarrow$.

Vamos a considerar ahora cómo puede introducirse en el funcionamiento de HB la distinción entre premisas y teoremas, de modo que el funcionamiento mismo de HB tenga un trato especial para las premisas. Esto puede conseguirse imponiendo algunas restricciones especiales a las demostraciones en HB que sean demostraciones a partir de premisas (más teoremas, naturalmente, por lo común). Para averiguar qué restricciones nos interesa imponer a las demostraciones en HB a partir de premisas, vamos a intercalar ahora una breve reflexión acerca de qué son premisas en el sentido relevante para este contexto.

'Premisa' puede tener uno de dos significados principales: 1.º) fórmula a la cual se aplica una regla de transformación; en este sentido es claro que también puede llamarse 'premisa' a un teorema o axioma; no es esto, naturalmente, lo que entendemos cuando hablamos de una 'demostración a partir de premisas', pues en este sentido de 'premisa' todas las demostraciones lo son a partir de ellas. 2.º) Fórmula cuya verdad no consta en el interior de un determinado sistema (en nuestro caso: HB). Tal vez se sepa, por datos procedentes de fuera del sistema, que la fórmula es verdadera. Tal vez se ignore si lo es o no, y tal vez no interese siquiera saberlo (éste

será siempre el caso en lógica pura). En todo caso, si consta su verdad, ello será empíricamente, por así decirlo: no formalmente, no porque sea un teorema del sistema formal. En este sentido hablamos aquí de premisas.

Que esa es la significación de 'premisa' que nos interesa puede quedar claro por lo siguiente: si las fórmulas que se usan como premisas en una demostración (en HB o en CDN) son teoremas, bastará intercalar en la demostración las demostraciones de dichos teoremas para que la demostración lo sea exclusivamente a partir de axiomas.

Así pues, una premisa es una afirmación factual sin justificación formal. Si es verdadera, lo será como una afirmación empírica, por ejemplo, como el enunciado 'este agua está a 35° C'.

De esta naturaleza de las premisas se desprende una interesante consecuencia para el cálculo con ellas: que no puede tener ninguna utilidad en HB sustituir variables de una premisa, o cuantificar sus variables individuales. Pues esas operaciones, practicadas sobre premisas, son en realidad introducciones disfrazadas de nuevas premisas. Transforman, en efecto, enunciados factuales; y transformar un enunciado factual es enunciar otro enunciado factual (afirmar un hecho en vez de uno primero no es explicitar una consecuencia, sino, simplemente, afirmar otro hecho). Por tanto, en vez de practicar transformaciones sobre una premisa, lo correcto es tomar desde el principio la que haga falta. Por ejemplo: supongamos que alguien cuenta con la premisa ' r ' y quiere aplicarla a un teorema ' $p \rightarrow q$ '; y supongamos que decide para ello someter ' r ' a $R\alpha_1$ con r/p ; y que luego afirma ' q ' a partir de la premisa ' r '. La operación ha sido un autoengaño: la persona en cuestión ha cometido una falacia si ' $p \rightarrow q$ ' era un teorema empírico; y ha perdido el tiempo si ' $p \rightarrow q$ ' era un teorema formal. En efecto:

Si ' $p \rightarrow q$ ' era un teorema empírico, como

Juan viene $\rightarrow$ Luis se va

y ' r ' un enunciado empírico cualquiera (distinto de ' p ', naturalmente), como

Pedro viene,

entonces la operación $R\alpha_1$ con r/p es un sinsentido.

Y si ' $p \rightarrow q$ ' era un teorema formal, entonces, por $R\alpha_1$ con p/r , también es un teorema formal ' $r \rightarrow q$ ', y con ' r ' se puede obtener correctamente ' q ' por $R\beta$. La transformación de la premisa era pues inútil, además de ser contraria a la naturaleza de premisa de ' r ', patente en el caso de la argumentación por $R\beta$ en un lenguaje empírico (en nuestro ejemplo de Juan, Luis y Pedro).

La anterior reflexión nos sugiere cuáles son las restricciones que debemos imponer a las demostraciones en HB a partir de premisas para que realmente traten a las premisas como lo que son, como afirmaciones factuales desde el punto de vista de HB. Definiremos del modo siguiente la

noción de demostración en HB a partir de premisas (para abreviar: 'P-demostración en HB'):

La demostración d es una P-demostración de X en HB a partir de la premisa Y cuando d es una demostración de X en HB a partir de Y , e Y no contiene variables de enunciado, variables predicativas o variables individuales que sean objeto de sustitución o cuantificación en d .

Sustituciones y cuantificaciones puede haber, naturalmente, en d , sin que por ello d deje de ser una P-demostración. Pero estarán practicadas sólo sobre axiomas, teoremas o fórmulas ya demostradas a partir de la premisa, axiomas y teoremas.

Adoptaremos además la convención, bastante natural y siempre seguida hasta ahora, de que la primera fórmula de una P-demostración de X a partir de Y es la premisa Y , o un axioma o un teorema de HB.

En este punto puede suscitarse una duda. Puede temerse, en efecto, que el concepto de P-demostración (es decir, las restricciones impuestas a las demostraciones a partir de premisas en HB) sea mucho más estrecho que la demostración bajo premisas en CDN, lo que haría inútil toda comparación entre ambos sistemas. Pues en CDN se pueden practicar cuantificaciones de las variables individuales de una premisa y , consiguientemente, también sustituciones de variables individuales (no de enunciados, ni predicativas).

Pero esta diferencia no supone una diversa concepción de la premisa, sino sólo dos modos distintos de precisar su naturaleza de premisa: en una demostración en CDN no hay ninguna línea afirmada, salvo las que no se escriben bajo premisa (es decir, las que se escriben bajo 0 premisa). Por eso las operaciones que se hacen bajo premisa no comprometen, por así decirlo. En cambio, en una P-demostración en HB cada línea se presenta como una afirmación, y por eso debemos tomar medidas restrictivas para que se presente de otro modo. Pues bien: en las fórmulas de CDN que son afirmaciones sucede con la premisa exactamente lo mismo que en las líneas de una P-demostración en HB, a saber: que la premisa no está transformada en nada. Ejemplo de esto puede ser la siguiente demostración en CDN:

L1 (1)	Py	RP.
L2 (1)	$\exists x Px$	$RI\exists; L1.$
L3 (0)	$Py \rightarrow \exists x Px$	$RE\exists; L2.$

En la única fórmula afirmada, ' $Py \rightarrow \exists x Px$ ', la premisa ' Py ' está exactamente igual que se introdujo. En CDN, más que transformar la premisa, que en realidad se va "arrastrando", lo que se hace es construir nuevas fórmulas bajo premisas.

En resolución: prohibir en CDN afirmar fórmulas todavía bajo premisas es lo mismo que prohibir en HB practicar sustituciones o cuantificaciones sobre premisas. El concepto de P-demostración no se separa, pues, del uso de premisas en cualquier cálculo, CDN u otro.

Una vez fijado el concepto de P-demostración, es posible mostrar que HB tiene una propiedad que, en la práctica, para la cuestión del rendimiento del sistema, equivale a la presencia de $RI \rightarrow$ en CDN. La propiedad en cuestión puede formularse así:

Si en HB existe una P-demostración (PD) de X a partir de la premisa Y, entonces en HB existe también una demostración (D) de $Y \rightarrow X$.

Esta afirmación en cursiva es una versión bastante intuitiva (por la noción de P-demostración) y simplificada (por no tomar en cuenta más que una premisa) de un metateorema llamado 'teorema de deducción'. (Un metateorema es un teorema del metalenguaje del sistema considerado, en este caso del metalenguaje de HB; también puede decirse 'teorema metalógico'.) Pero la llamaremos por brevedad 'teorema de deducción'. Para nuestra tarea — terminar la comparación de HB con CDN — esa versión va a ser suficiente, como veremos ahora, antes de proceder a demostrarla.

El teorema de deducción no nos va a legitimizar la operación de $RI \rightarrow$ en HB. En HB no se puede introducir una nueva regla (una regla auxiliar) más que construyéndola con las reglas primitivas. Y eso no es lo que hará la demostración del teorema de deducción. Pero el teorema de deducción nos va a garantizar que para todo teorema, $Y \rightarrow X$, obtenido en CDN por $RI \rightarrow$, hay en HB una demostración de $Y \rightarrow X$. Más precisamente: supongamos una fórmula, X, obtenida en CDN (sin usar $RI \rightarrow$) aún bajo la única premisa Y, y en HB a partir de la única premisa Y, usada auténticamente como premisa (o sea, en una P-demostración). Entonces, igual que en CDN es posible demostrar el teorema $Y \rightarrow X$ (fórmula sin ninguna premisa) por medio de $RI \rightarrow$ aplicada a la línea en que está X, así también existe en HB una demostración de $Y \rightarrow X$. La diferencia de estructura entre los dos sistemas sigue en pie, y a la luz del teorema de deducción esa diferencia puede expresarse así: en CDN no hay más que una demostración, en la cual, tras obtenerse X bajo la premisa Y, se obtiene $Y \rightarrow X$ bajo ninguna premisa, o sea, como teorema; en cambio, en HB hay dos demostraciones: una P-demostración de X a partir de Y y una demostración de $Y \rightarrow X$. Pero desde el punto de vista del rendimiento, que es el que aquí nos interesa, esa diferencia es irrelevante: lo esencial es que si $Y \rightarrow X$ es un teorema de CDN (en el que está obtenido con una aplicación de $RI \rightarrow$), entonces es también teorema de HB (pese a que en este sistema está excluida la aplicación de esa regla). Precisamente porque lo que nos interesa es comparar el rendimiento de los dos sistemas, aceptamos en 59 que CDN rinde todo lo que rinden las reglas $R\alpha_1$ y $R\alpha_2$ de HB, a pesar de que en CDN no hay ninguna regla que autorice a pasar, por ejemplo, de ' $p \vee p \rightarrow p$ ' a ' $r \vee r \rightarrow r$ ', cosa que en cambio es posible en HB. Pero para la comparación de rendimiento basta con saber que si ' $p \vee p \rightarrow p$ ' y ' $r \vee r \rightarrow r$ ' son teoremas en HB, entonces también lo son en CDN (aunque por dos demostraciones). Esto es

lo que vimos entonces. Y lo que nos indica el teorema de deducción es que si $Y \rightarrow X$ es teorema de CDN, entonces también lo es de HB (aunque con otro esquema demostrativo: con dos demostraciones en vez de una como en CDN).

Para justificar el teorema de deducción mostraremos que, dada una P-demostración, (PD) de X a partir de Y en HB, puede construirse otra demostración, (D), de $Y \rightarrow X$, la cual presenta las siguientes características: (D) conserva todas las líneas de (PD), pero tiene intercaladas entre ellas otras líneas nuevas, entre otras, para cada fórmula, V_n , de una línea, L_n , de (PD), una línea con una fórmula nueva, $Y \rightarrow V_n$. Para la última línea de (PD), que es X, la nueva línea de este género será $Y \rightarrow X$. Por tanto, (D) será una demostración de $Y \rightarrow X$.

De lo que se trata, pues, es de mostrar la posibilidad de esas intercalaciones para todas las líneas y casos posibles. El esquema de la argumentación es el de las demostraciones "por inducción", en este caso, por inducción sobre el número, n, de las líneas de (PD). Mostraremos: 1.º que para la línea L_1 de (PD) es posible una intercalación en (D) con la que se demuestra $Y \rightarrow V_1$; 2.º que, si la posibilidad de la intercalación se ha demostrado hasta la línea L_{n-1} , es decir, si se han podido intercalar fórmulas $Y \rightarrow V_1, Y \rightarrow V_2, \dots, Y \rightarrow V_{n-1}$, entonces la intercalación es también posible para la línea L_n , con una nueva fórmula demostrada $Y \rightarrow V_n$. Con esto, por el "principio de inducción" cuyo esquema es el 5.º axioma de Peano (cfr. 37), quedará establecida la posibilidad de intercalación para todas las líneas de (PD), o sea, la posibilidad de construir (D). Por simplicidad tipográfica usaremos letras distintas para las fórmulas, en vez de 'V' con subíndices.

Línea L_1 de (PD). — La fórmula de la línea L_1 de (PD), V, no puede ser, por la construcción de (PD), más que Y, o un axioma o un teorema.

Caso a. V es Y. Entre las líneas L_1 y L_2 de (PD) puede intercalarse entonces una demostración del teorema $Y \rightarrow Y$. (D) quedará construida en este primer tramo con el siguiente aspecto (comparado con el inicial de (PD), a la izquierda):

(PD)			D		
L_1	Y	Premisa	L_1	Y	Premisa
L_2			L_1	$p \rightarrow p \vee q$	A2.
.			L_{II}	$p \rightarrow p \vee p$	$R\alpha_1: q/p; L_1.$
.			L_{III}	$p \vee p \rightarrow p$	A1.
.			L_{IV}	$p \rightarrow p$	$RA\beta_1; L_{II}, L_{III}.$
.			L_V	$Y \rightarrow Y$	$R\alpha_1: p/Y; L_{IV}.$
.			L_2		
.			.		
.			.		
.			.		

(En la notación admitimos la mezcla de símbolos del lenguaje objeto y símbolos metalingüísticos — sintácticos — a la que ya antes aludimos y que no tiene riesgos. Para cada (PD), X sería, naturalmente, una fórmula del lenguaje objeto, y lo mismo Y .)

Caso b. V es un axioma o un teorema, en ambos casos una fórmula válida en HB. Puede intercalarse tras la línea L1 de (PD) las líneas que demuestran $Y \rightarrow V$. Aspectos inicial de (PD) y de la demostración resultante, (D), en este tramo en este caso:

(PD)			(D)		
L1	V	Axioma o teorema	L1	V	Ax. o teor.
L2			L1	$p \rightarrow p \vee q$	A2.
.			LII	$p \vee q \rightarrow q \vee p$	A3.
.			LIII	$p \rightarrow q \vee p$	$R\alpha_1$; L1, LII.
.			LIV	$V \rightarrow q \vee V$	$R\alpha_1$; p/V ; LIII.
.			LV	$q \vee V$	$R\beta$; L1, LIV.
			LVI	$\sim Y \vee V$	$R\alpha_1$; $q/\sim Y$; LV.
			LVII	$Y \rightarrow V$	(D2); LVI.
			L2		
			.		
			.		
			.		

La intercalación para la construcción de (D) es pues posible para la línea L1 de (PD).

Hipótesis de inducción (HI). — Suponemos ahora que la intercalación se ha conseguido hasta la línea $L_n - 1$ de (PD).

Línea L_n de (PD). — La fórmula de la línea L_n de (PD), V , puede haberse conseguido aplicando a una fórmula de alguna línea $L_n - r$, Z , alguna de las reglas $R\alpha_1$, $R\alpha_2$, $R\alpha_3$, $R\gamma_1$, $R\gamma_2$, o $R\delta$, o bien aplicando a dos fórmulas, Z y W , de dos líneas anteriores, $L_n - s$ y $L_n - t$, la regla $R\beta$. Como es natural, V puede ser también un axioma o un teorema. Pero entonces la intercalación que da $Y \rightarrow V$ es la misma que en los casos respectivos para la línea L1 de (PD). Atenderemos pues a las otras situaciones posibles.

a) V procede de Z por $R\alpha_1$. — Por la hipótesis de inducción, como la línea de Z es $L_n - 1$ o anterior a $L_n - 1$, hay en (D) una fórmula $Y \rightarrow Z$. Procedemos en (D) a la siguiente intercalación, con la misma visualización que antes:

(PD)			(D)		
.			.		
$L_n - r$	Z		$L_n - r$	Z	
.			.		
L_n	V	$R\alpha_1$; $L_n - r$	L_R	$Y \rightarrow Z$	(HI)
$L_n + 1$			.		
.			.		
.			.		
.			L_n	V	$R\alpha_1$; $L_n - r$.
.			L_{R+1}	$Y \rightarrow V$	$R\alpha_1$; L_R .
.			$L_n + 1$		
.			.		
.			.		

El subíndice R indica sólo que L_R es una línea nueva, una línea de (D), pero no de (PD) al principio, y va numerada con una cifra romana. El paso de Z a V en (PD) es por $R\alpha_1$, pongamos por p/q . Como (PD) es una P -demostración, Y no contiene ninguna variable sustituida en (PD) (en este caso del ejemplo, no contiene ' p '). Por eso la misma aplicación de $R\alpha_1$ que en (PD) permite pasar de Z a V , en (D) permite pasar de $Y \rightarrow Z$ a $Y \rightarrow V$. (No hemos supuesto que en (PD) Z fuera la fórmula inmediatamente anterior a V (no hemos supuesto $r = 1$). Esto es, naturalmente, irrelevante, pues la hipótesis inductiva nos dice que las intercalaciones son posibles hasta $L_n - 1$. Para simplificar tipográficamente los esquemas renunciaremos a partir de ahora a ese "realismo" y supondremos que Z está en la línea $L_n - 1$.)

b) V procede de Z por $R\alpha_2$. — Por la hipótesis de inducción, sabemos que en (D) se tiene ya, como fruto de una intercalación en (PD), la fórmula $Y \rightarrow Z$. El paso de Z a V en (PD) es por $R\alpha_2$, sustitución de una variable individual por otra, pongamos, x/y . Por construcción de la P -demostración (PD), Y no contiene a ' x '. Consiguientemente, como en el caso anterior, la misma sustitución que en (PD) lleva de Z a V , lleva en (D) de $Y \rightarrow Z$ a $Y \rightarrow V$. La intercalación es pues:

(PD)	(D)
.	.
.	.
Ln - 1 Z	Ln - 1 Z
.	.
Ln V R α_2 : x/y; Ln - 1.	L _R Y $\rightarrow$ Z (HI)
Ln + 1	Ln V R α_2 : x/y; Ln - 1.
.	L _{R+1} Y $\rightarrow$ V R α_2 : x/y; L _R .
.	Ln + 1
.	.
.	.

c) y d) *V* procede de *Z* por R α_3 o por R δ . — La intercalación se construye como en los dos casos anteriores, con la regla correspondiente.

e) *V* procede de *Z* por R γ_1 . — Por la hipótesis de inducción, en (D) se tiene ya una línea $Y \rightarrow Z$. De *Z* se pasa a *V* en (PD) por generalización posterior, R γ_1 . Esto supone: que *Z* es un condicional, $W \rightarrow S$, que la variable generalizada, digamos '*x*', está libre en *S* y no aparece en *W*; y que *V* es $W \rightarrow (x)S$. Además, como (PD) es una P-demostración, *Y* no contiene '*x*'. La intercalación entre fórmulas de (PD) para construir el correspondiente tramo de (D) es pues:

(PD)	(D)
.	.
.	.
Ln - 1 $W \rightarrow S$	Ln - 1 $W \rightarrow S$
Ln $W \rightarrow (x)S$ R γ_1 ; Ln - 1.	.
Ln + 1	.
.	L _R $Y \rightarrow [W \rightarrow S]$ (HI)
.	Ln $W \rightarrow (x)S$ R γ_1 ; Ln - 1
.	L _{R+1} $Y \rightarrow (x)[W \rightarrow S]$ R γ_1 ; L _R
.	L _{R+2} $Y \rightarrow [W \rightarrow (x)S]$ (TP18).
.	Ln + 1
.	.
.	.

f) *V* procede de *Z* por R γ_2 . — Por la hipótesis de inducción, hay en (D) una línea $Y \rightarrow Z$. Como en el caso anterior, *Z* es un condicional, $S \rightarrow W$, '*x*' está libre en *S*, '*x*' no se presenta en *W*, y *V* es $\exists x S \rightarrow W$. Por ser (PD)

una P-demostración, *Y* no contiene a '*x*'. La intercalación que necesitamos tras la línea Ln es:

(PD)	(D)
.	.
.	.
Ln - 1 $S \rightarrow W$	Ln - 1 $S \rightarrow W$
Ln $\exists x S \rightarrow W$ R γ_2 ; Ln - 1	.
Ln + 1	.
.	.
.	L _R $Y \rightarrow [S \rightarrow W]$ (HI).
.	Ln $\exists x S \rightarrow W$ R γ_2 ; Ln - 1
.	L _{R+1} $Y \rightarrow (x)[S \rightarrow W]$ R γ_1 ; L _R
.	L _{R+2} $Y \rightarrow [\exists x S \rightarrow W]$ (TP19).
.	Ln + 1
.	.
.	.

Queda por considerar el último caso:

g) *V* procede de *Z* y *W* por R β . — Por la hipótesis de inducción, hay en (D) una línea $Y \rightarrow Z$ y una línea $Y \rightarrow W$. *Z* tiene que ser en (PD) el condicional $W \rightarrow V$. La intercalación en (PD) para formar el correspondiente tramo de (D) es:

(PD)	(D)
.	.
.	.
Ln - 2 $W \rightarrow V$	Ln - 2 $W \rightarrow V$
Ln - 1 W	.
Ln V R β ; Ln - 2, Ln - 1	.
Ln + 1	.
.	L _R $Y \rightarrow [W \rightarrow V]$ (HI)
.	Ln - 1 W
.	.
.	.
.	L _{R+1} $Y \rightarrow W$ (HI)
.	L _{R+2} $Y \rightarrow V$ R β ; Ln - 2, L _{R+1}
.	Ln + 1
.	.
.	.

Con esto queda completada la demostración de que es posible realizar en la P-demostración (PD) de X a partir de Y las intercalaciones que la convierten en la demostración (D) de $Y \rightarrow X$. Lo cual es una demostración del teorema de deducción.

La anterior demostración es incorrecta en los pasos e) y g), porque en ellos se usan dos teoremas, (TP18) y (TP19), que no han sido previamente demostrados en HB. Pero los dos son teoremas de HB. He aquí, por ejemplo, una demostración de (TP18) en HB:

Empezaremos por justificar una nueva regla auxiliar de HB:

$$\text{RA}_{\gamma_4} \quad \frac{X \rightarrow [Y \rightarrow Z]}{X \wedge Y \rightarrow Z}$$

Justificación:

- | | | |
|-----|-----------------------------------|--|
| L1. | $X \rightarrow [Y \rightarrow Z]$ | Premisa de la regla. |
| L2. | $\sim X \vee [\sim Y \vee Z]$ | (D2); L1. |
| L3. | $[\sim X \vee \sim Y] \vee Z$ | paso meramente notacional, por la asociatividad de ' $\vee$ '. |
| L4. | $\sim [X \wedge Y] \vee Z$ | (D1); L3. |
| L5. | $X \wedge Y \rightarrow Z$ | (D2); L4. |

Rehaciendo la demostración a la inversa, se obtiene otra regla auxiliar más para HB:

$$\text{RA}_{\gamma_5} \quad \frac{X \wedge Y \rightarrow Z}{X \rightarrow [Y \rightarrow Z]}$$

Con estas dos reglas auxiliares podemos proceder a demostrar (TP18) en HB:

- | | | |
|------|---|---|
| L1. | $(x)Px \rightarrow Py$ | A5. |
| L2. | $(x)[\sim Qz \vee Px] \rightarrow \sim Qz \vee Py$ | R_{α_2} : $P/\sim Qz \vee P$; L1. |
| L3. | $(x)[Qz \rightarrow Px] \rightarrow [Qz \rightarrow Py]$ | (D2); L2 |
| L4. | $(x)[Qz \rightarrow Px] \wedge Qz \rightarrow Py$ | RA_{γ_4} ; L3. |
| L5. | $(x)[Qz \rightarrow Px] \wedge Qz \rightarrow (y)Py$ | R_{γ_1} ; L4. |
| L6. | $(x)[Qz \rightarrow Px] \wedge Qz \rightarrow (x)Px$ | $R\delta$; L5. |
| L7. | $(x)[Qz \rightarrow Px] \rightarrow [Qz \rightarrow (x)Px]$ | RA_{γ_5} ; L6. |
| L8. | $(x)[Py \rightarrow Px] \rightarrow [Py \rightarrow (x)Px]$ | R_{α_3} : Qz/Py ; L7. |
| L9. | $[p \rightarrow q] \rightarrow [s \vee p \rightarrow s \vee q]$ | A4. |
| L10. | $[(x)Px \rightarrow Py] \rightarrow [s \vee (x)Px \rightarrow s \vee Py]$ | R_{α_1} : $p/(x)Px$; q/Py ; L9. |
| L11. | $s \vee (x)Px \rightarrow s \vee Py$ | $R\beta$; L1, L10. |
| L12. | $s \vee (x)Px \rightarrow (y)[s \vee Py]$ | R_{γ_1} ; L11. |
| L13. | $s \vee (x)Px \rightarrow (x)[s \vee Px]$ | $R\delta$; y/x ; L12. |
| L14. | $\sim Py \vee (x)Px \rightarrow (x)[\sim Py \vee Px]$ | R_{α_1} : $s/\sim Py$; L13. |
| L15. | $[Py \rightarrow (x)Px] \rightarrow (x)[Py \rightarrow Px]$ | (D2); L14. |
| L16. | $(x)[Py \rightarrow Px] \leftrightarrow [Py \rightarrow (x)Px]$ | $\text{RA} \leftrightarrow_1$; L8, L15. |

Análogamente se demuestra en HB (TP19).

El papel del teorema de deducción en la teoría lógica no es el que tiene aquí de ilustrar la igualdad de rendimiento de HB y CDN; el teorema de deducción es parte de la metateoría de los sistemas axiomáticos y es, como el principio de éstos, históricamente anterior al cálculo de la deducción natural. Precisamente éste puede entenderse como inspirado en la idea de convertir el teorema de deducción, que es un teorema meta-lógico, un teorema sobre el cálculo y no del cálculo, en una regla de transformación del cálculo.

Pero incluso en su papel sustantivo para el estudio de los sistemas axiomáticos, el teorema de deducción está relacionado con el tema para el cual nos ha sido útil aquí. Pues el teorema de deducción se aduce en general, entre otros motivos, por mostrar que la noción de demostración a partir de premisas recoge "naturalmente" en el sistema axiomático la idea intuitiva de la relación de consecuencia entre enunciados, una esquematización extensional de la cual es ' $\rightarrow$ '. Este teorema tiende pues a documentar la "naturalidad" de la demostración en el sistema axiomático. Ahora bien: reproducir la "naturalidad" de la deducción es precisamente el motivo inspirador de CDN.

La presentación de la demostración inductiva del anterior teorema simplificado se basa en la construcción del teorema de deducción por A. Church para un sistema de lógica de enunciados y de lógica de predicados de primer orden con esquemas axiomáticos y axiomas. Difiere de ella por el uso hecho aquí de la noción de P-demostración, y por algunos expedientes de exposición. La idea del teorema de deducción es de J. Herbrand (1928).

PARTE TERCERA

LIMITACIONES Y ALCANCE
DEL CALCULO LOGICO

Sección Primera. — LAS LIMITACIONES DEL CALCULO LÓGICO

CAPÍTULO XI

RENDIMIENTO DEL CALCULO LÓGICO ELEMENTAL

61. Los metateoremas sobre el rendimiento. — En el capítulo X, con motivo de la comparación entre CDN y HB, hemos hablado repetidamente de 'rendimiento' sin precisar mayormente esa palabra. Sin embargo, contamos ya con conceptos que nos permiten precisar lo que quiere decir 'rendimiento de un cálculo'. Esos conceptos son los de consistencia, completud y decidibilidad (cfr. 18). Estudiar el rendimiento de un algoritmo es estudiar su comportamiento respecto de esas tres propiedades, especialmente respecto de las dos primeras, que son las de interés más básico. Un estudio de esa naturaleza es metalógico: sus resultados no son afirmaciones del cálculo, sino afirmaciones sobre el cálculo, metateoremas. (Un ejemplo de metateorema es el teorema de deducción, estudiado en 60.)

El estudio del cálculo lógico elemental ha establecido las siguientes propiedades del mismo:

- 1.ª: el cálculo lógico elemental es consistente;
- 2.ª: el cálculo lógico elemental es completo;
- 3.ª: no todo el cálculo lógico elemental es decidible, pero sí lo es parte de él.

En el presente capítulo vamos a estudiar la consistencia y la completud del cálculo lógico elemental. En el capítulo XIII estudiaremos cuestiones relativas a la decidibilidad de alguna de sus partes.

Los metateoremas sobre el rendimiento son, en general, de demostración sencilla en cuanto a su principio, pero algo laboriosa por su técnica. La complicación técnica se debe a la necesidad de suponer un universo de infinitos individuos, o sea, a la necesidad de referirse en las demostraciones a la totalidad de los símbolos individuales; esos símbolos pueden ser infinitos, y como no se les puede tener materialmente presentes a todos, es

necesario arbitrar procedimientos que garanticen, por así decirlo, que se les puede pasar revista a todos sin confundirlos. Esos procedimientos consisten en reglas estructurales sobre la formación de las fórmulas. En la formulación de dichas reglas se encuentran las principales dificultades técnicas de las demostraciones que nos interesa considerar.

Aquí deseamos evitar esas dificultades técnicas, pero sin renunciar por ello a tomar noticia de la naturaleza de las demostraciones metalógicas, con las cuales hemos empezado a familiarizarnos en 60. Para ello vamos a referir estas argumentaciones a un universo finito, esto es, a un universo del discurso con un conjunto finito de individuos, por ejemplo, el formado exclusivamente por los individuos a y b : $\Omega = \{a, b\}$. No por eso dejaremos de aludir también a universos cualesquiera. Pero entonces nuestras indicaciones serán meras ilustraciones o comentarios de procedimientos demostrativos. Nuestras reflexiones no serán demostrativas más que cuando se refieran a $\Omega = \{a, b\}$.

En cierto sentido, el expediente de limitarnos a un Ω finito no es tan grave desde el punto de vista práctico o aplicado cuanto lo es desde el punto de vista teórico. Pues, como sabemos (cfr. 37), la lógica elemental no basta para formular las nociones fundamentales de la aritmética, primera ciencia en la cual tiene un papel esencial la infinitud del campo de individuos.

Los siguientes esquemas de argumentación se basan en las demostraciones dadas por W. Ackermann en la 3.^a edición de Hilbert-Ackermann, *Grundzüge der theoretischen Logik* (1949), aunque introducen en ellas simplificaciones intuitivas y un punto de vista semántico que les es inesencial.

62. Consistencia del cálculo lógico elemental. — La noción de consistencia — o de compatibilidad, como también se dice — es en general la de ausencia de contradicción: un cálculo es consistente cuando en él no son demostrables contradicciones.

Tal como queda formulado — con la palabra ‘contradicción’ — el concepto de consistencia es semántico. Pero sólo verbalmente: pues en vez de ‘contradicciones’ puede decirse: ‘fórmulas como $p \wedge \sim p$ ’, o ‘una fórmula y otra fórmula que es como la anterior con el signo $\sim$ antepuesto’. Y estas formulaciones son sintácticas. En realidad, los conceptos de consistencia, completud y decidibilidad son susceptibles de varias formulaciones, y aunque en general nos atenderemos en cada caso a una noción, vamos ahora a obtener otro concepto de consistencia, éste sí verdaderamente semántico, que nos será útil para hacer más intuitiva la argumentación siguiente.

Si un cálculo es inconsistente, no puede tener ningún modelo, ninguna interpretación que haga verdadero al sistema de sus teoremas; pues ningún objeto, a , puede satisfacer una fórmula como, por ejemplo, $Pa \wedge \sim Pa$.

Entonces, por las leyes de la contraposición (TE4), puede afirmarse que si un cálculo tiene al menos un modelo para el sistema de sus teoremas, ese cálculo es consistente. La demostración de la consistencia del cálculo lógico elemental puede consistir pues en mostrar que tiene al menos un modelo. Vamos a construir un sistema de interpretaciones $\mathfrak{I}$, del cálculo lógico elemental, y a mostrar que toda interpretación particular, $\mathfrak{I}_a$, perteneciente al sistema $\mathfrak{I}$ es modelo de dicho algoritmo. $\mathfrak{I}$ será como sigue:

1.º Los símbolos de enunciado (atómicos) — como ‘ p ’, ‘ q ’, etc. — se interpretarán en el campo de objetos constituido por los valores V y F : $\{V, F\}$.

Las fórmulas moleculares — como ‘ $p \rightarrow q$ ’, etc. — quedan entonces interpretadas sin más, según las tablas de los funtores veritativos (cfr. 30).

2.º En cuanto a las fórmulas predicativas atómicas, establecemos las siguientes reglas de interpretación:

a) el campo de individuos para las variables individuales es $\Omega = \{a, b\}$;

b) las letras predicativas — como ‘ P ’, ‘ Q ’, etc. — no se interpretan, sino que se toman como parámetros fijos.

La sencillez del universo del discurso elegido nos permite tratar simultáneamente todas las interpretaciones, $\mathfrak{I}_a$, pertenecientes a $\mathfrak{I}$, como si se tratara de una sola. Por eso hablaremos (laxamente) de ‘la interpretación $\mathfrak{I}$ ’.

Como consecuencia del punto 2.º, las fórmulas cuantificadas se reducen a conjunciones o disyunciones (cfr. 55), del modo ejemplificable como sigue:

$$(x)Px \leftrightarrow \mathfrak{I} Pa \wedge Pb.$$

$$\exists x Px \leftrightarrow \mathfrak{I} Pa \vee Pb.$$

$$(x)(y)Rxy \leftrightarrow \mathfrak{I} Raa \wedge Rab \wedge Rbb \wedge Rba.$$

$$\exists x \exists y Rxy \leftrightarrow \mathfrak{I} Raa \vee Rab \vee Rbb \vee Rba.$$

$$(x)\exists y Rxy \leftrightarrow \mathfrak{I} [Raa \vee Rab] \wedge [Rba \vee Rbb].$$

$$\exists y(x)Rxy \leftrightarrow \mathfrak{I} [Raa \wedge Rab] \vee [Rba \wedge Rbb].$$

Etcétera.

‘ $\leftrightarrow \mathfrak{I}$ ’ debe leerse: ‘equivale por $\mathfrak{I}$ a’.

La interpretación $\mathfrak{I}$ es modelo del cálculo lógico elemental, que tomaremos en la presentación axiomática HB. Para justificar esa afirmación tenemos que mostrar dos cosas: primera, que $\mathfrak{I}$ es modelo de todos los axiomas de HB; segunda, que las reglas de HB transmiten esa propiedad de los axiomas de tener el modelo $\mathfrak{I}$; o sea: que aplicadas a fórmulas que (como los axiomas) sean verdaderas para $\mathfrak{I}$, las reglas del cálculo producen nuevas fórmulas que también son verdaderas para $\mathfrak{I}$.

Por lo que hace a lo primero: el axioma A1 toma por la interpretación $\mathfrak{I}$ el valor de

$$V \vee V \rightarrow V$$

o el de

$$F \vee F \rightarrow F,$$

que es siempre V, como puede verse por las tablas de los funtores 'v' y '→'. De este mismo modo puede comprobarse la verdad de los axiomas A2-A4 para $\mathfrak{I}$.

El axioma A5,

$$(x)Px \rightarrow Py,$$

se reduce a

$$Pa \wedge Pb \rightarrow Pa,$$

o a

$$Pa \wedge Pb \rightarrow Pb;$$

esos dos enunciados son verdaderos para $\mathfrak{I}$, como puede apreciarse por las tablas de '∧' y '→'.

El axioma A6,

$$Py \rightarrow \exists x Px,$$

se reduce por $\mathfrak{I}$ a

$$Pa \rightarrow Pa \vee Pb$$

o a

$$Pb \rightarrow Pa \vee Pb,$$

enunciados válidos ambos para $\mathfrak{I}$ según las tablas de 'v' y '→'.

Con esto queda resuelto el primer punto. La interpretación $\mathfrak{I}$ es modelo de todos los axiomas de HB. Ahora veremos que las reglas de transformación de HB transmiten a los teoremas la propiedad de los axiomas de ser verdaderos para $\mathfrak{I}$.

Las reglas de sustitución, $R\alpha_1$, $R\alpha_3$ y $R\delta$, no presentan ningún problema. Esas reglas no permiten más que dos cosas: o cambiar de nombre a un hecho (como cuando se escribe 'p' por 'q') o a un individuo (como cuando se escribe 'x' por 'y'), y éstos son casos de las reglas $R\alpha_1$, $R\alpha_2$ y $R\delta$; o bien sustituir un valor (un símbolo de enunciado) por un complejo de valores (por una fórmula molecular); éste es un caso de las reglas $R\alpha_1$ y $R\alpha_3$. En este caso no es posible la introducción de valores nuevos, pues el complejo de valores sustituyente está él mismo compuesto exclusivamente por los valores V y F, y tiene a su vez el valor V o el valor F, y ningún otro. Por ejemplo, dada la fórmula

$$p \wedge q \rightarrow q,$$

la sustitución $q/r \vee s$,

$$p \wedge [r \vee s] \rightarrow r \vee s,$$

no puede alterar el valor de la fórmula inicial ('principio de extensionalidad').

También la regla $R\beta$ transmite de los axiomas a los teoremas la propiedad que nos interesa. Esta regla permite pasar de unas fórmulas

$$X \rightarrow Y$$

y

$$X$$

a la fórmula

$$Y.$$

Supongamos que $X \rightarrow Y$ y X sean ambas verdaderas para $\mathfrak{I}$, y que Y , en cambio, no lo sea, porque la regla $R\beta$ no transmite dicha propiedad. Entonces $X \rightarrow Y$ o X tendrán que ser falsas para $\mathfrak{I}$, contra la hipótesis. Luego Y no puede ser falsa para $\mathfrak{I}$.

Queda por ver lo que ocurre con las reglas de cuantificación $R\gamma_1$ y $R\gamma_2$. $R\gamma_1$ permite pasar de

$$X \rightarrow Y$$

a

$$X \rightarrow (x)Y$$

siempre que 'x' esté libre en Y y no se presente en X . Tal es, por ejemplo, el caso de

$$(1) \quad q \rightarrow Px,$$

siendo 'q' tal que no contiene a 'x'.

Por $R\gamma_1$ se pasa a

$$(2) \quad q \rightarrow (x)Px.$$

Para el Ω de la interpretación $\mathfrak{I}$, (2) es lo mismo que

$$(3) \quad q \rightarrow Pa \wedge Pb.$$

$\mathfrak{I}$ es por hipótesis modelo de (1), lo que quiere decir que los dos enunciados siguientes son verdaderos para $\mathfrak{I}$:

$$(1a) \quad q \rightarrow Pa;$$

$$(1b) \quad q \rightarrow Pb.$$

Ahora bien: (1a) y (1b) juntas es lo mismo que ' $q \rightarrow Pa \wedge Pb$ ', que es (3). Luego $\mathfrak{I}$ es también modelo de (3). La regla $R\gamma_1$ transmite pues también la propiedad que nos interesa.

Por último, la regla $R\gamma_2$ permite pasar de

$$Y \rightarrow X$$

a

$$\exists x Y \rightarrow X,$$

con las mismas condiciones que $R\gamma_1$ sobre 'x', Y y X. Este es el caso, por ejemplo, de

$$(4) \quad Px \rightarrow q,$$

si 'q' no contiene a 'x'.

Por $R\gamma_2$ obtenemos de (4)

$$(5) \quad \exists x Px \rightarrow q.$$

Para el Ω de la interpretación $\mathfrak{I}$, (5) es lo mismo que

$$(6) \quad Pa \vee Pb \rightarrow q.$$

$\mathfrak{I}$ es modelo de (4). Pero (4), al no tener cuantificada la 'x', no precisa cuáles son los valores de 'x' para los cuales es verdadero el condicional ' $Px \rightarrow q$ '. Puede ser el valor *a* o el valor *b*. Y esto es precisamente lo expresado por (6). Consiguientemente, si $\mathfrak{I}$ es, como presupone la hipótesis modelo de (4), tiene que ser también modelo de (5). Así pues, también $R\gamma_2$ transmite la propiedad de ser verdadero para $\mathfrak{I}$.

En definitiva: $\mathfrak{I}$ es modelo de todos los axiomas del sistema y de todas las fórmulas obtenidas a partir de los axiomas por las reglas del sistema. $\mathfrak{I}$ es modelo de todos los teoremas del sistema. Este tiene, pues, al menos un modelo, y es por tanto consistente.

(En realidad: todas las interpretaciones $\mathfrak{I}_n$, pertenecientes al sistema de interpretaciones $\mathfrak{I}$, son modelos de HB.)

63. Completud del cálculo lógico elemental. — También la noción de completud se entiende en varios sentidos. En sentido semántico, que es el que ya conocemos (cfr. 18), se dice que un cálculo es completo cuando todas las verdades formales formulables en su lenguaje son teoremas del cálculo. En sentido sintáctico, se dice que un cálculo es completo cuando, al añadirse una fórmula que no sea un teorema suyo, el cálculo se hace inconsistente. (El añadido se entiende hecho al sistema de los teoremas del cálculo.) El sentido sintáctico de 'completud' incluye en cierto modo al semántico: pues si el añadir al cálculo una fórmula no demostrable en él le hace inconsistente, es que él demuestra todas las fórmulas verdaderas, o que, por decirlo gráficamente, es completo porque "no se le puede añadir nada".

Empezaremos por ver que la parte de la lógica elemental a la que hemos llamado 'lógica de enunciados' es completa en el sentido más fuerte (el sintáctico) de los dos dichos; luego veremos un esbozo de cómo se demuestra que la lógica de predicados de primer orden — o sea, el sistema de la lógica elemental — es completa (semánticamente).

Completud del cálculo de enunciados. — Supongamos una fórmula del lenguaje de la lógica de enunciados que no sea demostrable en el cálculo y que, sin embargo, se pretenda añadir al mismo como teorema. Sea esa fórmula, por ejemplo

$$(7) \quad p \wedge q \leftrightarrow r.$$

(8) es una forma normal conjuntiva de (7):

$$(8) \quad [\sim p \vee \sim q \vee r] \wedge [\sim r \vee p] \wedge [\sim r \vee q].$$

Si hay que admitir a (7) como teorema, hay que admitir también a (8). Ahora bien: (8) es una conjunción. Para que una conjunción sea verdadera tienen que ser verdaderos todos sus componentes principales; en este caso, los tres de que consta (8). Tomamos entonces uno de ellos, el segundo, por ejemplo, e iniciamos con esa fórmula la siguiente demostración:

L1	$\sim r \vee p$	Supuesta verdad.
L2	$\sim \sim s \vee s$	$R\alpha_1: r/\sim s; p/s; L1.$
L3	$s \vee s$	Ley doble negación.
L4	$p \vee p \rightarrow p$	A1.
L5	$s \vee s \rightarrow s$	$R\alpha_1: p/s; L4.$
L6	s	$R\beta; L3, L5.$

Al añadir (7) a los teoremas del cálculo, obtenemos pues que 's' es también un teorema del cálculo. Pero con 's' como teorema podemos establecer la siguiente demostración:

L1	s	Supuesto teorema.
L2	$\sim s$	$R\alpha_1: s/\sim s; L1.$

Así pues, si (7), que no es demostrable en el cálculo, se añade al sistema de teoremas del mismo, entonces el cálculo da también como teoremas, a la vez, 's' y ' $\sim s$ ', lo que quiere decir que se hace inconsistente (cfr. 59: la afirmación de una fórmula atómica es ya una inconsistencia). Por tanto, la parte del cálculo constituida por el cálculo de enunciados es completa.

Completud del cálculo de predicados de primer orden. — La argumentación cuyo núcleo vamos a considerar aquí en esbozo muestra la completud del cálculo de predicados de primer orden en sentido semántico. La tarea consiste en mostrar que, dada cualquier fórmula formalmente verdadera del lenguaje del cálculo, esa fórmula es un teorema del mismo, es demostrable en el mismo. La idea básica de la demostración (demostración debida originalmente a K. Gödel, 1930) es que para toda fórmula, X, del cálculo de predicados de primer orden es posible construir una fórmula, Y, del cálculo de enunciados (a la que llamaremos 'la reducida de X'), tal que si

X es una verdad formal, entonces también lo es Y , y tal que, además, X es demostrable en el cálculo a partir de Y . Ahora bien: puesto que el cálculo de enunciados es completo, si Y es una verdad formal, Y es demostrable, es un teorema. Mas como X es demostrable en el cálculo a partir de Y , entonces también X es demostrable en el cálculo, es un teorema. Con lo que queda demostrado que el cálculo de predicados de primer orden es completo: pues una verdad formal cualquiera, X , expresable en su lenguaje, resulta ser un teorema del cálculo.

La dificultad técnica principal está en la construcción de la reducida con todas las (infinitas) variables. La línea general del procedimiento puede ejemplificarse así: dada una fórmula, X , del cálculo de predicados de primer orden, puesta en forma normal prenexa, con todos los particularizadores delante de todos los generalizadores (forma normal prenexa de Skolem).

$$(9) (X =) \quad \exists y(x)Rxy,$$

se indica la construcción de una fórmula sin cuantificadores, Y , en la que, por tanto, no hay más estructura lógica relevante que la de los funtores veritativos, y que es una disyunción de todas las fórmulas predicativas atómicas resultantes de expresar X con todos los pares (en nuestro ejemplo diádico) de variables:

$$(10) (Y =) \quad Rx_1y_1 \vee Rx_1y_2 \vee \dots \vee Rx_ny_n \vee \dots$$

(La construcción de esta reducida tiene que basarse en medidas de precaución especiales para que no se mezclen las variables que estaban particularizadas con las que estaban generalizadas.)

Luego se razona así: si X es una verdad formal, entonces también tiene que serlo Y . Pero la lógica de enunciados es completa. Luego Y tiene que ser un teorema de HB.

El tercer paso de la demostración consiste en mostrar que X es demostrable en HB a partir de Y . Se procede del modo siguiente: Y es una disyunción. Para que sea formalmente verdadera basta con que lo sea alguno de los miembros de la disyunción. Tomemos uno de esos componentes principales verdaderos de (10), por ejemplo, Rx_ny_n . Entonces se construye la siguiente demostración:

L1	Rx_ny_n	Teorema.
L2	$s \rightarrow Rx_ny_n$	$RA_{\gamma_3}; L1.$
L3	$[p \vee p \rightarrow p] \rightarrow Rx_ny_n$	$RA_1: s/p \vee p \rightarrow p; L2.$
L4	$[p \vee p \rightarrow p] \rightarrow (x)Rxy_n$	$R_{\gamma_1}; L3.$

(x_n pierde el subíndice porque ya no tiene sentido: al ser generalizado el operando respecto de ' x ', la fórmula vale ya para cualquier variable que ocupe el primer lugar de argumento de ' R ').

L5	$[p \vee p \rightarrow p] \rightarrow (y)(x)Rxy$	$R_{\gamma_1}; L4.$
L6	$p \vee p \rightarrow p$	A1.
L7	$(y)(x)Rxy$	$R_{\beta}; L5, L6.$
L8	$(x)Px \rightarrow Py$	A5.
L9	$(y)(x)Rxy \rightarrow (x)Rxx$	$RA_3: P_1/(x)R_2; L8.$
L10	$Py \rightarrow \exists xPx$	A6.
L11	$(x)Rxx \rightarrow \exists y(x)Rxy$	$RA_3: P_1/(x)R_2; L10.$
L12	$(y)(x)Rxy \rightarrow \exists y(x)Rxy$	$RA_{\beta_1}; L9, L11.$
L13	$\exists y(x)Rxy$	$R_{\beta}; L7, L12.$

Según ese ejemplo se establece que cualquier fórmula del cálculo de predicados de primer orden que — como (9) en el ejemplo — se sepa formalmente verdadera por información ajena al cálculo, es también un teorema del cálculo. Éste es pues completo.

En el caso de nuestro universo finito, con $\Omega = \{a, b\}$, ese esquema de argumentación nos da inmediatamente una demostración de la completud del cálculo para el universo de Ω . La demostración es trivial, pues la reducida de (9) puede escribirse directamente y es, además, una mera forma de escribir (9); es (9) para Ω :

$$(11) \quad [Raa \wedge Rba] \vee [Rab \wedge Rbb].$$

Si (9) es una verdad formal, entonces también lo es (11). Pero (11) es una fórmula sin cuantificar, una fórmula del cálculo de enunciados. Luego (11) es un teorema puesto que el cálculo de enunciados es completo. Y como (11) es sólo una forma de escribir (9) para Ω , (9) es también un teorema.

Con las anteriores reflexiones, que sólo son ilustraciones plausibles de la cuestión, tendremos suficiente por lo que hace a la completud del algoritmo lógico elemental.

64. La lógica elemental y el programa algorítmico. — La completud del cálculo lógico elemental permite pensar que este algoritmo cumple en algún sentido el ideal algorítmico, la mecanización de la inferencia deductiva a la cual ha aspirado una larga tradición filosófica, desde Ramón Lull (cfr. 18). Por ser completo, este algoritmo puede suministrar todas las consecuencias deductivas de las verdades fundamentales de cualquier discurso temático al que sea aplicable. La intervención del trabajo intelectual creador se limitaría a suministrar a esa máquina completa de deducir (una vez construida) las verdades fundamentales del campo que interesara.

Pero también hemos visto (cfr. 37) que la lógica elemental no basta para formalizar, no ya las argumentaciones morales que quería algoritmizar Leibniz, sino ni siquiera las afirmaciones fundamentales de la aritmética. En la lógica elemental no es, en efecto, posible formalizar el 5.º de los axiomas aritméticos de Peano, el llamado 'principio de inducción', que requiere la cuantificación de un símbolo predicativo.

Esta limitación de la lógica elemental tiene importancia para la mayoría de las ciencias positivas, en la medida en que éstas utilizan como un instrumento más o menos esencial la matemática, cuya fundamentación se busca en la aritmética. Por eso es interesante considerar las propiedades algorítmicas del lenguaje lógico capaz de formular las afirmaciones fundamentales de la aritmética. Se trata de la lógica de predicados de orden superior, empezando por el segundo.

CAPÍTULO XII

LA LÓGICA DE PREDICADOS DE ORDEN SUPERIOR Y EL TEOREMA DE INCOMPLETUD DE GÖDEL

65. La lógica de predicados de orden superior. — La lógica de predicados de orden superior se introdujo como aquella en la cual se admite la cuantificación de símbolos predicativos (cfr. 35, 36, 37). Al admitirse su cuantificación, los símbolos predicativos empiezan a ser propiamente variables. Por otra parte, cuando un símbolo es cuantificable, ello significa que puede ser usado en posiciones sintácticas de sujeto. El ejemplo del 5.º axioma de Peano, que fue precisamente el enunciado que nos llevó a interesarnos por la lógica de predicados de orden superior, permite apreciar esta circunstancia. Dicho axioma puede entenderse como definición de la noción de propiedad de todo número natural. El axioma dice, en efecto: 'una propiedad, P , tal que (1.º) la posee el módulo elegido (cero o uno) y (2.º), si la posee cualquier número, x , también la posee el siguiente de x , es una propiedad de todos los números naturales'. El axioma en cuestión justifica, por tanto, frases como: 'la propiedad Q es una propiedad de todos los números naturales'. Usando ' M ' para simbolizar la propiedad ser-propiedad-de-todos-los-números-naturales, esa frase puede simbolizarse así:

MQ ,

expresión simbólica en la cual el símbolo predicativo ' Q ' ocupa una posición sintáctica de sujeto. Si admitimos como base de la discusión un campo de individuos constituido por los números naturales — es decir, si admitimos provisionalmente que los números naturales sean objetos individuales —, entonces el predicado ' Q ' es un predicado de individuos, el predicado ' M ' un predicado de predicados de individuos. También se dice que la propiedad Q y el predicado ' Q ' son de *orden lógico* 1, la propiedad M y el predicado ' M ' de orden lógico 2, y los individuos y sus símbolos — como ' a ', ' x ', etc. — de orden lógico 0.

Naturalmente, nada impide introducir también predicados de predicados de predicados, símbolos de orden lógico 3, 4, etc. Esta será en realidad

la estructura más naturalmente íntegra de la lógica de predicados como marco general de la forma del conocimiento: pues en éste no existen limitaciones en ese sentido de progresiva elevación del nivel de la predicación.

En la lógica de predicados de orden superior coexisten, según esto, símbolos de diversos órdenes lógicos, todos los cuales, excepto los de orden 0, pueden aparecer en posiciones de predicado y en posiciones de sujeto. Los de orden 0 sólo pueden usarse en posiciones de sujeto. Pero si se utiliza sin precauciones esa posibilidad, se producen paradojas como las ya conocidas. La razón de ello es que la reflexividad, que, según vimos, era un rasgo común a muchas formaciones lingüísticas paradójicas, se presenta fácilmente al usar un símbolo de orden o *tipo* lógico n como predicado de símbolos del mismo orden lógico n . Este es un uso tipológicamente reflexivo. Por ejemplo, si es lícito ese uso tipológicamente reflexivo, no estará prohibido por ninguna regla de formación el escribir una fórmula como

$$(1) \quad PQ,$$

' Q es P ', aunque ' P ' y ' Q ' sean del mismo tipo lógico. Pero entonces tampoco estará mal formada una expresión como (1) aunque ' P ' sea el predicado especial que consiste en la negación del símbolo al cual se aplica. ' P ' se definirá así:

$$'PQ' \leftrightarrow_{ar} \sim QQ'.$$

' $\sim Q$ ' y ' Q ', como ' P ' y ' Q ' en (1), pertenecen al mismo tipo lógico, al modo como pertenecen al mismo tipo lógico los adjetivos 'par' e 'impar', ambos aplicables a los mismos sujetos (nombres de números naturales).

Pero si esas formaciones son lícitas, entonces en base a una regla de sustitución que siempre será necesaria en el cálculo de predicados (el análogo de la regla $R\alpha_2$ de HB para la lógica de predicados de primer orden), se podrá obtener de la definición de ' P ' la siguiente fórmula:

$$(3) \quad PP \leftrightarrow \sim PP.$$

Esta contradicción, llamada 'paradoja de Russell', se debe a que hemos admitido como expresión bien formada, como fórmula de la lógica de predicados de orden superior, una expresión en la cual el predicado y el sujeto eran del mismo tipo lógico. Con eso hemos admitido la posibilidad de formaciones lingüísticas reflexivas del tipo de las paradojas.

Por esta razón Russell y Whitehead construyeron la lógica de predicados según la llamada '*teoría de los tipos lógicos*', la cual, en sus dos variantes principales, consiste sustancialmente en esto:

- (1.º) imponer una clasificación sistemática de los símbolos según sus tipos lógicos;
- (2.º) establecer la regla de formación según la cual una expresión predicativa no es una fórmula más que si los símbolos que ocupan

posiciones de sujeto respecto de símbolos de orden n son de tipo $n-1$. (Esto es lo que no se cumple en (1), pues ' P ' y ' Q ' se supusieron del mismo orden lógico.)

Debe observarse que la paradoja de Russell supone una comprensión lingüística de las fórmulas. Si (3) no se entendiera como una equivalencia entre enunciados o fórmulas, no sería necesario considerarla incorrecta. (H. B. Curry.) 'Tipo' y 'orden' no son sinónimos, pero aquí los usaremos simplificadoramente como tales.

La distinción entre los tipos puede hacerse mediante índices, que antepondremos a los símbolos en su parte superior, por ejemplo: 1P , 4Q .

La siguiente expresión no es una fórmula:

$$\exists {}^1P [{}^1P {}^1Q \vee {}^1Px].$$

Lo es en cambio la que sigue:

$$\exists {}^3P [{}^3P {}^2Q \vee {}^3P {}^2R].$$

Tal es la estructura de la lógica de predicados de orden superior. Y como en ella sí que parecen poder formularse las afirmaciones fundamentales de la aritmética, cuyo interés indirecto para las demás ciencias positivas hemos recordado ya, interesa plantearse la pregunta de si esa lógica puede o no formalizarse en un cálculo completo. El matemático y lógico antes citado a propósito de la completud de la lógica elemental, Kurt Gödel, demostró en 1931 que el cálculo de predicados de orden superior es incompleto. En 67 veremos un esbozo de su argumentación. Pero antes será necesario considerar en 66 una técnica introducida por el propio Gödel que legitima su demostración. Esta técnica, llamada 'gödelización' o 'aritmetización', permite representar la lógica de predicados en el lenguaje de la aritmética, es decir, mediante afirmaciones sobre números naturales.

66. La gödelización. — La gödelización se basa en el teorema fundamental de la aritmética que enseña que la descomposición de un número en factores primos es unívoca. La técnica consiste en representar toda fórmula de la lógica de predicados por un número que es un producto de factores primos afectados por exponentes.

Se empieza por asignar a cada símbolo elemental un número, llamado 'número de Gödel' del símbolo. Un ejemplo puede ser la tabla siguiente:

TABLA I. — Números de Gödel

Símbolo	Número de Gödel	Símbolo	Número de Gödel
p	1	1P	12
q	2	1Q	13
$\sim$	3	${}^1[$	14
$\vee$	4	${}^1]$	15
$\wedge$	5	E	16
$\rightarrow$	6	$($	17
$\leftrightarrow$	7	$)$	18
x	8	2P	19
y	9	.	.
a	10	.	.
A	11	.	.

El número de Gödel de una fórmula se obtiene del modo siguiente: dada una fórmula, por ejemplo.

$$(4) \quad p \vee q,$$

se numeran los lugares de sus símbolos, asignándoles números primos correlativamente a partir de 2:

$$\begin{array}{ccc} p & \vee & q \\ 2 & 3 & 5 \end{array}$$

Luego se afecta a cada uno de esos números primos con un exponente que es el número de Gödel del símbolo que ocupa el lugar correspondiente al número primo. El producto de las potencias así formadas es la expresión factorizada del número de Gödel de la fórmula. Así, el número de Gödel de (4) es:

$$2 \cdot 3^4 \cdot 5^2 = 4050.$$

Dado, a la inversa, el número de Gödel de una fórmula, el teorema fundamental de la aritmética garantiza que se hallará la fórmula representada por el número, ya que la descomposición de éste en factores primos es unívoca. Por ejemplo, dado el número de Gödel 24 y dada la Tabla I de los números de Gödel de los símbolos elementales, la descomposición de 24 en factores primos,

$$24 = 2 \cdot 2 \cdot 2 \cdot 3 = 2^3 \cdot 3,$$

indica que la fórmula con número de Gödel 24 es:

$$(5) \quad \sim p.$$

El número de Gödel de una demostración se obtiene del modo siguiente: dada la demostración, que es una sucesión de fórmulas, se obtiene el nú-

mero de Gödel de cada fórmula. Luego se numeran las fórmulas como antes los símbolos elementales de éstas para construir el número de Gödel de las mismas; o sea, se asigna a cada fórmula de la demostración un número primo, correlativamente a partir de 2, por el orden en que las fórmulas aparecen en la demostración. A continuación se forman las potencias cuya base es el número (primo) de orden de la fórmula y cuyo exponente es el número de Gödel de la fórmula. El número de Gödel de la demostración es el producto de esas potencias. Por ejemplo, dada la demostración en el cálculo de la deducción natural

$$\begin{array}{lll} L1 (0) & p & \text{Supuesto teorema} \\ L2 (0) & p \vee q & R\vee; L1, \end{array}$$

su número de Gödel, n , es

$$n = 2 \cdot 3^{4050}.$$

Dado, a la inversa, el número n , el teorema fundamental de la aritmética garantiza que obtendremos la expresión

$$2 \cdot 3^{4050}.$$

Ante ella sabemos inmediatamente que se trata de una demostración, y no de una sola fórmula, porque el número 4050 no corresponde a ningún símbolo elemental, sino a una fórmula. La Tabla I de los números de Gödel de los símbolos elementales nos da como primera fórmula de la demostración ' p '. La descomposición de 4050 en factores primos nos da como segunda fórmula de la demostración ' $p \vee q$ '.

En realidad, la gödelización es más complicada que el resumen que se acaba de leer: pues en ese resumen no se tiene en cuenta la necesidad de reservar números de Gödel para infinitas variables de cada tipo (que es la única manera de garantizar que 4050, en nuestro ejemplo, no es el número de Gödel de ningún símbolo elemental). Pero esto se consigue fácilmente de un modo que no altera en nada la esencia de la técnica tal como queda expuesta.

Lo esencial de la gödelización es esto: la lógica de predicados (de orden superior) permite expresar la aritmética. En lógica de predicados pueden construirse, pues, expresiones que hablan de números, los relacionan, etc. Esas expresiones constituirían (si dieran lugar a un cálculo completo) una formalización de la aritmética. Por otra parte, dada la infinitud de órdenes de la lógica de predicados, en ella pueden también expresarse afirmaciones acerca de las afirmaciones aritméticas: por ejemplo, que tal afirmación aritmética es verdadera o falsa, o que es demostrable, o que es indemostrable. Diremos, para abreviar, que en la lógica de predicados puede expresarse también la "metaaritmética". Pues bien: la gödelización de la lógica de predicados permite a su vez representar las expresiones de la lógica

de predicados — incluidas, naturalmente, las *metaaritméticas* — por expresiones aritméticas, de modo que una expresión metaaritmética sobre la verdad, falsedad, demostrabilidad o inde demostrabilidad de una expresión aritmética podrá representarse a su vez por una expresión aritmética. En esta circunstancia se basa la argumentación de Gödel, cuyo hilo nos interesa seguir aproximativamente en 67.

Antes, sin embargo, vamos a tener que introducir una convención sobre nuestro vocabulario básico, sobre nuestros símbolos elementales, para poder hablar de números naturales en el cálculo. Pues es claro que la gödelización nos va a obligar a referirnos a ellos. La convención va a ser la siguiente: entenderemos que la constante individual 'a' (que en nuestra Tabla I tiene el número de Gödel 10) es el nombre (la cifra) del módulo de la sucesión de los números naturales, tomando como tal a cero; y, además, que el símbolo predicativo constante 'A' es el functor 'siguiente'; la función siguiente-de, aplicada a un número como argumento, tiene como valor el siguiente de ese número en la sucesión de los números naturales; p. e.:

$$A5 = 6.$$

'A' puede entenderse también como un descriptor, y 'A5' como una descripción. 'A5' puede leerse:

$$(1x)(x \text{ es el siguiente de } 5).$$

La notación 'A' nos ahorra, simplemente, espacio, y hará más visualizables las fórmulas.

De este modo, todos los números naturales son expresables con el vocabulario básico del cálculo, y, más precisamente, con la parte de dicho vocabulario básico a la que hemos asignado convencionalmente números de Gödel en nuestra Tabla I. '3', por ejemplo, se expresará por

$$[A[A[Aa]]]$$

(el siguiente del siguiente del siguiente de a). Como no hay ninguna confusión que temer con estas expresiones, prescindiremos de los corchetes, escribiendo, por ejemplo, en el caso anterior:

$$AAAa.$$

Es claro que las cifras de números naturales tienen también cada una su número de Gödel. El de '3', por ejemplo, es:

$$2^{11} \cdot 3^{11} \cdot 5^{11} \cdot 7^{10}.$$

67. El teorema de incompletud de Gödel. — La argumentación de Gödel demuestra que ya la parte de la lógica de predicados necesaria para formular la aritmética es incompleta, en el sentido de que existe al menos

una fórmula de la teoría aritmética formulable en ella la cual es verdadera (cosa que se establece por un razonamiento metalógico, "metaaritmético") y que, sin embargo, no puede ser demostrada en la formulación de la aritmética misma en la lógica de predicados. El razonamiento procede en sustancia del modo siguiente:

Paso I: Consideración de la situación de partida. — a) Se empieza por admitir que el cálculo de predicados es consistente. En otro caso, como es natural, no tiene sentido discutir acerca de su completud, pues un cálculo inconsistente lo demuestra todo, tanto lo verdadero cuanto lo falso. Es, por así decirlo, hipercompleto.

b) La gödelización permite representar por fórmulas aritméticas las expresiones metaaritméticas, por ejemplo, expresiones que hablen de la demostrabilidad o inde demostrabilidad de fórmulas aritméticas.

Paso II: La relación de demostrabilidad. — Atendamos a una de esas relaciones entre fórmulas aritméticas, por ejemplo, la que existe entre una sucesión de fórmulas, X , y la fórmula Y cuando X es una demostración de Y . En la metaaritmética (sintaxis de la aritmética) esto podrá expresarse así, por ejemplo:

$$(6) \quad DXY.$$

Por la gödelización, X , que es una demostración, será representable por un número de Gödel, n , e Y , que es una fórmula, por otro, m . La relación metaaritmética D entre la sucesión de fórmulas X y la fórmula Y será entonces representable en la aritmética por una fórmula aritmética con un símbolo relacional, D , y los números n y m :

$$Dnm.$$

' Dnm ' es propiamente un enunciado; no tiene variables, y ' D ' es una constante, el nombre de una determinada relación aritmética. La fórmula que esquematiza enunciados de ese género (enunciados de demostrabilidad) es la fórmula aritmética

$$(7) \quad Dxy.$$

Hay que tener bien presente, porque ello es esencial para la demostración, que (7) y todas las fórmulas que obtengamos de ella en su mismo lenguaje son fórmulas aritméticas, no metaaritméticas. D es una relación aritmética entre números, como pueda serlo la divisibilidad. D es la relación aritmética que media entre dos números, x e y , cuando el primero es el número de Gödel de una demostración de la fórmula cuyo número de Gödel es y . Por eso al decir que ' Dxy ' es una 'fórmula de demostrabilidad' hablamos laxamente: ' Dxy ' representa más bien en la aritmética a una afirmación metaaritmética de demostrabilidad. Pero ' Dxy ' misma lo único que

afirma es la existencia de la relación aritmética D entre x e y . La distinción entre lenguaje y metalenguaje está pues respetada en fórmulas como ' Dxy '.

La siguiente fórmula

$$(8) \quad \sim \exists x Dxy$$

es entonces la fórmula aritmética que representa la afirmación metaaritmética: 'la fórmula cuyo número de Gödel es y es indemostrable'. Literal y aritméticamente, por sí misma, (8) dice: 'no hay ningún número, x , tal que x esté en la relación D con y '.

Paso III: Construcción de una fórmula aritmética indemostrable en el cálculo. — Sea b el número de Gödel de una fórmula. Supongamos que esa fórmula contiene la variable ' y ', cuyo número de Gödel es en nuestra Tabla I 9. Entonces, con una notación que ya hemos utilizado en otros contextos, escribiremos

$$[b/9] b$$

para indicar el número que resulta de b al sustituir, en la expresión factorizada de éste, la cifra del número de Gödel 9 por la cifra del número de Gödel de ' b '. En vez de 'la cifra del número de Gödel 9' diremos más brevemente: 'la cifra 9'. ' $[b/9] b$ ' puede entenderse también como una descripción.

Según esto, la fórmula

$$(9) \quad \sim \exists x Dx[y/9] y$$

puede parafrasearse por: 'no hay ningún número, x , tal que x esté en la relación D con el número que resulta de sustituir, en la expresión factorizada de y , la cifra '9' por la cifra del número de Gödel de ' y '. (9) es por tanto la representación aritmética de la siguiente afirmación metaaritmética: 'es indemostrable en el cálculo la fórmula cuyo número de Gödel resulta del número de Gödel y , al sustituir en la expresión factorizada de éste la cifra '9' por la cifra del número de Gödel de ' y '.

Dos observaciones para precisar el significado de esa notación: a) hablamos de 'cifras' porque una sustitución es una operación que se realiza sobre expresiones escritas; lo que se sustituye son pues símbolos, en este caso nombres de números, cifras, y no los números, que son conceptos, entidades abstractas; b) si en la expresión factorizada del número y no aparece la cifra '9', $[y/9]y$ es naturalmente el mismo y .

(9) es una fórmula de nuestra aritmética expresada en el cálculo de predicados. (En realidad, no lo es tal como está escrita, pues en nuestro vocabulario básico no está el símbolo '9', ni tampoco hemos introducido

a éste por definición. Pero su definición no ofrece ninguna dificultad con ' A ' y ' a '. '9' es una abreviatura de

$$AAAAAAAAAa,$$

expresión que efectivamente pertenece a nuestro cálculo de predicados, y que tiene el número de Gödel

$$g = 2^{11} \cdot 3^{11} \cdot 5^{11} \cdot 7^{11} \cdot 11^{11} \cdot 13^{11} \cdot 17^{11} \cdot 19^{11} \cdot 23^{11} \cdot 29^{10}.)$$

Puesto que es una fórmula de la aritmética formulada en el cálculo de predicados, (9) tendrá su número de Gödel. Sea n ese número. Para expresarlo en forma factorizada, llamemos ' g ' al número de Gödel de la cifra '9', ' d ' al número de Gödel de ' D ' y ' f ' al número de Gödel del símbolo '/'. Como vamos a necesitar usar estos números de Gödel no calculados, y aun otro, los reuniremos en la siguiente

TABLA II. — Números de Gödel

Símbolo	Número de Gödel
9	g
D	d
/	f
n	c

Entonces el número de Gödel de (9) es:

$$(10) \quad n = 2^3 \cdot 3^{10} \cdot 5^8 \cdot 7^2 \cdot 11^8 \cdot 13^{14} \cdot 17^9 \cdot 19^7 \cdot 23^9 \cdot 29^{15} \cdot 31^9.$$

Basándonos en la fórmula (9) vamos a construir ahora otra fórmula aritmética, (11), que sólo diferirá de (9) en que, en el lugar de la variable ' y ' que aparece en (9), tendrá una concreta cifra. Y esa cifra va a ser precisamente ' n ', la cifra del número n . La fórmula es pues:

$$(11) \quad \sim \exists x Dx[n/9] n.$$

(11) es una fórmula aritmética que dice: 'no hay ningún número, x , que esté en la relación D con el número que resulta de sustituir, en la expresión factorizada de n , la cifra '9' por la cifra del número de Gödel de ' n '. (11) representa por tanto la afirmación metaaritmética: 'es indemostrable en el cálculo la fórmula cuyo número de Gödel, $[n/9]n$, resulta de n , al sustituir, en la expresión factorizada de éste, la cifra '9' por la cifra del número de Gödel de ' n '.

Nos interesa ahora saber cuál es esa fórmula la afirmación de cuya indemostrabilidad en el cálculo está representada por (11). Llamemos, para

abreviar, X a esa fórmula. (11) representa pues en el cálculo la afirmación metalógica: 'X es indemostrable en el cálculo'.

Para averiguar cómo es X podemos construir su número de Gödel en expresión factorizada. Como sabemos, el número de Gödel de X resulta del número n , sustituyendo, en la expresión factorizada de éste, la cifra '9' por la cifra del número de Gödel de ' n '. En la expresión factorizada de n (que es (10)), '9' no aparece más que dos veces: como exponente de la cifra '17' y como exponente de '31'. Llamemos ' c ' a la cifra del número de Gödel de ' n '. Entonces la expresión factorizada del número de Gödel de X , m , es:

$$(12) \quad m = [n/9]n = 2^3 \cdot 3^{16} \cdot 5^3 \cdot 7^4 \cdot 11^8 \cdot 13^{14} \cdot 17^c \cdot 19^f \cdot 23^g \cdot 29^{15} \cdot 31^c.$$

Pero ahora resulta que m es el número de Gödel de una fórmula que ya conocemos, a saber, de (11), como debe comprobar el lector por inspección de (11) teniendo en cuenta las Tablas I y II de números de Gödel.

X es pues (11). Consiguientemente, (11) es la representación aritmética de la afirmación metaaritmética: 'la fórmula (11) es indemostrable en el cálculo'.

Es muy importante comprobar que (11) no tiene el defecto de estructura del enunciado causante de la paradoja de Epiménides, a pesar del parecido entre este enunciado y (11). Pues (11) no afirma de sí misma que es falsa, ni que es indemostrable. Lo que afirma, como fórmula aritmética que es, es que no hay ningún número, x , que esté en la relación aritmética D con el número m . Y como la relación D es aritmética, aquí no hay ninguna mezcla de niveles lingüísticos, como en la paradoja de Epiménides, cuyo enunciado, al afirmar falsedad, hace una afirmación metalingüística. Sólo en el metalenguaje sabemos que m es el número de Gödel de (11). En el cálculo mismo, no lo sabemos, ni podemos, por tanto, hacer ningún uso de esa circunstancia. En cambio, en la paradoja de Epiménides el mismo enunciado paradójico contiene—es—la afirmación de falsedad. (11) no es una afirmación de indemostrabilidad, sino que representa una afirmación de indemostrabilidad hecha en otro lenguaje.

Vamos a ver ahora que (11) es indemostrable en el cálculo, que no es un teorema de la aritmética formalizada en el cálculo de predicados. Supongamos que lo fuera. Entonces, es también un teorema del cálculo la afirmación aritmética que dice: 'hay al menos un número, x , que está en la relación D con el número m '; esa afirmación aritmética representa, en efecto, la supuesta verdad metaaritmética que acabamos de sentar: '(11) es demostrable en el cálculo'. Así pues, si (11) es un teorema, si (11) es demostrable en el cálculo, lo es también

$$\exists x Dxm,$$

o, más detalladamente

$$(13) \quad \exists x Dx[n/9]n.$$

Pero (11) y (13) juntas constituyen una contradicción. Esta contradicción se ha producido por admitir que (11) es demostrable. Por tanto, como el cálculo de predicados debe construirse de modo que sea consistente (ésa es nuestra condición de partida), (11) no puede ser un teorema del mismo, una fórmula demostrable en él (en la aritmética formulada en él).

Ahora bien: si a pesar de ello (11) es una fórmula verdadera, quedará demostrado que la formalización de la aritmética en el cálculo de predicados es incompleta.

Paso IV: La fórmula (11) es verdadera. — La verdad lógica de (11) está ya implícitamente establecida por las consideraciones del *Paso III*. Pero vale la pena volver a establecerla aquí explícitamente: como la formalización de la aritmética en el cálculo de predicados tiene que ser consistente, entonces vale la afirmación metaaritmética '(11) es indemostrable', o '(11) no es un teorema del cálculo', afirmación metalógica que prohíbe afirmar como teorema en el cálculo una fórmula que da lugar a una contradicción, a la inconsistencia del cálculo. Pero (11) es precisamente la representación aritmética, en el cálculo, de esa verdad metaaritmética. Luego (11) es verdadera a pesar de no ser un teorema. Consecuentemente, la formulación de la aritmética en el cálculo de predicados es incompleta.

Lo importante de este *Paso IV* es que la argumentación por la cual se establece la verdad de (11) es una argumentación metalógica respecto de la formulación de la aritmética en la lógica de predicados.

Este resumen de la marcha de la argumentación de Gödel es una variante de la esquematización de la versión de J. B. Rosser (1936) dada por E. Nagel y J. R. Newman en 1958. La principal diferencia está en la construcción y lectura de expresiones descriptivas como ' $[n/9]n$ '.

68. Sobre la significación del teorema de incompletud de Gödel para la teoría de la ciencia. — El teorema de incompletud de Gödel enseña por de pronto que *toda formalización de la aritmética en el cálculo de predicados es incompleta*. Como el cálculo de predicados, sin limitación de orden, es el algoritmo lógico más potente, puede decirse, de un modo más general, que *toda formalización de la aritmética es incompleta*.

Pero el intento de formalización de la aritmética se realiza con medios puramente lógicos. Vimos ya que desde la teoría de conjuntos de Cantor, y luego por obra de Frege y Russell y Whitehead, el concepto de número natural se construye con la idea lógica de clase o conjunto (cfr. cap. XIV); también veremos que puede construirse también con ayuda de la idea puramente lógica de relación (cfr. cap. XV). Consiguientemente, la incompletud de la formalización de la aritmética es una incompletud del instrumento formalizador mismo, o sea, del algoritmo lógico. De aquí que el resultado de Gödel pueda entenderse también así: *el cálculo de predicados es incompleto*. (Todos estos resultados se entienden con la condición de

nuestro punto de partida, a saber, que el cálculo de predicados es consistente.)

La lógica de predicados sin limitación de orden es aquella en la cual se intenta (sin éxito) formalizar la deducción para cualquier tipo de conocimiento que sea al menos de la complejidad de la aritmética. Y por debajo de la complejidad de la aritmética debe haber, puede pensarse, muy poco conocimiento teórico de interés. De aquí que, aún más laxamente, el teorema de Gödel haya podido entenderse también en el siguiente sentido filosófico: *la lógica es incapaz de formalizar la deducción necesaria para fundamentar definitivamente cualquier conocimiento de algún interés teórico.*

Por este camino de interpretación cada vez más laxa y vaga del teorema de incompletud de Gödel, algunos filósofos han llegado a afirmar que el resultado de Gödel demuestra "el fracaso de la lógica" o hasta "el fracaso de la razón". Estas afirmaciones carecen de fundamento, como puede verse por las siguientes consideraciones.

En primer lugar, lo único que demuestra el teorema de Gödel es que resulta imposible conseguir un conjunto de axiomas y un juego de reglas de transformación que suministren todas las verdades formales expresables en el lenguaje de la lógica de predicados. Esto, naturalmente, no excluye que los criterios logicoformales en particular, y los criterios racionales en general, sigan valiendo. Ellos son los que se aplican en las argumentaciones metalógicas que resultan necesarias por la incompletud del cálculo mismo. En principio, las fórmulas verdaderas indemostrables en un cálculo a partir de los axiomas del mismo pueden demostrarse metalógicamente con los mismos "criterios lógicos", como se hace, por ejemplo, en la argumentación de Gödel (cfr. *Paso IV* de 67). Claro que la demostración metalógica planteará a su vez el problema de la incompletud del razonar metalógico mismo: también en ese metalenguaje habrá verdades indemostrables, que serán demostrables en el metalenguaje de ese metalenguaje; y así sucesivamente. El proceso es, desde luego, ilimitado. Y ello muestra que el programa o ideal algorítmico, en el sentido de calculización de toda deducción, no es realizable. Pero eso sólo quiere decir que el sistema formal no se contiene ni se justifica nunca totalmente a sí mismo. Lo cual puede ser acaso molesto para filósofos idealistas como Platón, que crean en la autosuficiencia del mundo de los abstractos; pero no para un racionalismo como el aconsejado por la práctica de las ciencias, el cual debe ver en la experiencia con el mundo real la justificación de las formaciones abstractas, justificación siempre provisional y relativa.

En segundo lugar, el hecho de que la lógica misma haya descubierto y demostrado los límites o la inviabilidad de una realización universal del programa algorítmico en su forma clásica, es más bien un éxito que un fracaso de la actividad capaz de tal resultado. El resultado mismo significa que el pensamiento racional puede saber cuáles de sus actividades son

algoritmizables, ejecutables (en principio) mecánicamente, y cuáles no: cuáles son, como suele decirse, trabajo racional mecánico, y cuáles trabajo racional creador. Fracaso del pensamiento es más bien la situación en la cual el pensamiento no sabe cuál es el alcance de su actividad, como suele ocurrir, dicho sea de paso, a muchos filósofos.

En tercer lugar, debe observarse que la incompletud de un cálculo lógico tomado en toda su dimensión no excluye la completud de cálculos parciales contenidos por él. Es, por lo pronto, completo el cálculo lógico elemental, es decir, el formado por el cálculo de enunciados y el de predicados de primer orden. Pero además — y esto es lo más importante en la práctica — es posible construir en forma de cálculos completos partes de la lógica de predicados de orden superior.

En cuarto lugar, por lo que hace a la aritmética misma, debe observarse que los enunciados cuya indemostrabilidad establece la argumentación de Gödel no son del mismo estilo, por así decirlo, que los teoremas clásicos de la aritmética, los cuales se refieren a operaciones con números y son los realmente utilizados en la aplicación a otras ciencias o a la técnica. El enunciado (11) de 67 se parece, en efecto, muy poco al teorema de la descomposición de los números en factores primos, por ejemplo. Para estos teoremas de tipo "clásico" — o sea, para toda la parte "útil" de la aritmética (y de las disciplinas matemáticas basadas en ella, señaladamente el álgebra clásica y el cálculo infinitesimal) — se han construido cálculos (sistemas) que dan de sí todos los teoremas interesantes.

Por último, también puede conseguirse una cierta completud del cálculo de predicados en general, aunque pagando por ella el precio de una cierta ambigüedad semántica del cálculo, pues el sistema permite entonces interpretaciones no primariamente deseadas. Este último punto, establecido por L. Henkin (1947, 1950), no va a interesarnos aquí, pero debe tenerse en cuenta cuando se considera la significación del teorema de Gödel para la teoría de la ciencia.

69. El teorema de incompletud de Gödel y el programa de Hilbert. — En el capítulo III (cfr. 18) se advirtió que el programa de Hilbert, que hemos interpretado como versión moderna del viejo ideal algorítmico, tenía propiamente otra finalidad: conseguir una seguridad definitiva acerca de la solidez de los fundamentos de la matemática, empezando por los de la aritmética. El deseo de Hilbert era conseguir esa seguridad construyendo en el seno de la matemática clásica una demostración de su propia consistencia. El teorema de Gödel muestra que también eso es imposible: que es imposible demostrar la consistencia de la matemática (propiamente: de la aritmética) en su mismo lenguaje.

El teorema de Gödel, tal como lo hemos estudiado, consiste en la afirmación

(14) si la aritmética es consistente, entonces es incompleta.

(La condición de que fuera consistente fue nuestra condición de partida.) Vamos a buscar una representación de esa afirmación metalógica en el cálculo, en la formulación de la aritmética en la lógica de predicados: que un sistema es consistente, quiere decir que no toda fórmula de su lenguaje es un teorema del sistema; por ejemplo, que al menos la fórmula ' p ' o la fórmula ' $\sim p$ ' no es un teorema del sistema. De aquí que la consistencia de un cálculo pueda expresarse diciendo que no todas las fórmulas del lenguaje de ese cálculo son teoremas. Para una aritmética formalizada en el cálculo de predicados, esa afirmación metalógica puede *representarse* en el cálculo mismo, mediante la relación aritmética D , por la fórmula aritmética:

$$(15) \quad \sim(y)\exists xDxy,$$

que dice: 'no para todo número de Gödel, y , hay un número de Gödel, x , tal que x esté en la relación D con y '.

Por otra parte, la incompletud de la aritmética formulada en el cálculo de predicados (de la aritmética formalizada) puede representarse también en el cálculo por una fórmula aritmética, que es precisamente (11):

$$(11) \quad \sim \exists xDx[n/9]n.$$

Según esto, la representación de la verdad metalógica (14) en el cálculo es la fórmula aritmética

$$(16) \quad \sim (y)\exists xDxy \rightarrow \sim \exists xDx[n/9]n.$$

Podemos tomar (16) como una fórmula válida, pues es la representación aritmética del teorema de Gödel. Ahora bien: sabemos que (11) no es demostrable en el cálculo. Y como (11) es el consecuente del antecedente (15) en la fórmula válida (16), entonces tampoco puede ser demostrable (15), ya que si lo fuera lo sería también su consecuente (11), por una aplicación de $R\beta$ a (16). Luego (15) no es demostrable. Pero (15) es la representación aritmética de la afirmación de la consistencia de la aritmética formalizada en el cálculo de predicados. Por tanto, la afirmación que representa la consistencia de la aritmética no es demostrable en la formalización misma de la aritmética en el cálculo de predicados, con los meros instrumentos lógicos de esa formalización.

CAPÍTULO XIII

DECIDIBILIDAD EN LA LÓGICA ELEMENTAL

70. Decidibilidad de la lógica de enunciados. — Al tratar de la consistencia de la lógica elemental nos servimos de una interpretación de la lógica de enunciados basada en la noción aristotélica de enunciado apofántico (cfr. 62). Un enunciado apofántico es una formación lingüística que tiene que ser verdadera o falsa. De esa noción infirió Frege que los enunciados pueden ser considerados como nombres múltiples de dos únicos objetos: los valores verdad y falsedad, V y F. El conjunto $\{V, F\}$, constituido por esos dos valores, da según esto todas las interpretaciones posibles de un enunciado (considerado globalmente) en el universo del discurso que le es adecuado. De acuerdo con ello hemos trabajado ya con un sistema de interpretaciones, $\mathfrak{S}$, de las siguientes características:

- a) las fórmulas atómicas de la lógica de enunciados se interpretan en el campo finito de objetos $\{V, F\}$;
- b) el valor de cualquier fórmula molecular dada de la lógica de enunciados queda entonces determinado para cada interpretación por las tablas definitorias de los funtores veritativos.

He aquí un ejemplo de esa determinación: si la fórmula

$$(1) \quad p \vee q$$

se interpreta asignando a ' p ' el valor V y a ' q ' el valor F, la tabla de ' v ' da para (1) el valor V.

Ahora bien: una verdad formal es una fórmula verdadera para toda interpretación posible. Desde el punto de vista lógico, necesario para que tenga sentido hablar de 'verdad' y de 'verdad formal', las interpretaciones posibles quedan limitadas a un universo del discurso que contenga las nociones tradicionalmente lógicas, como las de verdad y falsedad. Consiguientemente, podremos decir que una fórmula de la lógica de enunciados es una verdad formal si y sólo si es verdadera para toda interpretación posible en el campo de valores $\{V, F\}$, que constituye el único universo del discurso relevante.

Apliquemos esta reflexión a la fórmula (1) para ver si es o no es una verdad formal. Las interpretaciones posibles de (1) son cuatro:

- 1.^a $p = V, q = V;$
- 2.^a $p = V, q = F;$
- 3.^a $p = F, q = V;$
- 4.^a $p = F, q = F.$

La tabla de 'v', nos indica que (1) es verdadera para las interpretaciones 1.^a, 2.^a, 3.^a, pero falsa para la 4.^a Consiguientemente, (1) no es una verdad formal.

Lo que acabamos de hacer ha sido *decidir* acerca de la fórmula (1), sin intentar siquiera demostrarla, o demostrar su negación. Esto permite preguntarse si será posible generalizar el procedimiento de decisión para cualquier fórmula de la lógica de enunciados, en cuyo caso ésta sería no sólo completa, como hemos visto, sino, además, decidible.

Tal es el caso, en efecto, como puede hacerse plausible por la siguiente consideración: ninguna fórmula atómica de la lógica de enunciados será una verdad formal, pues para una de sus dos interpretaciones posibles será falsa. En cuanto a las fórmulas moleculares, el número de símbolos contenidos en cualquier fórmula que se nos proponga será finito. Además, por la definición de 'fórmula de la lógica de enunciados', esa fórmula habrá sido compuesta enlazando fórmulas atómicas mediante funtores veritativos. Podremos examinar la fórmula propuesta partiendo de sus componentes moleculares mínimas, llamando así a aquellas que no contienen a su vez componentes moleculares. Las tablas definitorias de los funtores nos suministrarán los valores de esas componentes moleculares mínimas. Esos valores resultarán a su vez enlazados por funtores. Las tablas de éstos volverán a darnos unos valores para las componentes moleculares medias de la fórmula. Así llegaremos a una situación en la cual tendremos los valores a los que se aplica el functor principal de la fórmula propuesta (o los funtores principales de la fórmula propuesta). La tabla de ese functor (o de esos funtores) nos dará el valor de la fórmula en cuestión. Así pues, dada la finitud del número de símbolos de toda fórmula de la lógica de enunciados, y dada la finitud del campo de objetos {V, F}, único relevante para la interpretación, todas las fórmulas de la lógica de enunciados tienen que ser decidibles.

Esa misma argumentación de la decidibilidad de la lógica de enunciados resume la técnica más corrientemente usada para decidir acerca de fórmulas de la lógica de enunciados.

71. La técnica de las tablas veritativas. — Un ejemplo como la decisión sobre (1) en 70 suele presentarse del modo siguiente:

TABLA I

p	q	$p \vee q$
V	V	V
V	F	V
F	V	V
F	F	F

En las primeras columnas (empezando por la izquierda) se registran los valores posibles de las componentes atómicas de la fórmula que se estudia. En las columnas sucesivas se van registrando los valores resultantes para las componentes moleculares mínimas y medias de cada distribución de valores entre las componentes atómicas. En la columna final se registran los valores resultantes para la fórmula que se quiere decidir. En el ejemplo de la Tabla I no había fórmulas moleculares mínimas ni medias diversas de la fórmula a decidir. Pero sí las hay en el siguiente ejemplo:

$$(2) \quad [p \rightarrow q] \rightarrow [[q \rightarrow r] \rightarrow [p \rightarrow r]]$$

TABLA II

p	q	r	$p \rightarrow q$	$q \rightarrow r$	$p \rightarrow r$	$[q \rightarrow r] \rightarrow [p \rightarrow r]$	(2)
V	V	V	V	V	V	V	V
F	V	V	V	V	V	V	V
V	F	V	F	V	V	V	V
F	F	V	V	V	V	V	V
V	V	F	V	F	F	V	V
F	V	F	V	F	V	V	V
V	F	F	F	V	F	F	V
F	F	F	V	V	V	V	V

La fórmula (2), según nos informa la Tabla II, es una verdad formal, una fórmula verdadera para cualquier interpretación lógica — cosa que sabíamos, pues la teníamos demostrada como teorema (TE14).

El número de interpretaciones (valoraciones) posibles para cada fórmula dentro del sistema de interpretaciones $\mathfrak{I}$ es función del número de variables que contiene. Para una fórmula con n variables de enunciado hay 2^n composiciones de valores. Para no olvidarse de ninguna ni perder tiempo buscando las que faltan es conveniente, cuando hay más de dos letras de enunciado, seguir alguna regla práctica en la distribución de

valores; la siguiente, por ejemplo: en la primera columna (empezando por la izquierda) se alternan los valores V y F; en la segunda columna, un valor se escribe dos veces seguidas, y luego el otro también dos veces seguidas (2^1 veces), empezando por el mismo valor con que se haya empezado en la primera columna, por ejemplo; en la tercera columna, un valor se escribe cuatro veces seguidas, y luego el otro también cuatro veces seguidas (2^2 veces); en la cuarta, ocho veces seguidas cada valor (2^3 veces) y así sucesivamente hasta la columna n (2^{n-1} veces seguidas cada valor).

La técnica de las tablas veritativas sirve también, naturalmente, para decidir acerca de la consistencia o la contradictoriedad de fórmulas. La Tabla I, por ejemplo, muestra que la fórmula ' $p \vee q$ ' es (semánticamente) consistente, puesto que tiene algún modelo, aunque no sea una verdad formal. Y la Tabla III muestra que

$$(3) \quad [p \vee q] \wedge [p \vee \sim q] \leftrightarrow \sim p$$

es una contradicción.

TABLA III

p	q	$p \vee q$	$p \vee \sim q$	$[p \vee q] \wedge [p \vee \sim q]$	$\sim p$	(3)
V	V	V	V	V	F	F
F	V	V	F	F	V	F
V	F	V	V	V	F	F
F	F	F	V	F	V	F

72. Reducción del número de las funciones veritativas diádicas.—

En 31 (cap. V), al presentar las funciones veritativas diádicas, se definieron las 16 funciones de esa categoría, designadas entonces con las notaciones ' f_1 ' - ' f_{16} '. Pero sólo se estudiaron f_2 , que es $\vee$, f_4 , que es $\rightarrow$, f_{10} , que es $\leftrightarrow$, y f_{12} , que es $\wedge$. La razón por la cual se prescindió de las demás es que todas ellas son expresables mediante las cuatro funciones veritativas diádicas que hemos usado (y que son las de *Principia Mathematica*), junto con la función monádica $\sim$. Esto es cómodo de ver con la ayuda de tablas veritativas. Las funciones cuya reducibilidad a $\sim$, $\vee$, $\wedge$, $\rightarrow$ y $\leftrightarrow$ tenemos que ver son: f_1 , f_3 , f_5 , f_6 , f_7 , f_8 , f_9 , f_{11} , f_{13} , f_{14} , f_{15} , f_{16} . A continuación se indicará para cada una de esas funciones una posibilidad (no la única) de expresarla por algunas de nuestras cinco funciones. Se añadirá una equivalencia que expresa esa reducibilidad. Mediante una tabla puede decidirse del carácter de verdad formal de cada una de esas equivalencias.

Casos aparte constituyen las funciones f_1 y f_{16} . f_1 puede expresarse por cualquier verdad formal de dos variables, por ejemplo, por ' $[p \rightarrow q] \rightarrow$

$\rightarrow [\sim q \rightarrow \sim p]$ '. f_{16} puede expresarse por cualquier contradicción de dos variables, por ejemplo, ' $[p \vee q] \wedge [p \vee \sim q] \leftrightarrow \sim p$ '.

La función f_3 , utilizada por A. Church y llamada 'condicional inverso': mediante ' $\rightarrow$ ':

$$(4) \quad [p f_3 q] \leftrightarrow [q \rightarrow p].$$

La función f_5 , introducida por Sheffer: mediante ' $\sim$ ' y ' $\vee$ ':

$$(5) \quad [p f_5 q] \leftrightarrow \sim p \vee \sim q.$$

La función f_6 : mediante ' $\vee$ ', ' $\sim$ ' y ' $\wedge$ ':

$$(6) \quad [p f_6 q] \leftrightarrow p \vee [q \wedge \sim q].$$

La función f_7 : mediante ' $\wedge$ ', ' $\sim$ ' y ' $\vee$ ':

$$(7) \quad [p f_7 q] \leftrightarrow [\sim p \wedge q] \vee [\sim p \wedge \sim q].$$

La función f_8 : mediante ' $\sim$ ', ' $\vee$ ' y ' $\wedge$ ':

$$(8) \quad [p f_8 q] \leftrightarrow [p \wedge q] \vee [\sim p \wedge q].$$

La función f_9 : mediante ' $\sim$ ', ' $\vee$ ' y ' $\wedge$ ':

$$(9) \quad [p f_9 q] \leftrightarrow [p \wedge \sim q] \vee [\sim p \wedge \sim q].$$

La función f_{11} , que podría llamarse 'heterovalencia': con ' $\sim$ ', ' $\vee$ ' y ' $\wedge$ ':

$$(10) \quad [p f_{11} q] \leftrightarrow [p \wedge \sim q] \vee [\sim p \wedge q].$$

La función f_{13} : por medio de ' $\sim$ ' y ' $\wedge$ ':

$$(11) \quad [p f_{13} q] \leftrightarrow p \wedge \sim q.$$

La función f_{14} : por ' $\sim$ ' y ' $\wedge$ ':

$$(12) \quad [p f_{14} q] \leftrightarrow q \wedge \sim p.$$

La función f_{15} , usada, como f_5 , por Sheffer: por ' $\sim$ ' y ' $\wedge$ ':

$$(13) \quad [p f_{15} q] \leftrightarrow \sim p \wedge \sim q.$$

Reducción de las cinco funciones veritativas de Principia Mathematica unas a otras. — Al hablar de formas normales de la lógica de enunciados (cfr. 56, 57) vimos que, con la ayuda de las leyes de De Morgan (TE12)-(TE13), de la ley de doble negación (TE1), y de los teoremas (TE30)-(TE31), es posible prescindir de los funtores ' $\leftrightarrow$ ' y ' $\rightarrow$ ', en primer lugar, pues las funciones correspondientes pueden expresarse por medio de ' $\sim$ ', ' $\vee$ ' y ' $\wedge$ '. Y además, que no es necesario contar a la vez con ' $\vee$ ' y ' $\wedge$ ': basta uno de los dos, siempre que se disponga además de ' $\sim$ '.

Pero, a la inversa, también es posible prescindir de ' $\wedge$ ' y ' $\vee$ ' y expresarlos por medio de ' $\sim$ ' y ' $\rightarrow$ ', como puede verse decidiendo mediante una tabla acerca de las siguientes equivalencias, que son verdades formales:

$$(14) \quad p \vee q \leftrightarrow [\sim p \rightarrow q].$$

Por las leyes de De Morgan podemos obtener de (14)

$$(15) \quad \sim [\sim p \wedge \sim q] \leftrightarrow [\sim p \rightarrow q].$$

De (15):

$$(16) \quad \sim \sim [\sim p \wedge \sim q] \leftrightarrow \sim [\sim p \rightarrow q], \text{ o}$$

$$(17) \quad \sim p \wedge \sim q \leftrightarrow \sim [\sim p \rightarrow q].$$

En resolución: para construir todas las fórmulas de la lógica de enunciados con funciones tomadas de las cinco de *Principia Mathematica* basta cada uno de los tres pares:

a) $\sim$, $\vee$;

b) $\sim$, $\wedge$;

c) $\sim$, $\rightarrow$.

Reducción de todas las funciones veritativas diádicas a las funciones de Sheffer. — Por último, la discusión que hicimos a propósito de las funciones f_{15} y f_5 muestra que, igual que es posible reducir f_{15} a $\sim$ y $\wedge$, y f_5 a $\sim$ y $\vee$, es también posible reducir $\sim$ y $\wedge$ a f_{15} , y $\sim$ y $\vee$ a f_5 . Mas como los dos pares $\sim$ y $\wedge$, $\sim$ y $\vee$ bastan cada uno de ellos para expresar todas las demás funciones veritativas diádicas, eso quiere decir que tanto f_{15} cuanto f_5 bastan, cada una por sí sola, para expresar las cinco funciones veritativas de *Principia Mathematica*.

f_5 , conocida con los nombres de 'función trazo de Sheffer', 'negación alternativa', 'inconjunción' o 'incompatibilidad', se simboliza por ' $|$ ' (trazo de Scheffer) y ha tenido importancia en la teoría lógica. Su tabla es

$ $	V	F
V	F	V
F	V	V

La función trazo sólo da el valor F cuando sus dos argumentos son verdaderos, y en los demás casos da el valor V. Con ella se define la negación basándose en la siguiente equivalencia:

$$(18) \quad \sim p \leftrightarrow p | p.$$

(18) puede parafrasearse intuitivamente diciendo: ' p es falsa cuando es incompatible consigo misma'. La siguiente tabla justifica la equivalencia

TABLA IV

p	$\sim p$	$p p$	(18)
V	F	F	V
F	V	V	V

La fórmula

$$(19) \quad p \vee q \leftrightarrow [p | p] | [q | q]$$

da la reducción de $\vee$ a $|$, y puede comprobarse por la siguiente tabla:

TABLA V

p	q	$p p$	$q q$	$[p p] [q q]$	$p \vee q$	(19)
V	V	F	F	V	V	V
F	V	V	F	V	V	V
V	F	F	V	V	V	V
F	F	V	V	F	F	V

(19) puede parafrasearse intuitivamente así: 'disyunción es incompatibilidad de negaciones'.

La función $\wedge$ se reduce a $|$ según la equivalencia

$$(20) \quad p \wedge q \leftrightarrow [p | q] | [p | q],$$

que puede parafrasearse diciendo: 'conjunción es negación de incompatibilidad'.

Puesto que con la función trazo de Sheffer pueden construirse la negación, la disyunción y la conjunción, puede también construirse todas las demás funciones veritativas diádicas. Lo mismo vale de la función f_{15} , que simbolizaremos por ' $\vee$ ', y cuya tabla es

$\vee$	V	F
V	F	F
F	F	V

Se la suele llamar 'negación conjunta'. También 'negación simultánea' e 'indisyunción'. Es una buena esquematización de la conectiva 'ni' del lenguaje común. ' $p \vee q$ ' puede leerse: 'ni p ni q '.

Las funciones $|$ y $\vee$ son importantes en la teoría, porque reducen el número de conceptos primitivos de la lógica de enunciados. Pero en la

práctica son de uso incómodo, pues las fórmulas escritas con estos funtores son largas y poco legibles. Por ejemplo, una fórmula tan corta como

$$(21) \quad p \rightarrow q$$

entendida según la definición ' $p \rightarrow q \leftrightarrow_{df} \sim p \vee q$ ', requiere ya once manchas tipográficas (sin contar los corchetes) para su expresión con el trazo de Sheffer:

$$(22) \quad [[p|p] | [p|p]] | [q|q].$$

Por esta razón, cuando interesa más operar con el cálculo que hacer consideraciones teóricas, es costumbre conservar las cinco funciones veritativas de *Principia Mathematica*.

| y $\vee$ permiten interesantes reflexiones especulativas. Esas funciones, las únicas que son, cada una por separado, suficientes, representan formas de negación (alternativa o conjunta). La negación parece pues ser una operación lógica muy elemental. Por otra parte, son negaciones diádicas, a las que se puede reducir la monádica.

73. Uso de las formas normales como técnica de decisión en la lógica de enunciados. — En lo que precede se ha apuntado un parentesco entre la normalización de fórmulas de la lógica de enunciados y la técnica decisoria de las tablas veritativas. Con las tablas se justificaba la prescindibilidad de ciertos funtores de los que también se prescinde al establecer formas normales. Vamos a ver ahora que la normalización de fórmulas de la lógica de enunciados puede entenderse a su vez como un procedimiento de decisión.

La *forma normal conjuntiva* (de Schröder: mención que omitiremos para abreviar; cfr. 57) es útil para decidir si una fórmula dada es o no es una verdad formal. Si la forma normal conjuntiva de una fórmula, X , presenta en cada disyunción componente principal de la conjunción una componente atómica afirmada y negada, entonces la fórmula es una verdad formal, puesto que cada una de sus componentes es formalmente verdadera (es una versión del principio de tercio excluso). Por ejemplo, el teorema (TE14).

$$(23) \quad [p \rightarrow q] \rightarrow [[q \rightarrow r] \rightarrow [p \rightarrow r]],$$

una forma normal conjuntiva del cual, establecida en 57, es

$$(24) \quad [\sim p \vee r \vee p \vee q] \wedge [\sim p \vee r \vee p \vee \sim r] \wedge [\sim p \vee r \vee \sim q \vee q] \wedge [\sim p \vee r \vee q \vee \sim r],$$

es una verdad formal, pues cada una de las disyunciones componentes principales de la conjunción es una versión del principio de tercio excluso, con ' p ', con ' q ' o con ' r '. No es, en cambio, una verdad formal

$$(25) \quad p \vee q \rightarrow r,$$

cuya forma normal conjuntiva, tal como la establecimos en 57, es

$$(26) \quad [\sim p \vee r \vee q] \wedge [\sim p \vee r \vee \sim q] \wedge [\sim q \vee r \vee p] \wedge [\sim q \vee r \vee \sim p].$$

(26) no tiene ninguna componente principal que sea formalmente verdadera. Y basta con que una no lo sea para que tampoco lo sea la conjunción.

La *forma normal disyuntiva* es sobre todo útil para decidir acerca de si una fórmula dada es o no es consistente. Será inconsistente si todas las conjunciones componentes principales de su forma normal disyuntiva presentan una misma componente atómica afirmada y negada. Por ejemplo, la fórmula

$$(27) \quad \sim [[p \rightarrow q] \rightarrow [[q \rightarrow r] \rightarrow [p \rightarrow r]]],$$

tiene una forma normal disyuntiva

$$(28) \quad [p \wedge \sim r \wedge r \wedge q] \vee [p \wedge \sim r \wedge r \wedge \sim p] \vee [p \wedge \sim r \wedge \sim q \wedge q] \vee [p \wedge \sim r \wedge \sim q \wedge \sim p],$$

en todas cuyas conjunciones hay una componente atómica negada y no negada. Todas esas conjunciones son falsedades formales, contradicciones formales. Por tanto, también lo es (28) y, consiguientemente, (27). ((28) se obtiene de (27) por (TE31), (TE13), (TE8) y (TE11)).

En cambio, la fórmula (25), ' $p \vee q \rightarrow r$ ', que, como hemos visto por su forma normal conjuntiva, no es una verdad formal, no es tampoco una falsedad formal. En 57 hallamos también una forma normal disyuntiva suya:

$$(29) \quad [\sim p \wedge \sim q] \vee [r \wedge \sim q] \vee [\sim p \wedge r].$$

Ninguna de las conjunciones principales de la disyunción es una falsedad formal. Consiguientemente, ni (29) ni (25) son falsedades formales. (25) es pues consistente. (29) muestra cuáles son las interpretaciones para las cuales es verdadera (25): que ' p ' y ' q ' sean ambas falsas; que ' r ' sea verdadera y ' q ' falsa; que ' p ' sea falsa y ' r ' verdadera; que ' r ' sea verdadera, con independencia de lo que sean ' p ' y ' q '.

74. Decisión abreviada por "reducción al absurdo". — No siempre es necesario establecer toda una tabla veritativa o normalizar una fórmula de la lógica de enunciados para decidir acerca de ella. La práctica personal desempeña, naturalmente, cierto papel en esto. Entre las varias técnicas rápidas que se han propuesto para decidir acerca de fórmulas de la lógica de enunciados de un modo abreviado, la más común es la que consiste en reducir al absurdo la hipótesis de que una determinada fórmula no es una verdad formal. Al reducir esta hipótesis al absurdo, se muestra que la fórmula es una verdad formal. He aquí un ejemplo del uso de este expediente. Para mostrar que la fórmula

$$(30) \quad p \wedge q \rightarrow q$$

es una verdad formal, empezamos por admitir la hipótesis de que no lo sea. Anotamos 'F' debajo de su functor principal:

$$\begin{array}{c} p \wedge q \rightarrow q \\ F \end{array}$$

Si no es una verdad formal, la fórmula tiene que ser falsa, por lo menos, para una interpretación (si lo es para toda interpretación, entonces además de no ser una verdad formal es una falsedad formal). (30) no tiene más que una interpretación que la haga falsa: interpretar el antecedente como verdadero y el consecuente como falso. Hacemos (mentalmente por lo común; si hay que hacerlo por escrito, es bueno numerar los pasos del análisis) las anotaciones correspondientes:

$$\begin{array}{cc} p \wedge q \rightarrow q & \\ F & \text{Primer paso.} \\ V & F \quad \text{Segundo paso.} \end{array}$$

El resto del análisis consiste en ver si es posible esa interpretación. Por de pronto, ella nos exige dar el siguiente tercer paso:

$$\begin{array}{cc} p \wedge q \rightarrow q & \\ F & \text{Primer paso.} \\ V & F \quad \text{Segundo paso.} \\ F & \text{Tercer paso.} \end{array}$$

En efecto, 'q' ha quedado interpretada ya por F en el paso segundo.

Pero el antecedente de (30) es una conjunción. Tiene que ser verdadera según la hipótesis de que (30) no es una verdad formal. Y para que una conjunción sea verdadera, tienen que serlo todos sus componentes. Esto nos impone un cuarto paso y, con él, el final del análisis de (30). Lo repetiremos íntegro:

$$\begin{array}{cc} p \wedge q \rightarrow q & \\ F & \text{Primer paso.} \\ V & F \quad \text{Segundo paso.} \\ F & \text{Tercer paso.} \\ V & V \quad \text{Cuarto paso.} \end{array}$$

La hipótesis lleva a una contradicción en la columna de 'q'. 'q' tiene que ser a la vez verdadera y falsa para que valga la hipótesis. Consiguientemente, la hipótesis es inconsistente, está "reducida al absurdo". (30) es una verdad formal.

75. Decidibilidad de las expresiones monádicas de la lógica de predicados. — La decidibilidad de la lógica de enunciados se extiende a toda

expresión de la lógica elemental que sea reducible a una forma isomorfa de la de las expresiones de la lógica de enunciados. Este es el caso de las expresiones monádicas de la lógica de predicados. Nos haremos cargo de ello realizando una adecuada transformación de las fórmulas monádicas de la lógica de predicados de primer orden.

En primer lugar, prescindiremos del cuantificador 'E' con arreglo a los teoremas (TP24)-(TP27), escribiendo '(x)' en lugar de ' $\sim \exists x \sim$ ', '(x) $\sim$ ' en lugar de ' $\sim \exists x$ ', ' $\sim (x)$ ' en lugar de ' $\exists x \sim$ ' y ' $\sim (x) \sim$ ' en lugar de ' $\exists x$ '. Según esta primera transformación, el axioma A6 de HB, por ejemplo,

$$(31) \quad Py \rightarrow \exists x Px,$$

se escribirá

$$(31a) \quad Py \rightarrow \sim (x) \sim Px.$$

Una fórmula como el axioma A5 de HB,

$$(32a) \quad (x)Px \rightarrow Py,$$

no queda afectada por este primer paso de la transformación en curso.

En segundo lugar, y provisionalmente, haremos que en la fórmula no aparezca más que un símbolo variable individual. Más adelante nos liberaremos de esta limitación, pero, de todos modos, se observará que este resultado se puede obtener mediante operaciones permitidas en HB por las reglas R α_2 y R δ . Para (31a) y (32a) obtendremos, respectivamente, eligiendo convencionalmente a 'x' como variable única:

$$(31b) \quad Px \rightarrow \sim (x) \sim Px,$$

$$(32b) \quad (x)Px \rightarrow Px.$$

Luego prescindimos de la variable individual. Esta transformación podría tener consecuencias indeseables si la fórmula inicial (con las variables unificadas, es decir, en la forma b, con sólo 'x') no fuera inequívocamente reconstruible a partir de la nueva transformada sin 'x'. Pero es claro que (31b) y (32b) se reconstruyen inequívocamente a partir de (31c) y (32c), respectivamente, si se conviene, como resulta naturalmente impuesto por los pasos dados, que, siendo cualquier símbolo variable individual apto para el "relleno" de (31c) y (32c) cuando se quiera volver a las fórmulas b, todo relleno se hará con 'x':

$$(31c) \quad P \rightarrow \sim () \sim P,$$

$$(32c) \quad () P \rightarrow P.$$

El último paso que queremos dar en estas transformaciones consiste en prescindir de los generalizadores. Pero este paso sí que tendrá efectos no deseados, por lo que exige, para evitarlos, la introducción de alguna me-

dida de precaución. En efecto, si suprimimos simplemente los generalizadores obtenemos: de (31c)

$$(31d) \quad P \rightarrow P \text{ (pasando por } 'P \rightarrow \sim \sim P');$$

y de (32c)

$$(32d) \quad P \rightarrow P.$$

(31d) y (32d) muestran que hemos perdido la diferencia entre A5 y A6. Esa diferencia se ha perdido porque fórmulas como $'(x)Px'$ y $'\exists xPx'$ han quedado ambas transformadas en lo mismo en que quedan transformadas fórmulas sin cuantificar, como $'Px'$: es decir, han quedado ambas identificadas con esta última. Esta identificación es la que hay que evitar, porque ella comporta la identificación de $'(x)Px'$ con $'\exists xPx'$. Lo que hay que evitar es, más precisamente, esta última identificación. En cambio, la identificación de $'(x)Px'$ con $'Px'$ no es objetable en HB, pues ya varias veces hemos visto que si en HB vale una fórmula sin cuantificar, como $'Px'$, entonces vale también $'(x)Px'$. (Demostración por $RA\gamma_3$, RY_1 y $R\beta$ en cap. VII, 46, por ejemplo.) Lo que hay que evitar es, pues, la identificación de $'\exists xPx'$ con $'Px'$ y, consiguientemente, con $'(x)Px'$. Para conseguirlo vamos a introducir una convención muy poco natural, pero práctica. La convención se basa en la distinción entre negaciones que afectan a cuantificadores y negaciones que afectan a predicados. Para las segundas vamos a usar el símbolo que en aritmética se lee 'menos', o sea, $'—'$, antepuesto, como $'\sim'$, al símbolo negado. $'—'$ tendrá la misma tabla que $'\sim'$, y es nombre de la misma función: de $'—'$ valen los mismos teoremas que de $'\sim'$. Pero, con objeto de poder reconstruir fórmulas particularizadas, convendremos en no hacer nunca uso de la ley de doble negación para simplificar en fórmulas composiciones de $'\sim'$ con $'—'$. Esta convención es una restricción, pero *sólo notacional*: la ley de doble negación vale, naturalmente, para composiciones de $'\sim'$ y $'—'$, que nombran la misma función. Y la aplicaremos para computar valores de composición de $'\sim'$ y $'—'$. Pero no haremos uso de ello escribiendo $'P'$ en vez de $'\sim — P'$, como, en cambio, escribiremos $'P'$ en vez de $'\sim \sim P'$ o de $'— — P'$.

Con esta convención, la supresión del generalizador en (31c) y (32c) nos da, respectivamente:

$$(31e) \quad P \rightarrow \sim — P,$$

$$(32e) \quad P \rightarrow P.$$

La expresión (32e) recoge la identificación de $'(x)Px'$ con $'Px'$. Esto conlleva, naturalmente, la imposibilidad de reconstruir inequívocamente, por ejemplo, (32a) a partir de (32e). Esa imposibilidad puede considerarse irrelevante por la reflexión que ya hemos hecho, y por esta otra: la imposibilidad en cuestión impone el abandono de toda fórmula monádica que se presente sin cuantificar, y la adopción en su lugar de la misma fórmula, pero

generalizada. Esto es inesencial desde el punto de vista del cálculo, como vimos antes. Pero también lo es desde un punto de vista semántico, de significación: fórmulas como $'Px'$ son semánticamente imperfectas, por ambigüas. Si esa fórmula debe entenderse "al pie de la letra", es decir, si $'x'$ está de verdad tomada como variable, entonces se quiere decir que cualquier cosa es P ; y esto se expresa más perfecta e inequívocamente por $'(x)Px'$. Si, en cambio, la fórmula no puede tomarse "al pie de la letra", porque se quiere significar con ella que algo es P o que un determinado objeto innominado es P , entonces las expresiones inequívocas son, respectivamente, $'\exists xPx'$ y $'Pa'$, siendo $'a'$ una constante individual, un nombre, o una descripción, como por ejemplo: 'el objeto innominado en el que estoy pensando'.

Así pues, eliminar de entre las fórmulas afirmables con variables las fórmulas sin cuantificar, además de no ser nada peligroso desde el punto de vista sintáctico del cálculo, es, desde el punto de vista semántico, lo más razonable que puede hacerse.

De acuerdo con este abandono de fórmulas con variables sin cuantificar, $'P'$ significa en todo caso a partir de ahora (y en el presente contexto) $'(x)Px'$, o sea, $'\sim \exists x \sim Px'$; $'\sim P'$ significa $'\sim (x)Px'$, o $'\exists x \sim Px'$; $'— P'$ significa $'(x) \sim Px'$, o $'\sim \exists x Px'$; consiguientemente, la expresión $'\sim — P'$ de (31e) significa $'\sim (x) \sim Px'$, o sea $'\exists x Px'$. Las convenciones sentadas hasta aquí permiten pues la reconstrucción de las fórmulas iniciales a partir de las transformadas, con dos pérdidas inesenciales: la diversidad de símbolos variables individuales y la presencia de fórmulas con variables sin cuantificar.

Ahora bien: como fruto de ese juego de convenciones de transformación notacional y semántica, hemos llevado las fórmulas A5 y A6 a una forma isomorfa de la de expresiones de la lógica de enunciados: a una forma a la que pueden aplicarse las técnicas de decisión de ésta, por ejemplo, la de las tablas veritativas (con la anomalía notacional de dos funtores para la misma función, $'\sim'$ y $'—'$). La decisión sobre A5 y A6 es tan sencilla que no vale la pena establecer las tablas correspondientes. Veamos un ejemplo un poco más interesante: la decisión sobre el esquema silogístico que ya conocemos con el nombre de 'Darapti':

$$(33) \quad (x)[Px \rightarrow Qx] \wedge (x)[Px \rightarrow Rx] \wedge \exists xPx \rightarrow \exists x[Rx \wedge Qx].$$

Su transformación según las anteriores convenciones da sucesivamente:

$$(33a) \quad (x)[Px \rightarrow Qx] \wedge (x)[Px \rightarrow Rx] \wedge \sim (x) \sim Px \rightarrow \sim (x) \sim [Rx \wedge Qx].$$

$$(33c) \quad () [P \rightarrow Q] \wedge () [P \rightarrow R] \wedge \sim () — P \rightarrow \sim () — [R \wedge Q].$$

$$(33e) \quad [P \rightarrow Q] \wedge [P \rightarrow R] \wedge \sim — P \rightarrow \sim — [R \wedge Q].$$

Para (33e) puede construirse una tabla veritativa en la cual, como queda dicho, $'—'$ se trata como $'\sim'$:

$P, \sim P$	V	F	V	F	V	F	V	F
Q	V	V	F	F	V	V	F	F
R	V	V	V	V	F	F	F	F
$P \rightarrow Q$	V	V	F	V	V	V	F	V
$P \rightarrow R$	V	V	V	V	F	V	F	V
$[P \rightarrow Q] \wedge [P \rightarrow R] \wedge \sim P$	V	F	F	F	F	F	F	F
$[R \wedge Q], \sim \sim [R \wedge Q]$	V	V	F	F	F	F	F	F
(33e)	V	V	V	V	V	V	V	V

Por último, para liberarnos de la restricción de considerar una sola variable individual, podemos hacer lo siguiente: empezaremos por observar que las letras predicativas 'P', 'Q', etc. de la tabla anterior se han usado en realidad como letras de enunciado. 'P' significa el enunciado 'Px' (que equivale a su cuantificación universal), y análogamente las demás. Esto permite, dada una fórmula cualquiera de la lógica predicativa monádica, utilizar un diccionario o serie de correspondencias convencional, como el siguiente, por ejemplo:

$$\begin{aligned} Px &: p \\ Qx &: q \\ Rx &: r \end{aligned}$$

Representando así las fórmulas predicativas monádicas por letras de enunciado, las tablas tendrán el mismo aspecto que las de la lógica de enunciados. Ahora bien: esto sugiere una ampliación de este método de diccionarios para evitar la restricción de no usar más que una variable. La ampliación consiste en tomar en nuestro diccionario letras de enunciado diversas para representar fórmulas predicativas monádicas con variables individuales diversas, aunque estas fórmulas tengan el mismo símbolo predicativo. Todas las demás transformaciones se practican como antes. Así, por ejemplo, dada la fórmula

$$\exists y(x)[Py \rightarrow Qx] \rightarrow (x)\exists y [Py \rightarrow Qx],$$

empezamos por establecer el diccionario convencional

$$\begin{aligned} Py &: p \\ Qx &: q \end{aligned}$$

Por las demás transformaciones obtenemos sucesivamente:

$$\sim (y) \sim (x) [Py \rightarrow Qx] \rightarrow (x) \sim (y) \sim [Py \rightarrow Qx].$$

Aquí aplicamos el diccionario, y obtenemos:

$$\begin{aligned} \sim (y) \sim (x) [p \rightarrow q] &\rightarrow (x) \sim (y) \sim [p \rightarrow q], \\ \sim \sim [p \rightarrow q] &\rightarrow \sim \sim [p \rightarrow q] \end{aligned}$$

La tabla de esta última fórmula se computa de acuerdo con las reglas dadas para '—' y las normales de '~'.

Observaremos incidentalmente que la anterior transformación, con o sin eliminación de cuantificadores, suministra una nueva notación para las expresiones de la lógica de predicados de primer orden. 'Ex[Px ∧ Qx]', por ejemplo, podría expresarse por 'Ex[P ∧ Q]x', una vez que hemos admitido la composición de predicados mediante funtores veritativos. Esta notación, que expresa muy directamente el análisis veritativo-funcional de predicados compuestos, es poco frecuente, pero la usa a veces alguno de los principales autores, como R. Carnap.

Sección segunda. — EL ALCANCE ANALÍTICO DEL CÁLCULO LÓGICO

CAPÍTULO XIV

LÓGICA DE CLASES

76. **El álgebra de clases.** — Entre las varias razones por las cuales las limitaciones algorítmicas del cálculo lógico no anulan la utilidad del mismo se encuentra la siguiente: el algoritmo lógico que ya conocemos es susceptible de ciertas reinterpretaciones — dos de las cuales estudiaremos en esta sección, caps. XIV y XV — que hacen de él un instrumento apto para el análisis de contextos diversos, o sea, para la aplicación a diversos universos del discurso. En este capítulo vamos a considerar la primera de las dos reinterpretaciones aludidas, la cual da lugar a la llamada 'lógica de clases'.

La lógica de clases comprende, entendido en un nuevo universo del discurso, todo lo que comprende la lógica de predicados monádica. En ésta es posible distinguir dos partes: la lógica de enunciados y la lógica predicativa monádica propiamente dicha. Análogamente distinguiremos dos niveles de análisis en la lógica de clases: uno, el más elemental, al que llamaremos 'álgebra de clases', que es el nombre tradicional, dado a este cálculo por Boole; otro, más complejo, al que llamaremos 'lógica general de clases' o, simplemente, 'lógica de clases'.

El álgebra de clases es una reinterpretación extensional de la lógica de enunciados. Consideraremos fórmulas atómicas del álgebra de clases. En vez de las letras de enunciado usaremos para ellas letras minúsculas griegas, como ' α ', ' β ', ' γ ', etc. Esas letras representan clases, no valores veritativos o enunciados. 'Clase' es el término lógico que corresponde al matemático 'conjunto'. Una clase es una colección de objetos cualesquiera. En el tratamiento de una clase se prescinde de la naturaleza concreta completa de sus miembros, u objetos que la componen, y del orden en que se presenten dichos miembros.

Interesa ahora aclarar qué ocurre con los funtores en la nueva inter-

pretación. En la lógica de enunciados, los funtores son veritativos, componen enunciados moleculares a partir de enunciados dados, o, lo que es lo mismo, representan funciones de valores veritativos a valores veritativos.

En el álgebra de clases los funtores de clases compondrán símbolos moleculares de clases a partir de símbolos de clases, o lo que es lo mismo: representarán funciones de clases a clases. Pero es claro que si nos limitamos a eso, no podremos construir nunca enunciados sobre clases. Y un álgebra no contiene sólo símbolos para componer objetos con objetos, sino también otros para componer enunciados con nombres de objetos. El álgebra elemental, por ejemplo, no sólo permite construir nombres compuestos a partir de otros nombres, como el nombre ' $a + b$ ' a partir de los nombres ' a ' y ' b ', sino también enunciados sobre los objetos nombrados, como, por ejemplo, y con esos mismos nombres: ' $a = b$ ', o ' $a > b$ '.

En la lógica de enunciados que, como vemos, también puede entenderse como un álgebra, esto no plantea ningún problema por la excepcional circunstancia de que ya los objetos mismos son enunciados (o valores veritativos, que son valores de enunciados). Pero éste no es el caso en lógica de clases. Por eso será conveniente tener, además de funciones de clases a clases, de funciones que compongan clases con clases, también funciones de enunciado, de clases a enunciados, que compongan enunciados (sobre clases) con nombres de clases.

Los funtores de clases a clases comúnmente usados son:

- 1.º ' $\sim$ ', que corresponde a ' $\neg$ ' y, más precisamente, al ' $\neg$ ' que usamos en 75. Se lee 'complemento de'. ' $\sim \alpha$ ' se lee: 'la clase complemento de la clase α ', o, más brevemente, 'complemento de α '.
- 2.º ' $\cup$ ', que corresponde a ' $\vee$ ', y se lee: 'suma (o reunión) de'. ' $\alpha \cup \beta$ ' se lee: 'la clase suma (o reunión) de las clases α y β ', o, brevemente: 'suma (o reunión) de α y β '.
- 3.º ' $\cap$ ', que corresponde a ' $\wedge$ ', y se lee 'producto de' o 'intersección de'. ' $\alpha \cap \beta$ ' se lee: 'la clase producto (o intersección) de las clases α y β ', o, brevemente, 'producto (o intersección) de α y β '.

No es necesario que, para completar la analogía de esta álgebra de clases con la lógica de enunciados, busquemos funtores de clases correspondientes a ' $\rightarrow$ ' y ' $\leftrightarrow$ '. Pues tenemos los tres correspondientes a ' $\sim$ ', ' $\vee$ ' y ' $\wedge$ ', que son suficientes en lógica de enunciados para expresar cualquier función diádica. Igualmente son ' $\sim$ ', ' $\cup$ ' y ' $\cap$ ' suficientes para expresar cualquier función diádica de clases. ' $\sim \alpha \cup \beta$ ' puede, por ejemplo, entenderse como 'la clase-flecha de α y β ', y ' $[\sim \alpha \cup \beta] \cap [\sim \beta \cup \alpha]$ ' como 'la clase-doble-flecha de α y β ' (cfr. (TE30) y (TE31) del cálculo de enunciados, y las tablas de valores más adelante).

Los funtores de enunciado (sobre clases), o funtores de clases a enunciados, que se usan comúnmente son:

- 4.º ' $\subset$ ', que corresponde a ' $\rightarrow$ ' y se lee: 'está incluida en'. ' $\alpha \subset \beta$ ' se lee: 'la clase α está incluida en la clase β ', o ' α es una subclase de β '.

5.° ' $=$ ', que corresponde a ' $\leftrightarrow$ ' y leeremos 'es igual a'. ' $\alpha = \beta$ ' se lee: 'la clase α es igual a la clase β ', o ' α igual β '.

' $\subset$ ' y ' $=$ ' componen pues enunciados sobre clases con símbolos de clases, mientras que ' $\neg$ ', ' $\cup$ ' y ' $\cap$ ' componen símbolos de clases con símbolos de clases. Pero ' $\subset$ ' y ' $=$ ' no son los únicos funtores de enunciados utilizables en álgebra de clases. También son útiles los funtores veritativos de la lógica de enunciados, para componer enunciados sobre clases a partir — no de clases, como ocurre con los funtores ' $\subset$ ' y ' $=$ ' sino de enunciados sobre clases. En la siguiente fórmula, por ejemplo, se hace uso de funtores de clases a clases, funtores de clases a enunciados y funtores de enunciados (sobre clases) a enunciados (sobre clases):

$$(1) \quad [\alpha \subset \beta] \wedge [\beta \subset \gamma \cap \delta] \rightarrow [\alpha \subset \gamma \cap \delta].$$

(1) es una versión del "principio del silogismo": 'si la clase α esta incluida en la clase β , y la clase β está incluida en la clase intersección de las clases γ y δ , entonces la clase α está incluida en la clase $\gamma \cap \delta$ '.

Tenemos, en resumen:

A) símbolos de clase atómicos, que corresponden a las letras de enunciado de la lógica de clase. Son los símbolos variables: ' α ', ' β ', ' γ ' ...

B) símbolos constantes, funtores, que corresponden y duplican a los de la lógica de enunciados:

1.° funtores de clases a clases: ' $\neg$ ', ' $\cup$ ', ' $\cap$ '.

2.° funtores de clases a enunciados: ' $\subset$ ', ' $=$ '.

3.° funtores de enunciados a enunciados: ' $\sim$ ', ' $\vee$ ', ' $\wedge$ ', ' $\rightarrow$ ', ' $\leftrightarrow$ '.

La construcción del álgebra de clases como una reinterpretación de la lógica de enunciados deja sin alterar lo estructural o formal de ésta, pues la duplicación del uso de los funtores no introducirá ninguna posibilidad de composición de valores que no sea isomórfica de las de los funtores veritativos. Por eso el álgebra de clases debe tener las mismas propiedades lógicas que el cálculo de enunciados. Señaladamente, será decidible, y le serán aplicables las técnicas decisorias de la lógica de enunciados (todas ellas basadas en la de las tablas). Mas para que le sean aplicables esas técnicas, es necesario definir los funtores de clases a clases y de clases a enunciados mediante tablas de valores. Sólo así, en efecto, se podrán definir valores para los enunciados sobre clases. Por ejemplo: no sabemos qué valor veritativo (V o F) puede tener el enunciado

$$(2) \quad \alpha \cap \beta \subset \alpha$$

si no sabemos antes qué valores atribuir a $\alpha \cap \beta$ y a α y β .

Ahora bien: la definición de los funtores en la interpretación normal de la lógica de enunciados se establece mediante los valores V y F, atribuidos a los argumentos de aquellos funtores. Es claro que no tiene directamente sentido decir de una clase que es verdadera o que es falsa. Nuestra interpretación de la lógica de enunciados para hacer de ella una álgebra

de clases exige la elección de dos valores que se comporten respecto de los funtores de clases a clases como V y F respecto de los funtores veritativos, pero que tengan sentidos en el (pertenezcan al) universo del discurso de las clases. Convendremos en la siguiente reinterpretación: un símbolo de clase tendrá el valor positivo, que simbolizaremos por '+', cuando represente una clase a la cual pertenezcan todos los objetos, cuando sea una clase universal; ejemplo puede ser la clase de las cosas que tienen la propiedad Ser-o-no-ser-terrestres. Y una clase tendrá el valor negativo, que simbolizaremos por '0', cuando sea una clase nula, una clase sin ningún miembro, como, por ejemplo, la clase de las cosas que tienen la propiedad Ser-y-no-ser-terrestres.

Con esos valores, las definiciones de las funciones de clases a clases mediante tablas tienen el siguiente aspecto, isomorfo del de las tablas de las correspondientes funciones veritativas:

TABLA I

α	$\neg \alpha$
+	0
0	+

Lectura: 'si α es una clase universal, entonces su complemento es una clase nula; si α es una clase nula, entonces su complemento es una clase universal'.

TABLA II

$\cup$	+	0
+	+	+
0	+	0

Lectura: 'si α , o β , o ambas son clases universales, entonces $\alpha \cup \beta$ es una clase universal; si α y β son ambas clases nulas, entonces $\alpha \cup \beta$ es una clase nula'.

TABLA III

$\cap$	+	0
+	+	0
0	0	0

Lectura: 'si α , o β , o ambas son clases nulas, entonces $\alpha \cap \beta$ es una clase nula; si α y β son ambas clases universales, entonces $\alpha \cap \beta$ es una clase universal'.

Las funciones veritativas, o funciones de enunciados a enunciados, no pueden aplicarse a + y 0, sino sólo a valores veritativos. Pero las funciones $\subset$ y $=$, funciones de clases a enunciados, tienen aquí un papel mediador: ellas se aplican a + y 0 y dan valores veritativos a los que a su vez pueden aplicarse funciones veritativas:

TABLA IV

$\subset$	+	0
+	V	F
0	V	V

Lectura: 'el enunciado ' $\alpha \subset \beta$ ' es falso si y sólo si α es universal y β es nula; en cualquier otro caso de los cuatro considerados es verdadero'.

TABLA V

$=$	+	0
+	V	F
0	F	V

Lectura: 'el enunciado ' $\alpha = \beta$ ' es falso si α es universal y β nula, o si α es nula y β es universal; en los otros dos casos considerados es verdadero'.

Con estas tablas, las de las funciones veritativas y unas reglas de formación adecuadas, se puede decidir de cualquier fórmula del álgebra de clases:

- A) si representa (a) una clase nula, o (b) una clase universal, o (c) una clase ni universal ni nula, para toda posible distribución de valores,
 B) si es (d) una verdad formal, o (e) una falsedad formal, o (f) una fórmula consistente.

He aquí algunos ejemplos de decisión:
 la fórmula

$$(3) \quad \alpha \cup -\alpha$$

representa una clase universal para toda distribución de valores, como puede verse por la Tabla VI, en cuya última columna no aparece más que el símbolo '+'.

TABLA VI

α	$-\alpha$	$\alpha \cup -\alpha$
+	0	+
0	+	+

Diremos de una clase así que es 'formalmente universal', pero pronto abandonaremos este modo de hablar.

La fórmula

$$(4) \quad \alpha \cap -\alpha$$

representa una clase formalmente nula:

TABLA VII

α	$-\alpha$	$\alpha \cap -\alpha$
+	0	0
0	+	0

La fórmula

$$(5) \quad \alpha \cup \beta$$

representa una clase que no es formalmente universal ni formalmente nula:

TABLA VIII

α	β	$\alpha \cup \beta$
+	+	+
0	+	+
+	0	+
0	0	0

La fórmula

$$(6) \quad \alpha \cap \beta = -\alpha \cup -\beta$$

es una falsedad formal.

La fórmula

$$(7) \quad \alpha \cap \beta \subset \alpha \cup \beta$$

es una verdad formal:

TABLA IX

α	β	$\alpha \cap \beta$	$\alpha \cup \beta$	(7)
+	+	+	+	V
0	+	0	+	V
+	0	0	+	V
0	0	0	0	V

También es una verdad formal la fórmula

$$(8) \quad [\alpha \subset \beta] \wedge [\beta \subset \gamma] \rightarrow [\alpha \subset \gamma]$$

TABLA X

α	β	γ	$\alpha \subset \beta$	$\beta \subset \gamma$	$[\alpha \subset \beta] \wedge [\beta \subset \gamma]$	$\alpha \subset \gamma$	(8)
+	+	+	V	V	V	V	V
0	+	+	V	V	V	V	V
+	0	+	F	V	F	V	V
0	0	+	V	V	V	V	V
+	+	0	V	F	F	F	V
0	+	0	V	F	F	V	V
+	0	0	F	V	F	F	V
0	0	0	V	V	V	V	V

Considerando esas nociones elementales del álgebra de clases se aprecia su pobreza como instrumento para el análisis de un lenguaje científico en el cual se hable de clases: el álgebra de clases no permite referirse a la relación que media entre un objeto y una clase de la que es miembro. Esa relación, que se llama '*pertenencia*', se simboliza por ' ϵ '. Esta insuficiencia reproduce una que ya vimos en la lógica de enunciados respecto del universo del discurso propio de ella: del mismo modo que la lógica de enunciados no permite analizar la estructura de los enunciados atómicos, así tampoco el álgebra de clases permite analizar la estructura de las clases tomadas como sillares para componer otras. Para superar esa limitación del análisis, que tiene varias consecuencias, apelaremos a la lógica de predicados (monádica).

77. La lógica general de clases. — La relación entre la lógica de clases y la lógica de predicados monádica puede aclararse desde el punto de vista de la teoría del conocimiento: un enunciado predicativo monádico, como

$$(9) \quad \text{Juan es alto,}$$

puede entenderse de dos modos: uno intensional, según el cual (9) significa

$$(10) \quad \text{Juan tiene la propiedad de ser-alto;}$$

y otro extensional, según el cual (9) significa

$$(11) \quad \text{Juan es un miembro de la clase de las cosas llamadas 'altas'.$$

La distinción tradicional entre la comprensión y la extensión de los predicados (de los "conceptos") se relaciona directamente con nuestro tema (cfr. 15): todo predicado monádico tiene asociada a su intensión una clase. Por eso pueden reducirse las expresiones predicativas monádicas a símbolos que representan clases, que es en realidad lo que hicimos en 75 al mostrar la decidibilidad de la parte monádica de la lógica de predicados. La lógica de predicados monádica es, pues, la parte del cálculo lógico más adecuada para obtener, mediante una oportuna reinterpretación, una lógica general de clases. Utilizaremos en ella todo lo que necesitemos de la notación de la lógica de predicados, para definir los conceptos elementales de la lógica de clases. Volveremos incluso a definir los funtores de clases a clases y de clases a enunciados, con objeto de construirlos teniendo en cuenta la relación de pertenencia, cosa que no pudimos hacer en el álgebra.

Antes de proceder a esas definiciones con ayuda de los símbolos de la lógica de predicados, vamos a introducir una notación propia de la lógica de clases, a saber, la notación de *abstracción* de clases. El término '*abstracción*' nos es ya conocido en un contexto filosófico (cfr. capítulo I) y también en un contexto formal (cfr. cap. IV). Este último nos interesa precisar un poco ahora: en 28 hablamos de '*abstracción funcional*' para designar la operación por la cual se construye una función lógica a partir de determinadas formaciones lingüísticas. Análogamente pueden construirse clases, como acabamos de decir, a partir de expresiones monádicas de la lógica de predicados. Por ejemplo, dada la fórmula ' Pa ', puede construirse la clase de las cosas que son P . Llamaremos '*abstracción de clases*' a esta operación. A ella nos hemos venido en realidad refiriendo en este epígrafe hasta el momento. La notación es la siguiente: una expresión de la forma

$$(12) \quad \hat{x} [X],$$

en la cual ' $\hat{x}$ ' es una variable individual con un acento circunflejo y X es una fórmula de la lógica de clases (escrita con símbolos de la lógica de

predicados monádica también), se leerá: 'la clase de los objetos x , tales que satisfacen a X '. Los miembros de esa clase son los objetos cuyos nombres, usados como interpretación de la variable ' x ' en X , hacen verdadera a ésta. La justificación del término 'abstracción' en este contexto es obvia desde el punto de vista de la teoría del conocimiento: pues los objetos en cuestión serán siempre algo más que meros "satisfactores", por así decirlo, de X . Por tanto, al considerarlos sólo desde ese punto de vista, se hace abstracción de cualquier otro comportamiento suyo, de cualquier otra propiedad suya. Otra lectura más breve y frecuente de (12) es: 'la clase de los x tales que X '. Diremos que la variable ' x ' está ligada (como si estuviera cuantificada) en ' $\hat{x}[X]$ '. Estas expresiones se llaman 'abstracciones' o 'abstractos', y el símbolo " $\hat{}$ " puede llamarse 'abstractor'. Otra notación frecuente equivalente a ' $\hat{x}[X]$ ' es ' $x \varepsilon [X]$ ' (Es claro que si X no contiene a ' x ' la abstracción es "vacía": no define ninguna clase porque la clase en cuestión estará ya definida de otro modo. Pero es bueno no excluir la posibilidad de escribir la notación incluso en ese caso.)

Ejemplo: la expresión

$$(13) \quad \hat{x}[x \text{ es personaje del Quijote}]$$

representa la clase construida por abstracción de (o la clase asociada a) la propiedad Ser-personaje-del-*Quijote*; o, más simplemente, la clase de los personajes del *Quijote*.

Con esa notación de la abstracción de clases, podemos proceder a las siguientes definiciones:

$$(DC14) \quad -\alpha =_{df} \hat{x}[\sim x \varepsilon \alpha].$$

Lectura: 'el complemento, $-\alpha$, de α es la clase de los x que no son miembros de α '. Más literalmente: 'la clase complemento de α es la clase de los x tales que x no es un α '. Obsérvese que no es necesario escribir ' $x \varepsilon \alpha$ ' entre paréntesis para negarla, es decir, que no es necesario escribir ' $\sim (x \varepsilon \alpha)$ '. No hay confusión posible acerca de lo negado por ' $\sim$ ', pues ' $\sim$ ' es un functor veritativo, sólo aplicable a enunciados, y ni ' x ' ni ' $x \varepsilon$ ' son enunciados, tienen valor veritativo, sino sólo ' $x \varepsilon \alpha$ '. Observaciones análogas valen para las definiciones siguientes.

$$(DC15) \quad \alpha \cup \beta =_{df} \hat{x}[x \varepsilon \alpha \vee x \varepsilon \beta].$$

Lectura: 'la suma de α y β es la clase de los x que son miembros de α o son miembros de β o son miembros de ambas α y β '.

$$(DC16) \quad \alpha \cap \beta =_{df} \hat{x}[x \varepsilon \alpha \wedge x \varepsilon \beta].$$

Lectura: 'el producto de α y β es la clase de los x que son a la vez miembros de α y de β '.

$$(DC17) \quad '\alpha \subset \beta' \leftrightarrow_{df} '(x)[x \varepsilon \alpha \rightarrow x \varepsilon \beta]'$$

Lectura: 'el enunciado ' α está incluida en β ' equivale por definición al enunciado 'para todo x , si x es un α , entonces x es un β '.

$$(DC18) \quad '\alpha = \beta' \leftrightarrow_{df} '(x)[x \varepsilon \alpha \leftrightarrow x \varepsilon \beta]'$$

Lectura: 'el enunciado ' α igual β ' equivale por definición al enunciado: 'para todo x , x es un α si y sólo si x es un β '.

La identidad individual puede definirse del modo siguiente, que es una expresión del "principio de identidad de los indiscernibles", propuesto por Leibniz. Según este principio, puede llamarse 'idénticas' a las cosas que no se pueden discernir (separar, distinguir), porque pertenecen a las mismas clases:

$$(DC19) \quad 'x = y' \leftrightarrow_{df} '(\alpha)[x \varepsilon \alpha \leftrightarrow y \varepsilon \alpha]'$$

Con (DC18) y (DC19) tenemos definidas la identidad de clases y la identidad de individuos. Esto justifica el uso — que ya hemos practicado sin justificarlo — del símbolo '=' para definiciones ('=_{df}'). Cuando definamos enunciados, seguiremos usando, como hasta ahora, ' $\leftrightarrow_{df}$ ' entre nombres de enunciados, que generalmente se obtienen *entrecomillando* los enunciados. En cambio, cuando definamos clases nuevas a partir de clases ya conocidas — por ejemplo, una clase α diciendo que es la clase $\beta \cap \gamma$ —, usaremos ' $=_{df}$ ' entre los nombres de esas clases, que son precisamente esas letras *sin entrecomillar*.

Se observará que (DC18) y (DC19) son expresiones de segundo orden, es decir, que corresponden a fórmulas de la lógica de predicados de segundo orden. (DC18) tiene símbolos de clase en posiciones de sujetos (del predicado diádico '='). (DC19) tiene cuantificado un símbolo de clase. Y los símbolos de clases corresponden a símbolos predicativos. Así, por ejemplo, la expresión del "principio de identidad de los indiscernibles" en lógica de predicados es: 'pueden considerarse idénticas las cosas que tienen las mismas propiedades':

$$(20) \quad 'Ixy' \leftrightarrow_{df} '(P)[Px \leftrightarrow Py]'$$

La diversidad entre individuos puede definirse del modo siguiente:

$$(DC21) \quad '\sim x = y' \leftrightarrow_{df} '\exists \alpha[x \varepsilon \alpha \wedge \sim y \varepsilon \alpha]'$$

Y la diversidad entre clases de individuos:

$$(DC22) \quad '\sim \alpha = \beta' \leftrightarrow_{df} '\exists x[x \varepsilon \alpha \wedge \sim x \varepsilon \beta]'$$

(DC21) y (DC22) se traducen a la lógica de predicados por (23) y (24) respectivamente:

$$(23) \quad '\sim Ixy' \leftrightarrow_{df} '\exists P[Px \wedge \sim Py]'$$

$$(24) \quad '\sim IPQ' \leftrightarrow_{df} '\exists x[Px \wedge \sim Qx]'$$

También (DC21)-(24) son expresiones de orden superior (concretamente, tal como están escritas, de segundo).

La traducibilidad de las expresiones de la lógica de clases a expresiones de la lógica de predicados monádica, y viceversa, que acabamos de ver en los ejemplos (DC19)-(24), permite obtener teoremas de la lógica de clases con los métodos de cálculo de la lógica de predicados. En particular, para todos los teoremas obtenidos al estudiar la lógica de predicados hay teoremas análogos de la lógica de clases. Por ejemplo, el teorema de la lógica de predicados (TP18),

$$(x)[Py \rightarrow Px] \leftrightarrow [Py \rightarrow (x)Px],$$

suministra el teorema de la lógica de clases

$$(TC25) \quad (x)[y \in \alpha \rightarrow x \in \alpha] \leftrightarrow [y \in \alpha \rightarrow (x)[x \in \alpha]].$$

Otras transcripciones:

$$(TC26) \quad (x)[x \in \alpha] \rightarrow y \in \alpha \quad : \text{A5: } (x)Px \rightarrow Py.$$

$$(TC27) \quad y \in \alpha \rightarrow \exists x[x \in \alpha] \quad : \text{A6: } Py \rightarrow \exists xPx.$$

El teorema de la lógica de predicados (TP16) puede llevarse a una forma algebraica:

$$(TP16) \quad (x)[Px \rightarrow \sim Qx] \wedge (x)[Rx \rightarrow Px] \rightarrow (x)[Rx \rightarrow \sim Qx].$$

Una primera traducción nos da:

$$(x)[x \in \alpha \rightarrow \sim x \in \beta] \wedge (x)[x \in \gamma \rightarrow x \in \alpha] \rightarrow (x)[x \in \gamma \rightarrow \sim x \in \beta].$$

fórmula que, por la definición (DC17), nos da el teorema del álgebra de clases:

$$(TC28) \quad \alpha \subset \sim \beta \wedge \gamma \subset \alpha \rightarrow \gamma \subset \sim \beta.$$

78. Clase nula y clase universal. — La escasa capacidad analítica del álgebra de clases se nos puso de manifiesto en 76 por el hecho de que con aquellos medios algebraicos no era posible tomar en consideración los miembros de clases. Esto tiene una consecuencia importante: que para poder definir las funciones de clases por tablas de valores tuvimos que suponer que toda clase "atómica", sin analizar, es universal o nula. Este no es, naturalmente, el caso siempre en un normal discurso sobre clases. Por eso en el álgebra tuvimos que introducir una expresión que ahora vamos a abandonar: allí hablamos, en efecto, de clases 'formalmente nulas' y 'formalmente universales'.

Cuando una clase es lo que entonces llamamos 'formalmente nula', cuando en todos los lugares de la última columna de su tabla aparece el símbolo '0', es que carece de miembros "para cualquier interpretación" de las clases que la componen, es decir, carece de miembros por su estructura, cualesquiera que sean las clases que se utilicen para realizar esa estructura. Carece, pues, de miembros, por razones formales. (El ejemplo que vimos era el de las cosas que son a la vez terrestres y no-terrestres.) Pero una clase puede no tener miembros por razones no formales, sino empíricas. Por ejemplo: la clase de los reyes de Francia que han reinado en la primera mitad del siglo xx. Llamaremos 'vacías' a las clases sin miembros por razones no-formales. Y 'nula' a la clase sin miembros por razones formales, de estructura (de la fórmula o propiedad de que se abstrae).

En el álgebra de clases tuvimos también que considerar que una clase elemental, o sin analizar, que no fuera nula tenía que ser una clase miembros de la cual fueran todas las cosas. Pero éste tampoco es siempre el caso. Por eso distinguiremos también entre clase universal y clases simplemente no-vacías. Por ejemplo: la clase de los españoles es una clase no-vacía sin ser una clase universal. Clase universal es, en cambio, la de las cosas que son terrestres o no lo son.

Las clases universal y nula se definen, respectivamente, como sigue:

$$(DC29) \quad V =_{df} \hat{x} [x = x].$$

$$(DC30) \quad \Lambda =_{df} \hat{x} [\sim x = x].$$

La afirmación de que una determinada clase, α , es no-vacía se simboliza por ' $\exists ! \alpha$ '. Necesitamos esa afirmación, por ejemplo, para conseguir el teorema de la lógica de clases análogo del teorema (TP17) (Darapti) de la lógica de predicados:

$$(TP17) \quad (x)[Px \rightarrow Qx] \wedge (x)[Px \rightarrow Rx] \wedge \exists xPx \rightarrow \exists x[Rx \wedge Qx].$$

$$(TC31) \quad \alpha \subset \beta \wedge \alpha \subset \gamma \wedge \exists ! \alpha \rightarrow \exists ! \beta \cap \gamma.$$

La traducibilidad de las fórmulas de la lógica de clases a las de la lógica de predicados hace superflua la construcción de un cálculo especial para la lógica de clases. Pero la posibilidad de utilizar un mismo algoritmo para ambos lenguajes no debe hacer olvidar las diferencias semánticas (propriadamente lógicas) entre ambos. Así, por ejemplo, la notación ' $\exists ! \alpha$ ' significa que la clase α es no-vacía, que tiene al menos un miembro. Eso no equivale a la afirmación de la existencia de la clase α , sino que es una afirmación más fuerte: la afirmación de la existencia de cosas que son miembros de la clase α . Es una abreviatura de la que se puede prescindir por la siguiente definición:

$$(DC32) \quad '\exists ! \alpha' \leftrightarrow_{df} '\exists x[x \in \alpha]'.$$

Consiguientemente, la negación de $\exists! \alpha$, $\sim \exists! \alpha$, no es la negación de la existencia de la clase α , como sugeriría una precipitada traducción del cuantificador $\exists!$ de la lógica de predicados. $\sim \exists! \alpha$ significa que α no tiene miembros, que es vacía. Y esa negación no equivale a la negación de la existencia de la clase α . La clase α puede considerarse "existente" como clase vacía, pues una clase es un abstracto, no es simplemente sus miembros. Por eso debe distinguirse, por ejemplo, entre una clase α que no tenga más que un miembro, x , y este miembro mismo. Para indicar que α es la clase que contiene a x como único miembro suele escribirse

$$(33) \quad \alpha = \{x\}.$$

La definición de la clase $\{x\}$ es:

$$(DC34) \quad \{x\} =_{df} \hat{y} [y = x].$$

Análogamente se escribe, por ejemplo, si β es la clase de los dos únicos miembros z , u :

$$\beta = \{z, u\}.$$

La diferencia entre x y $\{x\}$ puede apreciarse observando que x tiene la propiedad de ser miembro de la clase $\{x\}$, propiedad que $\{x\}$ no tiene. Según el principio de identidad de los indiscernibles, x y $\{x\}$ no son, por tanto, idénticos. Una clase de un solo miembro es, por ejemplo, la clase de los satélites de la Tierra visibles a simple vista.

79. El principio de abstracción y la paradoja de Russell. — En 77 encontramos expresiones de la lógica de clases que pertenecen al segundo orden lógico. Esto ocurre también con el principio que expresa la formabilidad de clases a partir de expresiones predicativas, llamado 'principio de abstracción de clases'. Sabemos que una clase se constituye abstrayendo mediante una propiedad que permite reunir todos los objetos que la poseen. Esto es lo que tiene que expresar el principio de abstracción, a saber: que, para toda clase, α , hay una propiedad, P , tal que para todo objeto, x , el pertenecer a α es lo mismo que tener la propiedad P :

$$(35) \quad (\alpha) \exists P (x) [x \in \alpha \leftrightarrow Px].$$

Por nuestra experiencia con la lógica de segundo orden, o de orden superior en general, podemos temer que expresiones de la lógica de clases en las que, como en (35), se cuantifiquen símbolos de tipos lógicos diversos, como ' x ', que es de tipo 0, y ' α ' y ' P ', que son de tipo 1, den lugar a la paradoja de Russell (cfr. 65). Así es, en efecto: sustituyendo la variable predicativa ' P ' de (35) por la constante predicativa 'no ser x ', tenemos:

$$(36) \quad (\alpha) (x) [x \in \alpha \leftrightarrow \sim x \in x].$$

Si no hay ninguna restricción sobre el uso de símbolos de distintos tipos lógicos, (36) vale para todo α . Dicho de otro modo: 'para todo α ' quiere decir 'para todo símbolo de la lógica de clases, cualquiera que sea su tipo lógico'. Lo que permite afirmar (36) para el símbolo de la lógica de clases ' x ', o sea, sustituir en (36) ' α ' por ' x ', con lo que obtenemos:

$$(37) \quad (x) [x \in x \leftrightarrow \sim x \in x].$$

Esta contradicción es la paradoja de Russell escrita en el lenguaje de la lógica de clases.

Consiguientemente, también la lógica de clases tiene que construirse según la teoría de los tipos lógicos, de modo que se prohíban formaciones como ' $x \in x$ ', en la que los símbolos situados a ambos lados de ' $\in$ ' son del mismo tipo. Esto se consigue estableciendo una clasificación sistemática de los símbolos por sus tipos, e introduciendo en las reglas de formación de fórmulas y enunciados la prohibición de que sean fórmulas expresiones que afirmen la pertenencia entre un símbolo o expresión de tipo n y otro que no sea de tipo $n + 1$.

Con estas precauciones o restricciones no ha resultado posible deducir la paradoja de Russell. No lo es a partir del principio de abstracción de clases, puesto que la teoría de los tipos prohíbe la sustitución de ' α ' por ' x '. El principio se formulará, por ejemplo, para el orden n :

$$(35a) \quad ({}^n\alpha) \exists {}^nP ({}^{n-1}x) [{}^{n-1}x \in {}^n\alpha \leftrightarrow {}^nP {}^{n-1}x].$$

'(${}^n\alpha$)' no quiere decir ahora lo mismo que 'para todo símbolo de la lógica de clases, cualquiera que sea su tipo', sino: 'para todo símbolo de tipo n de la lógica de clases'. Y ' ${}^{n-1}x$ ' no es de tipo n . Por tanto, no puede sustituir a ' ${}^n\alpha$ ' en el operando al practicar la eliminación de ' $()$ ' respecto de ' α '.

La restricción, que equivale a prescribir que todas las fórmulas atómicas de pertenencia o identidad de la lógica de clases sean de las formas tipo-lógicas

$${}^{n-1}x \in {}^n\alpha, \quad {}^ny = {}^n\alpha,$$

no se impone a la inclusión: entre dos clases del mismo tipo puede haber inclusión. Sea, por ejemplo, la clase de los catalanes, γ . Entonces los símbolos de tipo lógico 0 representan individuos catalanes. La clase de los catalanes que viven en el distrito I de Barcelona se representa por un símbolo de tipo lógico 1, pues es una clase de individuos de tipo 0. Llamémosla clase ' α '. La clase de los catalanes que viven en Barcelona se representa también por un símbolo de tipo 1, pues es una clase de individuos de tipo 0. Llamémosla ' β '. Podemos definir α y β del modo siguiente:

$$(38) \quad \alpha =_{df} \hat{x} [x \text{ es catalán} \wedge x \text{ vive en el distrito I de Barcelona}].$$

$$(39) \quad \beta =_{\text{af}} \hat{x}[x \text{ es catalán} \wedge x \text{ vive en Barcelona}].$$

γ por su parte se define.

$$(40) \quad \gamma =_{\text{af}} \hat{x}[x \text{ es catalán}].$$

Con ' α ', ' β ' y ' γ ', todas de tipo 1, pueden escribirse las expresiones correctas (fórmulas), que además son verdaderas:

$$(41) \quad \alpha \subset \beta.$$

$$(42) \quad \alpha \subset \gamma.$$

$$(43) \quad \beta \subset \gamma.$$

$$(44) \quad \alpha \cap \beta \subset \gamma.$$

$$(45) \quad \alpha \cup \beta \subset \gamma.$$

Pero no son correctas las expresiones

$$(46) \quad \alpha \in \beta,$$

$$(47) \quad \alpha \in \gamma,$$

pues no tiene sentido decir que los catalanes que viven en el distrito I de Barcelona sean un individuo miembro de la clase de los catalanes que viven en Barcelona, ni un individuo miembro de la clase de los catalanes: sino que son subclases de esas clases.

Si en cambio construimos la clase de todos los grupos de catalanes, definidos con criterios locales o geográficos, y la llamamos ' δ ', entonces son expresiones correctas (y verdaderas):

$$(48) \quad \alpha \in \delta;$$

$$(49) \quad \beta \in \delta.$$

δ es de tipo lógico 2.

La teoría de los tipos introduce una considerable complicación en la lógica de clases. Pues exige que se definan de nuevo las constantes lógicas para cada tipo. En efecto: no es tipológicamente la misma relación ϵ la que media entre 0x y ${}^1\beta$ en

$$(50) \quad {}^0x \epsilon {}^1\beta,$$

que la que media entre ${}^1\alpha$ y ${}^2\gamma$ en

$$(51) \quad {}^1\alpha \epsilon {}^2\gamma.$$

La ' ϵ ' de (50) simboliza una relación entre individuos y clases de individuos. La ' ϵ ' de (51) simboliza una relación entre clases de individuos y clases de clases de individuos (clases cuyos miembros son clases de individuos). Así habrá

que distinguir entre ambas ' ϵ ', llamando a la primera, por ejemplo, ' ϵ' ', y a la segunda ' ϵ'' '. Análogamente habrá que volver a definir en cada tipo ' $\subset$ ', ' $\cup$ ', ' $\cap$ ', ' $\subset$ ', ' $\supset$ ', ' $\cap$ ' y ' Δ ', y formular el principio de abstracción para $n = 1, 2, 3, \dots$

Para evitar la complicación resultante se ha introducido alguna solución que permita prescindir de la teoría de los tipos sin caer en la paradoja de Russell. La solución más clásica es la introducción de la idea de elemento en la formulación del principio de abstracción (Zermelo). El principio de abstracción se formula como sigue:

$$(35b) \quad (\alpha) \exists P(x) [x \in \alpha \leftrightarrow \exists z [x \in z \wedge Pz]].$$

La existencia de z y la cláusula ' $x \in z$ ' representan el carácter de "elemento" de x , garantizan que x es una entidad "normal", por así decirlo, que ya pertenece a alguna otra clase z "antes" de ser abstraída la clase α . x no es una entidad paradójica. La idea de Zermelo, que es antigua, tiene varias versiones menos antiguas y alguna reciente.

80. Algunos conceptos fundamentales de la aritmética: números cardinales; pares ordenados.— Siguiendo a Cantor, Frege y Russell-Whitehead, es corriente concebir los números cardinales, que corresponden en la sistematización general de la aritmética a los números enteros positivos, como clases de clases de objetos cualesquiera. Se recordará (cfr. 15) que Cantor construyó su teoría de conjuntos, entre otras cosas, para definir el concepto de número cardinal. Cantor define el número cardinal de un conjunto como el conjunto de todos los conjuntos coordinables con dicho conjunto.

Según esa idea, el número cardinal 1, por ejemplo, es la clase de todas las clases coordinables con cualquier clase de un solo miembro, o sea, la clase de todas las clases que sólo tienen un miembro. La apariencia de petición de principio que tiene la definición del número cardinal 1 así construida se disipa si se tiene en cuenta que el adjetivo 'un' que aparece en el definiens no representa el concepto *técnico* del número cardinal 1, sino cierta noción intuitiva operatoria. Al escribir la anterior definición del número cardinal 1 en el lenguaje de la lógica de clases desaparece tipográficamente aquella apariencia de petición de principio:

$$(DC52) \quad 1 =_{\text{af}} \hat{\alpha} [\exists x [\alpha = \{x\}]].$$

Lectura: 'el número cardinal 1 es la clase, $\hat{\alpha}$, de todas las clases α , tales que hay un individuo, x , de modo que α es la clase la totalidad de cuyos miembros es x '.

Adoptando, para indicar una clase de n miembros, una notación

$$(53) \quad \{x_1, x_2, \dots, x_n\},$$

podemos definir cualquier número cardinal según el mismo procedimiento

usado antes para definir el número cardinal 1. He aquí, por ejemplo, la definición del número cardinal 3:

$$(DC54) \quad 3 =_{df} \hat{\alpha} [\exists x \exists y \exists z [\alpha = \{x, y, z\} \wedge \sim x = y \wedge \alpha \sim x = z \wedge \alpha \sim y = z]].$$

Otra noción de uso fundamental en aritmética es la de par ordenado, que sirve, por ejemplo, para definir nuevos números como pares ordenados de otros números conocidos. Un par ordenado es, por de pronto, una clase de dos miembros. Pero el orden en que se presentan los miembros es precisamente una de las dos características de que prescinde la lógica de clases, la abstracción de clases. Por eso el carácter ordenado de los pares (o, en general, de una clase de n miembros) tiene que ser definido y construido. Para simplificar la notación suelen usarse paréntesis angulares, ' $\langle \rangle$ ', para indicar que los miembros situados entre ellos están ordenados, según se presentan tipográficamente, en la clase de que se trate. Con esta notación escribiremos:

$$(DC55) \quad \langle x, y \rangle =_{df} (1z) [z = \{\{x\}, \{x, y\}\}].$$

Lectura: 'el par ordenado $\langle x, y \rangle$ es la clase z , la cual tiene dos miembros que son: la clase cuyo único miembro es x , y la clase cuyos miembros son x e y '.

Se observará que si x e y son de tipo 0, el par no ordenado $\{x, y\}$ es de tipo 1, y el par ordenado $\langle x, y \rangle$ es de tipo 2 (la clase z es una clase de clases de individuos del tipo de x y de y).

Según el mismo principio se define, por ejemplo

$$(DC56) \quad \langle {}^0x, {}^0y, {}^0z \rangle =_{df} (1^4u) [{}^4u = \langle \langle {}^0x, {}^0y \rangle, \langle {}^0y, {}^0z \rangle \rangle].$$

81. Morfología de la lógica de clases. — En este capítulo hemos prescindido de exponer ordenadamente la gramática de la lógica de clases. Nos hemos apoyado intuitivamente en la noción de fórmula ya adquirida en la lógica de predicados. La base morfológica de la gramática de la lógica de clases difiere, sin embargo, de la morfológica de la lógica de predicados de primer orden, única para la cual definimos propiamente la noción de fórmula. Aquí veremos brevemente lo esencial de la definición de 'fórmula de la lógica de clases', con lo que quedará indicado por analogía lo que es una fórmula de la lógica de predicados de orden superior. Salvo en la definición misma de fórmula, prescindiremos de repetir los elementos morfológicos que se tomen sin modificación de la lógica de predicados de primer orden. Por otra parte, en vez de usar minúsculas latinas para variables de tipo 0 y minúsculas griegas para variables de tipo $n > 0$, usaremos, puesto que nos proponemos indicar el tipo explícitamente, sólo minúsculas latinas.

Símbolos elementales o primitivos. — Además de los de la lógica de predicados de primer orden:

- a) Los símbolos constantes lógicos ' ϵ ', ' $=$ ' y ' $\wedge$ '. (Los símbolos ' $\neg$ ', ' $\cup$ ', ' $\cap$ ', ' $\subset$ ', ' $\supset$ ', ' $\vee$ ' y ' Δ ' no hacen falta como símbolos primitivos, puesto que se pueden definir por medio de ' ϵ ' y ' $\wedge$ ' y los símbolos de la lógica de predicados; tomamos de todos modos ' $=$ ' como símbolo elemental por comodidad.)
- b) Variables y constantes (si hacen falta) de cada tipo, clasificadas por ellos. Escribiremos las variables y las constantes con índices numéricos para indicar el tipo.
- c) Llaves, ' $\{ \}$ ', y paréntesis angulares, ' $\langle \rangle$ '.

Piezas morfológicas que no existen en la gramática de la lógica de predicados son los

Abstractos. — Si X es una fórmula, entonces $\hat{x}[X]$ es una expresión nominal que se lee 'la clase de los x tales que X '.

Otro elemento nuevo de la morfología es la definición de *Expresión de tipo n* :

- 1) Una variable ' nx ' es una expresión de tipo n .
- 2) Una constante ' na ' es una expresión de tipo n .
- 3) Si X es una fórmula y ' nx ' es una variable de tipo n , entonces ' ${}^n\hat{x}[X]$ ' es una expresión de tipo $n + 1$.

Definimos por último:

Fórmula de la lógica de clases:

- I. Si ' nx ' es una variable de tipo n , y ' na ' una constante de tipo n , y ' ${}^{n+1}y$ ' una variable de tipo $n + 1$, y ' ${}^{n+1}b$ ' una constante de tipo $n + 1$, entonces

$$\begin{aligned} {}^nx &\epsilon {}^{n+1}y, \\ {}^nx &\epsilon {}^{n+1}b, \\ {}^na &\epsilon {}^{n+1}y, \\ {}^na &\epsilon {}^{n+1}b \end{aligned}$$

son fórmulas (atómicas).

- II. Si ' nx ' e ' ny ' son variables de tipo n y ' na ' y ' nb ' son constantes de tipo n , entonces

$$\begin{aligned} {}^nx &= {}^ny, \\ {}^nx &= {}^na, \\ {}^na &= {}^nb \end{aligned}$$

son fórmulas (atómicas)

- III. El resultado de sustituir en una fórmula atómica una variable de tipo n por un abstracto de tipo n es una fórmula.
- IV. Si X es una fórmula, $\sim X$ también lo es.
- V. Si X, Y son fórmulas, $X \vee Y$ también lo es.
- VI. Si X, Y son fórmulas, $X \wedge Y$ también lo es.
- VII. Si X, Y son fórmulas, $X \rightarrow Y$ también lo es.
- VIII. Si X, Y son fórmulas, $X \leftrightarrow Y$ también lo es.
- IX. Si X es una fórmula, $({}^nx)[X]$ también lo es.
- X. Si X es una fórmula, $\exists {}^nx[X]$ también lo es.
- XI. Ninguna expresión es una fórmula sino en virtud de I-X.

CAPÍTULO XV

LÓGICA DE RELACIONES

82. *Las propiedades poliádicas como relaciones.* — Una expresión predicativa poliádica, como ' Pxy ', puede entenderse de un modo muy natural como simbolización de que entre x e y media la relación P . Si ' P ' simboliza la propiedad ser-padre-de, entonces ' Pxy ' podría leerse: ' x es padre de y ', o sea, ' x está en la relación padre-de respecto de y '. La lectura o interpretación relacional de las expresiones predicativas diádicas o poliádicas en general estaba ya implícita en las reglas de formación de fórmulas del cálculo de predicados. Una lectura de ' Pxy ' que consistiera en la atribución de la propiedad P a x e y por separado — por ejemplo: ' x es padre e y es padre' — correspondería en efecto más bien a una fórmula como ' $Px \wedge Py$ ', pues el predicado 'padre' en esa lectura es un predicado monádico.

La falta de una teoría general de las relaciones es una de las lagunas más importantes notadas en la lógica formal de la tradición. Los razonamientos matemáticos más elementales son razonamientos sobre relaciones, que no quedan formalmente recogidos con sólo predicados monádicos, como son los de la silogística categórica aristotélica. Sea, por ejemplo, la sencilla afirmación siguiente, para cuyo análisis bastaría la lógica de clases:

- (1) Si todos los andaluces son españoles, entonces todos los hijos de andaluces son hijos de españoles.

Al intentar esquematizar (1) con predicados monádicos, hay que hacer intervenir cuatro de éstos: 'andaluces', 'españoles', 'hijo de andaluces', 'hijo de españoles'. Esquematizándolos, respectivamente, por ' P ', ' Q ', ' S ' y ' T ', tendríamos:

- (2) $(x)[Px \rightarrow Qx] \rightarrow (x)[Sx \rightarrow Tx]$.

(2) no es demostrable, pues no es un teorema formal, no es universalmente verdadera, como puede verse construyendo su tabla, o aduciendo una de las muchas interpretaciones que la hacen falsa; por ejemplo, y en el mismo

universo del discurso de (1), interpretando ' P ' por 'andaluz', ' Q ' por 'español', ' S ' por 'asturiano' y ' T ' por 'extremeño'.

La situación cambia si se explicita la relación hijo-de, simbolizable, por ejemplo, mediante el predicado diádico ' R '. Entonces la esquematización de (1) procede así:

- a) esquematización de 'todos los andaluces son españoles':

$$(x)[Px \rightarrow Qx];$$

- b) esquematización de 'todos los hijos de andaluces son hijos de españoles':

$$(x)(y)[Px \wedge Ryx \rightarrow Qx \wedge Ryx].$$

Lectura: 'si x es andaluz e y es hijo de x , entonces x es español e y es hijo de x '.

Según eso, una aceptable esquematización relacional de (1) es:

- (3) $(x)[Px \rightarrow Qx] \rightarrow (x)(y)[Px \wedge Ryx \rightarrow Qx \wedge Ryx]$.

Y (3) sí que es demostrable:

Demostración de (3)

L1(1)	$(x)[Px \rightarrow Qx]$	RP
L2(1)	$Pu \rightarrow Qu$	RE (); L1
L3(3)	$Pu \wedge Rvu$	RP
L4(3)	Pu	RE $\wedge$; L3
L5(3)	Rvu	RE $\wedge$; L4
L6(1, 3)	Qu	RE $\rightarrow$; L2, L4
L7(1, 3)	$Qu \wedge Rvu$	RI $\wedge$; L6, L5
L8(1)	$Pu \wedge Rvu \rightarrow Qu \wedge Rvu$	RI $\rightarrow$; L7
v[u] L9(1)	$(y)[Pu \wedge Ryu \rightarrow Qu \wedge Ryu]$	RI (); L8
u L10(1)	$(x)(y)[Px \wedge Ryx \rightarrow Qx \wedge Ryx]$	RI (); L9
L11(0)	$(x)[Px \rightarrow Qx] \rightarrow (x)(y)[Px \wedge Ryx \rightarrow Qx \wedge Ryx]$	RI $\rightarrow$; L10.

83. *Relaciones y clases.* — Como hemos visto bajo el epígrafe anterior, la noción de relación puede obtenerse directamente de la noción de propiedad poliádica — o sea, mediante una reinterpretación de la parte poliádica de la lógica de predicados —, exactamente igual que obtuvimos el concepto de clase directamente del de propiedad monádica. Es corriente, sin embargo, fundar la noción de relación en la de clase, con objeto de referirse no a las intensiones, a los conceptos relacionales, sino a sus extensiones; en vez de a la intención o concepto padre-de, por ejemplo, a la extensión constituida por todos los pares de objetos tales que el uno es padre del otro.

Esta definición de la idea de relación por la de clase (que fue presentada sistemáticamente por Wiener y Kuratowski) no es, desde luego, inevitable.

En realidad, la idea de relación es, desde el punto de vista de la teoría de la ciencia, una de esas nociones fundamentales cuya definición no puede tener más que una finalidad técnica, a saber, fijar una significación a algo que intuitivamente fundamenta todo pensamiento. Todo el mundo sabe lo que es una relación, pues hablar es relacionar: afirmar un enunciado es relacionar de cierto modo algunos términos. En 82 demostramos una expresión relacional sin apelar a la idea de clase. Eso puede bastarnos para considerar que la lógica de relaciones es lógica fundamental y no tiene por qué basarse necesariamente en la lógica de clases. (Más adelante veremos incluso que las clases pueden definirse mediante relaciones.) Por eso, cuando, a partir de este punto, definamos relaciones por clases —siguiendo el uso tradicional—, acompañaremos a veces esas definiciones por otras que no apelan a clases, con objeto de recordar la independencia de la lógica de relaciones respecto de la de clases.

La definición del concepto de relación por el de clase se basa en la idea de pares ordenados. Limitándonos, como haremos frecuentemente, a las relaciones diádicas, se entiende así por relación (diádica) la clase de los pares ordenados cuyos miembros hacen verdadera cierta expresión predicativa diádica. Esto es lo que expresa la siguiente definición:

$$(DR4) \quad R =_{df} \hat{x} \hat{y} [Rxy].$$

Si nos inspiráramos fielmente en la notación de la lógica de clases, podríamos introducir minúsculas griegas para representar relaciones, dejando las mayúsculas latinas como predicados poliádicos "neutros", por así decirlo, o sea, sin interpretar relacionamente. Con ello (DR4) tendría el siguiente aspecto:

$$(DR4a) \quad \rho =_{df} \hat{x} \hat{y} [Rxy].$$

Pero la costumbre de usar las mayúsculas latinas para representar relaciones está ya definitivamente arraigada, y también la seguiremos aquí. Tanto (DR4) cuanto (DR4a) se pueden leer: ' R (o ρ) es la clase de los pares ordenados $\langle x, y \rangle$ tales que Rxy '; más brevemente: 'la clase de los x e y tales que Rxy '.

La definición (DR4) corresponde a un principio de abstracción relacional análogo al de las clases. Este principio, expresado ingenuamente (es decir, sin indicación de tipos, y sin evitar por tanto la paradoja de Russell) sería (utilizando por una vez minúsculas griegas):

$$(5) \quad (\rho) \exists R(x)(y) [\langle x, y \rangle \in \rho \leftrightarrow Rxy].$$

La abstracción de relaciones impone naturalmente, como la de clases, la distinción entre tipos lógicos. No nos detendremos aquí en ello. Pero convendremos en que al escribir, por ejemplo,

${}^n R_i$,

el índice ' n ' indica el tipo lógico de la relación, y el subíndice ' i ' el número de argumentos de la misma, o sea, que R es i -ádica. Al igual que hicimos

con las clases, prescindiremos por lo común de indicar el tipo de cada símbolo, aunque alguna vez valdrá la pena aludir a él para aclarar alguna situación. Lo mismo vale de los subíndices.

La notación ' Rxy ' indicará que la relación R media entre x e y . Análogamente para ' $R_{n, x_1 \dots x_n}$ ', con $n \geq 2$. Otra notación corriente es ' xRy ', pero tiene el inconveniente de ser menos apta para $n > 2$.

84. Functores de relaciones. Conversas. — Bajo el presente epígrafe estudiaremos funtores que, como los de clases, son susceptibles de tratamiento algebraico. Pero no nos atendremos a un tratamiento puramente algebraico de los mismos, por su relativa pobreza analítica, sino que utilizaremos todos los medios analíticos que nos suministran la lógica de clases y la de predicados en general.

$$(DR6) \quad \neg R =_{df} \hat{x} \hat{y} [\sim Rxy].$$

' $\neg R$ ' se lee: 'complemento de R '. Una definición que no recurriera a la idea de clase, sino que tratara simplemente a ' $\neg R$ ' como predicado (diádico en este caso), sería:

$$(DR6a) \quad \neg Rxy \leftrightarrow_{df} \sim Rxy.$$

Ejemplo: el complemento de padre-de es la relación que media entre x e y cuando x e y no están en la relación padre-de.

$$(DR7) \quad R \cup S =_{df} \hat{x} \hat{y} [Rxy \vee Sxy].$$

' $R \cup S$ ' se lee: 'suma lógica de R y S '.

Ejemplo: la suma lógica de padre-de y madre-de es padre-o-madre-de, o sea, progenitor-de.

$$(DR7a) \quad 'R \cup S xy' \leftrightarrow_{df} 'Rxy \vee Sxy'.$$

$$(DR8) \quad R \cap S =_{df} \hat{x} \hat{y} [Rxy \wedge Sxy].$$

' $R \cap S$ ' se lee 'producto lógico de R y S '.

Ejemplo: el producto lógico de colega-de y más-viejo-que es colega-más-viejo-que, o colega-mayor-de.

$$(DR8a) \quad 'R \cap S xy' \leftrightarrow_{df} 'Rxy \wedge Sxy'.$$

$$(DR9) \quad 'R \subset S' \leftrightarrow_{df} '(x)(y) [Rxy \rightarrow Sxy]'$$

' $R \subset S$ ' se lee: ' R está incluida en S '.

Ejemplo: padre-de está incluida en familiar-de.

$$(DR10) \quad 'R = S' \leftrightarrow_{df} '(x)(y) [Rxy \leftrightarrow Sxy]'$$

' $R = S$ ' se lee ' R y S son idénticas'.

Ejemplo: tío-de y hermano-del-padre-de son idénticas.

(DR6)-(DR10) deben, naturalmente, generalizarse para relaciones n -ádicas con $n \geq 2$. Ello puede hacerse especificando para cada n esquemas de definiciones como el siguiente de ' $\cap$ ':

$$R_n \cap S_n =_{\text{def}} \hat{x}_1 \hat{x}_2 \dots \hat{x}_n [Rx_1x_2 \dots x_n \wedge Sx_1x_2 \dots x_n].$$

Suelen distinguirse estos funtores de los correspondientes funtores de clases añadiéndoles un punto superior. Aquí renunciaremos a esa precaución, que no nos resultará necesaria para evitar ambigüedades.

Las nociones de relación universal y relación nula se definen del modo siguiente (limitándonos a relaciones diádicas):

$$(DR11) \quad V =_{\text{def}} \hat{x} \hat{y} [x = x \wedge y = y].$$

El fondo intuitivo de (DR11) es que la relación universal es la que media entre cualesquiera objetos por el hecho de ser objetos. La expresión sin alusión a la idea de clase sería:

$$(DR11a) \quad 'Vxy' \leftrightarrow_{\text{def}} 'x = x \wedge y = y'.$$

(DR11a) puede leerse: 'dos objetos, x e y , están entre sí en la relación universal si y sólo si cada uno de ellos es idéntico a sí mismo'.

$$(DR12) \quad \Lambda =_{\text{def}} \hat{x} \hat{y} [\sim x = x \wedge \sim y = y].$$

Para expresar que una relación no es nula escribiremos ' $\exists ! R$ ', de acuerdo (para relaciones diádicas) con la definición siguiente:

$$(DR13) \quad '\exists ! R' \leftrightarrow_{\text{def}} '\exists x \exists y Rxy'.$$

Una relación, Q , se llama la *conversa* de otra, R , si y sólo si vale la equivalencia ' $(x)(y)[Rxy \leftrightarrow Qyx]$ '. Para indicar la conversa de R se escribe ' $\bar{R}$ '.

$$(14) \quad Rxy \leftrightarrow Ryx.$$

Esto nos lleva a la siguiente definición:

$$(DR15) \quad \bar{R} =_{\text{def}} \hat{x} \hat{y} [Ryx].$$

Ejemplo: la conversa de progenitor-de es hijo-o-hija-de.

85. Productos y potencias relacionales. — El producto lógico de dos relaciones, definido por (DR8) bajo el epígrafe anterior, es una nueva relación que suele expresarse por un predicado compuesto, como 'colega-mayor-de'. Pero es frecuente que la composición de relaciones interese para obtener nuevas relaciones que en el lenguaje común se encuentran como predicados simples. Eso ocurre cuando las dos relaciones compuestas lo están a través de un objeto que participa de ambas. Así, por ejemplo si, dadas las relaciones hijo-de y hermano-de hay un individuo que participa de ambas puede ocurrir (por ejemplo, si el individuo común es segundo argumento de la primera relación y primero de la segunda) que aparezca una nueva relación corrientemente nombrada por un predicado simple (en el caso del ejemplo, la relación sobrino-de entre el primer argumento de la primera relación y el segundo de la segunda). Esta composición de relaciones se llama 'producto relacional' o 'producto relativo'. Se simboliza por

un trazo vertical entre las relaciones afectadas, y se define (para dos relaciones, por de pronto) del modo siguiente:

$$(DR16) \quad R | S =_{\text{def}} \hat{x} \hat{y} [\exists z [Rxz \wedge Szy]].$$

Sin alusión a clases puede escribirse:

$$(DR17) \quad 'R | Sxy' \leftrightarrow_{\text{def}} '\exists z [Rxz \wedge Szy]',$$

lo cual puede leerse, con la interpretación anterior por ejemplo: ' x es sobrino de y ' equivale por definición a 'hay un z , tal que x es hijo de z y z es hermano de y '.

El producto relacional es asociativo, lo que nos exige de definirlo para más de dos relaciones:

$$(TR18) \quad [R | S] | P = R | [S | P].$$

Interpretando ' P ' por padre-de y ' R ' y ' S ' según el ejemplo anterior, ' $[R | S] | P$ ' es 'sobrino-del-padre-de', y ' $R | [S | P]$ ' es 'hijo-del-tío-de', o sea, 'primo-hermano-de' en ambos casos.

(TR18) es fácilmente demostrable en su versión o traducción lógico-predicativa que, teniendo en cuenta (DR17), es:

$$(19) \quad (x)(y) [\exists z [\exists u [Rxu \wedge Suz] \wedge Pzy] \leftrightarrow \exists z [Rxz \wedge \exists u [Szu \wedge Puy]]].$$

Esta traducción se ha conseguido a través de los siguientes pasos (a)-(g):
Fórmula inicial (TR18): $[R | S] | P = R | [S | P]$.

Por (DR17):

$$[R | S] | Pxy \leftrightarrow \exists z [R | Sxz \wedge Pzy] \quad (\text{paso (a)}).$$

También por (DR17):

$$R | Sxz \leftrightarrow \exists u [Rxu \wedge Suz] \quad (\text{paso (b)}).$$

Por (a) y (b):

$$[R | S] | Pxy \leftrightarrow \exists z [\exists u [Rxu \wedge Suz] \wedge Pzy] \quad (\text{paso (c)}).$$

Por (DR17):

$$R | [S | P]xy \leftrightarrow \exists z [Rxz \wedge S | Pzy] \quad (\text{paso (d)}).$$

También por (DR17):

$$S | Pzy \leftrightarrow \exists u [Szu \wedge Puy] \quad (\text{paso (e)}).$$

Por (d) y (e):

$$R | [S | P]xy \leftrightarrow \exists z [Rxz \wedge \exists u [Szu \wedge Puy]] \quad (\text{paso (f)}).$$

La introducción de ' $\leftrightarrow$ ' en (19) en vez del ' $=$ ' de (TR18) está justificada, finalmente, por (DR10), que exige además la generalización de ' x ' e ' y '.
(paso (g)).

He aquí, a título de ejemplo de demostración de teoremas de la lógica de relaciones en el algoritmo de la lógica de predicados, la demostración de (TR18) o, más propiamente, de su traducción (19), en el cálculo de la deducción natural:

Demostración de (19)

L1(1)	$\exists z [\exists u [R_{xu} \wedge S_{uz}] \wedge P_{zy}]$	RP.
$a [x, y]$ L2(1)	$\exists u [R_{xu} \wedge S_{ua}] \wedge P_{ay}$	RE E; L1.
$b [a, x, y]$ L3(1)	$[R_{xb} \wedge S_{ba}] \wedge P_{ay}$	RE E; L2.
L4(1)	$R_{xb} \wedge S_{ba}$	RE A; L3.
L5(1)	P_{ay}	RE A; L4.
L6(1)	R_{xb}	RE A; L4.
L7(1)	S_{ba}	RE A; L4.
L8(1)	$S_{ba} \wedge P_{ay}$	RI A; L7, L5.
L9(1)	$\exists u [S_{bu} \wedge P_{uy}]$	RI E; L8.
L10(1)	$R_{xb} \wedge \exists u [S_{bu} \wedge P_{uy}]$	RI A; L9.
L11(1)	$\exists z [R_{xz} \wedge \exists u [S_{zu} \wedge P_{uy}]]$	RI E; L10.
L12(0)	$\exists z [\exists u [R_{xu} \wedge S_{uz}] \wedge P_{zy}] \rightarrow \exists z [R_{xz} \wedge \exists u [S_{zu} \wedge P_{uy}]]$	RI $\rightarrow$; L11.
L13(13)	$\exists z [R_{xz} \wedge \exists u [S_{zu} \wedge P_{uy}]]$	RP.
$c [x, y]$ L14(13)	$R_{xc} \wedge \exists u [S_{cu} \wedge P_{uy}]$	RE E; L13.
L15(13)	R_{xc}	RE A; L14.
L16(13)	$\exists u [S_{cu} \wedge P_{uy}]$	RE A; L14.
$d [c, y]$ L17(13)	$S_{cd} \wedge P_{dy}$	RE E; L16.
L18(13)	S_{cd}	RE A; L17.
L19(13)	P_{dy}	RE A; L17.
L20(13)	$R_{xc} \wedge S_{cd}$	RI A; L15, L18.
L21(13)	$\exists u [R_{xu} \wedge S_{ud}]$	RI E; L20.
L22(13)	$\exists u [R_{xu} \wedge S_{ud}] \wedge P_{dy}$	RI A; L21, L19.
L23(13)	$\exists z [\exists u [R_{xu} \wedge S_{uz}] \wedge P_{zy}]$	RI E; L22.
L24(0)	$\exists z [R_{xz} \wedge \exists u [S_{zu} \wedge P_{uy}]] \rightarrow \exists z [\exists u [R_{xu} \wedge S_{uz}] \wedge P_{zy}]$	RI $\rightarrow$; L23.
L25(0)	$\exists z [\exists u [R_{xu} \wedge S_{uz}] \wedge P_{zy}] \leftrightarrow \exists z [R_{xz} \wedge \exists u [S_{zu} \wedge P_{uy}]]$	RI $\leftrightarrow$; L12, L24.
$y [x]$ L26(0)	$(y) [\exists z [\exists u [R_{xu} \wedge S_{uz}] \wedge P_{zy}] \leftrightarrow \exists z [R_{xz} \wedge \exists u [S_{zu} \wedge P_{uy}]]]$	RI ($\cdot$); L25.
x L27(0)	$(x)(y) [\exists z [\exists u [R_{xu} \wedge S_{uz}] \wedge P_{zy}] \leftrightarrow \exists z [R_{xz} \wedge \exists u [S_{zu} \wedge P_{uy}]]]$	RI ($\cdot$); L26.

El teorema (19), que equivale por definición a (TR18), permite escribir simplemente ' $R | S | P$ ' en vez de ' $R | [S | P]$ ' o de ' $[R | S] | P$ '.

En cambio, el producto relacional no es conmutativo. Si, con el ejemplo ya utilizado, ' $R | S$ ' es 'sobrino-de' ('hijo-del-hermano-de'), entonces ' $S | R$ ' es 'hermano-del-hijo-de', o sea, si por 'hermano' seguimos entendiendo — como hasta ahora tácitamente — 'hermano de padre y madre', ' $S | R$ ' es simplemente 'hijo-de'.

Las potencias de una relación se construyen con el producto relacional, del mismo modo que las potencias de números con el producto de números. La notación es análoga:

$$(DR20) \quad R^1 =_{\text{df}} R.$$

$$(DR21) \quad R^2 =_{\text{df}} R | R.$$

$$(DR22) \quad R^n =_{\text{df}} R^{n-1} | R.$$

La noción de potencia de una relación se generaliza para poder dar sentido al exponente '0' y a exponentes negativos. La generalización más natural, sistematizada por R. Carnap siguiendo la línea de *Principia Mathematica*, es como sigue:

$$(DR23) \quad R^0 =_{\text{df}} I.$$

R^0 es la relación de cualquier objeto, que pueda ser argumento de R , consigo mismo, o sea, la relación identidad.

$$(DR24) \quad R^{-1} =_{\text{df}} \bar{R}.$$

$$(DR25) \quad R^{-n} =_{\text{df}} \bar{R}^n.$$

En la intención de Carnap, esas definiciones fundamentan la introducción de los números negativos (enteros) como exponentes de relaciones.

86. Dominios y campo. Descripciones relacionales. — Interesa frecuentemente referirse a las clases de objetos que pueden intervenir en una relación, o a los objetos mismos. Por ejemplo, a la clase de los objetos que están en la relación descendiente-de con algún objeto dado, o con una clase de objetos, o con objetos en general. Las siguientes definiciones facilitan un modo de hacerlo:

$$(DR26) \quad D_1 R_2 =_{\text{df}} \hat{x} [\exists y R_2 xy].$$

' D_1 ', que llamaremos 'dominio anterior', es un functor que, aplicado a una relación (argumento), da como valor una clase. El dominio anterior de la relación R_2 es la clase de los objetos x tales que x está en la relación R_2 con algo. (Esto es, dicho sea de paso, la definición de una clase mediante una relación.) La notación tradicional de ' $D_1 R$ ' es ' $D'R$ '. (La notación funcional ' $D_1 R$ ' procede de Carnap. Pero para que de verdad sea funcional, (DR26)

no tiene que ser una definición, sino, como en Carnap, simplemente una equivalencia. Tal como se presenta en (DR26), ' D_1R ' es una descripción, exactamente igual que la notación tradicional.) En vez de 'dominio anterior' es corriente decir simplemente 'dominio'.

$$(DR27) \quad D_2R_2 =_{df} \hat{y} [\exists x R_2xy].$$

Leeremos ' D_2 ' 'dominio posterior'. Es frecuente también la lectura 'dominio converso'.

Notaciones predicativas de (DR26) y (DR27) pueden ser las siguientes, sin alusión a la idea de clase:

$$(DR28a) \quad 'D_1R_2x' \leftrightarrow_{df} ' \exists y Rxy '.$$

$$(DR28b) \quad 'D_2R_2y' \leftrightarrow_{df} ' \exists x Rxy '.$$

De (DR26) y (DR27) y la definición de $\tilde{R}_2$ (DR15), se obtiene

$$(TR29) \quad D_1R_2 = D_2\tilde{R}_2.$$

Campo de una relación R es la suma lógica de su dominio anterior y su dominio posterior (tratándose de relaciones diádicas). Dicho sin apelar a clases: un objeto x tiene la propiedad de ser un objeto relacionado por R cuando está en la relación R con algún otro objeto o cuando algún otro objeto está en la relación R con él:

$$(DR30) \quad CR_2 =_{df} D_1R_2 \cup D_2R_2.$$

$$(DR30a) \quad 'CR_2x' \leftrightarrow_{df} ' \exists y [R_2xy \vee R_2yx] '.$$

Por (TR29) y (DR30):

$$(TR31) \quad CR = C\tilde{R}.$$

La generalización de las nociones de dominio y campo a relaciones n -ádicas con $n \geq 2$ no presenta más novedad que la necesidad de prescindir de las ideas de anterior y posterior, sustituyéndolas por la idea de dominio del lugar m ($m \leq n$) de la fórmula relacional. Nuestro uso de subíndices numéricos afectando a ' D ' sugiere ya esta generalización. Así podemos escribir el siguiente esquema de definiciones, a partir del cual pueden obtenerse definiciones propiamente dichas para $n = 1$, $n = 2$, etc.

$$(DR32) \quad D_mR_n =_{df} \hat{x}_m [\exists x_1 \exists x_2 \dots \exists x_{m-1} \exists x_{m+1} \dots \exists x_n R x_1 \dots x_m \dots x_n].$$

$$(DR33) \quad CR_n =_{df} D_1R_n \cup D_2R_n \cup \dots \cup D_nR_n.$$

Las anteriores nociones son útiles, entre otras cosas, para formalizar las nociones de 'primer miembro' y 'último miembro' de órdenes u ordenaciones

realizadas por una relación diádica. Primeros miembros de una tal ordenación por una relación, R_2 , serán aquellos que pertenezcan a D_1R_2 y no pertenezcan a D_2R_2 ; y últimos miembros de dicha ordenación serán aquellos que pertenezcan a D_2R y no pertenezcan a D_1R :

$$(DR34) \quad \text{primeros-}R =_{df} \hat{x} [x \in D_1R \wedge \sim x \in D_2R].$$

En la sucesión de los números naturales (con módulo 0) basada en la relación 'precedente-inmediato', 0 es un primer miembro, pues vale (simbolizando 'precedente inmediato' por ' R '):

$$0 \in D_1R \wedge \sim 0 \in D_2R, \text{ o sea} \\ 0 \in \hat{x} [x \in D_1R \wedge \sim x \in D_2R].$$

$$(DR35) \quad \text{últimos-}R =_{df} \hat{x} [x \in D_2R \wedge \sim x \in D_1R].$$

En el conjunto de los números naturales menores que diez, la ordenación menor-que tiene un último miembro que es, naturalmente, nueve. Pues vale (simbolizando 'menor-que' por ' R '):

$$9 \in D_2R \wedge \sim 9 \in D_1R, \text{ o sea} \\ 9 \in \hat{x} [x \in D_2R \wedge \sim x \in D_1R].$$

Las nociones de primer miembro y de último miembro de una ordenación pueden generalizarse para obtener la noción de miembro de lugar m :

$$(DR34a, 35a) \quad \text{miembro-}m\text{-}R_n =_{df} \hat{x} [x \in D_mR_n \wedge \sim x \in D_1R_n \wedge$$

$$\wedge \sim x \in D_2R_n \wedge \dots \wedge \sim x \in D_{m-1}R_n \wedge \\ \wedge \sim x \in D_{m+1}R_n \wedge \dots \wedge \sim x \in D_nR_n].$$

Para expresar que una ordenación por R no tiene ningún primer miembro se escribe:

$$(36) \quad (x) [x \in D_1R \rightarrow x \in D_2R],$$

en cuyo caso D_1R está incluida en D_2R (cfr. cap. XIV (DC17)), y, por tanto

$$(36a) \quad D_2R = CR.$$

Para expresar que una ordenación por R no tiene ningún último miembro se escribe

$$(37) \quad (x) [x \in D_2R \rightarrow x \in D_1R],$$

lo que significa que $D_2R \subset D_1R$, y, por tanto,

$$(37a) \quad D_1R = CR.$$

Hemos visto que, a pesar de la notación funcional, ' $D_n R$ ' es una descripción (de una determinada clase). Por el hecho de que, como vimos en 36, es deseable posibilitar la eliminación de descripciones, hemos introducido también una notación predicativa de las ideas de dominio y campo, sin alusión a clase determinada, sino mediante un predicado, ' $D_n R$ ' en el caso de ' D ', que representa la propiedad $D_n R$ (es la definición (DR28a)). Esta posibilidad debe tenerse en cuenta al estudiar las numerosas descripciones cuyo uso es frecuente en lógica de relaciones.

Así, por ejemplo, puede interesar referirse al concreto objeto que está en la relación R con el objeto y . Para ello puede empezarse por describirlo mediante R por ejemplo, como 'el R de y ' ('el padre de Cervantes'). La notación corriente es

$$(DR38) \quad R'y =_{\text{at}} (\lambda x) Rxy.$$

Con una notación sin descripciones podemos escribir:

$$(DR38a) \quad 'R'y \ x \leftrightarrow_{\text{at}} \exists x [Rxy \wedge (z) [Rzy \rightarrow z = x]].$$

$$(DR39) \quad \bar{R}'x =_{\text{at}} (\lambda y) Rxy.$$

Otras veces, la entidad a la que interesa referirse será no el x individual que está en la relación R con y , sino, porque haya más de un tal x , la concreta clase de éstos; por ejemplo, la clase de los objetos que están con Cervantes en la relación antecesor-de (o sea, la clase de los antecesores de Cervantes). Clases así son las definidas mediante descripciones como la siguiente:

$$(DR40) \quad \vec{R}'y =_{\text{at}} \hat{x} [Rxy].$$

$\vec{R}'y$ puede leerse: 'los R de y '.

La clase de los objetos con los cuales x está en la relación R se expresa por la descripción

$$(DR41) \quad \overleftarrow{R}'x =_{\text{at}} \hat{y} [Rxy].$$

$\overleftarrow{R}'x$ puede leerse: 'los y cuyo R es x '.

Ejemplo: si ahora ' x ' es 'Cervantes' y ' R ' sigue siendo 'antecesor', entonces ' $\overleftarrow{R}'x$ ' es 'la clase de los objetos cuyo antecesor es Cervantes'.

La siguiente definición describe el dominio anterior de una relación diádica, como en (DR38)-(DR41), para un dominio posterior reducido a objetos de una clase determinada, α . Por ejemplo, tratándose de la relación antecesor, como antes, α puede ser la clase de los descendientes que han vivido más de veinte años.

$$(DR42) \quad R''\alpha =_{\text{at}} \hat{x} [\exists y [y \in \alpha \wedge Rxy]].$$

Puesto que $\alpha \subseteq D_2 R$, de (DR42) se desprende

$$(TR43) \quad R''\alpha \subseteq D_1 R.$$

Restricciones parecidas de los dominios de una relación expresan las siguientes definiciones descriptivas, que suelen llamarse 'limitaciones' o 'especificaciones':

$$(DR44) \quad \alpha \upharpoonright R =_{\text{at}} \hat{x} \hat{y} [x \in \alpha \wedge Rxy].$$

(DR44) define la limitación de la relación R a un dominio anterior constituido sólo por la clase α .

$$(DR45) \quad R \upharpoonright \alpha =_{\text{at}} \hat{x} \hat{y} [y \in \alpha \wedge Rxy].$$

(DR45) define la limitación de la relación R a un dominio posterior constituido sólo por la clase α .

$$(DR46) \quad \alpha \upharpoonright R \upharpoonright \beta =_{\text{at}} \hat{x} \hat{y} [x \in \alpha \wedge y \in \beta \wedge Rxy].$$

$$(DR47) \quad R \upharpoonright \alpha =_{\text{at}} \alpha \upharpoonright R \upharpoonright \alpha.$$

La siguiente notación se usa para expresar la relación que media entre dos objetos, x e y , por el hecho de pertenecer x a la clase α e y a la clase β :

$$(DR48) \quad \uparrow \alpha \beta =_{\text{at}} \hat{x} \hat{y} [x \in \alpha \wedge y \in \beta].$$

87. Clases de relaciones diádicas. — Por el paralelismo entre las expresiones predicativas monádicas y las de clases, atribuir una propiedad, P , a una relación es lo mismo que decir que esa relación pertenece a la clase, α , abstraída de una fórmula predicativa del predicado monádico ' P '. Por ejemplo: decir que una relación, R , es reflexiva, es lo mismo que decir que R pertenece a la clase de las relaciones reflexivas. Esta propiedad (y clase) de relaciones es la primera que vamos a definir:

Reflexividad

$$(DR49) \quad \text{'Reflex } R' \leftrightarrow_{\text{at}} '(x) [x \in CR \rightarrow Rxx]'$$

(DR49) es una definición del predicado 'Reflex'. Define la expresión ' R es reflexiva'. Si se quiere una expresión más puramente predicativa puede sustituirse la notación de clases ' $x \in CR$ ' (' x pertenece a la clase CR ') por ' $\exists z [Rxx \vee Rzx]$ ' (' x está en la relación R con algo o algo está en la relación R con x ');

$$(DR49a) \quad \text{'Reflex } R' \leftrightarrow_{\text{at}} '(x) [\exists z [Rxx \vee Rzx] \rightarrow Rxx]'$$

Hablando, en cambio, de la clase de las relaciones reflexivas, en vez de hablar de la propiedad reflexividad, escribiremos:

$$(DR50) \quad \text{Reflex} =_{\text{at}} \hat{R} [(x) [x \in CR \rightarrow Rxx]],$$

y

$$(DR51) \quad 'R \in \text{Reflex}' \leftrightarrow_{\text{at}} '(x) [x \in CR \rightarrow Rxx]'$$

La comparación de (DR49) con (DR51) muestra el paralelismo de la notación de clases con la predicativa. A partir de ahora la notación será la de clases.

Se dice que una relación es reflexivo-total, o totalmente reflexiva (RefIt), cuando es reflexiva y su campo es la clase universal:

$$(DR52) \quad \text{RefIt} =_{\text{df}} \hat{R} [(x) Rxx].$$

$$(DR53) \quad 'R \in \text{RefIt}' \leftrightarrow_{\text{df}} '(x) Rxx'.$$

La clase de las relaciones no-reflexivas (NRefI) se define del modo siguiente:

$$(DR54) \quad \text{NRefI} =_{\text{df}} \hat{R} [\exists x [x \in CR \wedge \sim Rxx]].$$

Una relación no-reflexiva es una relación que no es reflexiva para todo x (puesto que hay al menos un x para el cual no lo es), pero puede serlo para algún otro x . En cambio, una relación irreflexiva (IrrefI) es aquella que no es reflexiva para ningún x :

$$(DR55) \quad \text{IrrefI} =_{\text{df}} \hat{R} [(x) \sim Rxx].$$

Ejemplos: La relación de identidad es reflexiva; la relación mayor-que es irreflexiva; la relación admirador-de es no-reflexiva.

Simetría

$$(DR56) \quad \text{Sim} =_{\text{df}} \hat{R} [(x)(y) [Rxy \rightarrow Ryx]].$$

Una relación simétrica está contenida en su conversa. Por (DR56) se tiene en efecto:

$$(57) \quad R \in \text{Sim} \leftrightarrow (x)(y) [Rxy \rightarrow Ryx].$$

Pero Ryx es $\check{R}xy$ (cfr. (DR15)). Por (DR9):

$$(TR58) \quad R \in \text{Sim} \leftrightarrow R \subset \check{R}.$$

Por otra parte, como $\check{R}$ será también simétrica si lo es R , se tendrá

$$(TR59) \quad R \in \text{Sim} \leftrightarrow \check{R} \subset R.$$

De (TR58) y (TR59)

$$(TR60) \quad R \in \text{Sim} \leftrightarrow R = \check{R}.$$

La demostración del teorema (TR60) utiliza un teorema auxiliar

$$S \subset Q \wedge Q \subset S \rightarrow S = Q$$

cuya demostración en el cálculo de predicados es muy sencilla. Por (DR9) y (DR10) se trata de demostrar la fórmula

$$(x)(y) [Sxy \rightarrow Qxy] \wedge (x)(y) [Qxy \rightarrow Sxy] \rightarrow (x)(y) [Sxy \leftrightarrow Qxy].$$

A continuación definimos la clase de las relaciones no simétricas: (NSim):

$$(DR61) \quad \text{NSim} =_{\text{df}} \hat{R} [\exists x \exists y [Rxy \wedge \sim Ryx]].$$

Por último, la clase de las relaciones asimétricas (ASim):

$$(DR62) \quad \text{ASim} =_{\text{df}} \hat{R} [(x)(y) [Rxy \rightarrow \sim Ryx]].$$

Ejemplos: La relación antípoda-de es simétrica; la relación mayor-que es asimétrica; la relación primo-de es no-simétrica (en cambio la relación primo-o-prima-de es simétrica).

Transitividad (Trans).

$$(DR63) \quad \text{Trans} =_{\text{df}} \hat{R} [(x)(y) [\exists z [Rxz \wedge Rzy] \rightarrow Rxy]].$$

Una relación transitiva contiene su segunda potencia. Por (DR63) se tiene, en efecto:

$$(DR64) \quad 'R \in \text{Trans}' \leftrightarrow_{\text{df}} '(x)(y) [\exists z [Rxz \wedge Rzy] \rightarrow Rxy]';$$

de donde, por (DR16) y (DR9):

$$R \in \text{Trans} \leftrightarrow R \mid R \subset R, \text{ o sea}$$

$$(TR65) \quad R \in \text{Trans} \leftrightarrow R^2 \subset R.$$

((TR65) puede también considerarse una definición.)

Ejemplo: El cuadrado de la relación transitiva mayor-que (llamémosle: 'mayor² que') está contenido en la misma relación mayor-que. 'x es mayor² que y' significa: 'x es mayor que un z que es mayor que y'. Vale pues 'x es mayor que y'.

A continuación definimos la clase de las relaciones no-transitivas (NTrans):

$$(DR66) \quad \text{NTrans} =_{\text{df}} \hat{R} [\exists x \exists y \exists z [Rxz \wedge Rzy \wedge \sim Rxy]].$$

Por último, la clase de las relaciones intransitivas (InTrans):

$$(DR67) \quad \text{InTrans} =_{\text{df}} \hat{R} [(x)(y)(z) [Rxz \wedge Rzy \rightarrow \sim Rxy]].$$

Ejemplos: La relación descendiente-de es transitiva; padre-de es intransitiva; colega-de es no transitiva.

Conexividad

Una relación diádica se llama 'conexa' cuando se da entre todo par de objetos distintos de su campo.

$$(DR67) \quad \text{Conex} =_{\text{df}} \hat{R} [(x)(y) [x \in CR \wedge y \in CR \wedge x \neq y \rightarrow Rxy \vee Ryx]].$$

Si además de ser conexa la relación es simétrica y transitiva, se tendrá en el consecuente de (DR67) ' $Rxy \wedge Ryx$ ' en vez de ' $Rxy \vee Ryx$ '. Por tanto, la relación será reflexiva. (Siendo simétrica, no hará falta escribir ' $Rxy \wedge Ryx$ '; bastará ' Rxy ', o ' Ryx '.)

Las relaciones entre las diversas propiedades o clases de relaciones son de bastante interés. Expresaremos algunas por unos cuantos teoremas.

(TR68) $\text{Trans} \cap \text{Sim} \subset \text{Refl}$.

Veamos, a título de ejemplo, la demostración de este teorema en un algoritmo general de la lógica de predicados, indicando antes los pasos que llevan de esa expresión algebraica a su traducción por una expresión puramente predicativa.

Eliminaremos, primero, los elementos algebraicos de la fórmula, o sea, ' $\cap$ ' y ' $\subset$ ', sirviéndonos de las definiciones correspondientes: (DC8) y (DC9) de la lógica de clases. Paso a paso, esta operación cubre (68a) y (68b):

(68a) $(R) [R \in \text{Trans} \cap \text{Sim} \rightarrow R \in \text{Refl}]$.

(68b) $(R) [R \in \text{Trans} \wedge R \in \text{Sim} \rightarrow R \in \text{Refl}]$.

En este punto prescindimos de la cuantificación de ' R ', con lo que este símbolo quedará reducido de variable a parámetro, como lo son todos los símbolos predicativos en la lógica de predicados de primer orden. Este mismo procedimiento hemos venido usando hasta ahora (cfr. p. e. (DR51)) para evitar expresiones de orden superior. Tendremos:

(68c) $R \in \text{Trans} \wedge R \in \text{Sim} \rightarrow R \in \text{Refl}$.

De aquí, por las definiciones (DR64), (DR57) y (DR49a), obtenemos:

(68d) $(x)(y) [\exists z [Rxx \wedge Rzy] \rightarrow Rxy] \wedge (x)(y) [Rxy \rightarrow Ryx] \rightarrow$
 $\rightarrow (x) [\exists z [Rxx \vee Rzx] \rightarrow Rxx]$.

Esta expresión de (TR68) está ya libre de signos de clases y es sometible a un cálculo de predicados. He aquí su demostración en el cálculo de la deducción natural.

Demostración de (68d)

L1(1)	$(x)(y) [\exists z [Rxx \wedge Rzy] \rightarrow Rxy] \wedge$	
	$\wedge (x)(y) [Rxy \rightarrow Ryx]$	RP
L2(1)	$(x)(y) [\exists z [Rxx \wedge Rzy] \rightarrow Rxy]$	RE $\wedge$; L1
L3(1)	$(x)(y) [Rxy \rightarrow Ryx]$	RE $\wedge$; L1
L4(1)	$(y) [\exists z [Rxx \wedge Rzy] \rightarrow Rxy]$	RE (); L2
L5(1)	$\exists z [Rxx \wedge Rzx] \rightarrow Rxx$	RE (); L4.

(Este expediente — aprovechar la sustitución autorizada por la regla RE () para sustituir ' y ' por ' x ' — se usa repetidamente; p. e., L7.)

L6(1)	$(y) [Rxy \rightarrow Ryx]$	RE (); L3
L7(1)	$Rxx \rightarrow Rxx$	RE (); L6
L8(8)	$\exists z [Rxx \vee Rzx]$	RP
$z[x]$ L9(8)	$Rxx \vee Rzx$	RE $\exists$; L8
L10(10)	Rxx	RP
L11(1, 10)	Rxx	RE $\rightarrow$; L7, L10
L12(1, 10)	$Rxx \wedge Rxx$	RI $\wedge$; L10, L11
L13(1)	$Rxx \rightarrow Rxx \wedge Rxx$	RI $\rightarrow$; L12
L14(1)	$(y) [Rzy \rightarrow Ryz]$	RE (); L3
L15(1)	$Rxx \rightarrow Rxx$	RE (); L14
L16(16)	Rxx	RP
L17(16, 1)	Rxx	RE $\rightarrow$; L15, L16
L18(1, 16)	$Rxx \wedge Rxx$	RI $\wedge$; L16, L17
L19(1)	$Rxx \rightarrow Rxx \wedge Rxx$	RI $\rightarrow$; L18
L20(1, 8)	$Rxx \wedge Rxx$	RE $\vee$; L9, L13, L19
L21(1, 8)	$\exists z [Rxx \wedge Rzx]$	RI $\exists$; L20
L22(1, 8)	Rxx	RE $\rightarrow$; L5, L21
L23(1)	$\exists z [Rxx \vee Rzx] \rightarrow Rxx$	RI $\rightarrow$; L22
x L24(1)	$(x) [\exists z [Rxx \vee Rzx] \rightarrow Rxx]$	RI (); L23
L25(0)	$(x)(y) [\exists z [Rxx \wedge Rzy] \rightarrow Rxy] \wedge$ $\wedge (x)(y) [Rxy \rightarrow Ryx] \rightarrow$ $\rightarrow (x) [\exists z [Rxx \vee Rzx] \rightarrow Rxx]$	RI $\rightarrow$; L24

La fórmula de la línea L25 es (68d).

No siempre se presenta el contraste de la anterior demostración entre la sencillez de los principios usados y la pesadez de su puesta en práctica. Pero lo que sí es común a gran número de demostraciones de fórmulas relacionales es el "truco" que consiste en sustituir unas variables individuales por otras. Así ocurre por ejemplo con el teorema, de demostración muy sencilla,

(TR69) $R \in \text{Sim} \leftrightarrow \check{R} \in \text{Sim}$,

que es, en formulación predicativa de primer orden,

(69a) $(x)(y) [Rxy \rightarrow Ryx] \leftrightarrow (x)(y) [Ryx \rightarrow Rxy]$.

La demostración se reduce a un mero cambio de variables, de nombres de lugares individuales, por así decirlo. He aquí su primera mitad; la segunda es completamente análoga:

L1(1)	(x)(y) [Rxy → Ryx]	RP
L2(1)	(y) [Ray → Rya]	RE (); L1
L3(1)	Rab → Rba	RE (); L2
a[b]	L4(1) (y) [Ryb → Rby]	RI (); L3
b	L5(1) (x)(y) [Ryx → Rxy]	RI (); L4
L6(0)	(x)(y) [Rxy → Ryx] → (x)(y) [Ryx → Rxy]	RI →; L5.

(Al demostrar el condicional inverso, '(x)(y) [Ryx → Rxy] → (x)(y) [Rxy → Ryx]', hay que usar, naturalmente, dos nuevos símbolos individuales, 'c' y 'd', por ejemplo, para los pasos de RE () análogos a L2 y L3 de la primera mitad de la demostración. Si no se hace así se incurre en doble anotación.)

Los dos ejemplos anteriores bastarán para ilustrar los procedimientos de demostración con fórmulas poliádicas. A partir de ahora nos limitaremos por lo general a argüir mediante consideraciones no-calculísticas la verdad de los teoremas afirmados, salvo en algunos casos de mayor interés. Por lo demás, la completud del cálculo de predicados de primer orden garantiza la demostrabilidad en él de las verdades formales de primer orden.

$$(TR70) \quad R \in \text{ASim} \leftrightarrow R^2 \in \text{Irrefl.}$$

Supongamos que R fuera asimétrica, pero R^2 no fuera irreflexiva. Si R^2 no es irreflexiva, entonces hay al menos un x para el cual vale R^2xx , y hay un y tal que valen ' Rxy ' y ' Ryx '. Pero entonces R no es asimétrica. Supongamos, a la inversa, que R^2 fuera irreflexiva, pero R no fuera asimétrica. Si R no es asimétrica, entonces hay al menos un y para el que valen ' Rxy ' y ' Ryx '. Pero entonces vale ' R^2xx ', luego R^2 no es irreflexiva.

$$(TR71) \quad \text{ASim} \subset \text{Irrefl.}$$

Supongamos que una relación R sea asimétrica, pero no sea irreflexiva. Si no es irreflexiva, vale entonces, para al menos un x , ' Rxx '. Pero ' Rxx ', de valer, vale, por así decirlo, "en las dos direcciones", puesto que sus dos argumentos son idénticos. Dicho de otro modo: si vale ' Rxx ', vale también ' Rxx '. Consiguientemente, R no es asimétrica en ese supuesto.

$$(TR72) \quad \text{Trans} \cap \text{ASim} = \text{Trans} \cap \text{Irrefl.}$$

(TR71) justifica la afirmación de que el producto lógico de la izquierda de (TR72) está contenido en el de la derecha. Para demostrar la identidad queda por justificar la afirmación inversa. Supongamos que R sea transitiva e irreflexiva. Entonces no pueden valer a la vez ' Rxy ' y ' Ryx ' para ningún x y ningún y , ya que, por ser R transitiva, valdría entonces también, para un x al menos, ' Rxx ', cosa que no puede ser porque se supone a R irreflexiva. Pero si no pueden valer para ningún x y ningún y a la vez ' Rxy ' y ' Ryx ', es que R es asimétrica.

$$(TR73) \quad R \in \text{Irrefl} \wedge Q \subset R \rightarrow Q \in \text{Irrefl.}$$

Si Q no es irreflexiva, hay al menos un x para el cual vale ' Qxx '. Pero, por $Q \subset R$, vale ' $Qxx \rightarrow Rxx$ '. Luego R tampoco sería irreflexiva. Este mismo esquema de argumentación vale para (TR74) y (TR75):

$$(TR74) \quad R \in \text{ASim} \wedge Q \subset R \rightarrow Q \in \text{ASim.}$$

$$(TR75) \quad R \in \text{InTrans} \wedge Q \subset R \rightarrow Q \in \text{InTrans.}$$

$$(TR76) \quad R^2 \in \text{Irrefl} \rightarrow R \text{ Irrefl.}$$

Si R no fuera irreflexiva, entonces habría al menos un x tal que valdría ' Rxx '. Pero entonces valdría ' $Rxx \wedge Rxx$ ', o sea, ' R^2xx '. Luego R^2 no sería tampoco irreflexiva.

$$(TR77) \quad R \in \text{Irrefl} \cap \text{Trans} \rightarrow R^2 \in \text{Irrefl.}$$

Si R^2 no es irreflexiva, entonces vale ' $Rxy \wedge Ryx$ ' para al menos un x . R tiene entonces que ser al menos no-reflexiva o al menos no-transitiva (y puede ser incluso reflexiva o intransitiva).

88. Relaciones y clases de equivalencia. — Se llama 'relaciones de equivalencia', o, simplemente — pero con cierto riesgo de ambigüedad — 'equivalencias', a las relaciones que son simétricas y transitivas. Estas relaciones son también, naturalmente, reflexivas (cfr. (TR68)). La igualdad entre números es una relación de equivalencia, pues se cumplen '(x)(y) [$x = y \rightarrow y = x$]' (simetría) y '(x)(y) [$\exists z [x = z \wedge z = y] \rightarrow x = y$]' (transitividad). También lo es, por ejemplo, la relación de convecindad, entendida en un sentido jurídico según el cual nadie pueda estar empadronado en dos municipios o vecindades a la vez, y admitiendo que tiene sentido decir que una persona es convecina de sí misma, es decir, que figura en el mismo padrón que sí misma.

Podemos establecer la siguiente definición de la clase (Eq) de las relaciones de equivalencia:

$$(DR78) \quad \text{Eq} =_{\text{df}} \hat{R}[R \in \text{Sim} \wedge R \in \text{Trans}].$$

O, lo que es lo mismo,

$$(DR78a) \quad \text{Eq} =_{\text{df}} \text{Sim} \cap \text{Trans.}$$

Por tanto:

$$(DR78b) \quad 'R \in \text{Eq}' \leftrightarrow_{\text{df}} 'R \in \text{Sim} \cap \text{Trans}'.$$

Ciertas clases que constituyen el campo de una relación de equivalencia, R , se llaman 'clases de equivalencia respecto de R '. Son las clases de los individuos x y los individuos y para los cuales vale ' Rxy '. Esas clases no

tienen ningún miembro común. Lo que sí puede ocurrir es que se trate de una sola clase dentro de un determinado universo del discurso. Por ejemplo, tratándose de la relación ser-de-la-misma-especie, si el universo del discurso está reducido a un campo de objetos Ω = los seres humanos, entonces dicha relación no tiene (en ese universo del discurso) más que una clase de equivalencia. Pero si una relación de equivalencia tiene varias clases de equivalencia, como ocurre con la relación de convecindad, entonces ninguna de ellas tiene un miembro común con otra: una clase de convecinos no puede tener ningún miembro en común con otra clase de convecinos, porque si tuviera alguno, entonces la transitividad y la simetría de la relación harían que las dos clases con miembro común fueran una sola y la misma clase: todos los miembros de la primera clase son convecinos del miembro que es común a la primera y a la segunda, y como este último es convecino de todos los miembros de la segunda clase, todos los de la primera lo serían de todos los de la segunda, y no habría más que una clase en vez de las dos supuestas. Por tanto, si el campo de una relación de equivalencia, R , contiene varias clases de equivalencia realmente distintas, ninguna de éstas puede tener ningún miembro común con otra de ellas. En esta discusión hemos usado sin explicitar las dos propiedades que caracterizan a las clases de equivalencia:

- 1.^a dada una relación de equivalencia, R , esa relación media entre cada par de miembros de cada una de las clases de equivalencia respecto de R . (R es conexas en cada una de sus clases de equivalencia);
- 2.^a dada una clase de equivalencia respecto de la relación R , por ejemplo, la clase α , todo objeto, y , con el cual algún miembro, x , de α esté en la relación R , pertenece también a la clase α .

Con nuestra relación de convecindad ocurre (1.^o) que para dos individuos cualesquiera de una clase de convecinos media entre ellos la relación de convecindad, y (2.^o) que si x es convecino de y , entonces y pertenece a la misma clase de equivalencia respecto de la convecindad (al mismo municipio) que x .

Las condiciones 1.^a y 2.^a pueden formularse del modo siguiente:

$$1.^a: (x)(y) [x \in \alpha \wedge y \in \alpha \rightarrow Rxy].$$

(Como R es una relación de equivalencia y, por lo tanto, simétrica, no es necesario escribir ' $Rxy \vee Ryx$ ' como consecuente del condicional para expresar la conexidad de R .)

$$2.^a: (x)(y) [x \in \alpha \wedge Rxy \rightarrow y \in \alpha].$$

Por tanto, la clase de las clases de equivalencia respecto de una relación R ('Eq— R ') —que corresponde a la propiedad de ser una clase de equivalencia respecto de R — podría definirse mediante la conjunción de esas dos condiciones, o sea, como sigue:

$$(DR79) \quad \text{Eq—}R =_{\text{df}} \hat{a} [(x)(y) [[x \in \alpha \wedge y \in \alpha \rightarrow Rxy] \wedge [x \in \alpha \wedge Rxy \rightarrow y \in \alpha]]].$$

Esta definición suele utilizarse en la siguiente forma "comprimida":

$$(DR79a) \quad \text{Eq—}R =_{\text{df}} \hat{a} [(x)(y) [x \in \alpha \rightarrow [y \in \alpha \leftrightarrow Rxy]]].$$

Para expresar que la clase α es una clase de equivalencia respecto de R escribiremos ' $\alpha \in \text{Eq—}R$ ', de acuerdo con la definición siguiente:

$$(DR80) \quad '\alpha \in \text{Eq—}R' \leftrightarrow_{\text{df}} '(x)(y) [x \in \alpha \rightarrow [y \in \alpha \leftrightarrow Rxy]]'.$$

La equivalencia entre la definición "larga" y la "condensada" de las clases de equivalencia puede establecerse demostrando en el cálculo de predicados la equivalencia de las dos fórmulas predicativas que traducen los dos definientia, o sea, demostrando

$$(81) \quad (x)(y) [[Px \wedge Py \rightarrow Rxy] \wedge [Px \wedge Rxy \rightarrow Py]] \leftrightarrow (x)(y) [Px \rightarrow [Py \leftrightarrow Rxy]].$$

Una demostración en el cálculo de la deducción natural, con todos los pasos explícitos, no debe requerir mucho más de 40 líneas.

La noción de relaciones y clases de equivalencia tiene importancia teórica para la definición de ciertas clases que desempeñan un papel fundamental en algunas ciencias. Haremos uso de ella con este fin. Por ahora, la posibilidad general de definir clases por relaciones de equivalencia puede quedar ejemplificada por nuestra relación de convecindad: con ella podríamos definir los municipios como las clases de equivalencia respecto de la relación de convecindad.

89. Univocidad. Isomorfía. Estructuras. — Bajo este epígrafe y bajo el siguiente vamos a precisar en alguna medida varios conceptos. Cuatro de ellos han sido usados constantemente en lo que precede, porque son conceptos muy básicos para la lógica misma: son los de isomorfía, estructura, orden y función.

Los conceptos de isomorfía y estructura están íntimamente emparentados con el concepto de forma, la abstracción básica de la lógica formal. El concepto de orden, representado por el de serie en 90, y el concepto de función han sido también de uso frecuente en los capítulos anteriores. El hecho de que sólo ahora dispongamos de medios para darles cierta exactitud formal permite una nueva observación filosófica, de teoría de la ciencia, que se sumará a las varias hechas en lo que precede de este manual. Se trata de lo siguiente: la marcha real de la adquisición de conocimientos es un constante vaivén entre conceptos a niveles más intuitivos y conceptos a niveles más definidos o tecnificados. La precisión técnica de un concepto suele suponer en el trabajo científico real el uso de resultados que se han conseguido realmente (psicológica e históricamente) mediante

aplicaciones menos claras, más intuitivas, de conceptos emparentados con el que finalmente se aclara o tecnifica. La presentación teórica de esos conocimientos — la que es propia de un tratado más que de un manual — esconde esa marcha real, y empieza con las nociones técnicas. Esa presentación teórica dispone ya de la fundamentación, en vez de buscarla, como es el caso en la marcha real de la ciencia.

Una relación n -ádica, R_n , se llama 'unívoca en el lugar m ', con $m \leq n$, o, simplemente, ' m -unívoca' (' m -Un'), si y sólo si para cada conjunto ordenado de $n-1$ argumentos de R_n hay un argumento y sólo uno que pueda ocupar el lugar de argumento m cuando la expresión relacional es verdadera. Esta noción puede sugerirse mediante el siguiente esquema de definiciones, del que, sustituyendo ' n ' y ' m ' por '1', '2', '3', ..., se obtienen definiciones concretas:

$$(DR82) \quad m\text{-Un}_n =_{df} \hat{R}_n [(x_1)(x_2) \dots (x_{m-1})(x_{m+1}) \dots (x_n)(x)(y) \\ [R_n x_1 x_2 \dots x_{m-1} x x_{m+1} \dots x_n \wedge R_n x_1 x_2 \dots x_{m-1} y x_{m+1} \dots x_n \rightarrow x = y]].$$

Tratándose de relaciones diádicas, m no puede ser más que 1 o 2. Llamaremos 'relaciones preunívocas' ('Preun') a las relaciones diádicas 1-unívocas, y 'relaciones postunívocas' ('Postun') a las relaciones diádicas 2-unívocas:

$$(DR83) \quad \text{Preun} =_{df} \hat{R}_2 [(x)(y)(z) [R_2 xy \wedge R_2 zy \rightarrow x = z]].$$

$$(DR84) \quad \text{Postun} =_{df} \hat{R}_2 [(x)(y)(z) [R_2 xy \wedge R_2 xz \rightarrow y = z]].$$

Una relación preunívoca es la relación padre-de. La relación natural-de es postunívoca.

Una relación diádica se llama 'biunívoca' ('Biun') cuando es a la vez preunívoca y postunívoca. Definición de la clase de relaciones Biun:

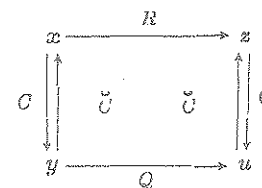
$$(DR85) \quad \text{Biun} =_{df} \text{Preun} \cap \text{Postun}.$$

La relación mitad-de es biunívoca (entre números naturales, por ejemplo).

Una relación *correlatora*, o, simplemente, una correlatora, es una relación diádica biunívoca que correlaciona los campos de dos relaciones n -ádicas ($n = 2, 3, \dots$) del modo siguiente:

1.º: cuando la correlatora, C , media entre dos argumentos, x e y , es que x pertenece al campo de una de las dos relaciones n -ádicas correlatadas, R por ejemplo, e y pertenece al campo de la otra, Q , por ejemplo. Además:

2.º: si valen ' Rxx ', ' Cxy ' y ' Czu ', entonces vale ' Qyu ', y recíprocamente. Esta explicación, que se restringe a relaciones correlatadas diádicas, puede representarse por la siguiente figura:



Las flechas de la figura representan las "direcciones" en que ligan las relaciones, o sea, el hecho de que los símbolos individuales situados por el lado de la punta de las flechas son segundos argumentos. x, y, z, u pertenecen los cuatro al campo de la correlatora C . x y z pertenecen además al campo de R ; y y u pertenecen al de Q .

La definición de la clase de las relaciones correlatorias (Corr) depende del número de argumentos de las relaciones correlatadas. El siguiente es un esquema de definiciones del que pueden obtenerse definiciones propiamente dichas sustituyendo ' n ' por '2', '3' ... Pero una correlatora es siempre una relación diádica, cualquiera que sea el número de argumentos de las relaciones correlatadas.

$$(DR86) \quad n\text{-Corr} =_{df} \hat{R}_2 [R_2 \in \text{Biun} \wedge \exists P_n \exists Q_n [(x) [x \in CP_n \rightarrow x \in D_1 R_2] \wedge \\ \wedge (x) [x \in CQ_n \rightarrow x \in D_2 R_2] \wedge \\ \wedge (x_1)(x_2) \dots (x_n)(y_1)(y_2) \dots (y_n) [R_2 x_1 y_1 \wedge \dots \wedge R_2 x_n y_n \rightarrow \\ \rightarrow [P x_1 \dots x_n \leftrightarrow Q y_1 \dots y_n]]]].$$

La definición de las correlatorias entre relaciones triádicas, por ejemplo, será (dejando de indicar que las correlatorias son diádicas y las correlatadas triádicas):

$$(DR86a) \quad 3\text{-Corr} =_{df} \hat{R} [R \in \text{Biun} \wedge \exists P \exists Q [(x) [x \in CP \rightarrow x \in D_1 R] \wedge \\ \wedge (x) [x \in CQ \rightarrow x \in D_2 R] \wedge \\ \wedge (x_1)(x_2)(x_3)(y_1)(y_2)(y_3) [R x_1 y_1 \wedge R x_2 y_2 \wedge R x_3 y_3 \rightarrow \\ \rightarrow [P x_1 x_2 x_3 \leftrightarrow Q y_1 y_2 y_3]]]].$$

Frecuentemente interesa referirse a una relación, R , correlatora entre dos relaciones dadas, P_n y Q_n . La definición de R , o de la expresión ' R es n -correlatora entre P_n y Q_n ', simbolizando por ' $n\text{-Corr-}P_n\text{-}Q_n$ ' la clase correspondiente, es:

$$(DR86b) \quad 'R \in n\text{-Corr-}P_n\text{-}Q_n' \leftrightarrow_{df} 'R \in \text{Biun} \wedge (x) [x \in CP_n \rightarrow x \in D_1 R] \wedge \\ \wedge (x) [x \in CQ_n \rightarrow x \in D_2 R] \wedge \\ \wedge (x_1) \dots (x_n)(y_1) \dots (y_n) [R x_1 y_1 \wedge \dots \wedge R x_n y_n \rightarrow \\ \rightarrow [P_n x_1 \dots x_n \leftrightarrow Q_n y_1 \dots y_n]]'.$$

Ejemplos de 2-correlatorias:

Suponiendo una elección por votación entre dos únicos candidatos, con

un número impar de electores, sin votos nulos ni abstenciones, la relación votar-*a* es correlatora entre la relación mayor-que, cuyo campo son los números de votantes en favor de cada uno de los candidatos, y la relación vencer-*a*, cuyo campo está constituido por los dos candidatos.

Suponiendo que todos los jugadores de un equipo de fútbol tienen una madrina, y sólo una cada uno, distinta de todas las de los demás jugadores, y que las madrinas forman una asociación *A*, entonces la relación biunívoca madrina-de es correlatora entre las relaciones diádicas coequipier-de y consocia-de- en-*A*.

El concepto de correlatora permite definir una noción frecuentemente usada de un modo intuitivo; la de *isomorfía* entre relaciones. La noción se introduce a continuación de acuerdo con la construcción de la misma por R. Carnap.

Se dice que dos relaciones son isomorfas cuando existe una correlatora de ambas. Igual que la de correlatora, la definición de isomorfía dependerá del número de argumentos de las relaciones consideradas. Y, también como en aquel caso, nos será de utilidad un esquema de definiciones para sugerir la noción de la clase de las relaciones *n*-ádicas isomorfas (*n*-Is):

$$(DR87) \quad n\text{-Is} \equiv_{\text{def}} \hat{P}_n \hat{Q}_n [\exists R [R \in n\text{-Corr-}P_n\text{-}Q_n]].$$

La expresión 'las relaciones triádicas *P* y *Q* son isomorfas', por ejemplo, se definirá:

$$(DR88) \quad '3\text{-Is } PQ' \leftrightarrow_{\text{def}} \exists R [R \in 3\text{-Corr-}P\text{-}Q].$$

La isomorfía, tal como ha quedado definida, es una relación diádica entre relaciones. En el definiendun de (DR88), *P* y *Q* son los argumentos del símbolo predicativo diádico (del predicado diádico) '3-Is'.

La relación de isomorfía es reflexiva, pues toda relación es isomorfa consigo misma: la correlatora puede ser la relación de identidad. La isomorfía es también simétrica, pues si hay una correlatora, *R*, entre *P* y *Q*, de modo que pueda afirmarse '*P* es isomorfa con *Q*', entonces hay también una correlatora, *R*, entre *Q* y *P*, de modo que puede afirmarse '*Q* es isomorfa con *P*' (cfr. la anterior figura). Por último, la isomorfía es transitiva: pues si *P* y *Q* son isomorfas y *Q* y *T* son isomorfas, o sea, si hay entre *P* y *Q* al menos una correlatora, *R*, y entre *Q* y *T* al menos una correlatora, *S*, entonces hay al menos una correlatora entre *P* y *T*, a saber, *R* | *S*. Luego *P* y *T* son isomorfas.

La demostración de esta última propiedad de la isomorfía en el cálculo de predicados es de principios muy sencillos, pero puede ser un buen ejemplo de "trucos" o expedientes permitidos por la regla RP del cálculo de la deducción natural y que son muy frecuentemente necesarios en la demostración de fórmulas poliádicas. Hay que demostrar, a partir de dos premi-

sas, que *R* | *S* es correlatora entre *P* y *T*. Por la definición de correlatora eso quiere decir, tomando, por ejemplo, el caso diádico, que hay que demostrar:

$$(89) \quad (x)(y)(z)(u) [R | S \ xy \wedge R | S \ zu \rightarrow [Pxx \leftrightarrow Tyu]].$$

Por la definición de '*R* | *S*' (DR16), (89) equivale a

$$(x)(y)(z)(u) [\exists v [R xv \wedge S vy] \wedge \exists w [R zw \wedge S wu] \rightarrow [Pxx \leftrightarrow Tyu]],$$

o sea, con los cuantificadores en prefijo:

$$(90) \quad (x)(y)(z)(u) \exists v \exists w [R xv \wedge S vy \wedge R zw \wedge S wu \rightarrow [Pxx \leftrightarrow Tyu]].$$

(90) tiene que demostrarse a partir de las dos premisas (X) e (Y):

$$(X): \quad (x)(y)(z)(u) [Rxy \wedge Rzu \rightarrow [Pxx \leftrightarrow Qyu]],$$

$$(Y): \quad (x)(y)(z)(u) [Sxy \wedge Szu \rightarrow [Qxz \leftrightarrow Tyu]],$$

que expresan, respectivamente, que *R* es correlatora entre *P* y *Q*, y que *S* es correlatora entre *Q* y *T*.

La demostración es como sigue:

L1(X)	(x)(y)(z)(u) [Rxy ∧ Rzu → [Pxx ↔ Qyu]]	RP
L2(Y)	(x)(y)(z)(u) [Sxy ∧ Szu → [Qxz ↔ Tyu]]	RP

(Ahora damos, para abreviar, en una sola línea, L3, cuatro eliminaciones de '()' L1, y en otra sola línea, L4, cuatro eliminaciones de '()' en L2.)

L3(X)	Rav ∧ Rcw → [Pac ↔ Qvw]	4RE(); L1
L4(Y)	Sob ∧ Swd → [Qvw ↔ Tbd]	4RE(); L2

(A continuación tomamos una nueva premisa que luego eliminaremos:)

L5(5)	Rav ∧ Sob ∧ Rcw ∧ Swd	RP
L6(5)	Rav	RE A; L5
L7(5)	Sob	RE A; L5
L8(5)	Rcw	RE A; L5
L9(5)	Swd	RE A; L5
L10(5)	Rav ∧ Rcw	RI A; L6, L8
L11(5)	Sob ∧ Swd	RI A; L7, L9
L12(X,5)	Pac ↔ Qvw	RE →; L3, L10
L13(Y,5)	Qvw ↔ Tbd	RE →; L4, L11

(Ahora utilizaremos, para abreviar, el teorema (TP16) de la lógica de predicados — la transitividad de '↔' — como regla auxiliar:)

L14(X,Y,5)	$Pac \leftrightarrow Tbd$	(TP16); L12, L13
L15(X,Y)	$Rav \wedge Sob \wedge Rcw \wedge Swd \rightarrow [Pac \leftrightarrow Tbd]$	RI $\rightarrow$; L14
L16(X,Y)	$\exists w [Rav \wedge Sob \wedge Rcw \wedge Swd \rightarrow [Pac \leftrightarrow Tbd]]$	RI $\exists$; L15
L17(X,Y)	$\exists v \exists w [Rav \wedge Sob \wedge Rcw \wedge Swd \rightarrow [Pac \leftrightarrow Tbd]]$	RI $\exists$; L16
d[a,b,c] L18(X,Y)	$(u) \exists v \exists w [Rav \wedge Sob \wedge Rcw \wedge Swu \rightarrow [Pac \leftrightarrow Tbu]]$	RI (); L17
c[a,b] L19(X,Y)	$(z)(u) \exists v \exists w [Rav \wedge Sob \wedge Rzw \wedge Swu \rightarrow [Paz \leftrightarrow Tbu]]$	RI (); L18
b[a] L20(X,Y)	$(y)(z)(u) \exists v \exists w [Rav \wedge Sov \wedge Rzw \wedge Swu \rightarrow [Paz \leftrightarrow Tyu]]$	RI (); L19
a L21(X,Y)	$(x)(y)(z)(u) \exists v \exists w [Rav \wedge Sov \wedge Rzw \wedge Swu \rightarrow [Pxz \leftrightarrow Tyu]]$	RI (); L20

La fórmula de la línea L21, que queda demostrada bajo las premisas (X) e (Y), es precisamente (90), es decir, la afirmación de que $R \mid S$ es correlatora entre P y T . Sus premisas son que R sea correlatora entre P y Q (premisa (X)) y que S sea correlatora entre Q y T (premisa (Y)).

La isomorfía es, pues, una relación simétrica y transitiva. Por tanto, es una relación de equivalencia. Entonces pueden existir clases de equivalencia respecto de la isomorfía o, más propiamente, respecto de cada relación de isomorfía, n -Is, con $n = 1, 2, \dots$

El campo de una relación de isomorfía está constituido por relaciones: no los individuos, sino las relaciones pueden ser o no ser isomorfas. Por tanto, una clase de equivalencia respecto de una relación de isomorfía es una clase de relaciones. Entre dos cualesquiera de esas relaciones (de una clase de equivalencia respecto de la isomorfía) media la relación de isomorfía. Dentro de una clase de equivalencia, todas las relaciones son isomorfas unas de otras.

Cuando dos relaciones son isomorfas, se dice que tienen la misma *estructura*. Una estructura puede pues concebirse como una clase de equivalencia respecto de una relación de isomorfía o, en términos de predicados, como la propiedad común a relaciones isomorfas por el hecho de ser isomorfas. Los siguientes esquemas definitorios sugieren un modo de fijar estas ideas (en lenguaje de clases). (DR91) define esquemáticamente la clase de las clases de equivalencia respecto de una relación de isomorfía, n -Is (o sea, la clase de las estructuras para n):

$$(DR91) \quad \text{Eq-}n\text{-Is} =_{\text{df}} \hat{a} [(R_n)(Q_n) [R_n \varepsilon \alpha \rightarrow [Q_n \varepsilon \alpha \leftrightarrow n\text{-Is } R_n Q_n]]].$$

(Compárese con la definición de la clase de las clases de equivalencia respecto de una relación R (DR79a)).

Que α es una clase de equivalencia respecto de n -Is (o, lo que es lo

mismo: una clase de relaciones n -ádicas isomorfas, una n -estructura) se expresará según el esquema de definiciones

$$(DR92) \quad '\alpha \varepsilon \text{Eq-}n\text{-Is}' \leftrightarrow_{\text{df}} '(R_n)(Q_n) [R_n \varepsilon \alpha \rightarrow [Q_n \varepsilon \alpha \leftrightarrow n\text{-Is } R_n Q_n]]'.$$

Es claro que la definición de la noción de estructura (que simbolizaremos por 'Str', el símbolo usado por Carnap) depende, como las de correlatora e isomorfía, del número de argumentos de la relación considerada. La clase de las estructuras diádicas (estructuras de las relaciones diádicas ('2-Str')), por ejemplo, se define:

$$(DR93) \quad 2\text{-Str} =_{\text{df}} \hat{a} [(R_2)(Q_2) [R_2 \varepsilon \alpha \rightarrow [Q_2 \varepsilon \alpha \leftrightarrow 2\text{-Is } R_2 Q_2]]].$$

(DR94) define ' α es una estructura diádica':

$$(DR94) \quad '\alpha \varepsilon 2\text{-Str}' \leftrightarrow_{\text{df}} '(R_2)(Q_2) [R_2 \varepsilon \alpha \rightarrow [Q_2 \varepsilon \alpha \leftrightarrow 2\text{-Is } R_2 Q_2]]'.$$

Sugerimos definiciones concretas de la clase de las estructuras (de la noción de estructura) mediante los dos siguientes esquemas de definiciones:

$$(DR95) \quad n\text{-Str} =_{\text{df}} \text{Eq-}n\text{-Is}.$$

$$(DR96) \quad '\alpha \varepsilon n\text{-Str}' \leftrightarrow_{\text{df}} '\alpha \varepsilon \text{Eq-}n\text{-Is}'.$$

Por (DR91) y (DR92), esos dos esquemas equivalen a los dos siguientes:

$$(DR95a) \quad n\text{-Str} =_{\text{df}} \hat{a} [(R_n)(Q_n) [R_n \varepsilon \alpha \rightarrow [Q_n \varepsilon \alpha \leftrightarrow n\text{-Is } R_n Q_n]]].$$

$$(DR96a) \quad '\alpha \varepsilon n\text{-Str}' \leftrightarrow_{\text{df}} '(R_n)(Q_n) [R_n \varepsilon \alpha \rightarrow [Q_n \varepsilon \alpha \leftrightarrow n\text{-Is } R_n Q_n]]',$$

respectivamente.

90. Algunos conceptos matemáticos fundamentales: número cardinal; serie; función. — *Número cardinal.* En la lógica de clases vimos posibles definiciones de los diversos números cardinales. Una definición (la de 1) era, por ejemplo:

$$(DC52) \quad 1 =_{\text{df}} \hat{a} [\exists x [\alpha = \{x\}]].$$

La lógica de relaciones permite una sencilla definición explícita del concepto de número cardinal en general, o de la clase de los números cardinales. Para llegar a esa definición empezaremos por observar lo que ocurre cuando la noción de correlatora, que tenemos definida para fórmulas predicativas n -ádicas con $n \geq 2$ (o sea, para fórmulas relacionales), se generaliza de modo que pueda aplicarse también a fórmulas monádicas (fórmulas de las cuales pueden abstraerse clases). La definición tendrá el siguiente aspecto:

$$'R \varepsilon 1\text{-Corr-P-Q}' \leftrightarrow_{\text{df}} 'R \varepsilon \text{Biun} \wedge (x) [x \varepsilon CP \rightarrow x \varepsilon D_1 R] \wedge (x) [x \varepsilon CQ \rightarrow x \varepsilon D_2 R] \wedge (x)(y) [Rxy \rightarrow [Px \leftrightarrow Qy]]'.$$

Puesto que 'P' y 'Q' son predicados monádicos, representan clases de individuos. Llamémosles, respectivamente, ' α ' y ' β '; con estos símbolos escribimos de nuevo la anterior definición:

$$(DR97) \quad 'R \in 1\text{-Corr-}\alpha\text{-}\beta' \leftrightarrow_{\text{df}} 'R \in \text{Biun} \wedge (x) [x \in \alpha \rightarrow x \in D_1 R] \wedge \\ \wedge (x) [x \in \beta \rightarrow x \in D_2 R] \wedge (x)(y) [Rxy \rightarrow [x \in \alpha \leftrightarrow y \in \beta]]'.$$

A la vista de (DR97) podemos afirmar que decir que dos clases tienen una correlatora es lo mismo que decir que son coordinables. Esto permite ampliar la noción de isomorfía a las clases: por analogía con las relaciones, diremos que dos clases son isomórficas si y sólo si tienen una correlatora. O, lo que es lo mismo: dos clases son isomorfas si y sólo si son coordinables. Esta reflexión se precisa en las siguientes definiciones:

$$(DR98) \quad '1\text{-Is } \alpha \beta' \leftrightarrow_{\text{df}} '\exists R [R \in 1\text{-Corr-}\alpha\text{-}\beta]'.$$

$$(DR99) \quad '\text{Coord } \alpha \beta' \leftrightarrow_{\text{df}} '1\text{-Is } \alpha \beta'.$$

¿Qué son las clases de equivalencia respecto de 1-Is, es decir, qué son 1-estructuras, estructuras monádicas? No serán clases de relaciones, sino clases de clases: clases de clases isomorfas, o, lo que es lo mismo, clases de clases coordinables. Y clases coordinables — conjuntos coordinables — son los que tienen el mismo número cardinal (cfr. 15). Según la definición (DR80a), si α es una clase de equivalencia respecto de 1-Is, podemos escribir:

$$(DR100) \quad '\alpha \in \text{Eq-1-Is}' \leftrightarrow_{\text{df}} '(x)(y) [x \in \alpha \rightarrow [y \in \alpha \leftrightarrow 1\text{-Is } xy]]'.$$

Puesto que x y y de (DR100) tienen que ser 1-isomorfas, es que son clases; α es entonces una clase de clases, y Eq-1-Is es una clase de clases de clases (coordinables). Eq-1-Is es pues de orden lógico 3. (En la definición (DC52), x es de orden 0, $\{x\} = \alpha$ es de tipo 1, y el número 1 es de orden 2.)

Atendamos a la clase Eq-1-Is, la clase de las clases de equivalencia respecto de 1-Is:

$$(DR101) \quad \text{Eq-1-Is} =_{\text{df}} \hat{a} [(x)(y) [x \in \alpha \rightarrow [y \in \alpha \leftrightarrow 1\text{-Is } xy]]].$$

Según (DC52), el número cardinal 1 es la clase de las clases coordinables con la clase $\{x\}$. Es entonces claro que es un miembro de la clase de las clases de clases coordinables, Eq-1-Is. Lo mismo vale de cualquier número cardinal definido como lo está el número 1 en (DC52). Consiguientemente, podemos decir que la clase Eq-1-Is es la clase de los números cardinales, y definir del modo siguiente que α es un número cardinal, un miembro de la clase de los números cardinales, a la que llamaremos ' N '

$$(DR102) \quad '\alpha \in N' =_{\text{df}} '\alpha \in \text{Eq-1-Is}'.$$

En general, la clase, N , de los números cardinales se definirá:

$$(DR103) \quad N =_{\text{df}} \text{Eq-1-Is}.$$

Ahora bien: las estructuras quedaron definidas (cfr. (DR95), (DR96)) como las clases de equivalencia respecto de isomorfías. Que α es una 1-estructura (una estructura monádica) se definirá pues:

$$(DR104) \quad '\alpha \in 1\text{-Str}' \leftrightarrow_{\text{df}} '\alpha \in \text{Eq-1-Is}'.$$

Una 1-estructura es una clase de clases 1-isomorfas. De (DR104) obtenemos

$$(DR105) \quad 1\text{-Str} =_{\text{df}} \text{Eq-1-Is}.$$

Y de (DR103) y (DR105) la siguiente definición explícita de la clase de los números cardinales:

$$(DR106) \quad N =_{\text{df}} 1\text{-Str}.$$

Dada la correspondencia entre clases y propiedades monádicas, podemos considerar a (DR106) como una definición explícita de la propiedad de ser un número cardinal: los números cardinales son las estructuras monádicas.

La ampliación del concepto de isomorfía a las clases equivale a admitir que la potencia de una clase, su número de miembros, es una estructura: dos clases tienen la misma estructura (son isomorfas) cuando tienen el mismo número de miembros (el mismo número cardinal). Esto sugiere una reflexión epistemológica: la generalización de los conceptos de isomorfía y estructura da lugar a una consideración cualitativa de la cantidad. La generalización de que se trata, y cuyo desarrollo explícito, tal como aquí lo hemos estudiado, se debe a R. Carnap, no es pues sólo interesante porque suministra una definición explícita del concepto de número cardinal, sino también porque da un fundamento elemental a complicados pasos de formación de conceptos y de razonamientos sobre cantidades y cualidades que son frecuentes en las ciencias, especialmente en las sociales.

Series. — Se llama 'serie' ('Ser') a toda relación irreflexiva, transitiva y conexa:

$$(DR107) \quad \text{Ser} =_{\text{df}} R [R \in \text{Irrefl} \wedge R \in \text{Trans} \wedge R \in \text{Conex}],$$

o bien

$$(DR107a) \quad \text{Ser} =_{\text{df}} \text{Irrefl} \cap \text{Trans} \cap \text{Conex}.$$

Las series son, naturalmente, asimétricas.

Obsérvese que, al igual que las isomorfías, una serie es primariamente una relación. Si no se amplía el concepto, no puede decirse que una clase, por ejemplo la clase $\{a, b, c\}$, sea una serie, pues una clase no está por sí misma seriada (ordenada serialmente), y puede dar de sí varias "series"

(en el caso del ejemplo: $\{a, c, b\}$, $\{b, a, c\}$, etc.). Sólo en un sentido ampliado o generalizado puede llamarse 'serie' a $\{a, b, c\}$. En sentido primario, la serie aludida al hablar así es la relación inmediatamente-anterior-en-el-alfabeto, con un campo, $\Omega = \{a, b, c\} = \{c, a, b\}$, etc. Por eso podría ser útil acostumbrarse a llamar 'seriación' o 'clase seriada' a lo que corrientemente llamamos 'serie', y 'seriadora', por ejemplo, a la relación correspondiente, abandonando la palabra 'serie', ambigua por su referencia a clases y a relaciones. Russell y Whitehead usaban la expresión 'relación serial'.

Ejemplo: la relación menor-que es una serie — o "seriadora" — en el campo de los números naturales.

Funciones. — El carácter fundamental de la noción de relación — y de la lógica de relaciones, por consiguiente — puede apreciarse al observar su íntimo parentesco con el concepto de función, igual que antes hemos visto el que tiene con el de estructura.

Intuitivamente es claro que una función es un esquema relacional, puesto que es un sistema de correspondencias unívocas entre valores. Por eso una determinada aplicación de una función puede expresarse mediante un enunciado relacional. Por ejemplo, a propósito de la función seno, podemos escribir, en vez de

$$(108) \quad \text{sen } 90 = 1,$$

$$(109) \quad R \ 1 \ 90,$$

simbolizando 'R' el predicado diádico 'ser-el-seno-de'.

Análogamente, si la función es la que conocemos por el functor ' $\sim$ ', podemos escribir, en vez de la expresión funcional

$$(110) \quad \sim p \leftrightarrow q,$$

la expresión relacional

$$(111) \quad Nqp,$$

simbolizando por 'N' 'ser-la-negación-de'. 'Nqp' se leerá: 'q es la negación de p'.

Esta traducción de expresiones funcionales a expresiones relacionales sugiere la idea de una traducción inversa, es decir, de expresiones relacionales a expresiones funcionales. Tal traducción inversa está sometida a ciertas limitaciones, como veremos. Su realización equivale a abstraer de una expresión relacional una función: por ejemplo, de (111) (110), de la relación N la función $\sim$. Para indicar esta *abstracción funcional* suele usarse el operador ' λ ' (lambda), introducido en este sentido por A. Church. Si 'X'

es una fórmula, entonces ' $\lambda x [X]$ ' se lee: 'la función de x tal que X'. Con esta notación puede escribirse:

$$(112) \quad \sim =_{\text{df}} \lambda y [\exists x Nxy].$$

De (112) obtenemos la definición de la aplicación de la función a un enunciado p:

$$(113) \quad \sim p = \lambda y [\exists x Nxy] p.$$

Una lectura de (113) puede ser: 'la negación funcional de 'p' es la aplicación a 'p' de la función que consiste en ser segundo argumento de la relación diádica N'. La abstracción funcional permite pues definir funciones con ayuda de relaciones.

Como primer resultado de la anterior discusión tenemos lo siguiente: toda expresión funcional de n argumentos (y un valor) puede traducirse por una expresión relacional de n + 1 argumentos:

$$(TR114) \quad (f_n) \exists R_{n+1} [f_n(x_2, x_3, \dots, x_{n+1}) = x_1 \rightarrow R_{n+1} x_1 x_2 \dots x_n x_{n+1}].$$

La notación de (TR114) es bastante grosera, porque deja cosas implícitas, pero nos bastará aquí, porque no ofrece dificultad a la intuición. ' $f_n(x_2, x_3, \dots, x_{n+1})$ ' debe entenderse según la corriente lectura matemática.

La afirmación inversa de (TR114), a saber,

$$(R_{n+1}) \exists f_n [R_{n+1} x_1 x_2 \dots x_{n+1} \rightarrow f_n(x_2, x_3, \dots, x_{n+1}) = x_1],$$

que podría creerse segundo resultado de nuestra discusión (lo que llamábamos 'traducción inversa'), no es en cambio válida: no toda expresión relacional puede traducirse por una expresión funcional, sino sólo las expresiones de relaciones para cuyo primer lugar de argumento haya siempre un individuo (para cada clase de argumentos de los lugares 2 a n + 1) y sólo uno (de modo que se trate de relaciones 1-unívocas). Sólo con estas condiciones hay valor en cada caso y éste está unívocamente determinado para cada conjunto dado de argumentos. Pues si la relación no es unívoca en el lugar 1, la función correspondiente tendrá varios valores para los mismos argumentos, lo cual es incompatible con la noción de función. Y si para una serie de argumentos para los lugares 2 a n + 1 la relación no tiene argumento para el lugar 1, entonces la función correspondiente no estará definida para esos n argumentos. Por eso sólo podemos afirmar, en vez de la "traducción inversa", lo siguiente:

$$(TR115) \quad (R_n) [R_n \in 1\text{-Un} \wedge (x_2)(x_3) \dots (x_n) \exists x_1 R_n x_1 x_2 \dots x_n \rightarrow \rightarrow \exists f_{n-1} [R_n x_1 x_2 \dots x_n \rightarrow f_{n-1}(x_2, x_3, \dots, x_{n-1}) = x_1]].$$

Por último, (TR114) y (TR115) sugieren un modo de definir la idea de función (la clase, F, de las funciones) por la de relación: funciones serán

abstractos de relaciones de las características dichas. Cuando hay una función, hay pues también una relación de esa naturaleza, y de la cual ha sido abstraída la función. El siguiente esquema de definiciones sugiere esa idea. De él pueden obtenerse efectivas definiciones para $n = 1, 2, \dots$

$$(DR116) \quad 'f_n \in F' \leftrightarrow_{df} [\exists R_{n+1} [R_{n+1} \in 1\text{-Un} \wedge \\ \wedge (x_2) \dots (x_{n+1}) \exists x_1 [R_{n+1} x_1 x_2 \dots x_{n+1}] \wedge \\ \wedge [R_{n+1} x_1 x_2 \dots x_{n+1} \leftrightarrow f_n(x_2, \dots, x_{n+1}) = x_1]]].$$

Como resultado general de esta discusión puede decirse: todas las funciones son relaciones, pero no todas las relaciones son funciones (sino sólo aquellas que son unívocas en el lugar de argumento 1 y tienen siempre un argumento para dicho lugar).

(Ejemplo: de la relación diádica *padre-de* puede abstraerse una función de un solo argumento. Pero no puede abstraerse una función de la relación diádica *hijo-de*, que no es 1-equívoca).

PARTE CUARTA

LOGICA FORMAL Y METODOLOGIA

CAPÍTULO XVI

LA DIVISIÓN Y LA DEFINICIÓN

91. **Lógica formal y metodología.** — En los viejos tratados de lógica solían aparecer cuestiones que no pueden estudiarse de un modo completo desde el punto de vista formal, porque pertenecen a la teoría del método de descubrimiento o de exposición de verdades materiales en las ciencias. Puesto que esas cuestiones se refieren a la materia del conocimiento, no pertenecen a la teoría lógica, razón por la cual no se recogen hoy día en los manuales y tratados de lógica. Pero esos temas son, naturalmente, susceptibles de consideración formal, y también es posible construir una teoría lógica de sus aspectos formales. La consideración formal de esos temas es una interesante aplicación analítica de la lógica, digna de atención desde el punto de vista del científico positivo (cfr. 14).

'Método' es palabra usada con diversas significaciones. Es corriente hablar de método a propósito, por ejemplo, de la deducción — "método deductivo" —, y también a propósito de cadenas de operaciones industriales, como el "método de las cámaras de plomo" en la industria química. Común a todos los usos de la palabra 'método' es la referencia a una serie de operaciones ordenadas y encaminadas a obtener un resultado. Cuando el resultado que se busca es la adquisición de un conocimiento o su transmisión, se habla de "métodos teóricos" ("heurísticos" o "didácticos", respectivamente). Estos son los únicos que nos interesan aquí.

Los tres temas metodológicos que vamos a considerar desde el punto de vista lógico en este capítulo y en el siguiente son los que solían aparecer en los viejos manuales de lógica: la división, la definición y la inducción. Los tres están relacionados entre sí en el método de las ciencias: la división (o clasificación) suministra muchas veces los elementos de la definición. Esto ocurre de modo prototípico en la sistemática biológica, en la que los principios de la definición de las unidades sistemáticas (especies, por ejemplo), son los mismos principios con los que se clasifican los individuos. A su vez, la definición precisa la extensión, por ejemplo, de un nombre de fenómenos, y por eso puede ser punto de partida de una división de esos

fenómenos, y también un requisito previo a cualquier afirmación general (obtenida por inducción) acerca de esos fenómenos. A la inversa, las inducciones conseguidas afinan las definiciones de que parten, al enriquecer el conocimiento que tenemos de los fenómenos estudiados.

92. Noción metodológica de la división. — Tradicionalmente se define la división como una expresión (correspondiente a una operación) que distribuye una cosa o un nombre por sus partes. Por 'cosa' hay que entender 'concepto' (normalmente: 'clase'), es decir, algo abstracto; pues la división en sentido metodológico no es un desmembramiento como el que puede operar un carnicero con una res, sino una diferenciación entre elementos significativos, entre subclases de una clase dada, o entre usos de una expresión.

Los elementos de una división son: el todo que se divide, o clase dividida; las partes en que se divide, o clases dividentes; y el punto de vista según el cual se practica la división, al cual se llama 'principio' o 'fundamento' de la división. Por ejemplo, cuando en 1 tomábamos como propiedad característica del vegetal el ser autótrofo, lo estábamos distinguiendo del animal, como heterótrofo que éste es. El todo de la división eran los seres vivientes; las partes o subclases dividentes eran la clase de los animales y la clase de los vegetales; y el fundamento de la división era el modo de nutrición. Este ejemplo ilustra también la relación entre división y definición en ciencias en las que la descripción y la clasificación (la sistemática) de objetos tengan mucha importancia.

No nos interesa aquí recoger las clasificaciones de la división misma que aparecen en los viejos tratados, pues esa clasificación no tiene mucha relevancia lógica, sino sobre todo metodológica. Interesantes desde un punto de vista formal son, en cambio, las llamadas 'leyes de la división':

- 1.ª: el fundamento de la división debe mantenerse constante durante toda la operación;
- 2.ª: la suma lógica de las subclases dividentes debe ser igual a la clase dividida;
- 3.ª: las subclases dividentes deben excluirse mutuamente.

Expresemos formalmente esas leyes.

93. Disyuntividad y exclusión mutua. — Lo primero que podemos hacer es expresar, con ayuda de la lógica de clases, las situaciones en que se cumplen las leyes 2.ª y 3.ª. Dada la clase dividida, α , y las subclases dividentes, $\alpha_1, \alpha_2, \dots, \alpha_n$, tiene que valer, según la ley 2.ª,

$$(1a) \quad \alpha_1 \cup \alpha_2 \cup \dots \cup \alpha_n = \alpha,$$

y, según la ley 3.ª,

$$(2a) \quad (\alpha_i)(\alpha_j) [\sim i = j \rightarrow \alpha_i \cap \alpha_j = \Lambda].$$

A veces pueden ser más cómodas otras expresiones de esas leyes:

$$(1b) \quad (x) [x \in \alpha \leftrightarrow x \in \alpha_1 \vee x \in \alpha_2 \vee \dots x \in \alpha_n].$$

$$(2b) \quad (\alpha_i)(\alpha_j) [\sim i = j \rightarrow (x) [x \in \alpha \rightarrow \sim [x \in \alpha_i \wedge x \in \alpha_j]]].$$

$$(1c) \quad [\alpha \cap \alpha_1] \cup [\alpha \cap \alpha_2] \cup \dots \cup [\alpha \cap \alpha_n] = \alpha.$$

$$(2c) \quad (\alpha_i)(\alpha_j) [\sim i = j \rightarrow [\alpha \cap \alpha_i] \cap [\alpha \cap \alpha_j] = \Lambda].$$

Las notaciones (c) se basan en que de (1a) y (2a) se desprende:

(3) Si α_i es cualquier subclase dividente de α , entonces:

$$(\alpha_i) [\alpha \cap \alpha_i = \alpha_i].$$

(3) recuerda el teorema de la lógica de clases

$$(\alpha) [\vee \cap \alpha = \alpha].$$

Esto puede interpretarse diciendo que en toda división la clase dividida se toma como clase universal.

94. El fundamento de una división. — Cuando se pasa de considerar el todo dividido y las partes dividentes a considerar el principio de la división, se abandona el lenguaje de clases para adoptar el de predicados. Pues es claro que el principio de la división está compuesto por nociones que se usan para construir con ellas las subclases dividentes, examinando los individuos para ver si presentan o no alguna de aquellas propiedades. La operación de dividir está pues relacionada con el principio de abstracción de clases a partir de fórmulas predicativas monádicas (cfr. 79). Cuando una división se contempla desde el punto de vista de su fundamento o principio, es pues adecuado expresarse en el lenguaje de la lógica de predicados, o, por lo menos, utilizar también éste en combinación con el de la lógica de clases. Aprovechemos esta reflexión para tratar la cuestión siguiente: ¿cómo puede expresarse el fundamento de una división, o la ley 1.ª de la división, que se refiere al fundamento?

Por de pronto observaremos que la relación entre predicados y clases — expresada por el principio de abstracción — se da también, como es natural, para la clase dividida. En el caso del ejemplo anterior, la clase (en sentido lógico, no de sistemática biológica) de los seres vivientes corresponde a un esquema del predicado monádico 'ser-viviente'. El fundamento de la división de esa clase en animales y vegetales es, dijimos, el modo de alimentarse. Pero aunque esta manera de hablar sea intuitivamente clara, no tiene, sin embargo, la precisión suficiente para expresar la situación; pues 'tener un modo de alimentarse' no puede ser el predicado que exprese la división a que estamos refiriéndonos: tanto los animales cuanto los ve-

getales tienen algún modo de alimentarse. Este predicado 'tener un modo de alimentarse' corresponde por tanto también a la clase de los seres vivos, y con él no pueden obtenerse las clases de los animales y de los vegetales.

En el caso de nuestro ejemplo, el principio de la división debería pues expresarse mediante una fórmula que contuviera los dos predicados — 'autótrofo' y 'heterótrofo' — de que se abstraen o construyen las clases de los vegetales y de los animales, respectivamente. Si esos predicados son 'A' ('autótrofo') y 'H' ('heterótrofo'), la clase de los vegetales podría definirse con el abstractor de clases a partir del esquema 'Ax':

$$(4) \quad \alpha = \hat{x} [Ax];$$

y la clase de los animales:

$$(5) \quad \beta = \hat{x} [Hx].$$

Basándose en (4) y (5), la clase, γ , de los seres vivos podría definirse por

$$(6) \quad \gamma = \hat{x} [Ax \vee Hx].$$

Pero lo normal es que la clase dividida (γ) se haya abstraído de un predicado menos analítico.

Entonces la fórmula predicativa monádica siguiente, consecuencia de (6), podría considerarse como expresión del principio de la división considerada ('G' es el predicado del que se ha abstraído la clase γ):

$$(7) \quad (x) [Gx \rightarrow Ax \vee Hx].$$

Prescindiendo ahora del ejemplo animales-vegetales, consideremos en general cómo debe expresarse, desde el punto de vista del fundamento de la división, la situación expuesta por las reglas 2.^a y 3.^a. Supongamos una clase, α , obtenida por abstracción a partir del predicado 'P' (propiamente: a partir del esquema 'Px') y dividida en n subclases, construidas por abstracción a partir de los predicados ' P_1 ', ' P_2 ', ..., ' P_n ' respectivamente. Entonces, según la ley 2.^a debe valer

$$(1d) \quad (x) [Px \rightarrow P_1x \vee P_2x \vee \dots \vee P_nx].$$

Y según la ley 3.^a debe valer, para cualquier par de predicados dividentes, ' P_i ', ' P_j ',

$$(2d) \quad \sim i = j \rightarrow [Px \rightarrow \sim [P_ix \wedge P_jx]].$$

Ahora vamos a ver que no es necesaria una formulación de la primera ley de la división, pues su contenido está ya presente en (1d)-(2d). En efecto:

aún supuesto (1d), si uno de los predicados ' P_1 ', ..., ' P_n ' no cumple con la ley 1.^a (referirse, dicho intuitivamente, al predicado 'P' como 'autótrofo' a 'modo de alimentarse'), entonces no va a haber ninguna garantía de construcción de que se cumpla (2d). Esto puede verse concretamente con el anterior ejemplo de división de los seres vivos, el cual, por ser dicotómico, o sea, del tipo más sencillo de división, nos ofrece la situación simplificada. En una división dicotómica, los predicados ' P_1 ' y ' P_2 ' ('autótrofo' y 'heterótrofo') se comportan como ' P_1 ' y ' $\sim P_1$ ', o como ' P_2 ' y ' $\sim P_2$ ', o, en general, como ' Q ' y ' $\sim Q$ '. De aquí que (1d), en el caso dicotómico, pueda expresarse por

$$(1e) \quad (x) [Px \rightarrow Qx \vee \sim Qx],$$

y (2d) por

$$(2e) \quad (x) [Px \rightarrow \sim [Qx \wedge \sim Qx]].$$

Pero ahora podemos observar que (2e) equivale a (1e):

$$\begin{aligned} (2e): \quad & (x) [Px \rightarrow \sim [Qx \wedge \sim Qx]] \\ & \Leftrightarrow (x) [Px \rightarrow \sim Qx \vee \sim \sim Qx] \\ & \Leftrightarrow (x) [Px \rightarrow \sim Qx \vee Qx] \\ & \Leftrightarrow (x) [Px \rightarrow Qx \vee \sim Qx]: (1e). \end{aligned}$$

Así pues, en el caso de la división dicotómica, las dos leyes 2.^a y 3.^a se funden en una sola, que puede expresarse por (1e); o, si se desea conservar la notación inicial, ' P_1 ' y ' P_2 ', de los predicados dividentes, la nueva ley conjunta que resume las dos leyes 2.^a y 3.^a para la división dicotómica puede expresarse mediante el functor veritativo de heterovalencia, ' f_{11} ' (cfr. 70, 72), al que podemos simbolizar con el aspa, 'X', propuesta por Touchais:

$$(I) \quad (x) [Px \rightarrow P_1x \times P_2x].$$

Supongamos ahora que se hubiera hecho, con violación de la ley 1.^a, una división dicotómica de los seres vivos en individuos autótrofos (' P_1 ') e individuos sin función clorofílica (' P_2 '). Esta división no cumple (I), pues ' P_1 ' y ' P_2 ' no están en la relación expresada por el functor 'X', o sea, no se comportan como ' Q ' y ' $\sim Q$ ' (y esto independientemente de que, además, es un hecho empírico que hay seres autótrofos sin función clorofílica propiamente dicha: la afirmación de que esa división viola (I) vale formalmente aunque no hubiera tales individuos autótrofos sin función clorofílica. Pues tampoco en este caso la exclusión mutua de las clases respectivas tendría razones formales, es decir, por mera significación y sin atender a los datos empíricos (salvo que se definiera un concepto por el otro)). Así

pues, el incumplimiento de la ley 1.^a implica el incumplimiento de la ley conjunta (I). Por tanto — según una ley de la contraposición (teorema (TE5)) — el cumplimiento de la ley conjunta (I) implica el cumplimiento de la ley 1.^a. Así pues, para la división dicotómica es superfluo sentar tres leyes, dos de ellas referentes a clases y una al fundamento.

Ahora bien: esta reflexión puede extenderse a cualquier división, aunque no sea dicotómica, a causa de que en una correcta división n -tómica un predicado dividente cualquier, ' P_i ', está en la relación dicotómica (I) respecto de la disyunción de los demás $n-1$ predicados dividendes. Es decir, que vale

$$(Ia) \quad (i)(x) [Px \rightarrow [P_i x \times [P_1 x \vee P_2 x \vee \dots \vee P_{i-1} x \vee P_{i+1} x \vee \dots \vee P_n x]]].$$

En toda correcta división n -tómica de una clase, α , vale, en efecto, que si un individuo es clasificado en la subclase β abstraída del predicado dividente ' P_i ' entonces no puede serlo también en la subclase de α complemento de β respecto de α , $\alpha - \beta$, abstraída de ' $P_1 \vee P_2 \vee \dots \vee P_{i-1} \vee P_{i+1} \vee \dots \vee P_n$ '. (Cfr. notación molecular de los predicados, 75.)

Ahora bien: toda división n -tómica puede obtenerse mediante una serie de $n-1$ divisiones dicotómicas al mismo nivel. Esa reducción puede proceder así: primero se divide la clase total, α , por

$$P_1 \vee [P_2 \vee P_3 \vee \dots \vee P_n]$$

Llamemos β a la clase abstraída de ' P_1 ' y γ al resto, que es la clase abstraída de ' $P_2 \vee P_3 \vee \dots \vee P_n$ '. γ se divide entonces por el fundamento

$$P_2 \vee [P_3 \vee \dots \vee P_n].$$

Y así sucesivamente. Cada vez valen:

$$P_i \vee [P_{i+1} \vee \dots \vee P_n]$$

y

$$\sim [P_i \wedge [P_{i+1} \vee \dots \vee P_n]]$$

si la división es correcta. O sea, cada vez vale, si la división es correcta:

$$P_i \times [P_{i+1} \vee \dots \vee P_n]$$

Por tanto, la situación lógica en una división n -tómica puede reducirse a la que es característica de la división dicotómica.

La anterior discusión debe completarse con un hecho que más adelante (95) valoraremos: lo más importante (desde el punto de vista formal) en una división es la exclusión mutua de los predicados dividendes. El agotamiento

de la clase dividida por las clases dividendes que está, desde luego, contenido en (Ia), se toma, en cambio, como un dato, y no es relevante para la discusión formal. Así, por ejemplo, la incorrecta división de los seres vivientes en individuos autótrofos e individuos sin función clorofílica (en sentido estricto) puede perfectamente ser exhaustiva porque todo ser vivo pueda incluirse en ella (ya que no hay heterótrofos con función clorofílica). Pero eso es una cuestión empírica, sin relevancia formal (como se ve por la anterior justificación entre paréntesis).

Mientras que la relación de exclusión mutua es formalmente definible y comprobable, vale por la pura significación de los predicados heterovalentes, sin necesidad de apelar a los datos empíricos. Y la ley correspondiente es precisamente la violada por la división incorrecta que acabamos de comentar.

Podemos pues concluir que la tradicional enumeración de tres leyes de la división es superabundante. A causa de un insuficiente análisis (por no distinguir entre lenguaje de clases y lenguaje de predicados), los lógicos tradicionales han dicho en la ley 1.^a, con lenguaje predicativo, lo mismo que decían en las leyes 2.^a y 3.^a con lenguaje de clases. Además, la ley 2.^a, más que una ley, es una descripción del punto de vista lógico ante cualquier división que se le somete a examen: admitir que el científico positivo autor de la división ha agotado con ella el universo del discurso.

Observemos, por último, que cuando una división se subdivide, la operación puede expresarse, desde el punto de vista del fundamento, por una fórmula análoga a (Ia), pero en la cual, en vez de encontrar predicados atómicos en disyunción, encontramos conjunciones de predicados atómicos en disyunción. Por ejemplo, supongamos que después de la división expresada por el fundamento (Ia), la clase abstraída de cada ' P_i ' ha sido subdividida en subclases (no necesariamente el mismo número de subclases para cada ' P_i '), abstraídas de los predicados ' P_{i_1} ', ' P_{i_2} ', ..., ' P_{i_m} '. Entonces el principio complejo de dichas división y subdivisión puede expresarse por

$$(Ib) \quad (i)(s)(x) [Px \rightarrow [P_i \wedge P_{i_s}] x \times \\ \times [[P_i \wedge P_{i_1}] \vee \dots \vee [P_i \wedge P_{i_{s-1}}] \vee [P_i \wedge P_{i_{s+1}}] \vee \dots \vee \\ \vee [P_i \wedge P_{i_m}] \vee \dots \vee [P_n \wedge P_{n_1}] \vee \dots \vee [P_n \wedge P_{n_l}]] x].$$

(Ib) puede desmembrarse para tener una formulación más detallada, que recoge la distinción entre dos leyes de la división (2.^a y 3.^a):

$$(Ib1) \quad (x) [Px \rightarrow [[P_1 \wedge P_{1_1}] \vee [P_1 \wedge P_{1_2}] \vee \dots \vee [P_1 \wedge P_{1_m}] \vee \\ \vee \dots \vee [P_n \wedge P_{n_1}] \vee \dots \vee [P_n \wedge P_{n_l}]] x].$$

$$(Ib2) \quad \sim i = j \vee \sim s = t \rightarrow [Px \rightarrow \sim [P_{i_s} \wedge P_{j_t}] x]$$

95. Un ejemplo de división.— Estudiando ahora un caso concreto, aunque muy sencillo, de división efectivamente utilizada en la ciencia, podremos apreciar a la vez dos cosas: que el análisis lógico de la operación refleja bien su estructura; y que es, en cambio, insuficiente para tomar las decisiones básicas en la práctica del procedimiento.

El orden *Eubacteriales* comprende la mayor parte de las bacterias, las bacterias unicelulares no ramificadas. El predicado 'unicelular-no-ramificada' es el que permite la abstracción de ese orden. Los biólogos dividen el orden *Eubacteriales* según un fundamento o principio complejo, basado a la vez en la forma y el modo de división de los individuos según dimensiones espaciales. A pesar de haber escogido este ejemplo por su sencillez, puede apreciarse que ya en él no va a haber simples predicados atómicos ' P_n ' ni en el primer estadio de la división de ese orden, sino predicados moleculares (en conjunción). Los predicados atómicos con conjunciones de los cuales se componen los predicados de que se abstraen las subclases (familias) de *Eubacteriales* son, en la división más breve de este orden:

P_1 : ser esférico.

P_2 : ser cilíndrico no encorvado ni arrollado.

P_3 : ser cilíndrico arrollado o encorvado.

Q_1 : dividirse según una sola dimensión espacial.

Q_2 : estar indeterminado en cuanto a dimensiones espaciales de la división.

Con esos predicados atómicos (aquí algo simplificados) los biólogos han compuesto los siguientes predicados moleculares para la abstracción de las subclases (familias) del orden *Eubacteriales*:

$P_1 \wedge Q_2$, del cual se abstrae la familia *Cocáceas* (a la que pertenecen, p. e., los cocos).

$P_2 \wedge Q_1$, del cual se abstrae la familia *Bacteriáceas* (a la que pertenecen las bacterias y los bacilos).

$P_3 \wedge Q_1$, del cual se abstrae la familia de las *Espiriláceas* (a la que pertenecen los vibriones y los espirilos).

Con nuestra notación (desmembrada), la división de *Eubacteriales*, si ' P ' es el predicado 'eubacterial', debe expresarse, desde el punto de vista del fundamento o principio, como sigue:

(Ib1a) $(x) [Px \rightarrow [(P_1 \wedge Q_2) \vee (P_2 \wedge Q_1) \vee (P_3 \wedge Q_1)] x]$,

que expresa el agotamiento del orden, y

(Ib2a) $(x) [Px \rightarrow [\sim (P_1 \wedge Q_2 \wedge P_2 \wedge Q_1) \wedge \sim (P_1 \wedge Q_2 \wedge P_3 \wedge Q_1) \wedge \sim (P_2 \wedge Q_1 \wedge P_3 \wedge Q_1)] x]$,

que expresa la exclusión recíproca de los predicados dividentes de los que se abstraen las familias.

Este ejemplo nos permite ilustrar nuestra discusión de 94. (Ib2a) es formalmente (semánticamente) verdadero por la definición de los predicados, ya que en cada conjunción negada en (Ib2a) hay dos predicados que se comportan como ' Q ' y ' $\sim Q$ ', a saber, los de misma letra y distinto subíndice. En cambio, (Ib1a) no es formalmente válido. No hay razón formal por la cual deba excluirse, por ejemplo, un predicado ' $[P_1 \wedge Q_2]$ ', que fuera 'forma triangular y división según dos dimensiones espaciales'. El que no exista una clase de bacterias abstraible de ese predicado es una cuestión empírica, no formal. Por esta razón, en vez de escribir ' $\leftrightarrow$ ' en las fórmulas que recogen la ley 2.ª de la división, hemos venido escribiendo ' $\rightarrow$ ', lo cual equivale a afirmar, en la primera versión en lenguaje de clases de esta ley, ' $\subset$ ' en vez de ' $=$ ', que fue lo escrito al traducir la formulación tradicional de la misma. Esto significa que lo formalmente exigible al examinar una división no es siempre la identidad entre la suma lógica de las clases dividentes y la clase dividida, sino la inclusión de la primera en la segunda. Cuando se exige la identidad, ello debe hacerse sin pretensión empírica, sino entendiendo entonces la situación como basada en una definición de la clase dividida por la suma de las clases dividentes, y reservando al juicio del científico positivo la decisión acerca de si ese modo de considerar la situación es correcto.

96. Noción metodológica de la definición.— Definir significa determinar o delimitar. La definición es la operación metódica que consiste en caracterizar suficientemente una noción para delimitarla y separarla de otras. A esta función de la definición en el método corresponde la noción tradicional de la misma, inspirada en Aristóteles: la definición es una expresión (*terminus complexus*, como la división) que expresa la naturaleza de una cosa o la significación de un término.

La clasificación tradicional de las definiciones comprende ante todo dos grandes clases, ya sugeridas por la definición misma de definición: definición real, que expone la naturaleza de una cosa, y definición nominal, que expone la significación de un término. Esta clasificación requiere ciertas observaciones necesarias para poder considerar formalmente el tema.

Por de pronto, la definición, incluso desde el punto de vista metodológico y no ya sólo desde el punto de vista formal, es una operación en el lenguaje, en un lenguaje. Otras operaciones de los métodos científicos — como pesar u obtener preparaciones microscópicas — no están totalmente contenidas en un lenguaje. Pero el definir, por su carácter más puramente teórico, tiene del todo lugar en un lenguaje. La distinción entre definiciones nominales y definiciones reales no debe, pues, entenderse en el sentido de que las primeras sean operaciones lingüísticas y las segundas no: unas y otras son expresiones, "términos complejos".

En segundo lugar, la definición real no debe entenderse como una determinación o delimitación de la "cosa" definida: la naturaleza real de al menos no puede ser delimitada o definida porque el botánico la define; la naturaleza en cuestión — si tiene algún sentido preciso hablar así — tendrá su identidad con independencia de cualquier actividad del botánico, incluida la de definir. Por tanto, lo determinado será a lo sumo el concepto, la noción, y no la cosa misma a que se refiera la expresión definidora.

Al tratar de la lógica de clases hemos establecido definiciones reales en este sentido. Por ejemplo, la definición (DC14) de 77,

$$-\alpha =_{df} \hat{x} [\sim x \in \alpha],$$

se presenta como una definición real: así puede interpretarse la ausencia de comillas; esta ausencia de comillas indica que las expresiones ' $-\alpha$ ' y ' $\hat{x} [\sim x \in \alpha]$ ' eran en (DC14) usadas, pero no mencionadas. Lo mencionado era la clase $-\alpha$. Pero es claro que el sentido de esa definición es decir que la clase $-\alpha$ es $\hat{x} [\sim x \in \alpha]$. $-\alpha$ y $\hat{x} [\sim x \in \alpha]$ no son dos "cosas" distintas, sino sólo dos conceptos distintos, dos sentidos en que se piensan unos mismos individuos, dos modos de referirse gráficamente a unos mismos individuos. En un caso, el de ' $-\alpha$ ', dando simplemente un nombre a su colección (que es un abstracto); en el otro, el de ' $\hat{x} [\sim x \in \alpha]$ ', pensándolos a través de la operación de abstracción que ha servido para coleccionarlos (para construir aquel abstracto).

La anterior reflexión nos enseña que sólo prudentemente puede hablarse de definiciones reales. Puede incluso admitirse que toda definición real es al menos traducible a una operación nominal de las que suelen llamarse 'definiciones ostensivas', o de 'palabra-cosa' (R. Robinson), si bien hay autores que no consideran que estas operaciones sean propiamente definiciones. Una definición ostensiva es una definición nominal que consiste en declarar la denotación de un término señalando materialmente la cosa que se desea nombrar con él. Por ejemplo: 'ése de la izquierda es Juan'. Pues bien: la definición (DC14) puede traducirse a una definición de la expresión ' $-\alpha$ ' parecida a las ostensivas; una definición que dijera: " $-\alpha$ es la de la expresión (materialmente indicada en la pizarra) ' $\hat{x} [\sim x \in \alpha]$ '. (Robinson distingue, sin embargo, entre definiciones reales y definiciones nominales de palabra-cosa.)

En resolución, puede decirse que la distinción entre definición real y definición nominal se refiere sobre todo a la intención y a las necesidades del que define. La intención del que define y sus necesidades son, dicho no subjetivamente, el lenguaje en que se formula la definición, o el estado del conocimiento científico en el momento de definir.

La gran importancia dada por los lógicos tradicionales a la idea de definición real se debe a motivos metafísicos. La metafísica aristotélico-escolástica pensaba que es posible captar la llamada "esencia" de una cosa

(que no sería lo que solemos llamar "esencial" de una cosa en el lenguaje cotidiano, sino algo que determinaría su ser tal o cual cosa) y expresarla mediante la conjunción del "género próximo" a que la cosa pertenece con la "diferencia específica" que distingue la especie o clase a que pertenece esa cosa dentro de aquel género (clase dividida). Así, la definición de la circunferencia como figura cerrada y plana cuyos puntos equidistan de otro llamado centro sería una definición real del tipo más perfecto ("esencial"), construida mediante la conjunción del género próximo — figura cerrada y plana (como también la elipse, por ejemplo) — con la diferencia específica — equidistancia de sus puntos respecto del centro (a diferencia de la elipse, por ejemplo). Esta noción suscita la siguiente crítica: la noción de género próximo es relativa al lenguaje en que se hace la definición, o, lo que es lo mismo, al estado del conocimiento científico en el momento de definir. En vez de figura plana cerrada, por ejemplo, los libros de geometría dieron luego, como "género próximo" de la circunferencia, sección cónica. (Esta observación fue hecha por el historiador de la filosofía Ueberweg.)

También pues por este camino llegamos a la conclusión de que cualquier definición, incluidas las de intención o sentido real, es relativa a un lenguaje. Precisamente por eso es el tema de la definición susceptible de consideración formal. Esta descubre en los sistemas o teorías otras clases de definiciones, de las que vamos a ocuparnos.

Pero antes interesa recordar dos de las varias leyes de la definición establecidas por los lógicos tradicionales (las demás que ellos establecieron son de interés metodológico y didáctico-psicológico). Esas dos leyes son:

- 1.^a Lo definido (*definiendum*) debe ser convertible con lo que define (*definiens*). O sea: entre la expresión definida y la definidora debe haber algún tipo de equivalencia. Esta exigencia está sugerida en toda definición formalmente propuesta en este libro mediante el uso de las notaciones ' $\leftrightarrow_{df}$ ' o ' $=_{df}$ ', sugeridoras de dicha relación de la familia de las equivalencias.
- 2.^a La definición no debe ser circular. O sea: el *definiendum* no debe estar contenido en el *definiens*. Esto no quiere decir que ningún símbolo de los situados a la izquierda de ' $\leftrightarrow_{df}$ ' o ' $=_{df}$ ' (*definiendum* en sentido amplio) pueda estar también entre los situados a la derecha de esos símbolos. El que no puede estar más que a la izquierda es el símbolo que realmente se desea definir (*definiendum* en sentido estricto). Por ejemplo, la siguiente definición no es circular si ' $\sim$ ' es ya previamente conocido y si se entiende como una definición de ' $\wedge$ ' por ' $\vee$ ':

$$\sim [p \wedge q] \leftrightarrow_{df} \sim p \vee \sim q.$$

97. Definiciones sintácticas. — Una definición sintáctica es una regla que permite sustituir (con la ventaja de una mayor brevedad o sin ella) una

expresión (definiens) por otra nueva (definiendum) y a la inversa, sin atender a la significación de una ni de otra. Las definiciones sintácticas son, pues, más nominales, por así decirlo, que las que la lógica tradicional llamaba de ese modo. Pues las definiciones sintácticas, que son definiciones en un cálculo o sistema, no pretenden dar la significación de una expresión, sino simplemente establecer la sustituibilidad recíproca entre dos expresiones del sistema. Una definición sintáctica es, por ejemplo, la siguiente:

$$'p \rightarrow q' \leftrightarrow_{\text{df}} '\sim p \vee q'.$$

si está dada en un cálculo sin interpretar.

Lo más frecuente en un cálculo es encontrarse no propiamente con definiciones sintácticas, sino con esquemas de ellas formulados en el metalenguaje, con variables sintácticas. En efecto, la utilidad de una definición sintáctica como la anterior consiste en que autoriza a sustituir cada fórmula de cierto tipo con ' $\sim$ ' y ' $\vee$ ' por una fórmula con ' $\rightarrow$ ' sin ' $\vee$ ' ni ' $\sim$ ', o a la inversa. Como seguramente se desea poder hacer eso con cualquier fórmula de ese género, y no sólo con ' $\sim p \vee q$ ', lo más oportuno es establecer dicha autorización con variables sintácticas, o sea, mediante un esquema de definiciones. Para el caso anterior se tratará del esquema

$$[X \rightarrow Y] \leftrightarrow_{\text{df}} \sim X \vee Y,$$

que es (D2) de 45.

Las definiciones recursivas que hemos usado varias veces, sobre todo para definir la noción de fórmula de un cálculo, pueden entenderse como definiciones sintácticas en el metalenguaje. Por ejemplo, la definición de la idea de fórmula de la lógica de enunciados puede servir para establecer que, en cualquier afirmación del metalenguaje del cálculo de enunciados, la expresión 'Z es una fórmula del cálculo de enunciados' puede sustituirse por una del tipo 'Z es $'p'$, o Z es $'q'$, o ..., o Z es $\sim X$ y X es una fórmula del cálculo de enunciados, o ..., o Z es $X \leftrightarrow Y$, y X e Y son fórmulas del cálculo de enunciados'. La efectiva traducción formal de una definición recursiva cualquiera a una corriente definición sintáctica es una tarea bastante complicada que ha ocupado a autores de primera importancia. Aquí sólo nos interesaba aclarar la naturaleza sintáctica de las definiciones recursivas en un cálculo. Las definiciones recursivas no hacen ninguna alusión al significado del definiendum, sino que exponen sólo su constitución gráfica.

98. Definiciones semánticas. — Una definición semántica es una atribución de significación (denotación o valor) a un símbolo o una expresión. Las definiciones semánticas son, por tanto, propias de los cálculos interpretados, o sea, de los lenguajes formalizados. Las tablas veritativas que determinan el uso de los funtores veritativos pueden considerarse definiciones semánticas de fórmulas con esos funtores. Así por ejemplo, la tabla de ' $\wedge$ '

podría traducirse por la siguiente definición en el metalenguaje semántico:

$$X \wedge Y \leftrightarrow_{\text{df}} X \text{ vale } V \text{ e } Y \text{ vale } V.$$

Una definición semántica es una definición nominal en el sentido tradicional, pues lo que determina es [el uso de] una expresión.

Cuando un cálculo está interpretado y es ya, por tanto, un lenguaje, las definiciones sintácticas que se hubieran hecho en él se convierten en semánticas. Pero eso no quiere decir que haya una correspondencia biunívoca entre definiciones sintácticas y definiciones semánticas correspondientes. A una definición semántica pueden corresponder varias definiciones sintácticamente distintas. Por ejemplo, supongamos un sistema axiomático de la lógica de enunciados sin más funtores primitivos que ' $\sim$ ', ' $\wedge$ ', ' $\rightarrow$ '. Después de los axiomas, la presentación del sistema puede contener la siguiente definición (entre otras):

$$(D1) \quad X \vee Y \leftrightarrow_{\text{df}} \sim [\sim X \wedge \sim Y].$$

Pero en vez de (D1) el sistema podría presentar

$$(D2) \quad X \vee Y \leftrightarrow_{\text{df}} [\sim X \rightarrow Y].$$

Normalmente no presentará las dos: primero, porque entonces las reglas primitivas del cálculo serían de aplicación ambigua a fórmulas con ' $\vee$ '; segundo, porque una de las dos definiciones basta calculísticamente. Supongamos que el sistema contiene (D2) y no (D1). Entonces, desde el punto de vista sintáctico, sólo (D2) es utilizable en demostraciones. (D1) y (D2) son sintácticamente dos definiciones.

En cambio si, como es corriente, la interpretación de ese cálculo por su semántica atribuye la misma significación, el mismo valor, a ' $\sim [X \wedge \sim Y]$ ' que a ' $[\sim X \rightarrow Y]$ ', la semántica de dicho cálculo contendrá propiamente una sola definición de ' $\vee$ ', la cual corresponderá a dos posibles definiciones sintácticas (D1) y (D2) (y a cualquier otra semánticamente equivalente). (D1) y (D2) son dos definiciones sólo desde el punto de vista sintáctico, que no atiende más que a los símbolos usados, al aspecto gráfico de las fórmulas. La fórmula con doble flecha entre las dos expresiones definidoras,

$$[\sim X \rightarrow Y] \leftrightarrow \sim [\sim X \wedge \sim Y],$$

será un teorema de la sintaxis, si las reglas de inferencia son las corrientes. En su semántica, la afirmación correspondiente a esa fórmula indicará que es lo mismo definir ' $\vee$ ' por (D1) que por (D2), o sea, que (D1) y (D2) son semánticamente la misma definición.

99. Definiciones por abstracción. Definiciones creadoras. — En un sentido amplio, puede considerarse definida por abstracción toda noción

construida mediante esa operación. Por ejemplo, la clase α puede en este sentido considerarse definida a partir del esquema ' $\sim x \in \alpha$ ' (o del esquema ' $\sim Px$ ') mediante la operación indicada por el abstractor de clases ' $\hat{x}$ '. (Cfr. 95.)

En un sentido más estrecho se considera definidas por abstracción las nociones construidas a partir de una relación de equivalencia (cfr. 88), o sea, de una relación simétrica y transitiva. Así, por ejemplo, la relación entre convecinos (en el sentido del ejemplo de 88), o relación de convecindad, permite abstraer de ella, como vimos, una serie de clases — las clases de equivalencia respecto de convecindad, como la clase de los convecinos de Juan, la clase de los convecinos de Pedro (no siendo Juan y Pedro convecinos), etc. —, y también ulteriormente, la propiedad de ser un municipio, propiedad que corresponde a la clase de todas las clases de equivalencia respecto de la relación convecindad.

Las definiciones por abstracción realizan, en el terreno formal, algo que recuerda la idea de definición real, más precisamente, lo que suele llamarse 'definición real genética', que expone la naturaleza de una cosa declarando cómo llega a ser o se engendra esa cosa. Una definición por abstracción expresa la creación o introducción de un concepto nuevo. Para seguir con nuestro ejemplo, si se dispone de la relación de convecindad, con ella puede obtenerse por abstracción el concepto de municipio, el cual es nuevo en el sentido de que permite pensar de un modo nuevo la realidad de los individuos dados. Pero debe observarse que la novedad que introduce una definición por abstracción es relativa al lenguaje en que se da, o sea, al estado del conocimiento, al repertorio de conceptos en el momento de definir. La definición de convecindad aporta algo nuevo en un lenguaje que, para el correspondiente universo del discurso, hubiera sido hasta entonces sólo relacional.

El parentesco que acabamos de observar con la idea de definición real se da también en las definiciones llamadas 'creadoras'. Dada una teoría — señaladamente en la forma de sistema axiomático — se dice que una definición es creadora en ella cuando hace posible demostrar fórmulas que, aunque no contienen más que símbolos primitivos, son sin embargo inde demostrables en el sistema si no se da aquella definición.

La definición creadora tiene pues en común con la definición por abstracción el introducir novedad y el hacerlo respecto de un cálculo o lenguaje determinado. Así vemos que incluso las más reales de las definiciones son operaciones lingüísticamente condicionadas.

Suponiendo una teoría de la división de números naturales que no tuviera símbolo para cero ni tuviera definida la división por cero, las definiciones

$$(D3) \quad (x) [x - x =_{\text{at}} 0],$$

$$(D4) \quad (x) [x/0 =_{\text{at}} n],$$

siendo ' n ' una constante, podrían servir para ilustrar la idea de definición creadora, pues con ellas se podría demostrar, por ejemplo, la fórmula

$$b/a - a = c/d - d,$$

que sólo contiene símbolos primitivos de la teoría y no sería demostrable en ella sin (D3) y (D4).

Las definiciones creadoras oscurecen la verdadera estructura de un sistema axiomático, por ejemplo, y por eso deben evitarse.

CAPÍTULO XVII

EL ANÁLISIS FORMAL DE LA INDUCCIÓN

100. El planteamiento clásico del problema de la inducción. — En 20 se recordó la definición tradicional de la inducción como razonamiento o inferencia que lleva de conocimientos menos generales a afirmaciones más generales. En el ejemplo de 20 se tiene el conocimiento de que unas cuantas masas de helio, que se han observado, presentan la propiedad de que el producto de la presión a que están sometidas en cada caso por el volumen que ocupan es 22,41; de ese conocimiento el químico pasa a la afirmación, más general, de que *toda* masa de helio (reduciendo el alcance del ejemplo a este gas) tiene dicha propiedad.

El paso de una afirmación a otra que se infiere de ella, paso que al formular reglas de cálculo representábamos con una raya horizontal, se expresa también frecuentemente por el llamado 'signo de aserción', '⊢', aunque éste no es el uso principal de dicho símbolo. A la izquierda de '⊢' se escriben las afirmaciones de partida (premisas del paso), y a su derecha aquella o aquellas a las que se pasa. Si simbolizamos por 'P' la propiedad del ejemplo recordado y por 'a₁', 'a₂', ..., 'a_n' las masas de helio observadas, el paso inductivo que estamos considerando puede representarse por

$$(1) \quad Pa_1, Pa_2, \dots, Pa_n \vdash (x)Px.$$

'x' es una variable cuyo campo está constituido por la clase de las masas de gas (en el ejemplo de 20) o de las masas de helio (en la reducción que acabamos de hacer de dicho ejemplo).

Cuando se trata de deducciones, un paso de 'p' a 'q' 'p ⊢ q', no da lugar a una afirmación deductivamente válida, 'q', más que si valen 'p → q' y 'p'. Si éste es el caso, una aplicación de la regla de inferencia o *modus ponens* permite la afirmación de 'q'. Es claro que esto no ocurre con (1). Pues si valiera

$$(2) \quad Pa_1 \wedge Pa_2 \wedge \dots \wedge Pa_n \rightarrow (x)Px,$$

como la conjunción 'Pa₁ ∧ Pa₂ ∧ ... ∧ Pa_n' puede construirse con el conocimiento que se tiene, con las premisas del paso (1), entonces '(x)Px' valdría deductivamente, y el paso (1) sería un paso deductivo. Pero (2) no es formalmente verdadero — no es una implicación — más que si a₁, a₂, ..., a_n son todas las masas de helio, es decir, todo el campo de 'x'. O sea, si vale:

$$(3) \quad (x) [Ixa_1 \vee Ixa_2 \vee \dots \vee Ixa_n],$$

lo cual, como es obvio, no es el caso: no se sabe que el número de las masas de helio sea precisamente n.

Como (3) no es verdadero, no vale la implicación (2), y el paso (1) no es un paso deductivo. Esto suele expresarse diciendo que el paso (1) no permite afirmar '(x)Px' "con certeza". El ejemplo de la ley de Boyle-Mariotte nos recuerda esta circunstancia, pues, como es sabido, ese enunciado no tiene más que una vigencia práctica, aproximativa y limitada, aunque en otro tiempo, a falta de experiencia suficiente sobre la licuefacción de los gases, se le tuvo por una verdadera ley científica.

Consideremos el caso de que (3) valiera. En vez de (2) podríamos escribir entonces:

$$(4) \quad Pa_1 \wedge Pa_2 \wedge \dots \wedge Pa_n \wedge (x) [Ixa_1 \vee Ixa_2 \vee \dots \vee Ixa_n] \rightarrow (x)Px.$$

(4) sí que es formalmente verdadero, y el paso correspondiente del antecedente al consecuente está, por tanto, formalmente justificado. Pero no responde ya al concepto de inducción, pues es un paso deductivo, y (4) una verdadera implicación. En realidad vale incluso:

$$(5) \quad Pa_1 \wedge Pa_2 \wedge \dots \wedge Pa_n \wedge (x) [Ixa_1 \vee Ixa_2 \vee \dots \vee Ixa_n] \leftrightarrow (x)Px.$$

(4) es deductivamente demostrable. Sin embargo, la tradición consideró inductivo el paso correspondiente a (4):

$$(6) \quad Pa_1, Pa_2, \dots, Pa_n, (x) [Ixa_1 \vee Ixa_2 \vee \dots \vee Ixa_n] \vdash (x)Px.$$

Se llamaba a este paso u operación 'inducción completa'. La razón por la cual se consideraba inductivo este tipo de argumentación era metafísica: se basaba en distinguir entre una generalidad meramente extensional, representada por los n primeros miembros de la conjunción que es el antecedente de (6), por ejemplo, y una universalidad intensional, la "captación de la esencia" del hecho, por así decirlo, que estaría representada por el consecuente de (4), que es la conclusión del paso (6). Es claro que, aun en el supuesto de que esa distinción tuviera algún valor filosófico, formalmente no puede tomarse en consideración: tanto el paso (6) cuanto su fundamento, la verdad formal de la implicación (4), son deductivos.

Podemos, por tanto, considerar que los casos propiamente inductivos son aquellos en los cuales el campo de la variable generalizada, 'x' en nuestro

ejemplo, tiene un número cardinal (si lo tiene) mayor que el número cardinal, n , de la clase de los casos observados. Ésta es la situación en el ejemplo de 20 y en su versión restringida al helio, (1) de esta sección. (1) es una instancia de lo que tradicionalmente se llamaba 'inducción incompleta'. Recoge una situación muy corriente en la ciencia empírica, y su significación es principalmente metodológica. La justificación del paso (1), si la tiene — es decir, la justificación de una afirmación del enunciado ' $(x) Px$ ' solo —, es una tarea de la teoría del método. Formalmente es injustificable, porque su supuesto deductivo o formal, (2), no es una verdad formal.

Pero eso no quiere decir que el problema de la inducción no sea susceptible de más análisis formal que el poco que nos ha sido necesario para llegar a la anterior conclusión negativa. Por una parte, es posible precisar más la situación desde un punto de vista lógico, no de teoría del método o teoría del conocimiento. Por otra parte, es posible encontrar un modo de construir una teoría formal de la inducción. Łukasiewicz ha ofrecido una interesante esquematización para aclarar el primer punto. Carnap ha suministrado el intento más importante de solución para el segundo. Las siguientes secciones 101-107 pueden leerse como una introducción a las ideas de esos dos importantes autores.

101. El esquema reductivo de Łukasiewicz. — Si los dos componentes principales de (1) de 99 se encontraran en posiciones invertidas,

$$(7) \quad (x) Px \vdash Pa_1, Pa_2, \dots, Pa_n,$$

se tendría un paso deductivamente correcto, basado en la verdad formal (implicación)

$$(8) \quad (x) Px \rightarrow Pa_1 \wedge Pa_2 \wedge \dots \wedge Pa_n,$$

y en la regla de inferencia, o modus ponens. Esto tiene, naturalmente, que ocurrir en todos los casos de "inducción incompleta" igual que en los de "inducción completa", en el sentido tradicional ya visto de esas expresiones. También ocurre en todos los pasos deductivos, es decir, de aplicación del modus ponens. La diferencia está en que en los pasos deductivos se tiene como premisa el antecedente de esos condicionales del tipo de (8), mientras que en los pasos inductivos se tiene como premisa el consecuente. El paso deductivo (7) puede esquematizarse así, con la representación para reglas que ya conocemos:

Esquema I

$$\frac{\begin{array}{l} (x) Px \rightarrow Pa_1 \wedge Pa_2 \wedge \dots \wedge Pa_n \\ (x) Px \end{array}}{Pa_1, Pa_2, \dots, Pa_n}$$

(Tanto (7) cuanto el Esquema I pueden entenderse como representación abreviada de n pasos.) En cambio, el paso inductivo (1) de 100, que recoge el razonamiento sobre la ley de Boyle-Mariotte en 20, tendría que representarse así:

Esquema II

$$\frac{\begin{array}{l} (x) Px \rightarrow Pa_1 \wedge Pa_2 \wedge \dots \wedge Pa_n \\ Pa_1, Pa_2, \dots, Pa_n \end{array}}{(x) Px}$$

La comparación de los esquemas I y II ilustra el hecho de que la diferencia principal entre pasos deductivos y pasos inductivos es, como vimos, que los segundos no están justificados formalmente, sino que, por así decirlo, se basan en un modus ponens al revés, carente de garantías formales. Ahora bien: si ésa es la diferencia principal, entonces la mayor o menor generalidad del antecedente respecto del consecuente no nos es necesaria para expresar lo que define la situación. Ésta puede quedar caracterizada en cada caso por los siguientes esquemas:

Esquema deductivo

$$\frac{p \rightarrow q}{p} q$$

Esquema reductivo

$$\frac{p \rightarrow q}{q} p$$

A la vista de esos dos esquemas puede admitirse la existencia en la práctica científica de dos grandes géneros de argumentaciones o pasos de unas afirmaciones a otras: pasos deductivos y pasos reductivos. Lo interesante de estos últimos desde el punto de vista formal no es que supongan una ganancia en generalidad. Y lo interesante formalmente de los pasos deductivos no es que supongan una pérdida de generalidad, como puede verse por el esquema deductivo, en el que no intervienen consideraciones de generalidad. La diferencia interesante es que el esquema deductivo es formalmente válido (válido para cualquier interpretación posible), mientras que el reductivo no lo es.

En relación con el anterior análisis de Łukasiewicz se ha desarrollado en la metodología la costumbre de entender por 'inducción' toda inferencia no-deductiva. La inducción en el sentido tradicional sería entonces aquel

caso especial de reducción en el cual el lugar de ' p ' en el anterior esquema reductivo está ocupado por un enunciado más general que el que ocupa el lugar de ' q '. Es la que Carnap llama 'inferencia inductiva universal'.

102. La relación de inducción. — El esquema reductivo es demasiado estrecho o exigente para poder recoger todas las situaciones inductivas, en el nuevo sentido amplio de esta palabra. Ese esquema sugiere muy adecuadamente la situación lógica característica en la interpretación de un fenómeno cuando ya se posee una ley o teoría sobre la clase de fenómenos a la que pertenece el observado. Supongamos que un biólogo observa, en una generación de n individuos, que sólo $n/4$ individuos presentan el carácter P , poseído por los abuelos. El biólogo posee una teoría con una ley estadística según la cual los caracteres recesivos aparecen en la cuarta parte de los individuos de la segunda generación. Esta es una ley del tipo ' $p \rightarrow q$ ', siendo ' p ' 'x es un carácter recesivo' y ' q ' 'x aparece en la cuarta parte de los individuos de segunda generación'. De ' $p \rightarrow q$ ' puede obtenerse el enunciado singular: ' P es un carácter recesivo $\rightarrow P$ aparece en la cuarta parte de los individuos de segunda generación'. Llamemos ' $s \rightarrow t$ ' a este enunciado singular deductivamente obtenido de la ley, de la que es una instancia. Nuestro biólogo ha hecho, pues, la observación expresada por ' t '. El paso reductivo consiste en afirmar ' s ', o sea, que el carácter P es recesivo.

El enunciado ' s ' (o ' p ') es en casos de este género una explicación del hecho expresado por ' t ' (o ' q '). Por eso podemos decir que el esquema reductivo recoge especialmente la inducción que busca hipótesis explicativas de los fenómenos observados. El científico puesto ante los fenómenos expresados por ' q ' busca, con el ejercicio racional de la imaginación, un enunciado que los explique o implique, esto es, un enunciado (por regla general, varios: toda una teoría) del que pudiera obtenerse deductivamente el enunciado ' q ' que describe los fenómenos observados. Pero el esquema reductivo no dice nada sobre ese paso imaginativo. Y, ciertamente, como tal paso imaginativo, este aspecto de la cuestión no es nada propio de la lógica, sino de la psicología. Mas el paso imaginativo hacia la hipótesis utiliza algo que sí es relevante para la lógica: una relación de ' q ' a ' p '; mientras que el esquema reductivo sólo atiende a la relación de ' p ' a ' q ', la cual es una relación deductiva, de implicación. Ahora bien: seguramente la relación que más importa aclarar para entender la inducción es precisamente la relación de ' q ' a ' p ' que subyace al paso imaginativo de lo conocido a lo hipotéticamente afirmado, de la información — como diremos en adelante — a la hipótesis.

En los primeros años del análisis moderno de la inducción, B. Russell indicó que el paso por una forma ' $p \rightarrow q$ ' (hipótesis implica información) es muchas veces ocioso y redundante, si es que realmente se amplía el concepto de inducción para incluir en él toda inferencia no-deductiva. Su-

pongamos, dice Russell, que se quiera justificar la afirmación 'Sócrates es mortal' (' r ') partiendo del conocimiento que tenemos acerca de hombres que han muerto. La situación no es deductiva; es por tanto inductiva. La concepción tradicional del uso de la inducción y la deducción en el método científico afirma ' r ', 'Sócrates es mortal', pasando por 'todos los hombres son mortales' (' p '). Pero ' p ' es un enunciado inductivamente obtenido a partir de un conjunto (o conjunción) de enunciados: ' a murió', ' b murió', etc., conjunción a la que llamaremos ' q ' y que no agota, como es natural, la clase de los hombres, porque si la agotara ya estaría dicho en la información que Sócrates es mortal, puesto que uno de los enunciados sería 'Sócrates murió'. Pero entonces, si ' p ' no sirve para aumentar la "certeza" de ' r ', porque ' p ' mismo no es "cierto", sino inductivo, ¿para qué dar el rodeo por ' p ' para inferir ' r '? Más sencillo es admitir que la verdadera inferencia es aquí de ' q ' a ' r ', una inferencia sin ningún momento deductivo, basada en una relación propiamente inductiva de la información, ' q ', a la hipótesis, ' r ', la cual es, por cierto, aún menos general que la información.

Esta discusión permite concluir: 1.º) el esquema reductivo, inspirado en la estructura deductiva del modus ponens, no pone suficientemente de manifiesto la relación inductiva de la información a la hipótesis; 2.º) aunque sin duda presente en la inducción que busca hipótesis explicativas, ' p ', de fenómenos observados y expresados en la información, ' q ', la relación deductiva $p \rightarrow q$ de la hipótesis a la información no parece un elemento necesario de la situación inductiva (en el sentido hoy generalmente admitido del término 'inductiva' (= 'no-deductiva')), mientras que en toda situación inductiva se encuentra, naturalmente, una relación propiamente inductiva de la información a la hipótesis.

Rudolf Carnap, aclarando nociones que habían aparecido anteriormente en la teoría del cálculo de probabilidades, ha abierto una nueva época en el análisis formal de la inducción al entender la relación inductiva de la información a la hipótesis como una *relación de confirmación*. La relación inductiva de ' q ' a ' p ' es la medida en la cual ' q ' confirma a ' p ', independientemente de que ' p ' implique a su vez a ' q ' o no lo implique. En efecto, el uso de la relación inductiva de ' q ' a ' p ' en la ciencia no supone siempre la existencia de la relación deductiva de ' p ' a ' q ', es decir, la verdad de ' $p \rightarrow q$ '. Esto puede verse examinando la situación que se da en los principales géneros de inferencia inductiva. Carnap distingue cinco de ellos, sin pretender que la división sea exhaustiva. Pero es suficiente para la consideración que tenemos que hacer:

1. *Inferencia inductiva universal.* — Es una de las dos únicas consideradas por la teoría tradicional de la inducción. Consiste en pasar de la información de que todos los miembros de una clase observada, γ (o sea, todos los individuos o hechos que tienen la propiedad de haber sido observados, C) pertenecen a la clase β (tienen la propiedad B), a la afirmación

de que todo individuo o hecho pertenece a β (todo individuo o hecho tiene la propiedad B). La situación es pues:

información, q : $(x) [x \in \gamma \rightarrow x \in \beta]$, o $\gamma \subset \beta$, o $(x) [Gx \rightarrow Bx]$;

hipótesis, p : $(x) [x \in \beta]$, o $\beta = V$, o $(x) Bx$.

Es claro que en este caso vale ' $p \rightarrow q$ '.

L1(1)	$(x) Bx$	RP
L2(2)	Gx	RP
L3(1)	Bx	RE (); L1
L4(1, 2)	Bx	RA2; L2, L3
L5(1)	$Gx \rightarrow Bx$	RI $\rightarrow$; L4
$\times$ L6(1)	$(x) [Gx \rightarrow Bx]$	RI (); L5
L7(0)	$(x) Bx \rightarrow (x) [Gx \rightarrow Bx]$	RI $\rightarrow$; L6.

Un ejemplo de esta inferencia es la obtención de la ley de Boyle-Mariotte según el ejemplo 20 (en un universo constituido por las masas de gas).

2. *Inferencia inductiva inversa*. — Es la inferencia por la cual se pasa de la información de que todos los miembros de una subclase, γ , de una clase, α , pertenecen a la clase β , a la hipótesis de que todos los miembros de α pertenecen a β . Esta inferencia, también contemplada por la teoría tradicional de la inducción, es una limitación de la inferencia inductiva universal a un universo más reducido, la clase α . En ella vale también la relación deductiva de la hipótesis, ' p ', a la información, ' q ', pues la situación es:

información previa, q' : $\gamma \subset \alpha$, o $(x) [Gx \rightarrow Ax]$;

información, q : $\gamma \subset \beta$, o $(x) [Gx \rightarrow Bx]$;

hipótesis, p : $\alpha \subset \beta$, o $(x) [Ax \rightarrow Bx]$.

Un ejemplo de esta inferencia es el anterior de la ley de Boyle-Mariotte limitado al helio y en el mismo universo constituido por las masas de gas.

3. *Inferencia inductiva directa*. — Es ésta una inferencia inductiva poco útil y frecuente en la práctica, pero bien caracterizable en la teoría. Supongamos que se tiene información estadística acerca de toda la clase (población) de individuos considerada; por ejemplo, que el 70 % de los miembros de la clase α pertenecen también a la clase β . De esa información puede pasarse a la hipótesis de que en una determinada subclase (muestra), γ , de la población α se encontrará un 70 % de individuos que son también miembros de β . La situación será en este ejemplo:

información, q : 70 % de los α son $\beta \wedge \gamma \subset \alpha$;

hipótesis, p : 70 % de los γ son β .

Aquí no vale ' $p \rightarrow q$ '.

4. *Inferencia inductiva predictiva*. — Es la inferencia inductiva más importante en la práctica científica. Consiste en pasar de una información sobre una muestra, β , de una población, α , a una afirmación hipotética acerca de otra muestra, γ , independiente de la primera, de la misma población α . Por ejemplo: la información puede ser que el 70 % de los inmigrantes que viven en el distrito III de Barcelona son levantinos; y la hipótesis puede ser que el 70 % de los inmigrantes que viven en el distrito VII de Barcelona son levantinos. Cuando la segunda muestra es un solo individuo o hecho, se tiene la situación del caso Sócrates discutido por Russell.

Tampoco en esta inferencia está implícita la relación de implicación de la hipótesis a la información. La situación puede describirse, con el anterior ejemplo, como sigue:

información, q : $\beta \subset \alpha \wedge \gamma \subset \alpha \wedge \beta \cap \gamma = \Delta \wedge$ 70 % de los β son δ ;

hipótesis p : 70 % de los γ son δ .

5. *Inferencia inductiva analógica*. — Tiene la misma estructura que la anterior, pero las dos muestras son de un solo individuo cada una. Tampoco en ella vale ' $p \rightarrow q$ '.

Para cada una de las tres inferencias inductivas en las que no vale ' $p \rightarrow q$ ' puede construirse una implicación ' $p' \rightarrow q$ ' modificando la hipótesis para hacerla universal. Pero entonces, naturalmente, se está en el caso de la inferencia inductiva universal o inversa. Ejemplo para la inferencia predictiva:

hipótesis modificada, p' : $(\gamma) [\gamma \subset \alpha \rightarrow$ 70 % de los γ son $\delta]$.

Este examen de las principales clases de inferencia inductiva confirma en la opinión de que lo principal del problema de la inducción es la relación inductiva, o de confirmación, que va de la información, ' q ' en los anteriores ejemplos, a la hipótesis, ' p '. Si simbolizamos esa relación por ' $\rightarrow$ ', podemos representar del modo siguiente el esquema de toda inferencia inductiva:

$$\frac{q \rightarrow p}{q}$$

p

La relación en que se basa un paso inductivo es entonces expresable por

$$(9) \quad q \rightarrow p,$$

análogamente a como la relación en que se basa un paso deductivo es expresada por ' $q \rightarrow p$ '. El que, además, valga ' $p \rightarrow q$ ', la implicación de la información por la hipótesis, es inesencial, salvo para caracterizar las inducciones inversas (universales o restringidas), únicas consideradas por la tradición con el nombre de 'inducción incompleta'.

103. La posibilidad de una lógica inductiva según Carnap. — A propósito de la deducción y de la inducción hemos distinguido entre pasos (u operaciones) y relaciones en que se basan esos pasos. Así, en el caso de la deducción, tenemos, por una parte, el paso de ' p ' a ' q ',

$$(10) \quad [p \rightarrow q], p \vdash q, \quad o$$

$$\frac{\begin{array}{c} p \rightarrow q \\ p \end{array}}{q},$$

y, por otra parte, la relación de implicación

$$(11) \quad [p \rightarrow q] \wedge p \rightarrow q.$$

El paso deductivo es una aplicación de la relación de implicación, mientras que la relación misma de implicación (11) es sobre todo de interés teórico. Cuando uno deduce en la práctica, usa la relación (11) en un esquema (10). Cuando uno estudia teoría lógica *se interesa por saber* si (11) es o no una verdad formal. Hablando restrictivamente, lo relevante para la teoría lógica es la relación (11), mientras que la operación (10) es de interés metodológico: interesa propiamente a la metodología de la deducción. Precisamente por eso el resultado de (10), que es ' q ', sólo es válido si valen ' $p \rightarrow q$ ' y ' p ', mientras que la afirmación (11) es válida en cualquier caso, aunque ' $p \rightarrow q$ ' y ' p ' sean empírica o formalmente falsas. El esquema operativo (10) tiene pues que ver con la experiencia de un modo directo. La relación (11) no, sino sólo del modo indirecto en que las verdades formales tienen que ver con la experiencia (cfr. 7).

Esa misma reflexión es aplicable al caso de la inducción. Para el paso inductivo de ' p ' a ' q ', por ejemplo, se tiene una relación

$$(12) \quad [p \rightarrow q] \wedge p \rightarrow q,$$

de la que podemos decir que ella será lo relevante para un estudio formal teórico de la inducción.

La conjunción ($\wedge$), que significa la copresencia de la relación y la información, no tiene ninguna peculiaridad: es la misma en la inducción que en la deducción. Por eso la relación que importa a la teoría lógica de la deducción es la expresada por ' $\rightarrow$ ' (en el sentido de implicación, o sea, de condicional universalmente válido), y la que importará a una posible teoría lógica de la inducción es la representada por ' $\rightarrow$ '. Las fórmulas relevantes son pues, respectivamente:

$$(13) \quad p \rightarrow q,$$

$$(14) \quad p \rightarrow q.$$

Ahora bien: si existe una teoría formal de la implicación, que es la teoría lógica de la deducción, ¿por qué no podría existir una teoría formal de ' $\rightarrow$ '? La lógica inductiva de Carnap, cuyo programa vamos a considerar ahora brevemente, es una teoría formal de ' $\rightarrow$ ', de la relación de confirmación entre enunciados.

La relación de confirmación se diferencia de la relación (deductiva) de implicación, ante todo, porque no puede dar lugar a un esquema de operación análogo al deductivo. Por ejemplo, un esquema

$$\frac{\begin{array}{c} p \rightarrow q \\ p \end{array}}{q}$$

representa una operación injustificable, pues la relación ' $\rightarrow$ ' no es la relación de implicación, de modo que ' p ' no "acarrea necesariamente" a ' q '; como suele decirse tradicionalmente, lo afirmado por inducción no puede afirmarse "con certeza". Podría pensarse en un esquema que reprodujera más fielmente las operaciones inductivas; por ejemplo, un esquema operativo la fórmula de cuya última línea estuviera afectada por algún símbolo que indicara "sin certeza":

$$\frac{\begin{array}{c} p \rightarrow q \\ p \end{array}}{\rightarrow q}$$

La comparación de este esquema (que más adelante abandonaremos, por poco adecuado aún) con el de las operaciones deductivas ilustra de un modo nuevo la diferencia principal que antes hemos visto entre uno y otro; inducción y deducción, como toda operación y a diferencia de la teoría lógica, parten de unos datos de hecho, a saber: que valen dos determinadas premisas. Pero, sobre esa base, la operación deductiva obtiene una afirmación "cierta", sin cualificar; mientras que la operación inductiva obtiene

afirmaciones "inciertas", cualificadas. Y como las operaciones se basan en las relaciones respectivas, esa diferencia sugiere lo siguiente: la relación deductiva (o implicación) es una relación de necesidad, por así decirlo, una relación sin excepciones y por la cual el miembro relacionante (antecedente) acarrea consigo el miembro relacionado (consecuente), de modo que puesto el primero está ya puesto el segundo. Esto no ocurre en la relación inductiva (confirmación). Mas, por otra parte, la relación inductiva tiene algún peso, puesto que a través de ella la información confirma la hipótesis. Podría, por eso, decirse que la relación inductiva rinde imperfecta o parcialmente lo que plena y perfectamente rinde la implicación: lleva de algún modo al consecuente inductivo, lo confirma, aunque "sin certeza". Por esta razón indica Carnap que la relación de confirmación puede entenderse como una relación de *implicación parcial*.

Como se recordará por la lógica de clases, una propiedad monádica corresponde a una clase, y la implicación entre propiedades monádicas corresponde a la inclusión entre las dos clases respectivas. Esto nos va a permitir aclararnos provisionalmente con unas figuras las situaciones de implicación y de implicación parcial. Sea la implicación

$$(15) \quad (x) [Px \rightarrow Qx].$$

Llamemos ' α ' a la clase asociada a ' P ' y ' β ' a la clase asociada a ' Q '. La situación recogida en (15) puede también expresarse por

$$(16) \quad \alpha \subset \beta, \text{ o por}$$

$$(17) \quad (x) [x \in \alpha \rightarrow x \in \beta].$$

La figura 1 la representa:

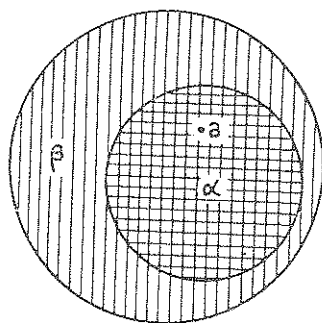


FIG. 1

Supongamos ahora que se tenga la afirmación adicional

$$(18) \quad a \in \alpha, \text{ o } Pa.$$

Entonces la inclusión de α en β , o la implicación de ' Qx ' por ' Px ', justifica el paso deductivo a

$$(19) \quad a \in \beta, \text{ o } Qa.$$

La situación quedó ya representada en la figura 1.

Pasemos ahora a la relación inductiva entendida como relación de implicación parcial. Que ' Px ' implica parcialmente a ' Qx ' puede entenderse

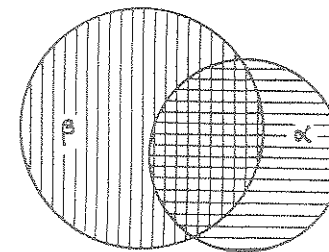


FIG. 2

en el sentido de que α está parcialmente contenida en β :

Esta situación puede expresarse en el lenguaje de la lógica de clases por

$$(20) \quad \alpha \cap \beta \neq \Lambda, \text{ o por}$$

$$(21) \quad \exists ! \alpha \cap \beta, \text{ o por}$$

$$(22) \quad \exists x [x \in \alpha \wedge x \in \beta].$$

Elegiremos (20).

Supongamos ahora que se tenga la misma información adicional que en el caso anterior, o sea

$$(18) \quad a \in \alpha, \text{ o } Pa.$$

Esta vez (18), junto con (20), no basta para legitimizar el paso a (19) pues el individuo a , que según (18) es un α , podría perfectamente ser uno de los miembros de α que no son miembros de β , es decir, un miembro de la clase $\alpha \cap \neg \beta$, que es una subclase de α , como puede apreciarse en la figura 2. Pero es claro que a podría ser también un miembro de la clase $\alpha \cap \beta$, que es también una subclase de α ; y en este caso valdría (19). La tarea de la lógica inductiva consiste, entre otras cosas, en precisar el sentido de esa expresión 'podría ser', al modo como la lógica deductiva aclara el sentido de expresiones como 'acarrear necesariamente' (el implicante al implicado) mostrando que 'necesariamente' no tiene en ese uso ninguna significación empírica o material. El programa de la lógica inductiva de Carnap aspira a precisar esa frase, o la "incerteza" de los resultados inductivos, abstra-

yendo de la relación de confirmación, $\rightarrow$), una función de confirmación, c . Se trata de obtener la teoría de un concepto cuantitativo de confirmación. En las siguientes secciones estudiaremos dicho programa, pero antes vamos a hacer una reflexión preliminar acerca de cómo puede, en general, cuantificarse un concepto lógico como el de confirmación.

Supongamos que en la situación de implicación parcial entre ' Px ' y ' Qx ' descrita por (18) y (20) y representada por la figura 2, se ignora la localización del punto que representa al individuo a , miembro de α , pero se cuenta en cambio con la información adicional de que la clase α tiene 100 miembros, 70 de los cuales lo son además de la clase β ; escribiremos abreviadamente:

$$(23) \quad 70 \text{ de los } 100 \alpha \text{ son } \beta.$$

Sea ahora la hipótesis

$$(19) \quad a \in \beta, \text{ o } Qa.$$

De acuerdo con el análisis de la inducción visto hasta aquí (19) no es una afirmación "cierta" en base a (18) y (20), sino "incierta", cosa que expresábamos por

$$(24) \quad \rightarrow a \in \beta, \text{ o } \rightarrow Qa.$$

Si la notación de (24) tuviera buen sentido, entonces, con la información adicional, (23), de que ahora disponemos, a saber, que 70 de los 100 miembros de α son también miembros de β , seguramente tendría sentido sustituir nuestro símbolo de "incerteza" por '70%', por ejemplo, o por '7/10', o por '0,7'. Pero en realidad una expresión como

$$(25) \quad 70\% a \in \beta, \text{ o } 7/10 a \in \beta, \text{ o } 0,7 a \in \beta, \text{ o } 70\% Qa, \text{ etc.,}$$

no hace buen sentido. Pues del mismo modo que, en cálculo de probabilidades, la afirmación de que la probabilidad de la hipótesis ' q ' respecto de una información ' p ' es 0,7 no tiene el absurdo significado de que de cada diez veces que afirmemos ' q ' siete veces vaya a resultar verdadera y tres falsa, así tampoco ' $0,7 a \in \beta$ ' puede querer decir el absurdo de que de cada diez veces que tomemos el individuo a siete va a resultar ser miembro de β y tres no. Mas como ese absurdo es lo que parece indicar la notación, ésta resulta poco adecuada y hay que abandonarla.

La fracción 0,7 no expresa la validez de la afirmación ' $a \in \beta$ ', sino el grado en el cual ' $a \in \beta$ ' está fundamentada por la información disponible, o sea, por la conjunción de (18), (20) y (23). Lo que debemos expresar mediante alguna notación adecuada es pues: la confirmación de la hipótesis (19) por la información (18) $\wedge$ (20) $\wedge$ (23) es 0,7. Carnap expresa esto con la notación funcional siguiente:

$$(26) \quad c((19), (18) \wedge (20) \wedge (23)) = 0,7,$$

o, llamando ' q ' a la hipótesis (19) y ' p ' a la información, conjunción de los enunciados (18), (20) y (23):

$$(27) \quad c(q, p) = 0,7.$$

En este punto son necesarias dos observaciones:

Primera: el concepto de confirmación, de relación inductiva, resulta coincidir con el concepto lógico de probabilidad. Debe distinguirse entre concepto empírico y concepto lógico de probabilidad. El primero se define como la frecuencia relativa, entre los miembros de una clase α , de una determinada propiedad P . El segundo puede definirse precisamente como el grado de confirmación de una hipótesis por una información. Razonamiento inductivo es pues razonamiento por probabilidad lógica.

Segunda: la discusión del ejemplo que nos ha llevado a (27) no puede servir para formular el concepto lógico de probabilidad, ni, por tanto, el concepto cuantitativo de inducción; pues ha sido precisamente un caso empírico de frecuencia relativa (de la propiedad de pertenecer a la clase β entre los miembros de la clase α). El partir de una discusión sobre clases es lo más fácil intuitivamente. Pero al hablar de clases, las frecuencias relativas — el 0,7 del ejemplo — son datos empíricos, no relaciones formales, y no pueden servir para definir la función c como función lógica, al modo, por ejemplo, como se define una función veritativa. Por eso el primer problema que debe resolver una lógica inductiva es el de encontrar relaciones medibles que no sean empíricas.

104. Descripciones de estado y ámbitos semánticos. — Acabamos de ver que las clases, aunque son entidades muy cómodamente medibles (por el número de sus miembros, por ejemplo), no son adecuadas para construir funciones lógicas de medición, sino sólo para construir funciones empíricas: pues la pertenencia de un individuo a una clase es un asunto empírico, como el de todo discurso sobre clases que tenga interés para las ciencias reales. ¿Qué entidades pueden, entonces, ser objeto de medición, sin ser ellas mismas empíricas?

El punto de vista lingüístico en lógica ofrece una posibilidad de resolver este problema, del mismo modo que también es útil para esclarecer la deducción. Dado un determinado lenguaje L , imaginemos una lista de todas sus constantes individuales y predicativas. Con esos sujetos y esos predicados pueden construirse todos los enunciados del lenguaje L (que será, naturalmente, un lenguaje con un número finito de constantes). Por de pronto podemos construir todos los enunciados atómicos y sus negaciones. Llamemos enunciados básicos a los enunciados atómicos y a sus negaciones. Imaginemos ahora una conjunción de enunciados básicos en el que estén representados todos los individuos del universo del discurso de L , es decir,

en el que aparezcan todas las constantes de L , y solamente una vez cada una. Una conjunción de esa naturaleza se llama en la terminología de Carnap una '*descripción de estado* para el lenguaje L ', pues describe una situación, un estado del universo del discurso de L . Por ejemplo, sea L un lenguaje que tiene tres predicados sólo, 'rojo', 'verde' y 'ámbar' (' R ', ' V ', ' A ') y sólo tres constantes individuales, ' a ', ' b ' y ' c ' (nombres, por ejemplo, de tres semáforos de un cruce complicado). El universo del discurso de L es un "mundo" compuesto por los tres individuos a , b y c ; sus estados son distribuciones entre esos tres individuos de las propiedades rojo, verde y ámbar. Los siguientes enunciados son algunas descripciones de estado para L (no todas las descripciones de estado para L):

- (28) $Ra \wedge Vb \wedge Ac$ (Z_I)
 (29) $Ra \wedge Rb \wedge Ac$ (Z_{II})
 (30) $Ra \wedge Rb \wedge Vc$ (Z_{III})
 (31) $Aa \wedge Vb \wedge Rc$ (Z_{IV})

etcétera.

Las descripciones de estado para un lenguaje L se excluyen unas a otras:

- (32) Para todo lenguaje L , si Z_n y Z_m son dos descripciones de estado distintas, vale

$$\sim [Z_n \wedge Z_m].$$

Y la disyunción de todas las descripciones de estado para un lenguaje L es exhaustiva, o sea, agota todos los modos posibles de presentarse el universo del discurso de L :

- (33) Para cualquier lenguaje L , si $Z_1, \dots, Z_n$ son todas sus descripciones de estado, entonces la disyunción

$$Z_1 \vee Z_2 \vee \dots \vee Z_n$$

es formalmente verdadera.

La letra ' Z ' usada por Carnap (en escritura gótica) para designar las descripciones de estado es la inicial de '*Zustand*', 'estado'.

Las descripciones de estado son enunciados que describen situaciones semánticas. Con su ayuda se define el ámbito de un enunciado cualquiera del lenguaje correspondiente. *Ámbito* de un enunciado X de un lenguaje L es la clase de las descripciones de estado que describen situaciones semánticas en las cuales X es verdadero. Por ejemplo, el enunciado de L

- (34) $Ra \wedge Ac$

tiene (suponiendo que Z_I - Z_{IV} fueran todas las descripciones de estado para L , cosa que no son), el ámbito $\{Z_I, Z_{II}\}$.

Decir que (34) tiene el ámbito $\{Z_I, Z_{II}\}$ es decir que vale cuando el universo del discurso se encuentra en la situación indicada por la descripción de estado Z_I , y también cuando se encuentra en la situación indicada por la descripción de estado Z_{II} . Por tanto, que vale cuando valen Z_I o Z_{II} , las cuales no valdrán a la vez porque las descripciones de estado son mutuamente excluyentes. Así pues, el ámbito de un enunciado, que hemos definido como una clase de descripciones de estado, debe entenderse como la disyunción (no la conjunción) de dichas descripciones de estado. Al hablar de ámbitos pueden por tanto usarse indiferentemente el lenguaje de clases y el de la lógica de enunciados. Así podemos escribir

$$\text{ámbito de } X = \text{ámbito de } Y,$$

y estaremos entonces concibiendo los ámbitos como clases de descripciones de estado; o

$$\text{ámbito de } X \leftrightarrow \text{ámbito de } Y,$$

en cuyo caso estaremos concibiendo los ámbitos como disyunciones de descripciones de estado. Análogamente, decir que el ámbito m está incluido en el ámbito n es lo mismo que decir que la disyunción formada por las descripciones de estado que componen el primero es parte de la disyunción de descripciones de estado en que consiste el segundo. A partir de ahora, en vez de decir 'el enunciado X vale cuando la situación del universo del discurso es la indicada por la descripción de estado Z_n ' diremos brevemente: 'el enunciado X vale para la descripción de estado Z_n '.

Los ámbitos son susceptibles de medición; por ejemplo, tiene un sentido bastante claro el atribuir a un ámbito con más descripciones de estado que otro un valor numérico mayor que el atribuido al segundo. Según esto diremos, atendiendo al anterior ejemplo, que el enunciado de L

- (35) $Ra \wedge Vb$

(si la tabla Z_I - Z_{IV} fuera completa) tendría un ámbito — $\{Z_I\}$ — menor que el de (34).

Pero a pesar de ser medibles, los ámbitos y las descripciones de estado no son entidades empíricas, sino lógicas (semánticas): no son afirmaciones de que el mundo sea de tal o cual manera, sino descripciones de estados posibles del universo del discurso de un lenguaje dado: para ellos no se plantea la cuestión de si son verdaderos o falsos, sino sólo la de si son (como enunciados) consistentes o no ('posibles'). A partir de estas entidades medibles, pero no empíricas, que son los ámbitos, pueden construirse funciones lógicas de confirmación de un enunciado por otro.

La realización de este programa para lenguajes de verdadero interés, es decir, que cuenten con infinitos símbolos individuales y con predicados poliádicos para expresar relaciones, es difícil y no puede aún considerarse

resuelta. Aquí vamos a estudiar un caso de lenguaje muy sencillo, sólo para poder apreciar las líneas generales del programa de lógica inductiva de Carnap y el tipo de análisis de la inducción que fundamenta dicho programa.

Antes de enfrentarnos con nuestro ejemplo conviene insistir en que el ámbito de un enunciado se define lógico-semánticamente, no empíricamente. No nos hace falta saber empíricamente que Z_I o Z_{II} son verdaderas o no — si describen o no el real estado del universo del discurso de L — para poder afirmar que $\{Z_I, Z_{II}\}$, o $Z_I \vee Z_{II}$ es el ámbito de (34) en L (si la tabla Z_I - Z_{IV} es completa). Por tanto, toda función que se base en mediciones de ámbitos será de naturaleza lógica, no empírica. Pasemos ahora al ejemplo.

Sea un lenguaje muy sencillo, l , que no consta más que de tres símbolos de enunciado, ' p ', ' q ', ' r ', y los funtores veritativos ' $\sim$ ', ' $\vee$ ', ' $\wedge$ '. Ese lenguaje es muy pobre y simplifica uno aparentemente menos pobre: un lenguaje con tres constantes individuales, ' a ', ' b ', ' c ', y un solo predicado, ' P '. ' p ' traduce y simplifica cómodamente ' Pa ', por ejemplo. El resto de la traducción puede establecerse así:

$$\begin{aligned} \sim Pa &\leftrightarrow \sim p \\ Pb &\leftrightarrow q \\ \sim Pb &\leftrightarrow \sim q \\ Pc &\leftrightarrow r \\ \sim Pc &\leftrightarrow \sim r. \end{aligned}$$

Un lenguaje como l puede servir, por ejemplo, para describir situaciones en un tablero de control de una estación del metro con tres vías y un único sentido de circulación de trenes por cada vía. Cada una de las vías tiene una señalización para indicar 'vía libre' o 'vía cerrada' a los conductores de tren. En el tablero de control hay tres bombillas, cada una correspondiente a una vía. Una bombilla está encendida cuando su vía correspondiente está libre, y apagada cuando su vía correspondiente está cerrada.

Las descripciones de estado, Z_I - Z_8 para l son

- (36) $p \wedge q \wedge r$ (Z_I) (las tres bombillas encendidas: $Pa \wedge Pb \wedge Pc$).
- (37) $\sim p \wedge q \wedge r$ (Z_2)
- (38) $p \wedge \sim q \wedge r$ (Z_3)
- (39) $\sim p \wedge \sim q \wedge r$ (Z_4)
- (40) $p \wedge q \wedge \sim r$ (Z_5)
- (41) $\sim p \wedge q \wedge \sim r$ (Z_6)
- (42) $p \wedge \sim q \wedge \sim r$ (Z_7)

- (43) $\sim p \wedge \sim q \wedge \sim r$ (Z_8) (las tres bombillas apagadas).

El enunciado

- (44) $p \wedge q$,

por ejemplo, tiene en l el ámbito

$$R = \{Z_I, Z_5\}.$$

Escribiremos:

$$R(43) = \{Z_I, Z_8\}, \text{ o también}$$

$$R(43) \leftrightarrow Z_I \vee Z_8.$$

' R ' es la inicial de 'range' (inglés) y 'Rang' (alemán).

Ahora debemos fijarnos en un sencillo metateorema que es fundamental para la construcción de las funciones inductivas de confirmación: para todo enunciado, X , de un lenguaje L hay un ámbito, $R(X)$, tal que X equivale formalmente a $R(X)$ en L (concebido como disyunción). O sea:

- (45) $(X) \exists R [X \leftrightarrow R(X)].$

En nuestro lenguaje l (45) es de comprobación muy trivial: dado un enunciado cualquiera de l , puede establecerse sencillamente en una tabla veritativa la equivalencia del mismo con la disyunción de descripciones de estado que es su ámbito. Sea, por ejemplo, el enunciado de l

- (46) $\sim p \vee q$, o $p \rightarrow q$, si ya se ha introducido ' $\rightarrow$ ' por definición.

Su ámbito es $Z_I \vee Z_2 \vee Z_4 \vee Z_5 \vee Z_6 \vee Z_8$.

Con una tabla veritativa puede comprobarse la equivalencia entre (46) y esa disyunción.

105. Funciones lógicas de medición. — Las funciones m de Carnap, o funciones de medición de ámbitos, se basan en la atribución de valores numéricos a las descripciones de estado.

Para que las funciones de medición resulten útiles en el análisis de los lenguajes de la ciencia, la atribución de valores numéricos a las descripciones de estado tiene que hacerse con bastante complicación y muchas precauciones que aquí pasaremos por alto. Por eso la función de medición m (' m ' latina) que aquí usaremos no debe confundirse con la función regular de medición de Carnap.

Una función de medición, m , para nuestro lenguaje l será una función de un solo argumento. Sus argumentos son los enunciados de l . Sus valores, por adecuarse con el uso corriente en el cálculo de probabilidades, serán números racionales de 0 a 1, ambos comprendidos.

La definición de m exige atribuciones de valores a las descripciones de estado. Estableceremos para nuestro ejemplo la siguiente regla de definición de m :

- (47) (1) Para toda descripción de estado, Z_n , de l ,
 $m(Z_n) = 0,125 (= 1/8)$.
 (2) Para toda disyunción de descripciones de estado distintas en l , $Z_i \vee \dots \vee Z_j$ ($1 \leq i, j \leq 8$),
 $m(Z_i \vee \dots \vee Z_j) = m(Z_i) + \dots + m(Z_j)$.

La atribución del mismo valor a toda descripción de estado, como aquí hacemos, es precisamente la simplificación que no puede hacerse si de verdad se quiere construir un sistema de lógica inductiva capaz de recoger la situación efectiva en los lenguajes de las ciencias reales.

La anterior definición de m nos permite calcular valores para ámbitos de l . Por ejemplo, un ámbito compuesto por dos descripciones de estado tendrá por m el valor 0,25 o $1/4$.

Ahora bien: según el teorema (45), todo enunciado equivale a su ámbito. Consiguientemente, la función m atribuye valores a los enunciados de l . He aquí alguna otra consecuencia de la caracterización de la función m :

- (48) Si $X \leftrightarrow Y$, entonces: $m(X) = m(Y)$.

En efecto: el ámbito de X equivale a X , y el de Y a Y . Por ser $X \leftrightarrow Y$ se tendrá: $R(X) \leftrightarrow R(Y)$. De aquí: $m(X) = m(Y)$.

- (49) Si $X \rightarrow Y$, entonces: $m(X) \leq m(Y)$.

Siempre que vale X , vale Y . Consiguientemente, el ámbito de X tiene que estar contenido, o equivaler al menos, al ámbito de Y .

- (50) Si X es formalmente falso, entonces: $m(X) = 0$.

Pues el ámbito de X no contiene ninguna descripción de estado.

- (51) Si X es formalmente verdadero, entonces: $m(X) = 1$.

Pues el ámbito de X contiene todas las descripciones de estado para l .

- (52) Si X no es formalmente verdadero ni formalmente falso, entonces: $0 < m(X) < 1$.

Por (50) y (51).

- (53) $m(\sim X) = 1 - m(X)$.

Por (2) de (47), $m(X \vee \sim X) = m(X) + m(\sim X)$. Por (51), $m(X \vee \sim X) = 1$. De aquí (53).

- (54) $m(X \vee Y) = m(X) + m(Y) - m(X \wedge Y)$.

Hay que restar $m(X \wedge Y)$ porque si no se suman dos veces los valores de las descripciones de estado que pertenecen a la vez al ámbito de X y al ámbito de Y .

- (55) Si $\sim [X \wedge Y]$, entonces: $m(X \vee Y) = m(X) + m(Y)$.

Por (54), (55) es siempre el caso si X e Y son descripciones de estado.

- (56) $m(X \wedge Y) = m(X) + m(Y) - m(X \vee Y)$.

Hay que restar $m(X \vee Y)$ porque si no se sumarían también los valores de las descripciones de estado que pertenecen sólo al ámbito de X (pero no al de $X \wedge Y$) y de las que pertenecen sólo al ámbito de Y (pero no al de $X \wedge Y$).

Utilizando la función de medición m se define la función de confirmación c . Antes de considerar esta función c puede ser útil parafrasear intuitivamente la definición de la función m , describir qué es lo medido por m . Lo medido por m es el ámbito de los enunciados de un determinado lenguaje, en nuestro ejemplo, del lenguaje l . El ámbito de un enunciado X es, por así decirlo, su alcance semántico, la probabilidad de ser verdadero, de quedar satisfecho o cumplido en el universo del discurso del lenguaje de que se trate: su "frecuencia relativa", podría decirse, pero no en el mundo real, sino en el mundo semántico de la consistencia ("posibilidad"). La marcha de la lógica inductiva de Carnap puede caracterizarse intuitivamente diciendo que en ella se construyen, en base a funciones de medición, m , que sirven para cuantificar probabilidades de ser satisfechos los enunciados de un lenguaje por referencia al universo del discurso de ese lenguaje, otras funciones, c , funciones de confirmación, que sirven para cuantificar la probabilidad que tiene un enunciado de un lenguaje de ser satisfecho no ya respecto del universo del discurso en general, sino respecto de la situación de ese universo del discurso descrita por un enunciado previo, llamado información: esta probabilidad es el grado de confirmación del primer enunciado por el segundo.

106. Funciones de confirmación. — Con ayuda de la función de medición, m , puede definirse una función de confirmación, c , función de dos argumentos (los cuales son enunciados del lenguaje de que se trate, l en nuestro caso) y cuyos valores son números racionales de 0 a 1, ambos comprendidos.

Por el hecho de haber escogido, de entre las varias posibles funciones de medición, una función m muy sencilla que no responde a las necesidades de una teoría capaz de recoger la real argumentación inductiva de las ciencias, tampoco la función c que vamos a definir será capaz de ello. Será, como m , útil sólo para ilustrar la naturaleza de una teoría lógica de la inducción.

Definición:

(57) Para todo par de enunciados de I , X e Y ,

$$c(X, Y) =_{df} \frac{m(X \wedge Y)}{m(Y)}.$$

En esta notación — que es la de Carnap — X es la hipótesis formulada en base a la información Y . $c(X, Y)$ tiene como valor un número racional de 0 a 1 que cuantifica la confirmación que Y presta a X .

Para hacer plausible la definición (57) nos fijaremos en algunas consecuencias de la misma que armonizan bien con la idea intuitiva de confirmación de un enunciado por otro. La definición quiere decir que el grado en que Y confirma a X se entiende como la razón entre la probabilidad de que X e Y valgan juntos en I y la probabilidad de que valga simplemente Y (en I). Dicho menos imaginativamente: como la razón del valor del ámbito de $X \wedge Y$ al valor del ámbito de Y . De aquí se desprende:

1.º) En el caso de que siempre que valga Y valga X , por darse

$$\begin{aligned} Y &\leftrightarrow X, \\ \text{o } Y &\rightarrow X, \end{aligned}$$

se tendrá:

$$c(X, Y) = 1.$$

Es que entonces el paso de Y a X es deductivo: se tiene entre Y y X la relación deductiva de implicación plena. Y es muy natural dar a esta relación el valor máximo de confirmación, o sea, 1.

2.º) Si X e Y nunca valen juntos, o, lo que es lo mismo, si nunca vale $X \wedge Y$, su ámbito es vacío y $m(X \wedge Y) = 0$. Si al mismo tiempo $m(Y) > 0$, se tendrá:

$$c(X, Y) = 0.$$

Es que en este caso X e Y son incompatibles, mutuamente excluyentes. Y también resulta natural tener en este caso que la confirmación de X por Y vale 0.

3.º) Por último, es claro que $X \wedge Y$ no puede tener un ámbito de valor superior al del ámbito de Y , sino, a lo sumo, igual. Por tanto, no ocurrirá para ningún X y ningún Y que $c(X, Y) > 1$. Y también es natural que no haya confirmación más fuerte que la implicación plena.

Esas tres observaciones pueden bastarnos para admitir la plausibilidad de la definición (57). Hay, sin embargo, una situación en la cual $c(X, Y)$ queda sin definir: es el caso de que Y sea formalmente falso, pues entonces $m(Y) = 0$. Para definir, como suele hacerse, mediante una convención la función para este caso, se presentan dos posibilidades intuitivamente plausibles: usando ' F ' en vez de ' Y '; para recordar la situación (que Y es formalmente falso), podemos optar por una de las dos convenciones siguientes:

$$c(X, F) = 0,$$

$$(58) \quad c(X, F) = 1.$$

(58) parece más conveniente, porque alude al hecho de que entre F y cualquier enunciado media una relación práctica de deductividad. Por lógica deductiva conocemos la regla de que "de lo falso se sigue cualquier cosa".

Análogamente a como hicimos con la función m , podemos también representarnos intuitiva o imaginativamente qué es lo que se mide con c . Un valor de c es un cociente de valores de m . Por tanto, un cociente de probabilidades respecto de un universo del discurso, determinado por un lenguaje que tiene que estar precisamente dado en cada caso. El dividendo de ese cociente es el valor del ámbito de la conjunción de la hipótesis y la información, o sea, la probabilidad de esa conjunción respecto del universo del discurso supuesto. El divisor, aquello respecto de lo cual se estima aquella probabilidad, es a su vez el valor del ámbito de la información, esto es, la probabilidad de la información respecto del universo del discurso. Por tanto, la relación de confirmación puede imaginarse como la probabilidad de [la conjunción de la información y] la hipótesis en aquella situación del universo del discurso en la cual queda satisfecha la información.

107. Idea de una lógica inductiva. — La cuantificación de la relación de confirmación por la función c puede ser objeto de un estudio formal, exactamente igual que la relación de implicación en lógica deductiva. No se hace, en efecto, recurso a nada empírico cuando se habla de confirmación en el sentido de Carnap, sino sólo a relaciones entre valores de ámbitos de enunciados. He aquí algunos ejemplos de teoremas de una teoría de la función c , que es una teoría formal de la relación de confirmación y, por tanto, de la inducción. Los siguientes teoremas valen para cualesquiera enunciados, X , Y , W de I .

$$(59) \quad 0 \leq c(X, Y) \leq 1.$$

$$(60) \quad \text{Si } Y \rightarrow X, \text{ entonces: } c(X, Y) = 1.$$

$$(61) \quad \text{Si } X \text{ es formalmente verdadero, entonces, para todo } Y: \\ c(X, Y) = 1 \text{ (por (60)).}$$

(62) Si $X \leftrightarrow W$, entonces: $c(X, Y) = c(W, Y)$.

Los anteriores teoremas se desprenden de las definiciones de m y de c .

$$(63) \quad c(X \vee W, Y) = c(X, Y) + c(W, Y) - c(X \wedge W, Y).$$

Demostración:

$$\begin{aligned} c(X \vee W, Y) &= \frac{m([X \vee W] \wedge Y)}{m(Y)}, \text{ por def. (57).} \\ &= \frac{m([X \wedge Y] \vee [W \wedge Y])}{m(Y)}, \text{ por lógica deductiva.} \\ &= \frac{m(X \wedge Y) + m(W \wedge Y) - m([X \wedge Y] \wedge [W \wedge Y])}{m(Y)}, \text{ por (54).} \\ &= \frac{m(X \wedge Y) + m(W \wedge Y) - m(X \wedge W \wedge Y)}{m(Y)}, \text{ por lóg. deduc-} \\ &\quad \text{tiva.} \\ &= \frac{m(X \wedge Y)}{m(Y)} + \frac{m(W \wedge Y)}{m(Y)} - \frac{m(X \wedge W \wedge Y)}{m(Y)} \\ &= c(X, Y) + c(W, Y) - c(X \wedge W, Y), \text{ por def. (57).} \end{aligned}$$

$$(64) \quad c(X \wedge W, Y) = c(X, Y) \times c(W, X \wedge Y).$$

Demostración:

$$\begin{aligned} c(X \wedge W, Y) &= \frac{m(X \wedge W \wedge Y)}{m(Y)}, \text{ por def. (57).} \\ &= \frac{m(X \wedge W \wedge Y) \times m(X \wedge Y)}{m(X \wedge Y) \times m(Y)} \quad \text{multiplicación de dividendo} \\ &\quad \text{y divisor por } m(X \wedge Y). \\ &= \frac{m(X \wedge Y)}{m(Y)} \times \frac{m(X \wedge W \wedge Y)}{m(X \wedge Y)} \\ &= c(X, Y) \times c(W, X \wedge Y), \text{ por def. (57).} \end{aligned}$$

(63) y (64) corresponden a los teoremas generales de adición y multiplicación (respectivamente) de la teoría de probabilidades. La teoría de la relación de confirmación, o lógica de la inducción, suministra el fundamento lógico de la teoría de la probabilidad. Tal es, precisamente, el sentido de la investigación de Carnap.

Antes de terminar con unas últimas observaciones el tema de la lógica inductiva, apliquemos a un par de enunciados de nuestro lenguaje I lo visto

hasta aquí. Queremos averiguar el valor de la confirmación de la hipótesis ' $q \rightarrow r$ ' por la información ' $p \rightarrow q \wedge r$ '.

Información, Y : $p \rightarrow q \wedge r$, o sea, $\sim p \vee [q \wedge r]$.

Hipótesis, X : $q \rightarrow r$, o sea, $\sim q \vee r$.

$$\begin{aligned} c(X, Y) &= \frac{m([\sim q \vee r] \wedge [\sim p \vee [q \wedge r]])}{m(\sim p \vee [q \wedge r])} \\ &= \frac{m(\sim q \vee r) + m(\sim p \vee [q \wedge r]) - m([\sim q \vee r] \vee [\sim p \vee [q \wedge r]])}{m(\sim p \vee [q \wedge r])} \\ &= \frac{6/8 + 5/8 - 7/8}{5/8} = 4/5 = 0,8. \end{aligned}$$

El fundamento de la afirmación de esos 4/5 o 0,8 en I es que ' $q \rightarrow r$ ' queda satisfecho en 4 de las cinco descripciones de estado en que queda satisfecho ' $p \rightarrow q \wedge r$ '.

Recapitularemos lo principal de las comparaciones que hemos ido haciendo entre lógica deductiva y lógica inductiva:

Ante todo, las dos son teorías formales basadas en el concepto semántico de ámbito. Estudian relaciones entre ámbitos de enunciados: la de implicación plena, o deducibilidad (confirmación igual a 1) y la de implicación parcial, o relación inductiva (confirmación menor que 1). Las expresiones 'lógica deductiva' y 'lógica inductiva' son en rigor impropias; debería más bien decirse: 'análisis y teoría formal de la deducción' y 'análisis y teoría formal de la inducción'. Esta naturaleza formal es lo principal que tienen en común las dos teorías.

Sólo en segundo lugar debe apreciarse, desde un punto de vista lógico, su principal diferencia, que consiste en lo siguiente: el análisis y la teoría formal de la deducción justifican en la práctica del razonamiento la afirmación de enunciados plena o deductivamente implicados por otros enunciados que se admitan como verdaderos. El paso deductivo de $Y \rightarrow X$ e Y a X está formalmente justificado. En cambio, el paso práctico inductivo de $c(X, Y) = n$ e Y a X "con la certeza n " no lo está ni tiene sentido. Pues n no es un valor propio o una propiedad intrínseca de X , sino que su uso sólo tiene sentido respecto de la información Y . Por esta razón la lógica de la inducción no puede construir cálculos o algoritmos que suministren conclusiones que sean sólo hipótesis, sin alusión a ninguna información, mientras que un tipo de algoritmo análogo existe para la deducción. Esto constituye una diferencia calculística entre la deducción y la inducción (cfr. 20). Pero eso no excluye toda formalización o algoritmización de la teoría de la inducción: significa simplemente que los resultados que arroje un algoritmo

en la teoría de la inducción se basarán siempre en enunciados atómicos de la forma ' $c(X, Y) = n$ ', y no de la forma X , o Y .

Esta diferencia, aunque importante en la práctica, es secundaria desde un punto de vista de teoría lógica. Se presenta en la aplicación de la teoría lógica, o sea, en la doctrina del método. Desde el punto de vista teórico, lo esencial es que el objeto de la teoría formal de la inducción es una relación lógica (semántica) entre enunciados, igual que el objeto de la teoría formal de la deducción.

La naturaleza lógica de la teoría de la inducción no queda, naturalmente, perturbada por la cuantificación del concepto de confirmación. Pues los teoremas de la teoría de la inducción — por ejemplo, (63) o (64) — no tienen como valor ningún número racional entre 0 y 1, como las relaciones de confirmación de que hablan esos teoremas, sino que son formalmente verdaderos. Y un enunciado que se pretenda afirmar como teorema de la teoría de la inducción sin serlo, será formalmente no-verdadero. Los teoremas de la lógica de la inducción hablan de valores de confirmación; pero ellos mismos, si realmente son teoremas, son formalmente verdaderos en base a definiciones y a teoremas de la lógica deductiva. También puede decirse que, como éstos, tienen la confirmación 1.

La teoría lógica de la inducción no ha llegado aún, sin embargo, a alcanzar la solidez y el consenso universal con que cuenta la teoría lógica de la deducción. Incluso la construcción más adelantada, el sistema de Carnap, está aún en vías de perfeccionamiento y no es compartido por todos los autores, aunque son ya una minoría los que siguen considerando imposible o superfluo un tratamiento formal de la inducción.

Debe indicarse, por último, que, con posterioridad a la primera edición de su estudio sobre la lógica inductiva, Carnap ha preferido dejar de hablar de 'confirmación', para expresarse en términos más matemáticos (tomados de la teoría del cálculo de probabilidades).

RELACIÓN DE TEOREMAS Y REGLAS AUXILIARES

que se encuentran aludidos en lugares distintos de los de su introducción

La cifra en negritas que sigue a la sigla de referencia del teorema o la regla indica el número del epígrafe bajo el cual se introdujo el teorema o la regla.

Teoremas

Del (TE1) al (TE15)	: 54.
Del (TP16) al (TP29)	: 55.
(TE30) y (TE31)	: 57.

Reglas auxiliares

RA γ_8 : 46.	RA β_1 : 59.
RA1 : 54.	RA α_1 : 59.
RA2 : 54.	RA $\leftrightarrow_1$: 59.
	RA α_2 : 59.

BIBLIOGRAFIA

Para una bibliografía general debe verse el *Journal of Symbolic Logic* (órgano de la Association for Symbolic Logic, 190 Hope Street, Providence 6, Rhode Island, U.S.A.), empezando por el vol. I, 1936, págs. 121-128, y el vol III, 1938, págs. 178-212.

A continuación se da una selección de lecturas recomendadas, empezando (I) por obras disponibles en castellano. No se recogen artículos.

I

Manuales de introducción:

FERRATER, J., y LEBLANC, H., *Lógica matemática*, 2.^a ed., México 1962.

QUINE, W. v. O., *El sentido de la nueva lógica*, Buenos Aires 1958.

TARSKI, A., *Introducción a la lógica y a la metodología de las ciencias deductivas*, Buenos Aires 1951.

Manuales sistemáticos:

HILBERT, D., y ACKERMANN, W., *Elementos de lógica teórica*, Madrid 1962.

QUINE, W. v. O., *Los métodos de la lógica*, Barcelona 1962.

Ensayos y monografías:

MARTÍN, R. M., *Verdad y denotación*, Madrid 1962. (Semántica.)

MORRIS, C. W., *Fundamentos de la teoría de los signos*, México 1958.

MUÑOZ, V., *De la axiomática a los sistemas formales*. Madrid 1961. (Histórico. Trata también la idea de sistema formal de Curry.)

NAGEL, E., y NEWMAN, J. R., *La prueba de Gödel*, México 1959.

POPPER, K. R., *La lógica de la investigación científica*, Madrid 1962. (Tema de la inducción-probabilidad.)

QUINE, W. v. O., *Desde un punto de vista lógico*, Barcelona 1962. (Serie de ensayos de tema lógico-filosófico.)

SÁNCHEZ MAZAS, M., *Fundamentos matemáticos de la lógica formal*, Caracas 1963. (Desarrolla el punto de vista intensional.)

II

Manuales:

- BETH, E. W., *Formal methods. An introduction to symbolic logic and to the study of effective operations in arithmetic and logic*, Dordrecht (Holanda) 1962.
- BOCHENSKI, I. M., y MENNE, A., *Grundriss der Logistik*, Paderborn 1954.
- CARNAP, R., *Einführung in die symbolische Logik*, Viena 1954.
- DOPP, J., *Leçons de logique formelle*, 3 vols., Lovaina 1949-1950. (Muy útil para el lector educado en la lógica tradicional.)
- LEBLANC, H., *Formal logic*, Oxford 1952.
- PRIOR, A. N., *Formal logic*, Oxford 1955.

Tratados:

- CASARI, E., *Lineamenti di logica matematica*, Milano 1960.
- CHURCH, A., *Introduction to mathematical logic*, I, Princeton (NJ) 1956).
- CURRY, H. B., y FEYS, R., *Combinatory Logic*, I, Amsterdam 1958.
- HAO WANG, *A survey of mathematical logic*, Pekín-Amsterdam 1963. (No es propiamente un tratado, sino una colección de estudios; pero éstos cubren una temática muy amplia y de un modo bastante sistemático.)
- KLEENE, S. C., *Introduction to metamathematics*, Amsterdam, ed. de 1957. (Tampoco propiamente un tratado de lógica. Pero obra básica.)
- LORENZEN, P., *Einführung in die operative Logik und Mathematik*, Berlin 1955.
- SCHOLZ, H., y HASENJAECGER, W., *Grundzüge der mathematischen Logik*, Berlin 1961.

Semiótica:

- CARNAP, R., *Introduction to semantics and Formalization of logic*, Cambridge (Mass.) 1958 (ed. con las dos obras en un vol.).
- CARNAP, R., *Meaning and necessity*, Chicago 1958 (reimpr. de la 2.ª ed.).
- HERMES, H., *Semiotik*, Leipzig 1938.
- KLAUS, G., *Semiotik*, Berlín 1963.
- MARTÍN, R. M., *Toward a systematic pragmatics*, Amsterdam 1959.
- TARSKI, A., *Logic, Semantics, Metamathematics. Papers from 1923 to 1938*, Oxford 1956.

Lógica y metodología:

- BRAITHWAITE, R. B., *Scientific explanation*, Cambridge 1959.
- CARNAP, R., *Logical Foundations of Probability*, Chicago 1950.
- KNEALE, W., *Probability and induction*, Oxford, reimpr. de 1963.
- REICHENBACH, H., *The theory of Probability*, Berkeley-Los Angeles 1949.
- ROBINSON, R., *Definition*, Oxford 1950.

Textos clásicos de la lógica simbólica que pueden encontrarse en reediciones recientes:

- BOOLE, G., *An investigation of the laws of thought, in which are founded the mathematical theories of logic and probability*, New York 1951.
- BOOLE, G., *The mathematical analysis of logic, being an essay towards a calculus of deductive reasoning*, Oxford-New York 1948.
- BOOLE, G., *Collected logical works*, 3.ª ed., New York 1951.
- FREGE, G., *Funktion, Begriff, Bedeutung. Fünf logische Studien*, Göttingen 1962. (Los cinco estudios son: "Funktion und Begriff", "Sinn und Bedeutung", "Begriff und Gegenstand", "Was ist eine Funktion?", "Über die wissenschaftliche Berechtigung einer Begriffsschrift".)
- FREGE, G., *The foundations of arithmetic*, Oxford 1950.
- FREGE, G., *Translation of the philosophical writings*, Oxford-New York 1952.
- GENTZEN, G., *Recherches sur la déduction logique*, Paris 1955.
- RUSSELL, B., y WHITEHEAD, A. N., *Principia Mathematica*, Cambridge 1950.

Historia de la lógica

- (Los títulos no están ordenados por el orden alfabético de nombres de autores, como los anteriores)
- FRANKE, L., *Geschichte der Logik im Abendlande*, 4.ª ed., Leipzig 1927, reimpr. 1950. (Antiguado y con errores históricos, sigue siendo útil por la abundante documentación.)
- SCHOLZ, H., *Geschichte der Logik*, Berlin 1931. (Al revés que el anterior, es muy breve y no trae documentos históricos; pero su moderno punto de vista ha influido en la mayoría de los estudios históricos contemporáneos.)
- BOCHENSKI, I. M., *Formale Logik*, Friburgo 1956. (Muy amplia y detallada; de lectura incómoda; muchos documentos históricos, pero todos traducidos al alemán.)
- KNEALE, W. y M., *The development of logic*, Oxford 1962. (Concentrada en lo esencial; lectura muy agradable.)
- BOCHENSKI, I. M., *Ancient formal logic*, Amsterdam 1951. (Exposición técnica, con notación simbólica.)
- LUKASIEWICZ, J., *Aristotle's syllogistic from the standpoint of modern formal logic*, Oxford 1951. (Es el estudio formal más importante de la silogística aristotélica realizado con método moderno.)
- MATES, B., *Stoic logic*, Berkeley-Los Angeles 1953. (El estudio básico sobre el tema.)
- BOEHNER, PH., *Medieval logic. An outline of its development from 1250-c. 1400*, Manchester 1952.
- MOODY, E. A., *Truth and consequence in medieval logic*, Amsterdam 1953. (Más monográfico que el anterior. Importante para cuestiones de semántica en la Edad Media. Breve.)
- COUTURAT, L., *La logique de Leibniz d'après des documents inédits*, Paris 1901. (Siguiendo siendo un libro importante a pesar de los errores que posteriormente se han indicado en él.)

REGISTRO ALFABÉTICO

- Abelardo: 58.
- abstracción: 13 ss., 28 ss.
 - a. básica: 16, 35, 109.
 - a. básica de la lógica formal: 17 ss.
 - a. científica: 15 ss.
 - a. de Buridan: 21.
 - a. de clases: 223.
 - a. funcional: 67.
 - falacia de a.: 17, 22.
- abstractor: 224.
- abstractos: 233.
- Ackermann: 87, 110, 118, 178.
- álgebra de clases: 216 ss.
- algoritmos: 44 s., 79.
- Alkarismi, Mohamed ben Musa: 54.
- ámbito semántico: 296.
- análisis: 55.
- análisis y síntesis en lógica: 53 s., 114 ss., 119 ss.
- analítica: 119.
- apofántico: 18, 57.
- aporias: v. paradojas.
- a priori: 26, 27.
- Aristóteles: 18, 25, 34, 37, 50, 54, 55, 57, 61, 67, 68, 88, 102.
 - concepto aristotélico de verdad: 25.
- Arnauld y Nicole: 30, 37.
- axiomas: 103.

- Bacon: 31.
- base primitiva: 103.
- Bernays: 110.
- Beth: 51.
- bicondicional: 76.
- Boole: 38.
- Buridán: 21.

- cálculo: 45 ss.
- campo de una relación: 242.
- Cantor: 39 s., 231.
- Carnap: 35, 50, 51, 87, 108, 215, 241 s., 256, 286, 287, 300 ss., 303, 306.
- catagorema, catagoremático: 21.
- categorías: 55 ss.
- Celarent: 140 s.
- Cervantes: 43, 54.
- clase nula: 227.
- clase universal: 227.
- clases de equivalencia: 241 ss.
- comillas simples: 14 s.

- compatibilidad: v. consistencia.
- complemento de una clase: 217, 224.
- complemento de una relación: 237.
- completud: 47, 105, 182 ss.
- comprensión: 39.
- concatenación: 49.
- conceptualismo: 58.
- condicional: 74.
- Condillac: 44.
- conectivas: 65.
- conexividad: 247.
- confirmación: 287.
- conjunción: 77.
- conjunto: 39.
- conocer y conocimiento: 17.
- conocimiento y lenguaje: 17, 19.
- consistencia: 24, 47, 105, 178 ss.
- constantes: 58 s.
 - c. lógicas: 59.
- contradicciones formales: 50.
- convención e intuición: 19.
- correlatora: 254 s.
- Countur: 54.
- "crisis de fundamentos": 32 ss.
- cuantificadores: 61 s.
 - c. existencial: 61 ss.
 - alcance existencial de los c.: 62, 141.
 - c. universal: 61 ss.
- Curry: 45.

- Church: 56, 58, 173, 205.

- Darapti: 141.
- decidibilidad: 47, 105, 202 ss.
- deducción: 283 ss.
- deducción natural: 120 ss.
- definición: 46, 131, 275 ss.
 - d. creadora: 280 s.
 - d. implícita: 108 s.
 - d. por abstracción: 279 s.
 - d. reales y nominales: 275 s.
 - d. recursivas: 78, 278.
 - d. semánticas: 278.
 - d. sintácticas: 277 s.
 - leyes de la d.: 277.
- "de lo falso se sigue cualquier cosa": 131.
- demonstración: 104.
 - d. en sistema axiomático: 114 ss.
 - d. por inducción: 187.
- denotación: 50, 122.

- derivación: v. demostración.
 descripciones: 64, 94, 241 ss.
 descripciones de estado: 295 s.
 dilema: 131.
 Dirichlet: 66.
 disyunción: 73.
 disyuntividad: 268 s.
 división: 268 ss.
 d. dicotómica: 271.
 fundamento de la d.: 269 ss.
 leyes de la d.: 268.
 dominios de una relación: 241 ss.
- efectividad: 45 ss.
 elemento: 39.
 enunciado: 18, 29, 57.
 e. atómico: 20.
 e. esquemático: 24.
 e. formal: 24.
 e. molecular: 20.
 e. primitivos: 35.
 esquemas: 21.
 esquemas axiomáticos: 105 s.
 e. finales: 23.
 e. y forma: 21.
 e. y fórmulas: 57, 80 ss., 95 ss.
 estructura: 258 s.
 Euclides: 99.
 exclusión mutua: 268 s.
 expresión: 56.
 extensión de un término: 39.
- falsedad formal: 22, 24, 50.
 falsedad material: 22.
 forma lógica: 18, 19.
 forma normal: 149 ss.
 f. n. conjuntiva: 151 s.
 f. n. c. de Schröder: 152.
 f. n. disyuntiva: 152 s.
 forma y contenido: 22.
 formalismo: 48.
 formalización: 48.
 formalmente fundamentado: 18.
 formalmente válido: 18.
 fórmula: 56.
 Frege: 38 s., 48, 67, 231.
 función: 66, 262 ss.
 f. asociada a una conectiva: 64 ss.
 f. de confirmación: 301 ss.
 f. lógicas: 66 s.
 f. lógicas de medición: 299 ss.
 f. veritativas: 68.
 f. veritativas diádicas: 73.
 f. veritativas monádicas: 72.
 funtores: 65.
 f. de clases a clases: 217.
 f. de clases a enunciados: 217 s.
 f. de relaciones: 237.
 f. veritativos: 72.
 fundamentos de la ciencia: 30 ss.
- Gentzen: 120.
 Gödel: 183, 189.
 gödelización: 189 ss.
 Grelling: 40.
 Grzegoreyck: 35.
- Hasenjaeger: 132.
 Henkin: 199.
 Herbrand: 173.
 Hermes: 127, 128.
 Hilbert: 48, 87, 110, 118, 199 s.
 hipótesis e información: 286.
 homonimia: 41.
- identidad: 102.
 i. de clases: 218, 225.
 i. de los indiscernibles: 102, 225.
 i. de relaciones: 227.
 implicación: 74, 101 s.
 i. y $\rightarrow$: 84.
 i. parcial: 292 s.
 inclusión entre clases: 217, 224.
 i. entre relaciones: 237.
 independencia: 105.
 individuo: 58.
 inducción: 52, 282 ss.
 i. completa: 283.
 i. incompleta: 284.
 i. y probabilidad: 295.
 inferencia inductiva analógica: 289.
 i. i. directa: 288.
 i. i. inversa: 288.
 i. i. universal: 287.
 i. i. predictiva: 289.
 intensidad: 39.
 interpretación: 22, 24, 50, 107.
 intersubjetividad y lenguaje: 19.
 isomorfía: 108, 125, 253 s., 256 ss.
- Jaskowski: 120.
 juicio: 29.
- Kant: 27 s.
 Kuratowski: 235.
- Leibniz: 38, 54, 102.
 lenguaje formalizado: 48 s.
 letra esquemática: 24.
 ley de doble negación: 133.
 leyes asociativas: 136.
 leyes conmutativas: 135 s.
 leyes de De Morgan: 139.
 leyes de la contraposición: 135.
 leyes de la definición: v. definición.
 leyes de la división: v. división.
 leyes distributivas: 136-138, 143-144.
 lógica como arte o técnica: 27.
 lógica como lenguaje formalizado: 50.
 lógica como teoría: 24.
 lógica de clases: 216, 223 ss.

- lógica de enunciados: 61.
 lógica de predicados: 88.
 l. de p. de primer orden: 93.
 l. de p. monádica y poliádica: 147, 234.
 lógica de relaciones: 234 ss.
 lógica elemental: 92.
 lógica e investigación de fundamentos: 49, 252 ss.
 lógica y conocimiento empírico: 27 s., 33, 36, 119 s.
 lógica y juegos: 25.
 lógica y metodología: 267 s.
 lógica y psicología: 17.
 lógica y teoría del conocimiento: 18.
 logicismo: 39.
 "lo verdadero se sigue de cualquier cosa": 116.
 Lukasiewicz: 43, 87, 284.
 Lull: 44.
- Megara, escuela de: 42.
 Melanchton: 18.
 mención: v. uso.
 metalógica: 49.
 método: 267.
 m. didáctico: 267.
 m. heurístico: 267.
 miembro: 39, 242 ss.
 modelos isomorfos: 106 s.
 modelos semánticos: 51.
 modos silogísticos: 140 s.
 modus ponendo ponens: 111, 131.
 monomorfismo: 108.
- Nagel: 197.
 negación: 72.
 Newman: 197.
 nociones primitivas: 34.
 nombres: 57 s.
 nominalismo: 57.
 número cardinal: 231 s., 259 ss.
- objeto formal: 13, 16.
 objeto material: 13.
 Ockham: 58.
 operadores: 65.
 orden lógico: 187 s.
 Organon: 30, 57.
- Pablo apóstol: 43, 54.
 paradojas: 40 s., 41.
 p. de Epiménides: 42 s.
 p. de Grelling: 40.
 p. de Russell: 188, 228.
 parámetros: 60 s.
 paréntesis: 77, 92.
 pares ordenados: 232.
 Pavlov: 16.
 Peano: 99.
- Pedro Hispano: 27 s.
 pertenencia: 222.
 Platón: 58.
 polimerismo: 108.
 Popper: 28.
 potencias de una relación: 241.
 pragmática: 51.
 predicado: 90 s.
 premisa: 20, 120, 155, 163 ss.
 eliminación de p.: 121.
Principia Mathematica: 40, 204.
 principio de abstracción: 228 s.
 principio de extensionalidad: 180 s.
 principio de inducción: 99 s., 185.
 principio del silogismo: 140, 218.
 principios lógicos: 24, 130 s., 131 s.
 producto lógico de clases: 217, 224.
 producto lógico de relaciones: 237.
 producto relacional: 239 s.
 proposición: 29.
- Quine: 15, 58, 61, 120.
- realismo: 58.
 reducción: 284 ss.
 reducción al absurdo: 209 s.
 reflexividad: 245 s.
 regla de generalización posterior: 113.
 regla de particularización anterior: 114.
 regla de separación: 111.
 reglas auxiliares: 116.
 reglas de sustitución: 110 s., 125.
 reglas de formación: 46.
 reglas de transformación: 46, 103.
 relación nula: 238.
 relación universal: 238.
 relaciones de equivalencia: 251.
 relaciones y clases: 235 s.
 resolutoria: 119.
 Robinson: 276.
 Rosser: 197.
 Russell: 38 s., 40, 42, 58, 87, 231, 286.
- Scholz: 50.
 Schröder: 28, 152.
 semántica: 49.
 semiótica: 51.
 sentencia: 29.
 sentido: 50.
 'ser': 88 ss.
 seric: 261 s.
 Sheffer: 205 ss.
 simbolización: 48.
 simetría: 246 s.
 sincategorema, sincategoremático: 21.
 sinonimia: 41.
 sintaxis: 49.
 síntesis en lógica: v. análisis.
 sistema axiomático: 103.
 s. a. formalizado y no formalizado: 103.

- sistemas formales: 25, 45.
suasorio, razonamiento: 9.
subconjunto: 39.
subíndices: 92.
suma lógica de clases: 217, 224.
suma lógica de relaciones: 237.
- tablas veritativas: 202 ss.
Tarski: 50 s.
tautologías: 25.
teorema: 103 s., 122.
teorema de deducción: 166 ss.
teorema de incompletud de Gödel: 192 ss.
teoría de los tipos lógicos: 188 s., 229.
teorías: 32, 34 s.
tercio exclusivo: 131.
Tomás de Aquino: 27 s.
Touchais: 271.
transitividad: 247.
trazo de Sheffer: 206.
- Ueberweg: 277.
universo del discurso: 107.
u. d. d. empírico y lógico: 97.
univocidad: 254.
uso y mención: 14 s.
- variable: 59 s.
v. libre: 93.
v. ligada: 93.
v. sintáctica: 49 s.
verdad formal: 26 s. (origen y modo de validez), 22, 24, 50.
sentido de las v. f.: 26.
verdad material: 22.
verdad semántica: 50.
- Whitehead: 40, 231.
Wiener: 235.
Wittgenstein: 25.
- Zermelo: 231.