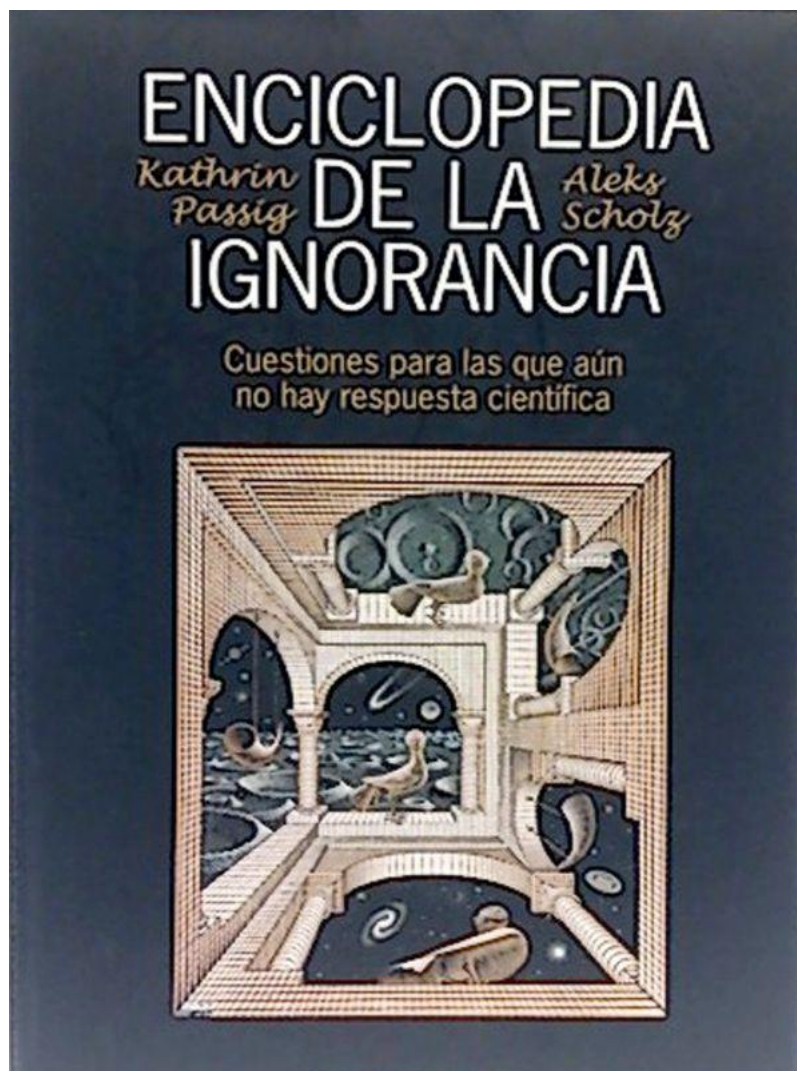




Reseña

¿Qué es la lluvia roja? ¿Por qué bostezamos? ¿En qué consiste la hipótesis de Riemann? ¿Qué es la materia oscura? ¿Cómo nos afectan los alucinógenos? ¿Existe la eyaculación femenina? ¿Por qué hay animales que necesitan dormir horas y otros tan sólo minutos? ¿Es más lo que sabemos sobre la realidad o lo que ignoramos? Cuestiones tan variopintas como las anteriores —unas próximas y cotidianas, otras sesudas y trascendentales, la mayoría sorprendentes y atractivas — siguen sin gozar de respuesta satisfactoria. La Enciclopedia de la ignorancia proyecta una mirada fresca sobre medio centenar de enigmas científicos sin resolver, mediante breves ensayos escritos con talento, conocimiento de causa, amplia documentación y mucho sentido del humor. Este magnífico libro de ciencia recreativa, que satisfará la curiosidad del más escéptico, tiene asimismo la virtud de mostrarnos hasta qué punto no se ha logrado explorar, descifrar ni demostrar todas las cosas. Porque resulta sencillo alardear de conocimiento, pero es más difícil determinar las islas de ignorancia, su magnitud y su ubicación. He aquí, pues, una parte —no la menos importante— de lo mucho que nos queda por saber.

*Cincuenta cuestiones que hasta ahora no
han tenido respuesta científica*

*Dedicado a los ratones de laboratorio y a
su incansable lucha contra lo logrado*

Índice

Lo que vale la pena conocer sobre lo desconocido

1. Agua
2. Alucinógenos
3. Americanos
4. Anestesia
5. Anguila
6. Bostezo
7. Ciempiés
8. Cinta autoadhesiva
9. Cúmulos globulares
10. Curva de Laffer
11. Dinero
12. Einemsen
13. Escritura del Indo
14. Estatura del ser humano
15. Estrella de Belén
16. Eyaculación femenina
17. Follaje otoñal
18. Goteo
19. Hawai
20. Hipótesis de Riemann
21. Lluvia roja
22. Manuscrito Voynich
23. Materia oscura
24. Miopía
25. Olfato
26. Parque Nacional Los Padres
27. Partículas elementales
28. Problema PN/P

29. Propina
30. Rayos globulares
31. Resfrío
32. Rey de ratas
33. Ronroneo
34. Rotación de las estrellas
35. Sonidos desagradables
36. Suceso de Tunguska
37. Sueños
38. Tamaño de los animales
39. Tectónica de placas
40. Tendencias sexuales
41. Túnel del Metro Norte Sur
42. Vida

Bibliografía

Lo que vale la pena conocer sobre lo desconocido

Hay conocimientos conocidos: hay cosas de las que sabemos que las conocemos.

Hay desconocimientos conocidos: es decir, hay cosas de las que ahora sabemos que no las conocemos. Pero también hay desconocimientos que están por conocer: hay cosas de las que no sabemos que las desconocemos. Y cada año descubrimos algunos más de esos desconocimientos por conocer.

DONALD RUMSFELD

Contenido:

1. *¿Qué es lo desconocido?*
2. *¿Por qué precisamente lo desconocido?*
3. *¿Cómo hubiera sido una enciclopedia de la ignorancia hace cien años?*
4. *¿Cómo se hace para encontrar lo desconocido?*
5. *¿Por qué se habla tanto sobre lo que se sabe, pero se habla mucho menos de lo desconocido?*
6. *¿Son más las cosas desconocidas que las conocidas? ¿Es todo quizá desconocido?*
7. *¿Cómo se ha hecho la selección de temas?*
8. *Aun así, falta mi tema favorito*
9. *¿No será que los extraterrestres tienen la culpa de todo?*
10. *¿Por qué hay uno o dos errores en el libro?*

1 ¿Qué es lo desconocido?

Las lagunas del conocimiento surgen habitualmente de la antigua técnica cultural del olvido.

De una manera que desde luego es menos bochornosa, este libro generará en cada lector 42 lagunas adicionales. Cada una de ellas es una laguna de buena calidad con la que no sólo nosotros nos devanamos los sesos, sino también el resto de la humanidad, incluidos muchos investigadores cuya inteligencia está por encima de la media. La *Enciclopedia de la ignorancia* es el primer libro tras cuya lectura se sabe menos que antes de leerlo, pero, eso sí, con un nivel muy alto.

Si nos imaginamos el estado de discernimiento de la humanidad como un gran mapa geográfico, el conjunto de todos los conocimientos constituye la masa terrestre de ese mundo imaginario. Lo desconocido se esconde en los mares y lagos. La tarea de la ciencia es reducir los lugares húmedos del mapa. No es cosa fácil, porque en ocasiones aparecen de nuevo charcos en lugares que creíamos secos desde hace tiempo. Un ejemplo es la pregunta que plantea quiénes fueron los primeros pobladores de América y cuándo llegaron allí: esta cuestión se consideraba aclarada desde hace más de medio siglo, pero tras el hallazgo de nuevos restos ha quedado totalmente abierta otra vez hace unos pocos años. Los investigadores no sólo consiguen aumentar los conocimientos de la humanidad, sino también lo desconocido. Así sucedía a finales del siglo XIX, cuando muchos físicos estaban convencidos de haber investigado ya todo lo que había en el mundo, y de que sólo les quedaban unos pocos detalles por aclarar. Esto fue así hasta que la mecánica cuántica y la teoría de la relatividad pusieron de manifiesto que en muchos aspectos se habían quedado demasiado cortos: un nuevo y enorme mar de desconocimiento se precipitó sobre ellos.

Lo desconocido sólo se puede describir perfilando sus bordes, mientras nos agarramos a las últimas certezas. Volviendo a la metáfora del mapa, diremos que cualquier entrada de esta *Enciclopedia de la ignorancia* es como el contorno de un lago: miramos hacia lo desconocido desde todas las perspectivas posibles, nos hundimos ocasionalmente en zonas pantanosas, encontramos un sendero que nos permite caminar un poco más lejos, pero, sin embargo, nunca podemos decir qué se esconde ante nosotros. La línea costera o fronteriza entre el conocimiento y el desconocimiento no se puede trazar de manera inequívoca, porque casi siempre hay varias teorías que compiten por la resolución de un problema determinado.

Lo desconocido que aquí nos ocupa debe satisfacer tres criterios: no puede existir ninguna solución que haya sido aceptada por una gran parte de los especialistas en la materia y que sólo precise todavía algunos retoques en cuestiones de detalle. Sin embargo, el problema debe haber sido tratado con la profundidad suficiente para que su perfil esté claramente definido. Además, debe ser un problema básicamente resoluble. Muchos interrogantes que la historia tiene abiertos son de tal naturaleza que nunca podremos darles respuesta (salvo que alguien invente una máquina del tiempo).

La descripción realizada por Donald Rumsfeld que hemos citado al principio, provoca a menudo risas, pero esto es injusto, ya que se trata de un hito en la forma de presentar públicamente lo desconocido. Según esta descripción, lo desconocido se puede clasificar en dos categorías: cosas de las que sabemos que no las conocemos, y cosas de las que ni siquiera sabemos que no las conocemos. En este libro sólo se puede tratar de las que figuran en la primera categoría, las «known unknowns», porque sobre las de la segunda no hay por ahora nada que decir.

2 ¿Por qué precisamente lo desconocido?

En el libro de Douglas Adams *Guía del autostopista galáctico*¹, unos seres pandimensionales e hiperinteligentes desarrollan el Computer Deep Thought, que debe dar respuesta a la pregunta relativa a «la vida, el universo y todo lo demás». Después de siete millones y medio de años está ya la respuesta calculada y es «42». Es entonces cuando los diseñadores de Deep Thought se dan cuenta de que no tienen ni idea de cuál es la pregunta. Hasta que lo averiguan pasan otros diez millones de años. A partir de esto se pueden aprender dos cosas: en primer lugar, se debería conocer la pregunta, si se quiere comprender la respuesta. Y, en segundo lugar, plantear la pregunta adecuada es a menudo más difícil que responderla (podemos observar el mismo fenómeno cuando miramos por encima del hombro a los usuarios inexpertos de Google). El físico Eugene Wigner recibió en 1963 la mitad del premio Nobel de física porque había planteado la pregunta adecuada, en concreto sobre el fundamento de los «números mágicos» del sistema periódico. La otra mitad fue para los dos investigadores que hallaron la respuesta.

¹ Anagrama, Barcelona, 2005. (*N. de la t.*)

El planteamiento de preguntas adecuadas para desvelar lo desconocido es una tarea importante en el ámbito científico. Porque lo desconocido siempre está ahí, sólo que no es evidente para todo el mundo; es como el animal negro que en un acertijo gráfico rellena el espacio que hay en torno al animal blanco, y que no solemos ver hasta que llevamos cierto tiempo contemplando la imagen.

Sin embargo, a partir de entonces no dejamos de percibirlo. Si este libro consigue captar mínimamente nuestra atención y dirigirla hacia el animal negro de lo desconocido, habrá alcanzado su objetivo. Después, el lector reconocerá lo desconocido incluso cuando tropiece con él en plena naturaleza.

3. ¿Cómo hubiera sido una enciclopedia de la ignorancia hace cien años?

Lo desconocido es una cosa huidiza. Desaparece y luego emerge de nuevo en otro lugar.

Resumiendo, es más difícil seguirle la pista a lo desconocido que a lo que conocemos. Por eso, una enciclopedia de la ignorancia no puede hacerse para siempre. Si comparamos este libro con un predecesor de hace cien años, que desgraciadamente nunca se escribió, nos encontramos con algo interesante: algunos temas desconocidos ni siquiera se conocían entonces (pensemos en la tectónica de placas o en la materia oscura). Otras preguntas sin respuesta estaban ya presentes en la historia del mundo, al alcance de cualquiera, pero por distintas razones no se abordaron en absoluto, o no se trataron aplicando medios racionales, como sucedió, por ejemplo, con el misterio de la eyaculación femenina. En cambio, otros problemas siguen estando sin resolver y, por eso, podrían de pleno derecho aparecer en ambas versiones de la enciclopedia, como es el caso, entre otros, de la hipótesis de Riemann o de la estructura de la materia. Sin embargo, nos sentimos optimistas al pensar en los campos del desconocimiento que no aparecen en este libro, aunque hace cien años eran grandes misterios. Por ejemplo, no se tenía ni idea de por qué brillan las estrellas. Aunque se sospechaba que el núcleo de la Tierra no era sólo de tierra, no se supo hasta dos décadas más tarde que se trataba de un fluido, lo cual era un hecho bastante inquietante.

No se sabía por qué los cítricos eran buenos para combatir el escorbuto. Ni siquiera eran conocidos los lugares de desove de las anguilas.

Si leyéramos ahora una enciclopedia de la ignorancia de hace cien años, pareceríamos probablemente muy listos. Es de esperar que les suceda lo mismo a nuestros bisnietos cuando tengan este libro en sus manos. Dirán que la materia oscura, por supuesto, la constituyen los superaxocuatrones que giran hacia la izquierda, como todo el mundo sabe, y ¿cómo se podía creer en aquellos tiempos que el sueño tiene alguna función? Evidentemente los gatos no ronronean, se trata de una ilusión acústica, y lo que son esos «reyes de las ratas»² figura claramente explicado en el manuscrito Voynich. Así, con el paso del tiempo, este libro contendrá cada vez menos temas verdaderamente desconocidos, hasta que no valga la pena hacer una nueva edición de las dos páginas que sigan siendo válidas. Por suerte, esto no lo verán los ojos de sus autores.

4. ¿Cómo se hace para encontrar lo desconocido?

¿Cómo se hace para encontrar agujeros? Pues basta con caminar lo suficiente como para llegar a un punto en que ya no hay suelo bajo los pies. Algo parecido sucede con lo desconocido: uno se hace más y más preguntas hasta que no pueda responder alguna, cosa que a menudo se consigue con gran rapidez. (No nos referimos a «no poder responder» en el sentido habitual de decir «Si entonces me hubiera estudiado mejor la química...», sino a que no exista realmente respuesta alguna).

El indicador más seguro para detectar un buen desconocimiento es que los expertos se dediquen en un congreso a hacer apuestas sobre la dirección en que se podría encontrar la solución de un problema determinado. Así pues, lo ideal sería torturar con preguntas a los expertos hasta que reconozcan de manera unánime que ya no saben más. Lástima que esto sólo sea factible en unos pocos casos. En vez de hacerlo así, tenemos que buscar lo desconocido con esfuerzo y de forma indirecta a través de las lagunas de los tratados científicos, que en la mayoría de los casos aclaran minuciosamente lo desconocido. Pocas veces encontramos indicaciones directas sobre lo desconocido. En cambio, casi todos los periódicos tienen una sección dedicada al conocimiento; con esos artículos, que informan orgullosamente

² Agrupaciones de ratas con las colas entrecruzadas y pegadas entre sí. Del alemán Rattenizónige. Aparece más adelante en la entrada correspondiente (N. de la t.)

sobre la resolución del problema X, se podría llenar innumerables archivadores, aunque suele tratarse casi siempre del mismo problema. Por ejemplo, sólo en los últimos diez años, se ha explicado de manera definitiva aproximadamente unas tres veces cómo funcionan los rayos globulares. Esas informaciones son un claro signo de la existencia de lo desconocido.

5. ¿Por qué se habla tanto sobre lo que se sabe, pero se habla mucho menos de lo desconocido?

Una razón para esto es con toda seguridad la forma de trabajar del científico: para no perderse en especulaciones infundadas, debe atenerse a lo que ya sabe, con lo que se sitúa, por decirlo de algún modo, de espaldas a lo desconocido. Sólo de vez en cuando se da la vuelta para no perder de vista lo que es realmente su función: el esclarecimiento de lo desconocido. Son esos momentos los que debemos investigar, si buscamos lo que todavía no es conocido.

Pero hay otras causas de ese abandono en que suele quedar lo desconocido dentro de la cobertura informativa: los periodistas prefieren informar sobre trabajos de investigación ya terminados y sobre los nuevos conocimientos. Está claro que un titular que diga «Nada nuevo sobre X» gusta menos que «El misterio de X finalmente resuelto». Además, no hay que esforzarse mucho para obtener información sobre resultados concretos a partir de los comunicados de prensa de las instituciones dedicadas a la investigación, mientras que lo desconocido requiere una investigación intensiva, lo cual hace que resulte más caro. Además, no hay que olvidar que es mucho más agradable fomentar la ilusión de que conocemos ya todo lo que es esencial. Pero, al mismo tiempo, esta idea puede resultar extraordinariamente contraproducente. En 1874 el profesor de física Philipp von Jolly intentó disuadir al joven Max Planck de que emprendiera el estudio de dicha disciplina, porque según él se trataba de una ciencia en la que casi todo estaba ya investigado. Por suerte Max Planck ignoró este consejo y, pocos años más tarde, dio el impulso inicial al desarrollo de la teoría cuántica, poniendo en marcha así una revolución de la física moderna.

Sin embargo, en unos pocos casos se investiga lo desconocido de una manera muy concreta.

Aunque pueda parecer lo contrario, lo «desconocido por conocer» no es simplemente un invento de Donald Rumsfeld. Se trata de un problema muy conocido dentro de la teoría militar, y el ejército de Estados Unidos lo ha bautizado con el nombre «*unk-unk*» (que viene de *unknown unknown*, o sea «lo desconocido desconocido» o «lo desconocido por conocer»). En una guerra hay muchas cosas que no pueden verse previamente, pero pueden existir, razón por la cual hay que tenerlas todas en cuenta, incluso lo imprevisible. Las omisiones pueden resultar penosas y muy caras. Por el mismo motivo, la NASA mantiene una base de datos de «*Lessons-Learned*» («lecciones aprendidas»), para que los errores debidos al desconocimiento de factores desconocidos se cometan como mucho una sola vez y no más. Esas instrucciones que nos recomiendan intentar controlar lo desconocido se las tenemos que agradecer al proyecto de investigación interdisciplinario llamado «*Nichtwissenskulturen*», que se desarrolló en la Universidad de Augsburgo desde 2003 hasta 2007.

6. ¿Son más las cosas desconocidas que las conocidas? ¿Es todo quizá desconocido?

Para responder a estas preguntas, Isaac Newton propuso, hace unos trescientos años, lo siguiente: «Lo que sabemos es una gota de agua. Lo que no sabemos es el océano». Es cierto que desde los tiempos de Newton ha cambiado un poco la situación, pero, sin embargo, la cantidad de cosas desconocidas no ha disminuido de una forma decisiva. En cuanto se llegaba a un punto en que se sabía más sobre algo, surgían al momento nuevas preguntas que quedaban abiertas. No obstante, de esto no debería deducirse inmediatamente que toda información es incierta y que en el futuro tendrá que ser sustituida por nuevos conocimientos. No hay que infravalorar las ciencias.

Más bien se trata de una situación que se puede resumir de la siguiente manera: sin duda disponemos de una cantidad considerable de conocimientos, la mayoría de los cuales no pueden ponerse en duda. (Quizá no todo el mundo dispone de ellos, pero sí la humanidad en su conjunto).

Por otra parte, hay además un gran número de problemas que, aun siendo importantes y de gran interés, están sin resolver, y que a algunos de nosotros nos

tienen ocupados a diario. Pero hay que ver también el lado positivo: a fin de cuentas, en el futuro habrá también personas que quieran cobrar un sueldo por estar en los laboratorios y rascarse la cabeza con cara de perplejidad.

7. ¿Cómo se ha hecho la selección de temas?

En julio de 2005, la revista *Science*, que es una autoridad en cuestiones desconocidas, publicó una lista de ciento veinticinco preguntas importantes a las que se tendría que dar respuesta en el siglo XXI. Sólo quince de esas preguntas aparecen en nuestra *Enciclopedia de la ignorancia*. Junto a éstas se presentan otras cuestiones que rara vez llegan a ver la luz del día, por ejemplo, el rey de las ratas. Esto se debe a que en este mundo son tantos los temas desconocidos, que no pueden caber en un solo libro, y es comprensible que la editorial no desee publicar ahora mismo una enciclopedia de veinticuatro tomos sobre dichos temas. Por otro lado, los temas se han elegido en parte por su importancia, y en parte también porque ilustran la habilidad con que lo desconocido se esconde dentro de lo que ya se conoce. Por ejemplo, investigar el origen del ronroneo de los gatos no es algo con lo que se pueda ganar el premio Nobel o cambiar de manera decisiva la faz del mundo. Sin embargo, es un tema de extraordinaria resonancia, sobre todo por la generalizada omnipresencia de estos animales. Por el contrario, para la mayoría de la gente es difícil comprender qué interés puede tener el bosón de Higgs, aunque quien realice tal descubrimiento puede tener asegurado el premio Nobel. Tampoco resulta en absoluto fácil averiguar cuáles de las preguntas abiertas serán decisivas a largo plazo para nuestra visión del mundo. Por lo tanto, el criterio de la importancia es tan subjetivo como cualquier otro. Porque, quién sabe si dentro de cien años no se construirá una central eléctrica ecológica con gatos ronroneantes. Hace doscientos cincuenta años nadie hubiera pensado que algún día seríamos capaces de pasear al aire libre, cuando ya ha oscurecido, sin usar velas, todo porque hemos llegado a comprender la razón por la que se contraen unas ancas de rana.

8. Aun así, falta mi tema favorito

Los temas filosóficos e históricos tienen una representación un tanto escasa en esta

Enciclopedia de la ignorancia, porque rara vez se corresponden con las cuestiones desconocidas que consideramos lógico o necesario investigar (véase el párrafo anterior). Aunque no esté clara la causa por la que en 1872 el velero *Mary Celeste* se encontró navegando a la deriva, sin tripulación, entre las islas Azores y Portugal, lo más probable es que no se intente jamás averiguar qué sucedió. Y también las preguntas relativas a las causas de determinadas enfermedades se quedan cortas, y hay motivos para ello. Al poco de que comenzáramos nuestra investigación, se nos amontonaron sobre la mesa los informes relativos a enfermedades cuya génesis no está clara, y tuvimos que reconocer que ningún desencadenante de enfermedades se comprende del todo, salvo quizá la causa de una fractura en la pierna (a saber, una acción mecánica violenta, a menudo ejercida por objetos que se encuentran donde no deberían estar). En la dirección de correo «vorschlag@lexikondesunwissens.de» recibiremos con agrado cualquier propuesta relativa a temas desconocidos, que podríamos incluir en una posible reedición de este libro.

9. ¿No será que los extraterrestres tienen la culpa de todo?

Es asombroso el gran número de preguntas de este libro que pueden responderse con ayuda de seres procedentes del universo exterior. Por ejemplo, los rayos globulares son naves espaciales que en ocasiones llegan de mundos extraños. Los extraterrestres entran en contacto con nosotros valiéndose de sustancias alucinógenas, según dice Rob McKenna. Además, la explosión de Tunguska tiene, por supuesto, algo que ver con seres de otros mundos, por no hablar de la estrella de los Reyes Magos. Christopher Chippendale, arqueólogo y autor de *Stonehenge Complete*³, se remite a ejemplos históricos para explicar las especulaciones sobre extraterrestres de los siglos XX y XXI, que son armas de múltiples usos contra los temas desconocidos de cualquier tipo.

Durante mucho tiempo, los habitantes de la Atlántida asumieron las funciones de los extraterrestres: poblaron América, construyeron Stonehenge, y el hundimiento de sus dominios resuelve el enigma de la reproducción de la anguila.

³ Christopher Chippendale, *Stonehenge*, Destino, Barcelona, 1989. (N. de la t.)

Antes de que la Atlántida se pusiera de moda, los fenicios desempeñaron un importante papel en la explicación de lo desconocido, ya que, según Chippendale, eran el prototipo de los extraterrestres de ahora. También los fenicios, un pueblo real, documentado históricamente, aparecen en algunas teorías como los primeros pobladores de América, y también podrían ser ellos quienes construyeron Stonehenge. « ¿Por qué no?», se preguntan algunos que buscan respuestas desesperadamente. « ¿Por qué no los fenicios?». Aquí es precisamente donde está el problema de las teorías de medio mundo: quieren parecer convincentes, pero no pueden probarse ni refutarse, sobre todo porque sobre los extraterrestres (como sobre la Atlántida o los fenicios) se sabe muy poco, o más bien nada. Están muy lejos, son poderosos y misteriosos. Se explica algo que es desconocido utilizando simplemente otra cosa totalmente desconocida. Por lo tanto, lo de los extraterrestres no tiene por qué ser falso, sólo que no se puede decir nada sobre lo que pueda haber de verdad en ello: como dijo el físico Wolfgang Pauli en un contexto parecido, esta teoría no es «ni siquiera falsa». Naturalmente esto es también una ventaja, pues así se puede adaptar sin dificultad cada solución individual a varias cuestiones diferentes sobre lo desconocido. Sin embargo, algún día también los extraterrestres pasarán de moda, como antes los fenicios, y serán sustituidos por algo aún más exótico e incomprensible. Por ejemplo, los erizos.

10. ¿Por qué hay uno o dos errores en el libro?

Esta *Enciclopedia de la ignorancia* contiene errores, porque los errores son muy valiosos para la experiencia humana. Hay dos tipos de faltas que no son evitables. Por una parte, se trata de errores por simplificación. Muchas circunstancias complicadas sólo se pueden expresar de una forma comprensible eligiendo comparaciones gráficas, que en sentido estricto son inexactas. Pero sin esos recursos el libro sería ilegible. Por otra parte, este libro contiene con toda seguridad suposiciones y afirmaciones que en un futuro cercano o lejano resultarán probadamente falsas.

Pero actualmente todavía no sabemos nada de esos errores, ni nosotros, ni los expertos. Dejando a un lado estos dos tipos de inexactitudes, en la *Enciclopedia de la ignorancia* también habrá auténticos errores, cosas que no es que sean

posiblemente falsas, sino falsas con toda seguridad. A pesar de los minuciosos controles y asesoramientos de los expertos, no siempre son evitables esos errores, y la culpa de esto es sólo nuestra. Por si acaso, pedimos disculpas y que se nos informe al respecto en «korrektur@lexikondesunwissens.de», para que podamos corregir esas faltas en futuras ediciones.

Capítulo 1

Agua

El agua no obedece tus «reglas». Va adonde quiere. Como yo, chaval.

BART SIMPSON

El agua es un producto químico abundante en el mundo, ya sea como líquido incoloro, como vapor incoloro o como trozo de hielo incoloro. En comparación con otras sustancias químicas como el aceite de marmota, el agua tiene una gran importancia para el ser humano. A pesar de ello, o precisamente por ello, el agua es una de las sustancias más enigmáticas del planeta.

Gran parte de su caprichoso comportamiento se debe a la estructura de la molécula de agua, que, como todos aprendimos de jóvenes, está formada por un átomo de oxígeno y dos de hidrógeno. Mientras que el oxígeno posee ocho electrones, los átomos de hidrógeno tienen tan sólo uno cada uno. Dos de los electrones del oxígeno podemos ignorarlos ya de entrada, pues están tan pendientes del núcleo del átomo que no prestan ninguna atención a su entorno. Así, nos quedan ocho electrones por cada molécula de agua que intentan alinearse por pares, pues estar solos no es el punto fuerte de los electrones. Lo que sucede es lo siguiente: cuatro electrones de oxígeno se aparean entre sí, y el resto lo hacen con los de los átomos de hidrógeno y mantienen la molécula cohesionada. El resultado es una figura con un tronco grueso (el átomo de oxígeno), dos brazos (los átomos de hidrógeno) y dos piernecitas (los dos pares de electrones de oxígeno) que sobresalen, desvalidas. Si en el mundo hubiera una sola molécula de agua y no muchas más, no haría falta añadir nada.

Sin embargo, la presencia de moléculas vecinas complica enormemente las relaciones sociales de los componentes del agua. Por una parte, los átomos de hidrógeno están en movimiento constante y cambian de molécula hasta mil veces por segundo. Por otra, la molécula de agua está muy polarizada eléctricamente: el grueso núcleo de oxígeno atrae los pares de electrones de modo que termina envuelto por una nube de electrones de carga negativa, y más rezagados quedan

los dos brazos de hidrógeno comparativamente desnudos, es decir, cargados positivamente. Como consecuencia, los brazos positivos de una molécula intentan agruparse con las piernas negativas de otra. El agua es un elemento de enlace fácil pero, por otro lado, inconstante en sus uniones, por lo que puede formar todo tipo de estructuras, de las que hablaremos más tarde. Además, estas propiedades hacen también que transmita diligentemente la electricidad y el calor, que la sal se disuelva en ella y que se pueda unir con sustancias orgánicas como el albumen. El agua se mezcla sin rechistar con todo, por eso los seres vivos están formados principalmente de agua.

Otra consecuencia de la estructura única de la molécula de agua son sus más de sesenta anomalías que parecen contradecir las leyes de la física molecular en vigor y que se comprenden tan sólo hasta cierto punto. La más conocida es la propiedad de alcanzar la menor densidad a cuatro grados Celsius, y no en el punto de congelación o por debajo de éste, como el resto de sustancias, que sí respetan el código de circulación. Lo normal sería que un bloque de hielo atara bien corto a las moléculas, mientras que en estado líquido éstas pudieran campar más a sus anchas. Por ello, en el caso de la mayoría de sustancias, caben más moléculas en un volumen concreto cuando el estado de ésta es sólido que cuando es líquido y, por lo tanto, la densidad es mayor. Pues con el agua sucede al contrario, algo que se puede explicar de la siguiente forma: en las estructuras de hielo que se forman en condiciones normales, las moléculas quedan ordenadas en mucho espacio y de forma muy poco inteligente o, dicho de otra forma, mucho menos apretadas de lo que les gustaría. Apenas se funde el hielo, se juntan con una densidad mayor. Por eso el hielo flota encima del agua, lo que permite que los ríos y lagos se empiecen a congelar por arriba.

Y muchos peces se alegran de ello.

Muchas peculiaridades del agua se deben al hecho de que el agua caliente se comporta de forma distinta que el agua fría en muchos sentidos. Una anomalía de la lista que de momento no ha sido aclarada es el llamado efecto Mpemba, que aborda la cuestión de por qué el agua caliente a veces se congela más rápido que el agua fría. El primero en describir el fenómeno fue Aristóteles; al parecer, en su época era una práctica común poner al sol el agua que luego había que congelar, porque el

calor acelera el posterior proceso de congelación. Como hay muchas bautizadas en honor a personajes griegos de la antigüedad pero sólo unas pocas en honor a personajes africanos, los científicos de todo el mundo acordaron olvidarse profesionalmente del tema un rato para que en el año 1963 Erasto Mpemba pudiera redescubrirlo en Tanzania. Mpemba tenía que preparar un helado para la clase de física con leche cocida, pero como quería ahorrar tiempo, metió la mezcla aún caliente directamente en la nevera. En realidad, con ello ganó el doble de tiempo, pues su mezcla caliente de leche y azúcar se congeló antes que las mezclas frías de sus compañeros.

Mpemba necesitó seis años, muchas lecciones de física y mucha tenacidad hasta que, finalmente, se le reconoció el hecho.

El que tenga conocimientos previos sobre física de sustancias frías y calientes creerá que el efecto Mpemba es una leyenda. El enfriamiento de líquidos en situaciones ideales funciona como una carrera de resistencia con velocidad constante donde el punto de congelación es la línea de meta. Si uno coloca dos recipientes, uno con agua fría y uno con agua caliente, uno junto a otro en la nevera, el agua fría debería congelarse antes, pues la distancia entre la temperatura de salida y el objetivo es menor. En determinadas circunstancias, sin embargo, se da el caso opuesto. ¿Por qué el agua es tan poco de fiar en este sentido? Al parecer, las dos aguas no se diferencian tan sólo por la temperatura, sino también por otros factores. En el efecto Mpemba son relevantes tanto la cantidad de agua como el contenido gaseoso y mineral de ésta, la forma y el tipo del recipiente, el tipo de nevera y, naturalmente, la temperatura. Para investigar el fenómeno hay que tener en cuenta todos esos parámetros.

Existen muchas teorías enfrentadas en relación con el efecto Mpemba. Por ejemplo, una dice que el agua caliente se evapora y pierde moléculas, por lo que al final quedan menos que enfriar.

Una segunda posibilidad se fija en los gases disueltos en el agua, como el anhídrido carbónico, que se escapan cuando ésta se calienta. Eso altera la disposición de las moléculas del agua de tal forma que se congela antes. (En ese sentido, el agua líquida tiene «memoria», pues la configuración de sus moléculas varía en función del manejo y el agua «percibe» ese cambio; en cualquier caso, basta con agitarla

con fuerza para que se olvide de todo). Hace poco, finalmente, se propuso una nueva variante sorprendentemente sencilla: el agua caliente contiene menos minerales porque, al calentarla, éstos se depositan en el fondo del recipiente (es algo que se puede apreciar en una cacerola) y por eso se enfría más rápido. Es cierto que la presencia de sales dificulta el enfriamiento del agua y por eso hoy en día se utiliza sal para eliminar la nieve de las calles. Todas estas explicaciones tienen una cosa en común: los expertos en agua no han aceptado ninguna.

Si bien el efecto Mpemba puede estudiarse con medios limitados, otros misterios del agua requieren todas las virguerías técnicas de que dispone la física actual. Por ejemplo, cuando se trata de determinar por qué los patines de hielo se deslizan por el hielo. Muchos creen que se debe tan sólo a la alta presión que efectúan las delgadas cuchillas. Eso, según se cree, hace que se forme una fina película líquida sobre la que se desliza el patinador. Eso, sin embargo, sucede tan sólo con patines de hielo, apenas con esquís y en absoluto si intentamos deslizarnos con zapatos normales. Una alternativa posible es la vieja teoría de la fricción. Según una hipótesis formulada por primera vez en la década de 1930, la fricción de las cuchillas y los zapatos encima del hielo genera suficiente calor para fundir algo de agua, que actúa como lubricante. Existen algunas pruebas experimentales que corroboran esa hipótesis y actualmente los científicos están casi seguros de que el calor de la fricción desempeña un papel importante a la hora de deslizarse encima del hielo. Por desgracia, sin embargo, el hielo es liso cuando nada se mueve sobre él y no hay ninguna fricción.

En caso necesario, la peculiar estructura de la molécula de agua nos serviría para salir del paso de todas estas situaciones. Si lanzamos electrones, protones o rayos X sobre una superficie helada veremos que las moléculas que hay en lo más alto se comportan como si fueran líquidas, una idea que ya apuntó el legendario físico Michael Faraday en 1850 sin la ayuda de cañones de electrones. Así pues, podría ser que el hielo fuera capaz de generar por sí solo, sin presión ni fricción, una capa pseudolíquida sobre la que puede uno patinar o deslizarse. No se conoce exactamente el proceso que se genera en el hielo, aunque la ciencia actual se dedica a estudiarlo.

Seguramente tiene que ver con el hecho de que, en la superficie, las moléculas, con su ansia por establecer enlaces, no saben ya adónde ir con sus nubes de electrones, y patalean y bracean desesperadas, exactamente lo mismo que hacen en estado líquido. Sin embargo, por bonito que suene, también existen escépticos: el físico Miguel Salmerón y su equipo han analizado la superficie del hielo utilizando un microscopio de fuerza atómica y han descubierto que, a pesar de su capa «líquida», el hielo en escala atómica sigue siendo muy áspero y nada resbaladizo. De momento no se sabe por qué entonces el hielo es liso.

Por cierto, que no existe un único tipo de hielo. En condiciones normales, cuando el agua se congela se convierte en el llamado «Hielo II» (abreviación de «hielo uno hexagonal»). Seis moléculas se encuentran en un hexágono y se sujetan con brazos y piernas al hexágono siguiente.

La formación tiene forma de colmena y, como ya hemos dicho, una estructura ligera. Con una mayor presión atmosférica y temperaturas más bajas, algo que en la tierra es bastante infrecuente por no decir inexistente, pueden formarse estructuras completamente distintas y en ocasiones muy moléculas presentan una densidad mayor que en el vulgar «hielo II». El «hielo III», por ejemplo, no se une con hexágonos, sino mediante pequeños tetraedros moleculares y se forma con una temperatura de -20°C y una alta sobrepresión. El «hielo IX» tiene un aspecto parecido, pero se forma si se enfría bruscamente el «hielo III». Por suerte, no tiene nada que ver con el «hielo 9» de la novela de Kurt Vonnegut *Cuna de gato*, que se hiela a $+46^{\circ}\text{C}$ y que pronto convierte toda la tierra en una bola de nieve.

Desconocemos muchas cosas de los diversos tipos de hielo, sus propiedades y su formación.

Tal vez existan también más tipos de hielo de los que se han descubierto hasta el momento: mañana mismo el Instituto de la esquina podría presentar un nueva forma de hielo, nunca vista antes. La existencia de tantas modalidades de hielo se debe una vez más a la peculiar estructura de la molécula de agua, que se puede organizar de muchas formas distintas sin mayores problemas. El análisis de los tipos de hielo más exóticos permite ver las propiedades de los enlaces del agua que, a su tiempo, podrían aportar datos relevantes sobre los numerosos juegos de sociedad biológicos y químicos en los que participa de buena voluntad. Otra obra de

arte que crea el agua al transformarse en hielo son los copos de nieve. Expresado de forma científica, los copos de nieve son pequeños cristales de hielo que se forman al helarse las nubes de vapor de agua. Los copos más sencillos son pequeños discos o prismas hexagonales (porque, como ya hemos dicho, la estructura normal del hielo que se forma en condiciones normales es hexagonal). Cuando se forma el primer hexágono de vapor de agua, a las moléculas de agua contiguas no les queda más remedio que seguir la estructura anterior. Sin embargo, si por culpa del viento que cruza la nube, por ejemplo, varía la temperatura o la humedad atmosférica alrededor de este «primer copo», pasan cosas muy emocionantes: el disco hexagonal aumenta de altura o de diámetro y le sale un agujero en el centro, o se forman brazos a ambos lados del disco que pueden proliferar en complejas estructuras. Así, cada copo de nieve arrastra con su forma una detallada historia vital.

La lástima es que nos cueste tanto descifrar el idioma. Uno de los expertos en copos de nieve más antiguos de los últimos años, Kenneth G. Libbrecht, del California Institute of Technology de Pasadena, no sólo pasa las vacaciones en lugares donde nieva seguro para poder recopilar copos de nieve, sino que también ha efectuado numerosos ensayos de laboratorio para descubrir cómo se forma la nieve. Su lema es: «Por lo menos una persona en este mundo debería saber cómo surgen los copos de nieve». Le queda trabajo antes de lograrlo, pues aún nadie sabe por qué bajo unas condiciones determinadas se forma un copo de nieve u otro, o de dónde sale todo el zoológico de tipos de nieve. ¿Cómo funciona el crecimiento de los copos de nieve bajo circunstancias distintas? ¿Qué parámetros son importantes además de la humedad y la temperatura? ¿Y cómo se puede comprender la formación a nivel microscópico partiendo de una molécula de agua, de apariencia tan sencilla? Libbrecht, como no podía ser de otra forma, recomienda realizar nuevos análisis de laboratorio a mayor escala. Rendirse, desde luego, no es una opción.



Capítulo 2

Alucinógenos

Cuando Dios toma LSD, ¿ve seres humanos?

STEVEN WRIGHT, humorista estadounidense

El cerebro humano se desorienta por cualquier cosa con tanta facilidad y tan a gusto, que hay que estar muy agradecidos a la evolución porque nos ha dado la posibilidad de conducir vehículos, al menos de vez en cuando. Finalmente, la naturaleza se esfuerza por enriquecer nuestro entorno con ilusiones ópticas y sustancias químicas que llevan nuestra percepción más allá de lo que es razonable. Si surgen ratones blancos allí donde nadie diría que los hay, eso se llama una alucinación. Los alucinógenos (o sustancias que producen alucinaciones) no llevan el nombre adecuado ya que se limitan a cambiar la percepción de lo que existe en el entorno, de tal modo que, si hay ratones blancos, hacen surgir unos de colores. Ésta es la razón por la que algunos expertos abogan por que se cambie el nombre de estas sustancias llamándolas «psicodélicas» (es decir, «sustancias que hacen que se manifieste el alma»), pero mientras no esté claro si el alma tiene realmente pelo y cuatro patas, nos quedaremos en principio con la denominación «alucinógeno». Tales sustancias están contenidas no sólo en varios cientos de plantas y en muchos hongos, sino también en algunas especies de sapos y peces; en cualquier caso, no se conoce hasta ahora la existencia de piedras alucinógenas. Junto a alucinógenos clásicos como el LSD, la psilocibina y la mezcalina, hay una gran cantidad de sustancias naturales y sintéticas que por distintos procesos farmacológicos producen un efecto relativamente parecido: entre las consecuencias físicas cabe citar mareo, debilidad, sopor y trastornos visuales. La percepción se transforma de tal modo que no se puede confiar en los colores y las formas que se ven, y se pueden producir sinestesias, es decir, la percepción de, por ejemplo, sonidos de colores u olores cuadrangulares. A esto se añade la sensación de estar soñando, una percepción del tiempo en parte drásticamente modificada y, en casos extremos, la impresión de

que toda la personalidad del individuo se diluye como un terrón de azúcar en el café.

Aunque los alucinógenos se utilizan desde hace muchos siglos en la religión y para el ocio, no sabemos gran cosa sobre lo que hacen en el cerebro. Sí es conocido el hecho de que todos ellos entran en los receptores de neurotransmisores que hay en el cerebro. Los neurotransmisores son sustancias mensajeras con las que se salva la distancia entre las prolongaciones de dos neuronas.

En lugar de estas sustancias, son los alucinógenos los que se conectan y se comportan como mensajeros que han olvidado su cometido, que abren las cartas y escriben otras cosas totalmente distintas con una caligrafía desfigurada. Desde la década de 1970 se han realizado algunos avances en la identificación de los receptores competentes, pero aún no se ha investigado muy a fondo el modo en que se llega mediante este procedimiento a los efectos antes mencionados.

Sin embargo, lo interesante en relación con los alucinógenos no es sólo el hecho de que sepamos poco sobre ellos, sino también por qué esto es así. Después de que Albert Hofmann en 1943 descubriera accidentalmente el LSD, llegaron dos décadas fructíferas en las que aparecieron varios miles de publicaciones científicas sobre el modo de actuación y las posibilidades de utilización de los alucinógenos. Desde mediados de la década de 1960, empeoró de manera drástica la fama de estas sustancias en la prensa, en buena medida porque su consumo creció hasta convertirse en un fenómeno de masas, y la media de las dosis de LSD que podían conseguirse en la calle era entonces unas diez veces mayor que en la actualidad. En consecuencia los consumidores eran sometidos a menudo a un lavado completo con centrifugado incluido. Además los políticos estadounidenses sospechaban la existencia de una relación entre el creciente consumo de drogas y las nuevas costumbres de sus ciudadanos, que de repente llevaban el pelo largo, quemaban banderas y decían que eran homosexuales.

A lo largo de la década de 1960, se regularon ante todo en Estados Unidos, y de una manera cada vez más estricta, los alucinógenos más consumidos y finalmente, en 1970, se decretó una prohibición total; la mayoría de los países occidentales siguieron la misma pauta más o menos voluntariamente. Los especialistas tuvieron que decidir si a costa de la carrera científica continuaban investigando los

alucinógenos o preferían cambiar de tema discretamente. También se podía ir eligiendo distintas actitudes, como hicieron los psiquiatras estadounidenses Jerome Levine y Arnold M. Ludwig, cuyos estudios dieron resultados favorables al LSD en la década de 1960, pero, tras el cambio que se produjo en la opinión pública, llegaron a conclusiones críticas con respecto a la misma sustancia. Hasta mediados de la década de 1990, apenas se concedieron autorizaciones para nuevos estudios, pero luego la investigación sobre alucinógenos cobró de nuevo algo de impulso. Hoy en día se acepta como algo seguro que los alucinógenos de uso corriente ni producen daños en los distintos órganos, ni llevan a una dependencia física o psíquica.

Esta difícil situación explica también por qué en las últimas décadas se ha investigado tan poco con personas y tanto con ratas. Esto es seguramente menos fatigoso que llevar a cabo experimentos con seres humanos, porque las ratas no están todo el tiempo haciendo risitas y hablando de Dios. El problema es que las ratas tampoco pueden dar información sobre cómo es el efecto de las drogas. Con respecto a la mayoría de los alucinógenos sintéticos que se conocen actualmente sí existen datos precisos sobre sus efectos, y esto se debe a que su descubridor, el químico estadounidense Alexander Schulgin, los comprobó mediante una larga serie de experimentos que realizó consigo mismo y describió los resultados.

Además, los animales que se utilizan en el laboratorio, a diferencia de muchos seres humanos, no se prestan de buena gana a tomar alucinógenos si se les da opción a elegir, a pesar de que no se asustan ante drogas que no tienen efectos alucinógenos, como la cocaína, la heroína, las anfetaminas, la nicotina y el alcohol. Se sospecha, por lo tanto, que hay que tener un cerebro altamente desarrollado para encontrar divertido lo que los alucinógenos desencadenan en el cerebro. Sin embargo, hemos de agradecer a las ratas el conocimiento de algunos datos nuevos sobre los receptores compartidos. Muchas sustancias alucinógenas se parecen en su estructura claramente a la serotonina, uno de los neurotransmisores más importantes del cerebro. Existen muchos receptores de serotonina diferentes, aunque el efecto alucinógeno probablemente se produce sobre todo mediante la activación del llamado receptor del tipo 2A para la serotonina.

No obstante, es un fastidio que precisamente el LSD deje más bien de lado a este receptor, a pesar de que en dosis menores desencadena enormes cambios en la

percepción y produce un efecto considerablemente más fuerte que otros alucinógenos. En este caso son probablemente otros los receptores que entran en juego, entre ellos los que lo son para el neurotransmisor llamado dopamina.

Aquellos lectores que no pueden comprender del todo cómo estos procesos de los receptores hacen que el cerebro pase de una percepción normal a otro estado de consciencia diferente, están en buena compañía, porque a los especialistas les sucede prácticamente lo mismo. Sin embargo, en los últimos años se han registrado ciertos avances: según estudios recientes los alucinógenos actúan principalmente sobre los lóbulos frontales y el tálamo, la «puerta de la percepción». Los lóbulos frontales y el tálamo se consideran los lugares más probables en los cuales, tras producirse estímulos exteriores, se genera la consciencia y se construye la realidad, de lo cual se puede concluir que no hay una «sede» de la consciencia claramente perfilada en el cerebro. Una explicación sería que los alucinógenos impiden al tálamo seleccionar las informaciones que nos llegan en masa. Las percepciones se agolpan todas ellas entrando sin filtros en los lóbulos frontales y se comportan allí como un saco de pulgas. Un estudio que se publicó en 2002 en la Universidad de Utah (por raro que pueda parecer, procedente de la especialidad de matemáticas) indica que las alucinaciones geométricas visuales, tales como diseños de ajedrez, telas de araña, túneles y espirales podrían producirse por la confusión generada en una zona específica del cerebro que en cualquier otro caso sería responsable del procesamiento de bordes y siluetas.

También las transformaciones de la percepción del tiempo plantean cuestiones interesantes, ya que la forma en que se percibe y se elabora el tiempo en el cerebro humano no está en absoluto investigada de forma concluyente, ni con drogas, ni sin ellas. En todo caso está claro que las regiones del cerebro más influidas por sustancias alucinógenas son al mismo tiempo las que más interesan a los investigadores que estudian la percepción. Si se consiguiera averiguar más sobre el modo en que actúan los alucinógenos sobre la red de conexiones de nuestro cerebro, se estaría probablemente más cerca de dar respuesta a la pregunta planteada sobre cómo a partir de ciertos estados del cerebro se puede llegar a la consciencia.

Sin embargo, no todos los especialistas consideran evidente que la consciencia surja de la actuación conjunta de distintas zonas de un órgano formado por sustancia gris y parecido a un flan.

¿Podría ser el cerebro sólo una especie de televisor, y la consciencia un programa televisivo que existe fuera de ese receptor y con independencia de él? Por otra parte, aquí tampoco puede faltar la consabida teoría de los extraterrestres: según el etnofarmacólogo y filósofo Terence McKenna no debemos esperar que la vida extraterrestre entre en contacto con nosotros con ayuda de los medios técnicos que nosotros tengamos previstos para ello. McKenna afirma que en el caso de los hongos de la psilocibina⁴ se trata más bien de una forma consciente extraterrestre que en contacto con las superficies de los planetas configura un micelio, pero se difunde por el resto de la galaxia en forma de esporas. Quien desee hablar con extraterrestres no necesita ningún costoso radiotelescopio, sino únicamente poner un par de frutos de ese micelio en el té.

Aunque estas teorías parezcan rollos *hippies*, cabe preguntarse por qué las personas que se encuentran bajo el influjo de sustancias alucinógenas se ven invadidas tan a menudo por las mismas ideas de las que ya han hablado las grandes religiones: la unidad mística de Dios y el universo, y la existencia humana como una ilusión. ¿Es que a los seres humanos siempre se les ocurren básicamente sólo ese par de ideas? ¿O es que la religión y las ideas que sugieren las drogas surgen de los mismos procesos en el cerebro? El neurólogo y etólogo Roland Griffiths, a propósito de un estudio sobre experiencias espirituales bajo el influjo de la psilocibina realizado en 2006 en la Universidad Johns Hopkins, dice lo siguiente: «Todavía no hemos recogido datos al respecto, pero hay buenas razones para suponer que las experiencias religiosas profundas se basan en mecanismos similares, con independencia del modo en que se realicen esas experiencias (mediante ayunos, meditación, control de la respiración, privación de sueño, experiencias en la proximidad de la muerte, enfermedades infecciosas o sustancias psicoactivas como la psilocibina). La neurología de la experiencia religiosa se llama actualmente neuroteología y está dando que hablar como nuevo campo de investigación». El estadounidense Herbert Kleber, profesor de psiquiatría y antiguo subdirector del

⁴ Hongos del género *Psilocybe*. (N. de la t.)

Office of National Drug Control Policy (ONDCP), justificó este estudio ante la prensa de la siguiente manera: con anterioridad se había procurado no inducir en la juventud unas ideas absurdas a las que podrían haber contribuido ciertas publicaciones científicas, pero en la era de Internet se dispone de tanta información sobre las drogas y su utilización, que un estudio más o menos no podría producir grandes perjuicios.

Es una suerte que exista Internet. El farmacólogo David E. Nichols, que trabaja en el Heffter Research Institute investigando las aplicaciones médicas de los alucinógenos, comunicó en 1998 lo siguiente: «Una cosa es segura: si conseguimos seguir disponiendo de financiación para nuestra investigación, tenemos en perspectiva los más emocionantes avances en la química médica de las sustancias psicodélicas». En cualquier caso, no faltan voluntarios dispuestos a ser sujetos de experimentación en esta investigación. Y eso es bueno, porque en última instancia hay que pensar también en las ratas de laboratorio, que después de un experimento con alucinógenos no suelen reconocerlo como uno de los acontecimientos más significativos de su vida —a diferencia de más de dos tercios de los sujetos de experimentación que participaron en el estudio de Roland Griffiths.



Capítulo 3

Americanos

América no es una tierra joven. Es vieja, sucia y mala. Eso ya era así antes de los colonos y antes de los indios. El mal está siempre allí y se mantiene al acecho.

WILLIAM S. BURROUGHS, *El almuerzo desnudo*

Desde que aprendieron a caminar en el clima cálido de África Central, los hombres primitivos empezaron a emigrar al mundo frío. Primero fueron a Oriente Próximo y luego, desde allí, a Europa, Siberia y el sudeste asiático, de donde partieron hace 50 000 años para establecerse en Australia. Cómo se produjo exactamente aquel éxodo de amplios espacios, quiénes, cuándo y cómo llegaron, y cómo les fue en tierras extrañas, para todo esto las explicaciones son cualquier cosa menos concordantes, pero eso no es algo por lo que debemos preocuparnos aquí. De cualquier manera, casi todos los investigadores están de acuerdo en que América se pobló muy tarde, a pesar de ser una zona muy bella del planeta. Sin embargo, lo que no está claro es cuándo pisaron el continente los primeros seres humanos; las teorías que actualmente se barajan sitúan este intervalo de tiempo entre hace 60 000 y 11 000 años. Aparte de esto, nadie sabe en qué dirección llegaron los primeros americanos, cómo se las arreglaron en las nuevas tierras y qué fue de ellos finalmente. En cualquier caso, no fue Colón quien descubrió América, sino algún otro. El que llegó tenía probablemente una lanza en la mano.

Hace 10 000 años América tenía un aspecto distinto del que tiene ahora. Después de 20 000 años de grandes fríos, acababa de terminar la última glaciación, cuyos glaciares en su máxima amplitud cubrían aproximadamente toda la masa de tierra de Canadá, una zona de grandes dimensiones. Los drásticos cambios climáticos del final del período glacial marcan también el final de una era geológica, que actualmente llamamos pleistoceno, y el principio de una nueva con un clima mucho mejor, llamada holoceno. En el pleistoceno predominaba en América una mezcla

ecléctica de animales gigantescos. Si uno se fija en los grandes animales de hoy en día y los contempla luego a través de una lente de aumento muy potente, se puede hacer una idea de cómo era la mega fauna de aquel período glacial. De esta manera podemos convertir unos elefantes que son como de juguete en los mastodontes y mamuts de tiempos remotos. Había alces gigantes, tortugas gigantes, castores gigantes, lobos gigantes, tigres con dientes como sables y, por si esto fuera poco, corrían libres por allí unos osos de nariz corta que sacaban varias cabezas a los modernos osos. Se cuenta la anécdota de un ruso al que le enseñaron en el museo de Utah el fémur de uno de esos osos gigantes y, al verlo, preguntó desesperado: «¿Por qué Estados Unidos ha de tener siempre lo más grande de todo lo que hay en este planeta?».

¿Cuándo aparecieron los primeros seres humanos en ese parque temático llamado América? Hasta hace pocos años había una respuesta ampliamente aceptada, que más o menos decía lo siguiente: los primeros americanos, llamados hombres de Clovis, emigraron de Siberia hace unos 12 000 años, desplazándose en dirección este. A causa de los glaciares, el nivel del mar estaba bastante más bajo que ahora, por lo que Rusia y Alaska estaban unidas por un puente de tierra. Así pues, los hombres de Clovis pasaron a América sin mojarse los pies. Por una afortunada coincidencia, el período glacial tocaba justo entonces a su fin. Las masas de los glaciares retrocedieron y dejaron libre un pasillo por el cual los recién llegados atravesaron Canadá para ir hacia el sur. En su tiempo libre se dedicaban a la caza mayor, matando mamuts, bisontes, camellos y caballos, para lo cual utilizaban unas lanzas hechas de una manera característica, de tal modo que sus puntas de piedra se convirtieron en la marca que dejaron aquellos primeros americanos.

Clovis es el nombre de la ciudad de Nuevo México donde los arqueólogos descubrieron las primeras de estas puntas de lanza en la década de 1930. Pronto se puso de manifiesto que cientos de puntas de lanza del mismo tipo estaban diseminadas por toda Norteamérica.

Al parecer, las puntas de lanza de Clovis y los cazadores correspondientes se extendieron rápidamente por todo el continente, todo lo «rápidamente» que se podía entonces, porque hablamos de una época en que se carecía de cosas tan básicas como carreteras y líneas ferroviarias. En menos de mil años, según esta

teoría, conquistaron los hombres de Clovis el continente, desplazándose desde Alaska hasta el extremo sur, hacia la Tierra de Fuego. Cada generación avanzó unos quinientos kilómetros más en dirección sur y al mismo tiempo se iban reproduciendo como es debido. En el nuevo continente había comida abundante, aunque por desgracia solía estar provista de grandes dientes. Coincidiendo con la expedición conquistadora de los hombres de Clovis, se produjo en América una extraña mortandad que hizo desaparecer todos los animales anteriormente mencionados. En la tradicional teoría de Clovis, la aparición de los seres humanos y la desaparición de los animales están estrechamente ligadas: en una especie de guerra relámpago avanzaron los hombres de Clovis a través de aquellas tierras y exterminaron en su camino toda la megafauna americana, de tal modo que al final sólo quedaron unos pequeños y bonitos animales.

Esta historia resulta espectacular, pero tiene enormes probabilidades de estar equivocada. El hallazgo decisivo que llevó finalmente a su refutación procede del sur de Chile, concretamente de un lugar llamado Monte Verde, donde el estadounidense Tom Dillehay y su equipo realizan excavaciones desde la década de 1970. Lo que sacaron a la superficie resultó revolucionario: lugares donde se había hecho fuego, vestigios de una especie de asentamiento, restos antiquísimos de carne de mastodonte, y herramientas fabricadas por la mano del hombre, que no sólo tenían el mismo aspecto que las ya conocidas como pertenecientes a los hombres de Clovis, sino que además tenían unos 12 500 años. Incluso se encontró la huella de un pie humano, lo cual es normalmente una prueba inequívoca de la presencia de seres humanos. Según el paradigma de Clovis, los hombres primitivos de Monte Verde tendrían que haber ido a América por el brazo de tierra antes de lo que se había supuesto, en una época en la que grandes zonas de Norteamérica se encontraban todavía cubiertas por el hielo. Los inmigrantes llegaron sólo hasta la zona de lo que actualmente es Fairbanks y luego tropezaron con el insalvable glaciar. Tuvo que haber además otro camino hasta Monte Verde, y por lo tanto una colonización anterior a la de Clovis.

Monte Verde no fue el primer lugar de América en el que se encontraron vestigios de los predecesores de los hombres de Clovis, pero en todos los demás casos los arqueólogos no llegaron a dar una opinión que fuera en alguna medida unitaria. Por

ejemplo, hubo décadas de controversias sobre la excavación de Meadowcroft, una vivienda de trogloditas en Pensilvania, en la que James Adovasio y sus colegas descubrieron puntas de lanza y otros útiles en la década de 1970. La antigüedad de estos objetos se estimó en parte en unos 16 000 años, por lo que eran claramente anteriores al umbral de la época de Clovis. Durante muchos años Adovasio intentó en vano convencer a los que criticaban su hallazgo. Sobre todo los datos cronológicos fueron considerados muy discutibles. Los métodos arqueológicos más importantes para la determinación de la antigüedad se basan en el contenido de carbono 14, un isótopo del carbono que es radiactivo y se desintegra con el paso del tiempo. Por lo tanto, la cantidad de carbono 14 que todavía exista en la actualidad puede utilizarse como una especie de reloj, siempre que se consiga mantener el control sobre todos los átomos en unos objetos desenterrados que tienen milenios de antigüedad.

Por ejemplo, es necesario asegurarse de que los huesos antiguos no se hayan contaminado en los vericuetos de su manipulación por haber estado en contacto con átomos de carbono más jóvenes.

A pesar de todos los problemas, en el caso de Monte Verde los especialistas llegaron a un acuerdo después de un debate que duró veinte años y aceptaron en general los datos relativos a los precursores de los hombres de Clovis. En 1997 un selecto consorcio de expertos, llamados por Adovasio la «paleopolicía», hizo un control del lugar de las excavaciones y confirmó las conclusiones de Dillhay. En palabras de otro experto del mismo grupo llamado David Meltzer y procedente de Dallas: «Monte Verde fue el punto de inflexión. Se había superado el listón de Clovis». Desde entonces reina de nuevo el desconocimiento sobre la historia de América.

Los hallazgos del sur de Chile no son el único problema de la teoría de Clovis. También es al menos dudoso que los primeros americanos fueran realmente capaces de eliminar la megafauna.

Hoy en día está más extendida la hipótesis de que, o bien los cambios climáticos extremos de finales del período glacial, o agentes patógenos traídos de fuera, ayudaran a los hombres primitivos en su lucha contra los grandes animales peludos. Otras reflexiones contrarias al paradigma de Clovis provienen de los lingüistas, que se quejan desde hace mucho tiempo de que 12 000 años no son suficientes para

desarrollar a partir del lenguaje de los hombres de Clovis las casi novecientas lenguas indias que se hablaban en América en tiempos de Colón. O bien se pobló América mucho antes, o la colonizaron distintos pueblos que fueron llegando sucesivamente. Los expertos en genética han llegado a unas conclusiones similares: en el caso de ciertos genes se sabe aproximadamente cada cuánto tiempo sufren una mutación.

Así, por comparación de los genes de norteamericanos y asiáticos se puede calcular el tiempo transcurrido desde que se separaron. También aquí se han obtenido datos que apuntan a una colonización de América hace entre 15 000 y 30 000 años, lo cual coincide con lo que las demás ramas de la ciencia han afirmado.

Sin embargo, en arqueología los argumentos decisivos vienen siempre del suelo. Por eso muchos se alegraron cuando en las últimas décadas los descubrimientos de Monte Verde quedaron confirmados por excavaciones realizadas en lugares con nombres tales como Cactus Hill, Topper y Taima-Taima, diseminados por toda América. Tras superarse por primera vez el listón de Clovis, el proceso no se detuvo ya: algunos de los nuevos hallazgos tienen posiblemente una antigüedad de entre 30 000 y 50 000 años. Está garantizado que en un futuro cercano continuarán las discusiones entre los arqueólogos americanos.

En el centro de estas discusiones figura desde el verano de 1996 el esqueleto de un hombre que vivió hace unos 9000 años en el noroeste de Estados Unidos, cerca de la ciudad de Kennewick. El «hombre de Kennewick» es casi el doble de antiguo que el hombre europeo del período glacial, llamado «Otzi». El pobre hombre tuvo seguramente una vida muy dura: sobrevivió a varias heridas en el cráneo y en las costillas, y llevó una punta de lanza dentro de su cadera. La representación de su rostro, que se publicó en los periódicos, tiene para los legos en la materia un aspecto sospechosamente europeo; en realidad tenía un mayor parecido con los primitivos habitantes del Japón. Lo cierto es que en absoluto tuvo parecido alguno con los indios modernos, que gustan de definirse a sí mismos como los «primeros americanos». Cualquiera que fuese su procedencia, se encontraba en (o más bien bajo) la tierra que en otro tiempo poblaron los indios umatilla y, como dice la creencia popular de éstos, «desde el principio de los tiempos» y no sólo desde hace unos pocos milenios, como sostienen los científicos. Una ley estadounidense

reconoce ahora el derecho de los indios a enterrar los restos de sus antepasados y además, si lo desean, sin dejar que los huesos sean previamente examinados por cualquier procedimiento. Los indios dicen que, si el hombre de Kennewick vivía en el territorio de los umatilla, entonces era un umatilla. Los arqueólogos replican que es altamente dudoso que ese esqueleto primitivo tenga algo que ver con los indios modernos, y les gustaría más tenerlo en el laboratorio que enterrado. Desde hace más de diez años, el señor Kennewick está tirado por ahí sin hacer nada y esperando a que acabe la vista de la causa. En última instancia, aquí no se trata de un par de huesos viejos, ni de peleas académicas, sino de la cuestión que plantea a quién pertenece América.

¿Quiénes eran pues los primitivos habitantes de América, y de dónde vinieron? Dado que las dudas sobre el paradigma de Clovis son inmensas, se discuten múltiples variantes. La mejor solución sería, por supuesto, que los habitantes de América descendieran de extraterrestres, pero rara vez se pronuncian los arqueólogos serios sobre esta posibilidad. Una teoría popular sostiene que la colonización de América fue un prolongado viaje en barco a lo largo de la costa del Pacífico; o, como lo llama Adovasio «el “club náutico” de finales del pleistoceno». Es posible que los japoneses primitivos durante el período glacial se desplazaran hacia América sin darse cuenta, en canoa o con embarcaciones de vela en dirección norte hasta el puente de tierra de lo que hoy es el estrecho de Bering, y siguieran después la costa americana hasta zonas muy lejanas del sur. Algo que les ayudó fue el hecho de que el continente nunca estuvo totalmente cubierto de hielo, porque en la costa hubo siempre una estrecha franja donde no lo había, y por la que supuestamente también los osos pardos avanzaron hacia el sur. Además, casi toda la línea costera estaba llena de selvas de algas, lo cual, por una parte, atraía animales marinos que podían servir como provisiones para el viaje, y por otra parte hacía que el mar estuviera menos agitado, y así el viaje por mar resultaba más fácil. Esta teoría de las barcas tenía otra ventaja: no era preciso hacer una marcha agotadora por desiertos y selvas y estar matando continuamente animales gigantes. Al final, es posible que la colonización de América fuera incluso placentera. Por desgracia, la teoría del club náutico es difícil de demostrar, porque todos los posibles lugares de asentamiento de la costa están actualmente inundados

a causa del ascenso del nivel del mar. Además, algunos antiguos yacimientos arqueológicos están al otro lado de Norteamérica y, para llegar allí, los hombres primitivos hubieran tenido que caminar a marchas forzadas.

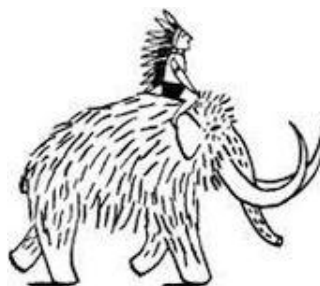
Otra idea, que desde hace años pulula por las revistas de arqueología, suena mucho más espectacular que un tranquilo viaje en velero a lo largo de la costa Oeste americana.

Evidentemente las puntas de lanza de Clovis se parecen a las que fabricaban los europeos de la cultura solutrense, al menos esto dicen algunos científicos, como Dennis Stanford del Smithsonian Institute. Los hombres de la cultura solutrense vivieron hace unos 20 000 años en las costas del sur de Europa y, según la teoría, zarparon de allí en pequeñas embarcaciones rumbo a América, una aventura extraordinariamente arriesgada para aquella época. Es verdad que los seres humanos habían utilizado embarcaciones desde mucho antes, pero ¿cómo se atrevieron a cruzar todo un océano, con vientos, mareo, tiburones y todos los demás imponderables posibles? ¿Será cierto que los americanos descienden de los europeos? Es ésta una hipótesis que en el ámbito de los especialistas en la materia suscita posturas dispares; algunos, aunque lo dicen con todas las precauciones, consideran que es una bobada, mientras que para otros es plausible. Por supuesto, la colonización de América podría haberse realizado en varias oleadas sucesivas, primero mediante botes desde Asia, luego en barcos desde Europa, y más tarde a pie desde Asia, o al revés, o de cualquier otra manera totalmente diferente.

Si se puede atravesar el Atlántico, entonces también será posible hacerlo con otros océanos.

En consecuencia, algunos científicos plantean otras situaciones en las que los primitivos americanos cruzarían el Pacífico desde Asia o desde Australia para llegar a su nuevo hogar. Para averiguar si las aventuras náuticas a gran escala vienen a cuento cuando se habla de hombres de la edad de piedra, el aventurero noruego Thor Heyerdahl, entre otros, navegó a vela en 1947 con la primitiva balsa *Kon-Tiki* desde Sudamérica a Oceanía. Al fin y al cabo, lo único que prueba esto es que Thor Heyerdahl puede atravesar los mares del mundo con unos barcos grotescos, pero sobre los hombres primitivos dice más bien poco. Porque no todo lo que es factible se ha llevado a cabo realmente. En cualquier caso, hay muchas posibilidades para

explicar la colonización de América. Finalmente, acaso vinieron de la Antártida, porque está claro que allí hace mucho frío, y además está a oscuras durante la mitad del año. ¿Quién no emigraría en tales circunstancias?



Capítulo 4

Anestesia

¿Podría usted decirme cuál de estos dos pañuelos huele más fuerte a cloroformo?
SAM & MAX, *Freelance Police*

En la práctica, para que algo funcione, afortunadamente no es imprescindible saber por qué funciona. Si no fuera así, mucha gente no podría utilizar un bolígrafo. Desde hace ciento cincuenta años se sabe que las anestésicas tienen un funcionamiento fiable, pero nadie sabe exactamente cómo funcionan (no siga usted leyendo esto si ha de ir dentro de poco tiempo al quirófano). Está claro que una anestesia total actúa sobre la médula espinal, el tronco encefálico y la corteza cerebral, produciendo así un estado de inconsciencia, insensibilidad al dolor y relajación muscular, del que los pacientes generalmente no pueden recordar nada cuando ha finalizado. Por lo que respecta al ajuste de la anestesia a los efectos deseados, y a la reducción de los efectos secundarios no deseados, se ha llegado muy lejos, aunque siga faltando una explicación rigurosa de lo que es la anestesia. Desde luego nos alegramos de que no sea al revés, pero hoy en día estamos como hace cien años ante las siguientes preguntas: ¿Cómo y en qué lugar de la célula actúan las sustancias anestésicas? ¿Cómo es posible que las sustancias más diferentes que uno pueda imaginar, tengan en el cuerpo unos efectos relativamente uniformes? ¿Y cómo se producen, a partir de esos efectos sobre las células, esas complejas consecuencias en la consciencia?

Hay una gran variedad de sustancias anestésicas: desde un simple gas noble hasta una intrincada maraña de moléculas, tenemos de todo. Debido a que la estructura química o física de estas sustancias es tan variada, prácticamente no puede haber unos receptores específicos para ellas. Sin embargo, esto no significa que no haya en absoluto características comunes. Hacia el año 1900, un farmacólogo de Marburgo llamado Hans Horst Meyer y el catedrático no numerario Charles Ernest Overton, de Zurich, descubrieron más o menos al mismo tiempo una relación curiosa: cuanto más liposoluble es un anestésico, más fuerte es su efecto. La

Llamada hipótesis de Meyer-Overton sostiene que toda sustancia liposoluble produce un efecto narcótico en las células de los seres vivos, especialmente en las células nerviosas, o neuronas, porque, como escribió Meyer, «la estructura química de éstas administra las sustancias parecidas a las grasas». Se supone que los anestésicos se disuelven en la membrana de las neuronas y cambian así sus características. No estaba claro cómo funcionaba esto, por lo que tuvieron que pasar unas cuantas décadas de intentos estériles hasta que se demostró esta «teoría de los lípidos relativa a la anestesia». Desde la década de 1970 quedó sustituida por la «teoría de las proteínas», según la cual los anestésicos afectan a las proteínas, es decir a las albúminas, en la membrana de las neuronas. Simplificando las cosas se puede decir que se interrumpe la comunicación entre las neuronas, de tal modo que éstas, por ejemplo, en vez de transmitir una sensación dolorosa debidamente dirigida, se quedan murmurando « ¿Qué? ¿Cómo era esto?». La teoría de las proteínas está bien apoyada en sus fundamentos básicos. Pero ¿qué es exactamente lo que hacen los anestésicos con las proteínas?

Hace pocos años, el farmacólogo alemán Uwe Rudolph, mediante experimentos realizados con cobayas genéticamente modificadas, consiguió demostrar que algunos anestésicos actúan de manera específica sobre determinados canales iónicos (una especie de válvulas que están en la membrana de las neuronas). Una sustancia que lleva el bonito nombre de ácido gammaaminobutírico (AGAB)⁵ gobierna la apertura y el cierre de esas válvulas y puede así paralizar la transmisión de señales nerviosas. La «hipótesis del AGAB» formulada por Rudolph afirma que los anestésicos se fijan en los mismos lugares que son utilizados habitualmente por el AGAB, por lo que pueden propiciar la misma paralización y la misma desconexión controlada del cerebro (de un modo que aún está por investigar). De todos modos, existe una multiplicidad de distintos receptores AGAB que realizan funciones diferentes y que pueden verse influidos de distinta manera por las sustancias en cuestión, de tal modo que hasta que se aclare definitivamente la cuestión del canal iónico, todavía habrá que agotar a un par de doctorandos más. Además, la hipótesis del AGAB se refiere sólo a unos cuantos anestésicos concretos que se inyectan

⁵ Al principio este nombre se utilizó en alemán: Gamma-Aminobuttersäure, pero las siglas correspondientes a su traducción al inglés, GABA, se utilizan frecuentemente también en textos en alemán y, por lo que respecta al castellano, unos escriben AGAB y otros GABA. (N. de la t.)

directamente por vía sanguínea: sobre los puntos en que inciden las sustancias que se toman por inhalación, siempre se ha sabido poco. Lo que sí se sabe es que producen algún efecto en muchas proteínas de las células, pero es tema de debate si este bombardeo superficial provoca los efectos de la anestesia, o si una gran parte de las proteínas afectadas es totalmente irrelevante con respecto a dichos efectos.

Con los nuevos modelos explicativos de la última década llegó también el final de la «hipótesis unitaria sobre la anestesia», cuya idea es que todos los anestésicos funcionan esencialmente de la misma manera. Ahora se sabe ya que el efecto de algunos anestésicos se basa en múltiples procesos que son independientes unos de otros, mientras que otras sustancias sólo parecen intervenir en determinados lugares. Por lo tanto, como tantas veces sucede, todo está dispuesto de un modo muy desordenado.

A pesar de sus diversas formas de proceder en el cerebro, los efectos de las distintas sustancias se parecen de una manera sospechosa. A veces lo importante es más bien la supresión del dolor, pero otras veces es más la profunda pérdida de consciencia o la relajación muscular, pero no se responde a la pregunta fundamental: cuál es la causa de que durante una anestesia las funciones corporales se vayan desconectando en el mismo orden en que lo hacen cuando uno se duerme o tiene un desmayo. En teoría se podría pensar que primero falla el sentido del gusto, luego la pierna derecha y finalmente el pensamiento lógico, aunque la capacidad de quejarse al médico se mantiene durante todo el tiempo. Overton ya había observado que la anestesia se parece tanto al sueño «que sin querer nos sentimos inducidos a preguntar si el sueño natural no podría estar causado por alguna sustancia de efecto narcótico producida por el propio organismo». Hasta ahora no se ha conseguido aislar ninguna sustancia de este tipo, pero bien puede ser que los anestésicos se limiten a poner en marcha un mecanismo, ya existente, surgido de forma evolutiva.

Las reflexiones sobre la clase de mecanismo que podría ser éste, no están por ahora mucho más avanzadas que en tiempos de Overton.

Los primeros indicios de lo que la anestesia hace realmente en el cerebro se han descubierto en los últimos años mediante distintos experimentos en los que se han

aplicado procedimientos tales como la tomografía por resonancia magnética y por emisión de positrones, en la que se obtienen imágenes del interior del cerebro sin tener que cortarlo antes en filetes. Si se coloca a una persona en un tomógrafo y se le suministra una anestesia, se puede observar en qué orden cronológico decaen las funciones cerebrales y en qué zonas del cerebro se reduce más el metabolismo. Así se observa también que los distintos anestésicos frenan la actividad de zonas diferentes del cerebro: el gas narcótico Halothan, por ejemplo, disminuye sobre todo la actividad del tálamo y del mesencéfalo, las zonas de conexión que distribuyen las informaciones a los elaboradores competentes que están en la corteza cerebral. En cambio el práctico anestésico inyectable Propofol actúa directamente y con mayor intensidad en la propia corteza cerebral. Parece, por lo tanto, que hay más de una posibilidad de desconectar la consciencia de forma controlada y, dado que los experimentos son relativamente nuevos, lo más que se puede hacer es formular una prudente hipótesis sobre el modo en que este metabolismo reducido de algunas partes del cerebro puede llevar a la pérdida de consciencia, la insensibilidad al dolor y a la amnesia.

El hecho de que la investigación haya progresado en ciento cincuenta años de una manera poco apreciable no se debe (o, al menos, no en primer lugar) a que los anestesiólogos dediquen demasiado tiempo a jugar al golf y demasiado poco a investigar. En realidad serían los especialistas de otras disciplinas quienes tendrían que averiguar primero cómo funcionan la consciencia, el dolor y el sueño. Pero quizá sea al revés y las investigaciones relativas a la cuestión de cómo perdemos la consciencia nos ayuden a explicar qué es esa consciencia. Ya veremos quien llega más rápido.



Capítulo 5

Anguila

ANGUILA (n.) Dícese de quien cambia su nombre para estar lejos del frente.

DOUGLAS ADAMS, *El sentido de la vida*

Las anguilas han conseguido siempre, desde hace siglos, mantener fuera de nuestra vista sus condiciones de vida. Sin embargo, todos sabemos que es posible verlas en muchos lugares (al menos ahumadas), y tampoco falta una gran cantidad de ambiciosos investigadores que se dedican a observar estos animales. Aristóteles, por ejemplo, se interesó mucho por estos peces, que en su época aún no se consideraban como tales, sino como una especie de gusanos que, según creía Aristóteles, se deslizaban por el barro del lecho de los ríos. Hasta muy avanzada la edad moderna estuvieron circulando otras teorías que no eran menos absurdas; así, en 1858 todavía se decía que las anguilas, para reproducirse, se enrollaban en forma de huso alrededor de un tallo de junco y se excitaban con el balanceo de esta planta. De todos modos, lo que sí se conocía desde tiempos remotos era la migración de la anguila: las anguilas adultas nadan bajando los ríos hasta el mar, y luego las crías vienen desde el mar, lo que induce a pensar que la reproducción se produce en aguas marinas. Dónde, cuándo y cómo sucede esto, son las preguntas a las que se enfrentan todos los que se interesan por las anguilas.

Las investigaciones sobre las anguilas han ido avanzando lentamente durante los últimos trescientos años. En 1777, el italiano Carlo Mondini descubrió los ovarios de la anguila y así supo que la anguila hembra, como cualquier otro pez razonable, pone huevos para contribuir a la conservación de su especie. Pasaron casi cien años hasta que se descubrieron los órganos sexuales masculinos. El biólogo Simon von Syrski, que era de Trieste, descubrió dos pequeños órganos en forma de lóbulo y los identificó acertadamente como los testículos de la anguila, pero surgió un misterio para los investigadores de entonces: no contenían tipo alguno de esperma. En la misma época, Sigmund Freud, que era todavía estudiante de zoología, estudió también los testículos de la anguila. Trabajando prácticamente a destajo, Freud hizo

la disección de unas cuatrocientas anguilas en busca de los órganos sexuales masculinos. Algunos creen que así hacía frente a sus propios problemas sexuales: al matar a las anguilas, que tienen forma fálica, Freud castraba simbólicamente no sólo a sus competidores, sino también (cuatrocientas veces) a su propio padre, convirtiéndose así la anguila en una víctima inocente del complejo de Edipo. Ahora bien, con su masacre de anguilas no le aportó Freud ningún nuevo conocimiento a la zoología.

Hacia finales del siglo XIX, se dio un importante paso hacia delante. Los biólogos Yves Delage y Giovanni Batista Grassi demostraron de forma concluyente que un ser marino plano y transparente, llamado *Leptocephalus brevirostris* y considerado hasta entonces como una especie independiente, no era sino la larva de la anguila de río. La explicación del origen de esas larvas se la tenemos que agradecer sobre todo al danés Johannes Schmidt. Durante las tres primeras décadas del siglo XX, este científico emprendió costosas expediciones que le llevaron a navegar por el Atlántico. Schmidt siguió el camino de las pequeñas anguilas en sentido contrario; fue cada vez más lejos hacia el continente americano y encontró larvas que eran cada vez más pequeñas, para dar finalmente con las más pequeñas en el mar de los Sargazos, al sur de las islas Bermudas. Esta zona de alta mar, conocida también como Triángulo de las Bermudas y tristemente famosa por los misteriosos naufragios y caídas de aviones, se considera desde entonces el lugar de nacimiento de las anguilas de río europeas. Lo raro es que nadie haya intentado nunca establecer una relación entre los desastres navales y la reproducción de las anguilas.

He aquí todo lo que creemos saber actualmente sobre la vida de la anguila: tras salir de los huevos en el mar de los Sargazos, las larvas de la anguila europea flotan y nadan a lo largo de la costa americana hacia el norte y cambian finalmente la dirección de su avance para seguir la corriente del Golfo hacia Europa. Al principio van acompañadas por las larvas de la anguila americana, que proceden de la misma zona, pero que luego prefieren renunciar a la fatigosa travesía del Atlántico. No se sabe por qué esas larvas europeas, en vez de quedarse en América, emprenden una penosa travesía del Atlántico, que dura varios años. Al llegar a las costas europeas, las larvas se convierten en las llamadas angulas, aunque no está claro qué es lo que

desencadena esa metamorfosis; puede que sea el alivio de ver por fin tierra otra vez. Las angulas son pequeños animales transparentes, parecidos a los gusanos, que están considerados como un bocado exquisito y se pescan en grandes cantidades. La experiencia de haber realizado una travesía de varios años por el Atlántico para acabar desembarcando en el plato de un refinado restaurante se podría calificar con toda seguridad de poco ecológica y decepcionante.

Todas las angulas supervivientes, y aquí comienza la parte de la vida de este animal que es conocida desde hace mucho tiempo, se desarrollan hasta llegar a su forma adulta, llamada anguila, y este proceso tiene lugar en los ríos de agua dulce de Europa. (Lo mismo sucede al otro lado del Atlántico con las pequeñas anguilas americanas). En este punto podemos saltarnos muchos años de la vida de la anguila, porque durante ellos no sucede nada especial. La anguila adulta vive como pez entre muchos otros peces, aunque algunas acaban entretanto ahumadas, pero a las que se libran de este destino les llama una voz misteriosa para que vuelvan a dirigirse al mar cuando tienen cinco, diez e incluso veinte años de edad. En su camino hacia la costa no hay quien las pare, ni los diques, ni la tierra. Para lo que las anguilas no han encontrado medio alguno es para protegerse de las turbinas de las centrales hidráulicas, que periódicamente las convierten en bocaditos de pescado. Durante el viaje hacia el océano se produce una interesante transformación: la piel de la anguila se vuelve plateada, sus ojos se agrandan y, lo más decisivo, el tracto digestivo se atrofia. Cuando la llamada anguila plateada llega al mar, su destino está ya predeterminado: se encuentra realizando una misión suicida cuya duración depende de las reservas de grasa que tenga en el cuerpo.

A continuación empieza la parte misteriosa de la vida de las anguilas. El experto en estos peces Friedrich-Wilhelm Tesch siguió a las anguilas hasta la dorsal Atlántica, pero luego se quedaron sin carga las baterías de los emisores con que las había equipado. Otras expediciones descubrieron anguilas plateadas en el mar de los Sargazos, pero pronto las perdieron de vista. Al parecer estos animales consiguen de algún modo regresar a su lugar de origen. De camino hacia allí, producen los machos por fin el esperma que Sigmund Freud y otros buscaron tan afanosamente. Tras llegar a su destino, las hembras tendrían que realizar el desove de la freza que las anguilas macho fecundarían, para que la reproducción se llevara a cabo

finalmente. En cualquier caso, ésta es sólo la teoría, ya que, a pesar de los enormes esfuerzos realizados, nadie ha podido observar todavía este importante proceso en el medio natural. Además, no está claro qué les sucede después a los progenitores, que en realidad no pueden hacer otra cosa que morir de hambre poco después de la travesía del Atlántico. Sin embargo, nunca hasta ahora se han encontrado sus esqueletos, y no se sabe de la existencia de cementerios de anguilas.

Hoy en día nadie cree ya que las anguilas salen simplemente del barro. Todas las investigaciones sobre la reproducción animal indican que debe existir una cierta proximidad en el espacio entre los padres y los hijos, por lo menos al principio. Sin embargo, en el caso de las anguilas no hay prueba de que se establezca relación alguna entre las distintas generaciones. Al parecer, las pequeñas larvas de anguila surgen de la nada. Al mismo tiempo los adultos desaparecen sin dejar rastro en el mar de los Sargazos. ¿Acaso los progenitores se convierten de nuevo sencillamente en larvas? ¿Es la anguila, por lo tanto, inmortal? Esto es un gran misterio. Por otra parte, en el caso de las anguilas asiáticas se plantea exactamente el mismo problema. Se conoce en muchos casos el origen de las anguilas, se puede deducir cómo llegan las larvas a las corrientes de agua dulce, se sabe que los adultos vuelven al mar y a los lugares de desove, pero queda por conocer el último paso, el que constituye en realidad la reproducción, así como la relación entre madre, padre y larvas. Los investigadores tienen que sentirse como esos niños pequeños que, aunque saben que la hipótesis de la cigüeña es falsa, no tienen, sin embargo, idea alguna sobre la procedencia de los bebés.

Una solución del problema que durante mucho tiempo ha estado ampliamente difundida pone en duda la existencia de las anguilas europeas. El zoólogo británico Denys W. Tucker aventuraba en 1959 la teoría de que el camino de regreso al mar de los Sargazos era demasiado largo, por lo que ninguna anguila que viviera en Europa conseguiría volver a los lugares de desove. En lugar de esto, las anguilas de Europa tendrían que descender de sus colegas americanas, que en cualquier caso sí se reproducirían en el mar de los Sargazos. Aunque esta hipótesis fue descartada tras largas discusiones (las anguilas europeas y las americanas muestran diferencias genéticas claras y, por lo tanto, han de ser consideradas como dos especies distintas), dejó tras de sí interesantes conclusiones: un grupo de

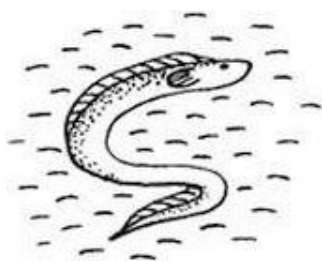
«ariósofos» partidarios de la tesis que sostiene que los arios proceden de la Atlántida, dedujo a partir de la teoría de Tucker que las anguilas antiguamente se adentraban por los ríos de la Atlántida, y no por los de Europa. Habría sido después del hundimiento de la Atlántida cuando las larvas llegaron por primera vez a Europa, pero nunca habrían podido acostumbrarse a hacer un viaje de regreso que sería el doble de la distancia habitual. En todo caso, esto es lo que decían los ariósofos, que querían recuperar su patria, la Atlántida, con ayuda de las anguilas, una audaz empresa aún más estéril que la búsqueda de las anguilas cuando se están reproduciendo.

Para averiguar si las anguilas tienen la capacidad de atravesar el Atlántico, un equipo holandés de investigadores organizó hace poco una prueba de natación: pusieron un grupo de anguilas a nadar en círculos en un depósito de agua durante medio año, sin alimento, sin pausas para la publicidad y sin bebidas energéticas. Aunque estas anguilas perdieron una quinta parte de su peso corporal, recorrieron, en una hazaña asombrosa, una distancia maratonia de 5.500 km.

Después de esta penosa prueba, en vez de subir al podio de los vencedores, los pobres animales acabaron en una mesa de disección. No se sabe cómo consiguen nadar durante tanto tiempo, pero con este experimento parece probado que pueden hacerlo. Además, hay quien dice que las anguilas pueden dejarse guiar por el campo magnético terrestre, lo cual les daría la posibilidad de orientarse en el mar. También se ha conseguido ya observar cómo las anguilas fecundan sus huevos en cautiverio, pero nunca ha sido posible verlo cuando están en libertad. Por otra parte, un nuevo trabajo de Tsukamoto y sus colegas japoneses demuestra que unas anguilas capturadas en el Atlántico habían vivido siempre allí, lo cual es todo un misterio, porque pone en entredicho la teoría de las migraciones de anguilas que hemos expuesto anteriormente. Por otra parte, han surgido unas ligeras dudas sobre si es correcto decir que todas las anguilas europeas de agua dulce pertenecen genéticamente a la misma especie, y si por consiguiente pueden aparearse entre sí y perseguir los mismos objetivos vitales que desde hace cien años se han dado por supuestos, incluidas las zonas de desove.

Desde siempre hay también mucho campo para los mitos sobre la fertilidad, las divinidades con forma de anguila y las especulaciones sobre la telequinesia en los

peces. La investigación sobre las anguilas podría mantener su impulso en el futuro a causa de la también misteriosa disminución de la población de angulas. Cada vez llegan menos anguilas jóvenes a las costas de Europa, lo cual podría deberse a la acción de los parásitos, al calentamiento de los mares, a la contaminación medioambiental y a otras cosas totalmente diferentes. Pero, dado que las angulas son un factor de la economía, hay esperanza de salvación. Aunque de las anguilas se podría esperar incluso que un día desaparecieran de la Tierra sin motivo aparente.



Capítulo 6

Bostezo

Cuando una persona está cansada y falta de sueño, el bostezo es una expresión espontánea que no resulta desagradable. [...] Cuando el cansancio es auténtico, la persona se rinde y se acuesta.

DOCTORA ANNA FISCHER-DÜCKELMANN,
La mujer como médico de familia

Hagamos una precisión anticipada: el bostezo no se debe forzosa y exclusivamente al aburrimiento o al cansancio. En los juegos olímpicos se observa a menudo que algunos deportistas de elite inmediatamente antes de la prueba decisiva bostezan ampliamente. Lo que está garantizado es que correr en competición ante millones de espectadores es todo menos aburrido, al menos eso se dice. Lo único en que coinciden todos los expertos es en que el bostezo es contagioso. Pero no se sabe por qué esto es así, ni qué es el bostezo. Todos aquellos lectores a los que la exposición que vamos a realizar aquí les produzca un bostezo, deberán considerarse como sujetos de un experimento en el marco de una investigación de ámbito mundial.

Bostezar, o más en general abrir la boca ocasionalmente, es una costumbre muy difundida. Se ha observado en muchos vertebrados, entre otros en peces, aves, serpientes, elefantes e innumerables animales de otras especies. Las personas comienzan a hacerlo cuando todavía están en el vientre materno, en una edad en que casi no pueden hacer otra cosa. El proceso de bostezar es desencadenado evidentemente por un cóctel de hormonas a merced de las cuales nos encontramos todos. Según algunos estudios, las personas que padecen enfermedades tales como la esquizofrenia bostezan considerablemente menos de lo normal, mientras que otras patologías inducen una tasa de bostezos elevada. Las personas sanas bostezan con mayor frecuencia poco después de despertarse y poco antes de dormirse. Otros seres vivos, como las ratas macho, tienen a veces una erección

mientras bostezan. Hay quienes bostezan preferentemente cuando están en compañía, al menos eso se dice en algunos textos. Entre los monos, el bostezo acompañado de un gesto de mostrar los dientes se considera una amenaza, cosa que llamó la atención de Darwin.

Como ya hemos mencionado, el bostezo es contagioso, no sólo entre los seres humanos, sino al menos también entre los chimpancés y los macacos, pero no entre los bebés. Los hombres bostezan probablemente con mayor frecuencia que las mujeres y, en casos excepcionales, un bostezo demasiado amplio puede producir una dislocación de la articulación de la mandíbula. En cualquier caso, la investigación sobre el bostezo rara vez es aburrida.

A menudo se dice que la causa de los bostezos es una escasez de oxígeno. Según la creencia popular, parece ser que las personas bostezan cuando están en espacios mal ventilados, con el fin de aspirar más aire al abrir ampliamente la boca. Esta idea es quizá demasiado simple para ser verdad. En todo caso, estas sencillas reflexiones sobre lo que es plausible suscitan ligeras dudas: ¿por qué bostezan los leones cuando están echados en la sabana? ¿Es que allí les falta oxígeno?

¿Por qué bostezan los no nacidos en el vientre materno, donde reciben oxígeno a través del cordón umbilical (y no por la boca)? ¿Por qué no se bosteza con mucha mayor frecuencia durante las actividades deportivas, en las que se consume mucho oxígeno? Esto tendría que comprobarlo la ciencia, y es precisamente lo que hizo el psicólogo Eichard Provine con sus colaboradores a finales de la década de 1980: los investigadores midieron si un aporte aumentado de oxígeno reducía la frecuencia del bostezo, y si un suministro de aire «viciado», es decir enriquecido con dióxido de carbono, induce más bostezos. Finalmente, comprobaron también si la actividad deportiva provoca bostezos. El resultado fue en todos los casos claramente negativo: el «aire viciado» y la escasez de oxígeno pueden ser causa de cansancio, pero probablemente no son los desencadenantes del bostezo. Así es como hay que investigar: teoría, experimento, refutación de la teoría, y listo.

Provine, que califica de «perverso» su interés por los bostezos, se mostraba sorprendido por los resultados. Desde entonces cree que el bostezo no tiene nada que ver con la respiración, sino con un cambio del estado de vigilia que se produce cuando las personas están cansadas o muy animadas. Su colega Ronald Baenninger

aporta algunos datos justificativos de la segunda hipótesis. Pidió a las personas sometidas a la prueba que llevaran durante todo el día unas pulseras sensibles al movimiento y que comunicaran cada bostezo a este aparato pulsando un botón. El resultado fue que los bostezos precedían generalmente a las fases de gran actividad.

Posiblemente algunos animales se comportan de una manera parecida: por ejemplo, entre los peces combatientes siameses, los machos abren la boca antes de atacar, en un gesto que al menos recuerda al bostezo, aunque represente también algo totalmente diferente. El problema es que los combatientes siameses no se dejan poner una pulsera.

Quizá este caso aclare la extraña frecuencia del bostezo en los corredores de maratón poco antes de la salida, o en los paracaidistas antes del salto, y en los estudiantes antes de un examen, cuando, sin embargo, están realmente despiertos. Puede ser que al abrir la boca se bombee al cerebro un aporte suplementario de sangre y de esta manera aumente la concentración. Si esto fuera cierto, el bostezo sería una especie de arranque de motores en el cerebro. Pero, entonces, ¿por qué bostezamos también cuando nos cansamos, como cada uno de nosotros podrá corroborar?

¿Es el bostezo en este caso una especie de señal de alarma para decir al cuerpo que no podemos seguir así?

Además, se puede producir un bostezo sólo con que nos lo imaginemos o reflexionemos laboriosamente sobre ello. Por eso es también muy normal que al leer sobre la investigación de los bostezos uno tenga que contenerlos frecuentemente. Desde hace algunos años, los que investigan este tema utilizan los modernos métodos de la neurología y toman imágenes de la actividad cerebral de la persona mientras ésta bosteza: se coloca a los sujetos del experimento en un tubo de metal (el escáner de resonancia magnética) y se les muestran vídeos de bostezos —no películas especialmente aburridas, sino películas donde se ve a gente bostezando—. Un grupo de trabajo germano-finlandés realizó esta prueba y halló 2005 indicios de la causa del bostezo contagioso: según sus resultados, no se trata sencillamente de una imitación de comportamientos, o sea, más concretamente, no es un proceso de aprendizaje en el sentido de que observamos lo que hacen otros y

luego hacemos lo mismo. Las partes del cerebro que se ocupan de este proceso de aprendizaje, llamadas neuronas espejo, no se activan más al ver vídeos de bostezos, que cuando ven películas del mismo tipo pero sin bostezos. Por lo tanto, no es preciso que comprendamos cómo bostezan otros, para poder hacerlo nosotros mismos; es algo que se produce de una manera totalmente automática. Por eso se especula con la posibilidad de que se trate de un mecanismo muy antiguo que eventualmente se pone en marcha para la comunicación entre los individuos de un grupo y para la sincronización de comportamientos colectivos. El aviso que lo desencadena podría ser « ¡Vamos, ataca! » o « ¡Cuidado, enemigo a la vista! » o simplemente « ¡Venga, vamos a dormir! ». Es posible que el bostezo fuera antiguamente un método rápido y efectivo para explicar las necesidades vitales sin tener que hablar mucho.

En cambio, el investigador de neurociencias Steven Platek y su grupo llegaron a otra conclusión después de realizar un experimento sobre el bostezo. Según su teoría, el contagio del bostezo es un acto de compasión o simpatía; una idea que ya se expresó en la década de 1970.

Para comprobar su teoría, en 2003 Platek y sus colegas compararon el comportamiento en cuanto a bostezos en personas que, según los correspondientes testes de personalidad, tenían una capacidad de empatía especialmente alta o especialmente baja. También en este caso se les mostraron vídeos de bostezos, y se produjo el efecto esperado: entre los sujetos del ensayo, a aquellos que tenían facilidad para ponerse en el lugar de otras personas era más fácil hacerles bostezar que a los que eran menos sensibles. Para comprobar el resultado, Platek y los suyos utilizaron el procedimiento de obtención de imágenes del cerebro. De hecho, al bostezar se activaban zonas del cerebro que supuestamente desempeñan un papel en la generación de empatía. En determinadas circunstancias, el bostezar conjuntamente no significa «A mí también me aburre», sino «Comparto tu sentimiento».

Queda por ver si esta empatía mitiga además el aburrimiento.

El gran problema lo plantea el hecho de que el bostezo se produce en contextos muy diferentes. Por una parte, se podría pensar que hay un componente físico: se bosteza para respirar hondo, estimular la circulación sanguínea, o simplemente

como ejercicio de relajación de los músculos faciales. Además el bostezo abre la comunicación entre la boca y el oído, las trompas de Eustaquio, con lo que consigue equilibrar la presión en el oído medio cuando, por ejemplo, estamos resfriados o dentro de un avión que aterriza. El bostezo es por lo tanto sano, en eso están de acuerdo los expertos. Pero, por otra parte, el bostezo parece tener un componente social y comunicativo, tanto si se produce por empatía, como si se trata de sincronizar acciones. En este contexto el bostezo contagioso desempeña probablemente un papel importante. Asusta pensarlo, pero es evidente que sostenemos conversaciones de manera inconsciente mediante los bostezos.

Por otra parte no podemos impedirlo, aunque nos pongamos la mano delante de la boca. A pesar de ello, el cerebro es lo bastante listo como para darse cuenta de que el otro bosteza. Con independencia de la razón por la que bostecemos, el hecho real de abrir la boca resulta a la vista siempre sorprendentemente igual, a pesar de las múltiples causas posibles. El complejo significado del acto de bostezar es único, sobre todo si pensamos que otros esfuerzos involuntarios del cuerpo, como estornudar, toser o reír, no son atribuibles a tantos motivos. El bostezo parece una h curiosa existencia un nuevo componente singular.



Capítulo 7

Ciempiés

El que tiene cien pies llega antes a la ruina.

Sentencia de Asia Oriental

Tener ciempiés como animales de compañía es seguramente un *hobby* poco corriente. No obstante, si tanto el anfitrión como el miriápodo se prestan a ello, nadie puede tener nada que objetar. Por desgracia los ciempiés, que habitan la tierra desde hace más de cuatrocientos millones de años y, en consecuencia, son sumamente testarudos, no se ciñen a esa regla de conducta. Desde hace muchos años, sufridos seres humanos de todo el mundo se quejan de que ejércitos enteros de ciempiés invaden cada año sus casas. El mundo desconoce por qué los ciempiés (nombre científico: *Diplopoda*) se comportan así y por qué por lo menos no piden permiso antes.

En realidad, la región de Rüns, en el estado de Vorarlberg, no es especialmente llamativa; como en muchos otros pueblos, hay un estanque, una iglesia, una zona industrial y una señal de entrada en el municipio. Sin embargo, por lo menos en tres casas (y la tendencia es ascendente) a finales de verano se puede observar algo curioso: cientos de ciempiés entran en las casas por debajo de las puertas y por las grietas, y trepan por las paredes buscando la platería. Con tantos pies, no costará imaginar la de polvo que dejan en la casa. Si uno los trata con descortesía, estos huéspedes no deseados no sólo no se marchan, sino que además secretan un líquido apestoso que deja unas manchas asquerosas. Cada día, los propietarios de la casa barren cientos de ciempiés de las paredes. No se les puede ni siquiera soltar un pájaro, pues a la mayoría de pájaros no les gustan los ciempiés; algo comprensible, por otra parte, pues tienen un caparazón duro y hacen cosquillas en el pico.

Inicialmente tiene una la esperanza de que se trate de una peculiaridad de la región de Vorarlberg, pero la pierde en cuanto le echa un vistazo al informe elaborado por el experto Klaus Zimmermann, de Dornbirn. Actualmente, el biólogo tiene entre

manos una veintena de casos de invasiones de miriápodos en Austria, Alemania y otros países, y ha comprobado que son hasta cinco las especies que se instalan en las casas. Pero eso es tan sólo la punta del iceberg, un iceberg marrón oscuro y hormigueante, pues se han documentado fenómenos similares en Suecia, Inglaterra, Chequia, Malasia y otros países. En algunos casos, los ciempiés han llegado incluso a interrumpir el tráfico ferroviario.

En la década de 1950, Hugh Scott realizó minuciosas observaciones de las invasiones de ciempiés. Scott vivió relativamente tranquilo en su casa de Inglaterra durante catorce años hasta que, en la primavera de 1953 encontró los primeros ciempiés en el invernadero. Durante los años siguientes llegaron siempre en esa misma época, sólo que cada vez eran más numerosos y su estancia en la casa se prolongaba cada vez más tiempo. Una noche se colaban por el resquicio de alguna puerta y a la mañana siguiente ya trepaban por paredes y escaleras, hasta que Mr. Scott los mataba y los archivaba. Pero en el año 1958 el asunto se salió de madre: entre febrero y junio encontró 567 ciempiés, de los cuales 325 eran hembras, 239 machos y tres cuyo sexo no se pudo determinar porque al matarlos quedaron demasiado dañados. Durante muchas noches de abril, entraban en la casa entre diez y treinta nuevos inquilinos. No sabemos qué sucedió más tarde; probablemente Hugh Scott entregó la casa a aquellos huéspedes indeseados y se mudó amargado a un país sin ciempiés.

Ni Scott ni el resto de mundo científico han sido capaces de encontrar la causa de este fenómeno. Partiendo del hecho de que en las casas afectadas se hallan sobre todo animales maduros, se supone que se trata de migraciones masivas de apareamiento. Otros científicos aseguran que los animales buscan un lugar apropiado para el desove, aunque eso no explicaría que los machos tomen también parte en la migración. Existen numerosas teorías que esgrimen causas climáticas: la humedad atmosférica podría ser un factor a tener en cuenta, pues los ciempiés, que generalmente viven en la capa terrestre más superficial y se alimentan de plantas muertas, buscan de repente un entorno seco. Y, sin embargo, ¿por qué iban a buscar refugio en casas modernas extremadamente secas? Sea como fuere, Hugh Scott no relacionó su presencia en modo alguno con la humedad, sino con la temperatura. Sus ciempiés lo invadían sobre todo en las noches frías, aunque no

cuando era demasiado frío, pues entonces no había ningún miriápodo que se colara tras la puerta. Otros expertos informan de una proliferación de ciempiés en los días cálidos. En realidad, casi todas las teorías fracasan ante la arbitrariedad con que los ciempiés eligen los días en que salen de excursión y sus objetivos.

Un biólogo inglés llamado John Cloudsley-Thompson, que en la década de 1950 publicó una monografía sobre arañas, ciempiés y cochinillas, observó algunos paralelismos entre la aparición masiva de diferentes insectos, por ejemplo las nubes de langostas: en un primer momento se produce una enorme proliferación motivada por unas condiciones favorables, pero pronto el mundo se vuelve hostil para los animales, que se ven obligados a retirarse al campo, un poco como sucede con la migración de los pueblos. Según Hugh Scott, para estudiar ese extremo con mayor detalle en el caso de los ciempiés convendría observar un lugar de forma sistemática durante un largo período y determinar bajo qué condiciones los animales comienzan a actuar de forma extraña. Para ello habría que encontrar varias personas de buena voluntad que se prestaran a compartir su casa con unos cientos de ciempiés unas semanas al año. Tampoco puede ser tan difícil.

En cualquier caso, en Róns, en Vorarlberg, la gente piensa de otra forma. Un primer intento de poner coto a la plaga con *mesostigmata*, ácaros ladrones, que se comen los huevos de los ciempiés, terminó fracasando tras un aparente éxito inicial. Desde otoño del 2006, Klaus Zimmermann aplica un nuevo método: tierra de diatomeas, un polvo no venenoso hecho con restos vegetales fósiles cuyos finos cristales se clavan en las articulaciones de los caparazones de quitina e impiden que éstos crezcan. Al verse así tratados, un gran número de ciempiés se desmigajaron sin oponer mayor resistencia y, con ellos, murió también la posibilidad de conocer la razón de sus migraciones.



Capítulo 8

Cinta autoadhesiva

Sólo por la presión del aire, el animal se adhiere al objeto que escala.

ALFRED EDMUND BREHM, «*La salamanquesa*»⁶

Si se pregunta a los expertos por qué se pega realmente la cinta autoadhesiva, se obtienen unas respuestas sospechosamente evasivas. Pocos son los que lo reconocen abiertamente, pero es evidente que esta pregunta tan esencial para la supervivencia de la civilización no tiene aún una respuesta definitiva. En la mayoría de los casos el problema se aborda desde el punto de vista práctico: lo principal es que pegue.

Entre dos superficies actúan distintas fuerzas adherentes, entre dos y siete, según el pegamento y las ganas de establecer diferencias que tenga el especialista en la cuestión. En el caso de los pegamentos de tubo que luego se endurecen, hay que tener en cuenta algunas variantes: se supone que la adherencia mecánica cimienta el pegamento en la superficie como si lo hiciera con los pequeños ganchos de un velcro. Cuando la adherencia se hace por difusión se mezclan los varios cientos de moléculas superiores del pegamento y de la superficie donde se ha de pegar. Otra posibilidad son los enlaces químicos entre los materiales que se pegan. El papel que desempeñan estas distintas fuerzas adherentes no está claro ni para superficies y materiales corrientes, ni para los que son especiales. Tampoco se comprenden todavía algunas adherencias que se producen en la naturaleza; hasta ahora no se ha sabido exactamente cómo se adhieren los moluscos a superficies mojadas o cómo se fijan directamente bajo el agua, lo cual es un desafío extremo para cualquier sustancia adhesiva.

Otras preguntas se refieren a los llamados materiales autoadhesivos, como es el caso de la cinta autoadhesiva o los papelitos del *post-it*. Se pegan al instante, sin necesidad de secarse o fraguar. Este efecto se atribuye sobre todo a las fuerzas de

⁶ Alfred Edmund Brehm, *Vida animal*, Plaza & Janés, Barcelona, 1997. (N. de la t.)

Van der Waals. Se trata de unas fuerzas muy débiles basadas en la atracción eléctrica que se produce entre las cargas positivas y negativas en átomos o moléculas individuales, por lo que sólo pueden actuar a distancias cortas. Están condicionadas por la necesidad de que los dos lados que se han de unir puedan conseguir un contacto muy estrecho, lo cual se consigue, por ejemplo, cuando las dos superficies son extremadamente lisas. Un pegamento que fluya llenando todas las oquedales hace que también las superficies desiguales se pongan en contacto de esta manera. Las fuerzas de unión de este tipo no se deterioran, por lo que las pegatinas que se ponen en las ventanas o el papel transparente para conservar alimentos se pueden pegar y despegar del cristal todas las veces que se desee.

Las fuerzas de Van der Waals son muy del gusto de los expertos en sustancias adhesivas, porque para investigarlas reciben una salamanesa en sus laboratorios. Este animal posee una indudable y envidiable habilidad para pegarse, despegarse y pegarse de nuevo; es capaz de colgarse del techo con un solo dedo y quedarse así mientras se balancea, y de agarrarse con un solo pie para detener su caída. Después de haber estudiado las patas de la salamanesa durante doscientos años, y los treinta últimos siguiendo una buena pista, se sabe actualmente con bastante seguridad que la salamanesa puede sujetarse en el techo sobre todo gracias a las fuerzas de Van der Waals y con un poco de ayuda de las fuerzas capilares. (En este caso las fuerzas capilares se basan en el hecho de que hay agua en los diminutos huecos que quedan entre las moléculas de la salamanesa y las de la pared; para ello basta con que el aire tenga una humedad algo elevada).

Dado que ambas fuerzas son tan débiles, nos reiríamos de cualquiera que tuviera la idea de inventar una salamanesa, pero por suerte este animal ya existe. Además, mediante una hábil ampliación de la superficie de las patas, utilizando para ello unos pelos que tiene en ellas, consigue aprovechar al máximo estas débiles fuerzas, más o menos como si hiciéramos que unas hormigas levantaran un camión, desmontado previamente, hasta quedarse en sus piezas más pequeñas.

Pero ¿basta con las fuerzas de Van der Waals para explicar la adherencia de la cinta autoadhesiva? En cualquier caso, un francés especializado en sustancias adhesivas, Cyprien Gay, pone en duda esta teoría. Si se mide la cantidad de energía que es necesaria para separar las partes que ya están pegadas, según Gay se pone de

manifiesto que el pegamento crea una adherencia diez mil veces mejor que la que se podría justificar mediante las fuerzas de Van der Waals. Una manera posible de salvar la situación es recurrir a lo que se llama viscoelasticidad: las grandes moléculas del pegamento no se comportan con buenas maneras, sino que, en muchos aspectos igual que los espaguetis, sólo se pueden separar unas de otras con gran esfuerzo y mucha paciencia. Cuando la viscoelasticidad y las fuerzas de Van der Waals actúan juntas, para separar dos objetos pegados hay que trabajar cien veces más que si actuaran sólo las fuerzas de Van der Waals. Ahora bien, con esto se explica únicamente el uno por ciento de la «energía adherente».

¿Cuál es la fuerza que se encarga del 99 por ciento restante? ¿Y por qué para separar una cinta autoadhesiva se necesita primero un fuerte tirón, y luego seguir tirando de manera constante? Al medir la fuerza que había que utilizar se obtuvo una curva cuya forma era más o menos parecida a la de un sillón: primero se necesita mucha fuerza durante un breve espacio de tiempo (respaldo), y luego poca fuerza durante un período más largo (superficie del asiento). La razón por la que la curva tiene este aspecto y no otro no se puede explicar recurriendo a las fuerzas de Van der Waals.

La teoría de la cavitación, que Cyprien Gay presentó junto con Ludvik Leibler en 1999, podría dar respuesta a ambas preguntas. Afirma que en el pegamento de la cinta autoadhesiva y de las hojas del *post-it* hay una gran cantidad de pequeñas burbujas que actúan como ventosas y se resisten cuando vamos a despegar la cinta o las hojas. Si es correcta la teoría de la cavitación, las pequeñas burbujas ofrecen resistencia en un principio a causa de la disminución de presión, hasta que se dilatan tanto que acaban deshaciéndose. Entonces sólo queda ya superar la resistencia de las fibras del pegamento. Si esto sucediera realmente así, la cinta autoadhesiva se despegaría con mayor facilidad en las cimas de las montañas, donde la presión atmosférica es más baja. La realización de este experimento se anunció el mismo año en que apareció la teoría, pero parece que hasta la fecha no se ha puesto en práctica. Quizá transportar un laboratorio a la cima de una montaña sea más difícil de lo que se piensa.

¿Por qué los fabricantes de pegamento no se deciden sencillamente a imitar a la salamanquesa? La pregunta resulta aún más interesante si tenemos en cuenta que

las salamanquesas utilizan al parecer sólo una fracción de la fuerza de adherencia que en teoría tienen. El biólogo estadounidense Kellar Autumn calculó que una salamanquesa tokay situada en una pared podía transportar 140 kilos de peso. Sólo que no quiere hacerlo. Pero las patas de la salamanquesa son nanoestructuras, y como tales un poco complicadas de fabricar. Además este animal se preocupa de mantener limpios sus pelos de contacto, mientras que las salamanquesas artificiales se ensucian enseguida y pierden su fuerza adherente. Además, sospecho que al consumidor no le agradaría que las notas autoadhesivas que utiliza en la oficina corretearan por las paredes y el techo y se pusieran a cazar moscas. En condiciones normales, ya es bastante difícil no perder de vista las cosas.



Capítulo 9

Cúmulos globulares

Quien quiera hacer una tarta de manzana sin utilizar ingredientes preparados, tendrá que crear primero el universo.

CARL SAGAN *Cosmos*

Supongamos que cada una de las estrellas que vemos en el cielo es una sola casa iluminada.

Entonces, tendremos que considerar que galaxias tales como nuestra Vía Láctea son grandes ciudades en las que varios cientos de miles de casas están ordenadas formando estructuras predominantemente lógicas. Siguiendo esta comparación, los cúmulos globulares serían los suburbios: se encuentran diseminados en el entorno próximo de la galaxia y cada uno de ellos consta de varios miles de casas. La Vía Láctea, por ejemplo, tiene un halo donde hay unos ciento cincuenta cúmulos globulares. Dado que todos ellos se encuentran muy lejos, al mirarlos con el telescopio, incluso los de mayor tamaño parecen sólo manchas de niebla desvaídas y esféricas.

Sin embargo, cuanto más potente sea el telescopio con el que se observa el cielo, más claramente se ve de qué se trata: una bengala de dimensiones extraordinarias llena de estrellas. A pesar de los grandes avances realizados en los últimos años, hasta ahora no se ha conseguido explicar por qué las galaxias están rodeadas de cúmulos globulares, ni cómo han surgido estas grandes aglomeraciones de estrellas. Cuando se logra el acercamiento necesario para poder reconocer cada una de las estrellas, se puede determinar la edad de los cúmulos globulares con mayor fiabilidad que la de los perros de un albergue para animales. Así, se ha averiguado que la mayoría de los cúmulos globulares surgieron hace entre diez mil y catorce mil millones de años, lo cual resulta sorprendente si se piensa que el universo en su totalidad no es más viejo, según los conocimientos actuales al respecto. Por lo tanto, los cúmulos globulares serían algo así como el Stonehenge del universo: restos de una época sobre la cual poco más sabemos. Son tan antiguos que gracias

a ellos se espera averiguar más sobre la infancia del universo, es decir, sobre la época en que surgieron las primeras estrellas y galaxias.

Las primeras teorías relativas a la aparición de cúmulos globulares partían de la idea de que éstos podrían ser los precursores de las galaxias; una especie de tímido intento del universo para organizar ordenadamente sus estrellas recién fabricadas. Se suponía que, después de existir estas situaciones «primarias», los cúmulos globulares serían anexionados por las galaxias que habían surgido posteriormente. Desde la década de 1980, se fueron descartando cada vez más estos modelos, porque se puso de manifiesto que los cúmulos globulares y las galaxias no se han agrupado casualmente, sino que evidentemente se conocen desde hace más tiempo. Así, se descubrió, entre otras cosas, una relación entre la composición química de las estrellas del cúmulo y la luminosidad que tenía en conjunto la correspondiente galaxia madre, lo cual apuntaba a la existencia de un pasado común. Desde entonces se cree que los cúmulos, o bien surgieron al mismo tiempo que las galaxias, o más tarde dentro de la galaxia ya formada. A partir de esto se espera poder resolver, con ayuda de los cúmulos globulares, otros dos grandes problemas de la astronomía: cómo nacieron las galaxias, y cómo pasaron su infancia.

Ni siquiera con los mejores telescopios puede verse claramente lo que sucede en ese hervidero que es el centro de los cúmulos globulares. Varios cientos de estrellas se hacían allí en un año luz cúbico, que desde un punto de vista astronómico es un volumen muy pequeño. Para hacernos una idea, el vecino más próximo al Sol está a una distancia de más de cuatro años luz. El principio general relativo al origen de estas agrupaciones de estrellas es, sin embargo, conocido: se forman a partir de una gigantesca nube de gas y polvo. Una nube así es estable cuando la propia fuerza de gravedad, que tiende a comprimir todo en el espacio más reducido posible, se ve compensada por otras fuerzas. Una de éstas es la que produce el calor almacenado en la nube; cuando la materia se calienta, experimenta una dilatación y actúa contra la fuerza de la gravedad.

Si se destruye ese equilibrio, por ejemplo, comprimiendo repentinamente la nube, entonces vence la gravedad y la nube se hunde bajo su propio peso; y surgen agrupaciones de estrellas. Por algún motivo, las galaxias espirales habituales, como

nuestra Vía Láctea, ya no consiguen actualmente generar de esa forma nuevos cúmulos globulares de gran tamaño; por el contrario, actúan de manera titubeante y prefieren crear sólo nuevas luces en el cielo. Pero ¿por qué sucede esto?

Una de las teorías relativas al origen de los cúmulos globulares habla de fusión de galaxias, accidentes de tráfico a gran escala, en los que de dos galaxias resulta una. Desde hace tiempo se sospecha que las grandes galaxias elípticas podrían originarse por la fusión de dos galaxias espirales. Si esto es cierto, según dijo el astrónomo Sidney van den Bergh en 1984, entonces ¿por qué no tienen las galaxias elípticas la suma del número de cúmulos globulares de dos galaxias espirales, sino un número claramente mayor? La respuesta que posiblemente sea acertada, la dieron los norteamericanos Keith M. Ashman y Stephen E. Zepf en 1992: Al producirse la colisión se generan condiciones que posibilitan el nacimiento de nuevos cúmulos globulares, de tal modo que la nueva galaxia posee al final más cúmulos que la suma de sus partes.

En la década de 1990, el modelo de Ashman-Zepf despertó las simpatías de muchos expertos, sobre todo por dos razones. Entre otras cosas predice que en las galaxias elípticas debe haber dos tipos de cúmulos globulares: unos viejos y con larga barba, y otros que se formaron más tarde, al fusionarse las galaxias. Entretanto, gracias principalmente al telescopio espacial Hubble y su poderosa visión del universo, se ha descubierto que existe realmente esa división en dos tipos de cúmulos globulares: unos son «bajos en metales», lo cual para los astrónomos es sinónimo de «muy antiguos», porque el universo joven se compone sólo de hidrógeno y helio, y todos los elementos pesados, como los metales, se van formando posteriormente. En cambio, los del otro tipo son «ricos en metales» y, por eso, quizá surgieron después, al chocar las galaxias, siempre según este modelo. El segundo descubrimiento importante de los últimos años, en este contexto, es el siguiente: cuando las galaxias chocan entre sí, se originan realmente unos cúmulos de estrellas dotados de una cantidad enorme de masa, que hoy en día se consideran en general como casos raros entre los cúmulos globulares.

Por otro lado, sin embargo, la teoría de Ashman-Zepf tiene que luchar con algunas dificultades. Por ejemplo, algunas investigaciones ponen de manifiesto que las galaxias poseedoras de una cantidad especialmente grande de cúmulos globulares

muestran una elevada proporción de cúmulos bajos en metales. Según el modelo de la colisión sería de esperar justo lo contrario: cuanto más a menudo se produzcan los choques, deberían originarse más cúmulos globulares ricos en metales. Una posibilidad alternativa para explicar los dos tipos diferentes de cúmulos globulares sería la siguiente: en principio cada galaxia lleva consigo su propio rebaño de cúmulos; con el paso del tiempo, y tiempo es lo que le sobra al universo, las galaxias se van desplazando y, cada vez que se encuentran, las más grandes absorben algunos cúmulos de las que son menores. Este planteamiento, entre otras cosas, ayuda a comprender una curiosa diferencia que se da entre la Vía Láctea y la nebulosa de Andrómeda, aunque ambas son galaxias espirales de estructura parecida: mientras que en la Vía Láctea todos los cúmulos globulares parecen ser muy antiguos, en los últimos tiempos proliferan los indicios de la existencia de algunos cúmulos globulares «más jóvenes» en la nebulosa de Andrómeda, es decir, cúmulos que no tienen más de cinco mil millones de años. Las pruebas no son todavía contundentes, pero es posible que la nebulosa de Andrómeda haya robado estos flamantes cúmulos nuevos a sus pequeñas galaxias vecinas.

Sin embargo, este modelo de galaxias caníbales tampoco está libre de problemas a la hora de explicar las complicadas relaciones entre los cúmulos globulares y sus galaxias madres. Otros expertos sospechan que, al menos por lo que respecta a algunos cúmulos globulares, no se trata en absoluto de cúmulos de estrellas, sino de los núcleos indigestos de antiguas galaxias enanas, cuyos componentes ha incorporado para sí misma la galaxia madre. Entretanto ha quedado claro que las galaxias en absoluto están predispuestas a llevar una existencia solitaria y aislada, sino que, por el contrario, tienen una intensa vida social. Chocan unas con otras, se despedazan, se fusionan entre sí y se devoran mutuamente; todos estos comportamientos pueden observarse actualmente con los telescopios adecuados. Actualmente se piensa que, con toda seguridad, estos acontecimientos desempeñan un papel en la formación de las galaxias y, por consiguiente, también en la de los cúmulos globulares. Así pues, mediante la observación precisa de los cúmulos globulares de una galaxia determinada, se puede averiguar mucho sobre la agitada vida de ésta.

Pero el problema fundamental, que es el origen de los cúmulos globulares, ni siquiera así ha podido resolverse.

Con todo esto, apenas queda otra posibilidad que suponer que una gran parte de los cúmulos globulares, al menos los que son pobres en metales, han tenido que surgir más o menos al mismo tiempo que las galaxias. Cualquier astrónomo que tenga una idea de cómo se forman las galaxias, y lo simule todo en un ordenador, debe producir al mismo tiempo, junto con la galaxia, cierto número de cúmulos globulares. Dado que es posible determinar, en cierto modo de manera fiable, las características de los cúmulos (distribución del espacio, cantidad, masa, composición química), los observadores, desde sus telescopios, tienen realmente posibilidades de comparar la realidad con las predicciones y así averiguar cuál de los modelos de sus colegas teóricos tiene sentido y cuál no. Los cúmulos globulares sirven como indicio claro para configurar nuestras representaciones del universo en épocas anteriores.

Las primeras estrellas del universo surgieron probablemente como unas grandes nubes de gas y materia oscura, que pudieron ser las precursoras de las galaxias. Este suceso, que se produciría algunos cientos de millones de años después del *big bang*, marca el final de los tiempos oscuros del universo. Desde «fuera» tuvo que verse como si alguien hubiera encendido la luz en aquella oscuridad. La pregunta que queda sin respuesta es qué fue exactamente lo que produjo luego la aparición de los cúmulos globulares, para la cual, como ya se ha dicho, tuvieron que comprimirse en breve tiempo grandes cantidades de gas. Probablemente fue la luz de la primera estrella la que calentó el hidrógeno, generando así ondas longitudinales que se propagaban a través de la nube de gas, como cuando se arroja una piedra al agua que antes estaba en reposo. Por otra parte, es un misterio la razón por la cual no se ha encontrado ni rastro de materia oscura en los cúmulos globulares, teniendo en cuenta que éstos han surgido en un entorno donde seguramente no se puede salir de casa sin hundirse hasta las rodillas en dicha materia oscura. De algún modo los cúmulos han conseguido librarse de sus sombras oscuras. Finalmente, se seguirá especulando sobre la razón por la cual los cúmulos globulares pasaron de moda unos pocos miles de años más tarde, y el universo redujo su producción de manera drástica, de tal modo que hoy en día sólo podemos

comprarlos de segunda mano, a menos que choquemos justo ahora con otra galaxia, como se ha dicho anteriormente.

¿Cómo podemos resolver estos problemas? Los astrónomos lo ven muy fácil. Se limitan a pedir mejores juguetes: telescopios más potentes, mejores cámaras fotográficas y ordenadores más rápidos. Si no se les niegan esos deseos, seguramente aclararán todos los misterios antes de las próximas vacaciones de verano.



Capítulo 10

Curva de Laffer

Pomperipossa lo oyó por primera vez cuando una amiga le preguntó un día: «¿Te has enterado ya de que tus impuestos marginales de este año ascienden a un 102 por ciento?». «¡Imposible! —dijo Pomperipossa—. ¡No existe un porcentaje tan alto!»

ASTRID LINDGREN *Pomperipossa en el mundo del dinero*

La curva de Laffer, bautizada con el nombre del economista estadounidense Arthur B. Laffer, expresa la relación entre el nivel de los tipos impositivos y los ingresos fiscales que recibe el Estado según dichos tipos. Es indiscutible la verdad de un dicho muy antiguo: «Cuando se exagera la explotación de los súbditos [en términos modernos: “cuando se exagera con los impuestos”], la hacienda pública va a la ruina». Los vampiros tienen un problema parecido. Si le chupan demasiada sangre a su víctima, ésta se les muere, y se acaba una fuente barata de recursos alimenticios. Hasta cuándo se puede chupar, es decir, dónde se sitúa el nivel óptimo del tipo imponible, o sea el punto máximo de la curva de Laffer, es algo que nadie sabe. Laffer dice que, cuando los tipos impositivos superan este máximo, el Estado puede obtener más bajando los impuestos que subiéndolos. No está claro si esto funciona, ni cómo lo hace. ¿La curva de Laffer es un recurso importante para controlar la economía, o es sólo una especie de vudú económico?

El principio básico puede explicarse mediante un ejemplo sencillo. Desde principios de la década de 1990, muchos países europeos, entre ellos Alemania, están subiendo lenta y continuamente el impuesto del tabaco. Hasta hace poco, esto producía el efecto deseado, dentro de lo posible: la gente seguía fumando a pesar de todo, pero aumentaban los ingresos del Estado por el impuesto del tabaco. En 2003, cuando el gobierno federal aumentó el precio de cada cigarrillo en un 1,2 por

ciento mediante una reiterada subida del impuesto, se preveían unos ingresos adicionales de unos mil millones de euros. Un año después estaba claro que había sucedido justo lo contrario. De repente, las arcas del Estado recibieron menos dinero por el impuesto del tabaco, y el Ministerio de Economía reconoció un déficit de varios cientos de millones de euros.

Posiblemente hubo muchos fumadores que dejaron su hábito a causa de los altos precios, pero también se podría pensar que lo que hicieron algunos fue recurrir a la mercancía de contrabando.

Fuera cual fuese la causa, la subida de impuestos produjo paradójicamente una pérdida de ingresos.

Para el total de las rentas fiscales del Estado, el tabaco tiene escasa importancia; mucho más significativas son otras fuentes de ingresos, como, por ejemplo, los impuestos sobre las rentas del trabajo. Pero también en este caso puede suceder que un tipo impositivo demasiado elevado produzca menos ingresos fiscales, aunque las causas son más complejas que en el tema del impuesto sobre el tabaco. Pero ¿qué significa en este contexto «demasiado elevado»? ¿Dónde está la frontera mágica? ¿Cuánto debe recibir el Estado por la renta de cada ciudadano, para alcanzar unos ingresos máximos?

Hay que empezar por hacer unas consideraciones sencillas. Si el tipo impositivo es cero, o sea, si no se recaudan impuestos, el Estado no ingresa cantidad alguna. Todos los ciudadanos estarán felices y contentos, porque sus ingresos brutos serán iguales a sus ingresos netos, pero el Estado se queda arrinconado y se siente ofendido, porque no recibe su parte. Éste es el primer punto de la curva de Laffer: cero por ciento en el tipo impositivo produce ingresos fiscales cero.

Sin embargo, está claro que el inicio de la curva de Laffer da lugar posteriormente a un ascenso: si se aumenta el tipo impositivo de cero a, por ejemplo, dos por ciento, siguen todos contentos, pero los ingresos del Estado suben un poco.

De seguir así, se supone que con unos tipos impositivos del cien por cien, es decir, si cada ciudadano entrega todos sus ingresos al Estado, no es de esperar recaudación tributaria alguna.

Sólo trabajarían un par de idealistas o de locos, bien porque no son capaces de hacer las cuentas, o porque les encanta bajar a la mina. Los ciudadanos preferirían

dedicarse a actividades ilegales, emigrar, mendigar, quedarse en la cama y morir de hambre, pero nadie querría trabajar legalmente y pagar impuestos, para acabar en cualquier caso muriéndose de hambre. Por lo tanto, si el tipo impositivo es el cien por cien, los ingresos fiscales del Estado serán nulos o alguna cantidad insignificante. El punto inicial y el punto final alcanzan el cero. Para unir ambos puntos hay que trazar una curva que al principio es ascendente, alcanza un máximo en algún lugar y luego desciende para llegar de nuevo al cero. Así es la curva de Laffer en un caso ideal. Si el Estado es listo, exigirá a sus ciudadanos la cantidad de impuestos que le lleve al máximo de recaudación, a la cima de la curva de Laffer. Si el Estado exige unos impuestos demasiado elevados, disminuyen sus ingresos. En consecuencia, según los argumentos de Laffer, es posible aumentar los ingresos cuando se reducen los impuestos.

La idea fundamental en que se basa la curva de Laffer no es en absoluto nueva. El propio Laffer se la atribuye al político árabe Ibn Chaldun, que en el siglo XIV escribió lo siguiente: «Se ha de saber que al principio de una dinastía el Estado obtiene grandes ingresos fiscales con pocos impuestos. Al final de una dinastía los impuestos son altos y el Estado recibe pocos ingresos fiscales». Esto suena inteligente, pero está expresado de una manera tan críptica que permite sospechar mucho contenido oculto; es un problema general expresado al modo de Chaldun, que las generaciones posteriores interpretarían a menudo justo al contrario. En cualquier caso, Laffer entiende esas palabras como un llamamiento a la reducción de impuestos.

El concepto de curva de Laffer surgió en 1974 como consecuencia de un encuentro, ahora casi legendario, entre Laffer, que entonces era profesor de la Universidad de Chicago, y representantes de Gerald Ford, que por aquel tiempo era presidente de Estados Unidos. Estaba presente, entre otros, Dick Cheney como representante de Donald Rumsfeld, que en aquel momento era jefe de Estado Mayor de la Casa Blanca, aunque hay quien dice que el propio Rumsfeld participó en la reunión. Laffer no se acuerda ya de los detalles de aquel encuentro, por lo que no hay más remedio que creer la versión de Jude Wanninski, que entonces trabajaba como coeditor del prestigioso *Wall Street Journal* y ha sido hasta ahora un potente defensor de la teoría de Laffer. Según Wanninski, Laffer dibujó la curva en una servilleta, para

convencer a Cheney de que había que bajar los impuestos, y no subirlos, con el fin de reactivar la economía y, como consecuencia de ello, sanear los ingresos fiscales. Aunque Cheney y Rumsfeld quedaron bastante impresionados, Ford no aceptó la propuesta. El propio Laffer considera como dudosa la historia de la servilleta que cuenta Wanninski, ya que su madre, la de Laffer, «le había educado para no profanar las cosas bellas». En cualquier caso, desde aquel episodio se llama «curva de Laffer» a la relación entre tipos impositivos e ingresos fiscales, y en cualquier libro de texto aparece dibujada como una hermosa y simétrica curva de campana que empieza y acaba en el cero, y tiene un máximo en el 50 por ciento del tipo impositivo.

Sin embargo, esta forma ideal de la curva no deja de ser un invento y posiblemente no tenga nada que ver con la realidad. El trazado exacto de la curva de Laffer, sobre todo la posición del máximo, es objeto de vivas controversias. En numerosas publicaciones se habla sobre los intentos de investigar la relación entre los tipos impositivos y los ingresos fiscales resultantes, utilizando para ello modelos matemáticos. Al mismo tiempo, se describen los procesos económicos mediante un sistema de ecuaciones, que están vinculadas entre sí. Por lo tanto, los ingresos del Estado procedentes del impuesto sobre la renta no sólo dependen de los tipos impositivos, sino también de la cantidad de personas que estén trabajando y de las remuneraciones de estas personas. El nivel de los salarios brutos depende de cómo le vaya al empresario, es decir, de cuántos productos consiga vender. Esto, a su vez, depende, entre otras cosas, de cuánto dinero les quede a los ciudadanos después de deducir los impuestos y, por lo tanto, de lo que puedan gastar, lo cual lógicamente se relaciona de nuevo con los tipos impositivos. En conjunto, lo que surge aquí es una estructura compleja e intrincada.

Veamos ahora unos cuantos resultados de los cálculos realizados según los distintos modelos matemáticos. La curva de Laffer que Peter Ireland obtuvo en un estudio del año 1994 alcanza el máximo con un tipo impositivo del 15 por ciento. Con otro modelo algo diferente, Paul Pecorino calculó en 1995 un tipo impositivo óptimo que se situaba entre el 60 y el 70 por ciento. En un trabajo del año 1982, Don Fullerton obtuvo un tipo impositivo óptimo del 79 por ciento. Es desconcertante la diferencia de cifras. En un trabajo más reciente, realizado en 2005, N. Gregory Mankiw y

Matthew Weinzierl intentaron encontrar un enunciado más general y llevaron a cabo numerosas ampliaciones con el propósito de lograr que el modelo se acercara más a la realidad.

Entre otras cosas, llegaron a la conclusión de que la economía reacciona de una forma muy sensible ante los cambios fiscales y que las bajadas de impuestos pueden en parte autofinanciarse.

Sin embargo, es difícil predecir en qué medida sucede esto. En definitiva, no está claro si realmente es posible, como dice Laffer, aumentar los ingresos exigiendo menos impuestos.

La simulación realista de una economía nacional es una tarea difícil. Incluso si los modelos son complejos y están muy bien madurados, sólo reproducen la realidad de una forma enormemente simplificada, porque lo real es aún más complejo y maduro. El impuesto sobre la renta, por ejemplo, es progresivo en muchos países (como sucede en Alemania), es decir, a las rentas más altas se les aplica tipos impositivos también más altos. Así, puede darse el caso de que con un determinado nivel de ingresos se produzca el efecto Laffer, porque los impuestos que le corresponden sean demasiado altos, pero puede ser el único nivel en el que se llegue a dicho efecto. Además, la gran cantidad de excepciones, normativas especiales y excepciones de esas normativas generan un barullo complicado e impenetrable en el sistema impositivo, por lo que resulta casi imposible predecir lo que puede suceder cuando se cambia algún pequeño detalle en cualquier lugar. Con nuestros sistemas impositivos hemos creado a través de los siglos un monstruo de muchas patas que ahora es difícil de domar.

En las publicaciones de Laffer relativas a su curva aparecen pruebas detalladas que justifican su teoría y que se refieren sobre todo a sucesos de la historia reciente de Estados Unidos. Entre otros, se discuten los programas de reducción de impuestos que se aplicaron en tiempos de John E Kennedy en la década de 1960 y de Ronald Reagan en la de 1980. En ambos casos se observan claros indicios de que la bajada de impuestos tuvo efectos positivos en el crecimiento económico, aunque no era éste el objetivo. En realidad, sí lo era el aumento de los ingresos fiscales del Estado, pero esto no se consigue justificar con claridad. Aún peor: incluso en los casos de éxito era prácticamente imposible demostrar una relación causa-efecto entre la

bajada de impuestos y variaciones en los ingresos del Estado, porque en las cuentas intervenían muchos factores desconocidos. Un ejemplo: la situación global de la economía nacional no depende sólo de la situación en el propio país, sino también, entre otras cosas, de la demanda de productos en el extranjero. Puede suceder que los consumidores alemanes compren menos *ketchup* americano porque en ese momento aparezca una nueva marca barata de *ketchup* en los comercios de Alemania. O porque se implante en este país un impuesto especial sobre el *ketchup*, y esto impulse a los consumidores a utilizar más la mostaza.

Finalmente, se puede dar un giro de 180° a la discusión planteando la pregunta de si debería ser obligación del Estado maximizar su recaudación tributaria. Porque no son sólo una fuente de ingresos para el Estado, sino que además pueden servir para influir en los modos de comportamiento y darles una nueva dirección. Un buen ejemplo es el caso del tabaco, que ya hemos mencionado anteriormente: ¿Se eleva este impuesto para forzar la entrada de más ingresos procedentes de la venta de tabaco, o se trata por el contrario de impulsar a las personas a dejar de fumar? En este último caso, también se valoraría como un éxito que hubiera menos ingresos fiscales, es decir, que se vendieran menos cigarrillos. Esto lleva inmediatamente a preguntarse hasta qué punto el Estado puede inmiscuirse en las vidas de sus ciudadanos. ¿Son los Estados simplemente una especie de entidades económicas que ofrecen determinadas prestaciones (por ejemplo, la construcción de carreteras) y a cambio reciben el pago en forma de impuestos? ¿O tienen además la obligación de preocuparse por el bienestar, la educación y la organización en general de la sociedad? Sea cual fuere la respuesta que se dé a esta pregunta, el debate sobre la curva de Laffer puede ser al final totalmente irrelevante.

Con independencia del valor que se le dé, la idea básica de la curva de Laffer es de una belleza y una claridad insuperables: si una cosa que en sí misma es buena y útil (por ejemplo, chupar sangre) se explota en demasía, llegará a actuar en contra del objetivo inicial y, por lo tanto, se volverá perjudicial. Esta idea nos lleva directamente a una crítica general del exceso. Comer es necesario para la vida y parece en principio una buena idea; sin embargo, cuando se practica en exceso es malo para la salud y produce una degeneración adiposa. Los medicamentos son beneficiosos cuando no se toman en sobredosis. Las sanciones tienen en el mejor

de los casos un efecto educativo, si se aplican en la medida adecuada, ni mucho, ni poco. No obstante, Alexander Solzhenitsin lo ve de otra manera y afirma que las penas de prisión que duran menos de veinticinco años sólo sirven para embrutecer al individuo, pero, sin embargo, todas las que superan ese tiempo contribuyen a hacer personas honestas. Esperemos que a nadie se le ocurra trasladar esta filosofía al sistema impositivo.



Capítulo 11

Dinero

El dinero nos hace ricos.

Lotería del sur de Alemania

El dinero es en realidad una cosa sencilla: si se tiene, se puede meter en la máquina del café o en el mercado de valores. Si no se tiene, hay que dedicarse a recoger chatarra. Entre los profanos en la materia apenas se plantean problemas con respecto a comprender qué es el dinero, dejando a un lado lo difícil que es explicar por qué se va tan rápido de las manos.

Para los expertos la cosa es diferente. El periodista francés Marcel Labordère, que trabajaba en prensa financiera, escribió lo siguiente en la década de 1920: «Es evidente que el ser humano nunca sabrá qué es el dinero, del mismo modo que nunca llegará a saber qué es Dios en el mundo espiritual». Quizá descubriremos ambas cosas algún día; en todo caso, hasta ahora sólo se puede hablar de modestos avances. Especialmente las preguntas básicas «¿Qué es el dinero?», «¿Cuánto dinero existe?» y «¿Qué efectos produce el dinero?» acaban siempre haciendo que los profesores de economía política se digan unos a otros cosas terribles.

Según la mayoría de los expertos en finanzas, el dinero tiene tres funciones: sirve como medio de canje, para la conservación del valor y como pauta de valor. Es decir, con el dinero podemos comprar, también podemos colocarlo, y lo necesitamos porque sin él no sabríamos cuánto vale un sello de 55 céntimos. El experto en economía política Hans-Joachim Jarchow en su libro de texto *Geldtheorie* explica lo siguiente: «En general se puede considerar dinero o medio de pago todo aquello que en el marco del servicio de pagos de una economía nacional es aceptado en general para el pago de bienes y servicios». Karl Kraus resume los mismos hechos de una manera aún más concisa: «Con dinero pueden comprarse mercancías, porque es dinero, y es dinero porque con él pueden comprarse mercancías».

¿Por qué existe el dinero? Ésta es una pregunta con la que los expertos hacen trampas para no tener que dar una respuesta. Pues el dinero está ahí. ¿Y a quién le importa por qué los seres humanos en algún momento se decidieron a cambiar

mercancías por piezas de metal y por papel impreso? De manera intuitiva, estamos dispuestos a aceptar que el dinero es sencillamente una mercancía especialmente práctica, y además fácil de transportar y de conservar, por lo que como medio de canje es claramente más adecuada que, por ejemplo, los pepinillos. Sin embargo, esto puede ser bastante discutible; algunos especialistas creen que el dinero (en forma de oro, u otras pompas y boatos) se inventó para cultos simbólicos como ofrenda a los dioses o a los sacerdotes, y fue más tarde cuando se estableció como recurso en la vida cotidiana. Otros parten de la idea de que el dinero surgió por cargas sobre la propiedad, es decir, habría sido al principio un pagaré que se tenía que presentar para recuperar una propiedad pignorada.

Especialmente desde el final de la cobertura en oro, las circunstancias se han vuelto aún más difíciles de entender intuitivamente. Antes cada billete de banco se correspondía de algún modo con una determinada cantidad de oro que poseía el Estado, y éste no podía acuñar arbitrariamente más dinero sólo porque quisiera comprar un par de portaviones nuevos. Después de la segunda guerra mundial, sólo dos países siguieron teniendo una divisa cubierta por el oro: Estados Unidos hasta 1968 y Suiza hasta 1999. Dado que la obligación de la cobertura en oro había tenido desde la década de 1930 una función más bien decorativa, en general nadie es muy consciente hoy en día de que fue abolida. Por supuesto, sólo se recuerda cuando se trata de comprender ese fenómeno llamado dinero.

Para complicar aún más un asunto que de por sí no es transparente, resulta que no hay una sola clase de dinero, sino muchas más: junto al dinero en efectivo solemos tener, en el mejor de los casos, el dinero que está en una cuenta bancaria. Este último puede transferirse, sacarse del cajero automático o cargarlo en la tarjeta de crédito; dado que en muchos aspectos se comporta como si fuera dinero en efectivo, tendrá que ser también dinero. Hasta aquí va bien, pero si el dinero de la cuenta corriente puede ser dinero, ¿por qué no lo puede ser también el dinero invertido a un mes?

¿Y por qué sólo un mes? Podríamos decir lo mismo del dinero invertido a un plazo más largo, de los valores y del dinero invertido en seguros. Y es así como sucede, por lo que surgen distintas masas monetarias que se designan desde MO hasta M3. Lamentablemente hay ahora otros asuntos que se sitúan en el margen más externo

de este espectro y que no tienen mucho que ver con el dinero que vemos en nuestro monedero. Un caso así son las acciones. El Deutsche Bundesbank dice: «Aquello que razonablemente consideramos como dinero no es [...] una cuestión que pueda explicarse de una vez por todas con precisión científica, sino una cuestión de utilidad. [...] Para el Banco Central Europeo la masa monetaria M3, ampliamente delimitada, ocupa un primer plano en su valoración de situaciones». Por el contrario, en Estados Unidos la M3, la masa monetaria que con mucho es más estable, está considerada como «un dato inservible» y desde 2006 ya no se da a conocer públicamente. «El intento de definir la masa monetaria —dice el experto en economía nacional Paul A. Samuelson— hace que los expertos rigurosos lleguen al borde de la desesperación, ya que no existe ninguna línea de separación claramente definida en el caleidoscopio de las inversiones que permita fijar el punto en el que el dinero se diferencia de otras inversiones».

Los críticos de las distintas masas monetarias plantean la objeción de que tiene tanto sentido contabilizar conjuntamente dinero y saldos activos, como sumar manzanas con dibujos de manzanas que uno ha prestado a otras personas. El dinero es solamente aquello que por ley ha de ser aceptado como medio de pago, es decir, monedas y billetes de banco. También se discute, por lo tanto, si los bancos, cuando prestan dinero, crean un dinero que antes no existía. Cuando alguien lleva 100 euros al banco, éste lo presta a otros clientes, y no sólo una vez, sino tantas que al final se produce un total de 900 euros. Esto es lo que se llama creación de depósitos bancarios, pero la pregunta sobre si realmente se crea dinero, como sugiere esta palabra, siempre da a los economistas el pretexto para iniciar aburridos debates del tipo «¡Que no! ¡Que sí! ¡Que no! ¡Que sí!».

La cuestión relativa a cuánto dinero existe, no es sólo una sutil pregunta definitoria, pues los billetes de banco pretenden controlar la masa monetaria, con el fin de que el poder adquisitivo permanezca estable. Alan Greenspan, que fue durante muchos años presidente del banco emisor estadounidense, explicaba en una entrevista: «El problema principal es definir qué parte de nuestra estructura de liquidez es realmente dinero. Hace años que intentamos hallar indicadores para esto. El criterio que aplicamos aquí es que, con ayuda de dichos indicadores, se podría predecir la dirección del desarrollo económico y financiero. Lamentablemente hasta ahora no lo

hemos conseguido con ninguno de los indicadores que hemos inventado [...]. Esto no significa que el dinero no nos parezca importante; se trata sólo de que nuestros métodos de medición eran insuficientes. [...] No se puede manejar aquello que uno no ha podido definir previamente».

Incluso si nos limitamos al dinero en efectivo, los bancos emisores saben cuánto han emitido, pero no tienen conocimiento de la cantidad que está realmente en circulación. El ensayista Helmut Creutz estima que en la década de 1990 sólo circulaba un tercio escaso de la cantidad total de marcos alemanes emitidos. Al parecer, el resto estaba metido en las huchas y en los cofres de dinero negro, o se había llevado al extranjero: en los tiempos del marco pudo haber en algunos momentos más dinero alemán en efectivo en Turquía que en la propia Alemania. Cinco años después de la implantación del euro seguían faltando catorce mil millones de marcos alemanes, que posiblemente no habían desaparecido en sumideros, lavadoras o pozos de la suerte.

Aunque diéramos por supuesto que se sabe cuánto dinero hay, qué es y qué parte de él está en circulación, quedaría sin respuesta la pregunta relativa al efecto que produce la disponibilidad de dinero. En el siglo XIX no había duda alguna de que el dinero sólo era un factor neutral, un «velo» ante la producción y el intercambio de mercancías. Sin embargo, más tarde se le ocurrió a alguien que había una relación entre el nivel de los intereses y la marcha de la coyuntura; no podía ser que el dinero careciera de influencia. A partir de 1936 adquirió máxima relevancia la tesis del economista británico John Maynard Keynes, según la cual no hay separación alguna entre la economía y el dinero. En el keynesianismo y el post-keynesianismo se supone que la política monetaria tiene influencias claras y duraderas en la economía real. Finalmente, en la década de 1950 Milton Friedman fundó el monetarismo: si acaso, el dinero ejerce sólo una influencia a corto plazo en la economía. Para ilustrar este hecho, Friedman ideó el lamentablemente hipotético «ejemplo del helicóptero», en el que llueve dinero del cielo, de tal modo que la masa monetaria en circulación se duplica de un día para otro. Como única consecuencia suben los precios a un nivel más alto, pero no cambia ninguna otra cosa. El efecto que esto tiene en la política monetaria es que habría que mantener todo lo más estable posible, mientras los keynesianos apuestan por una política

monetaria con un rumbo anticíclico. Hoy en día nadie cree ya en una total neutralidad del dinero, pero por desgracia esto responde a motivos diferentes y tiene distintas consecuencias.

Por una parte, hay que asombrarse de que, a pesar de la naturaleza misteriosa del dinero, en última instancia todo funciona muy bien, e incluso se consigue de vez en cuando pagar puntualmente las facturas. Por otra parte, los críticos de nuestro sistema monetario alegan que hemos perdido de vista las consecuencias perniciosas de las actuales prácticas monetarias y crediticias, porque las consideramos lógicas e inevitables. Con sólo dar un par de vueltas de tuerca al sistema, se conseguiría poner a raya la explotación, la distribución injusta de la riqueza e incluso la guerra. Quizá lo descubramos algún día. Siguiendo la vieja norma que dice «Primero limpiar la habitación y luego acabar con las costumbres guarras», seguramente no nos haría daño esforzarnos por conseguir una explicación de las cuestiones relacionadas con el dinero, antes de cambiar de arriba abajo la economía mundial.



Capítulo 12

Einemsen⁷

Lo que las ranas tienen por dentro es como un libro con siete sellos.

KATHRIN PASSIG

Por si no fuera ya suficientemente malo que los animales tengan a veces un aspecto tan absurdo que podemos pasar todo el día ante un cercado señalándolos con el dedo, suelen exhibir además modos de comportamiento descabellados, y probablemente lo hacen sólo para volvernos locos.

Por supuesto, les da igual que en consecuencia se tenga que malgastar a lo tonto un dinero que se destina a la investigación biológica. Por ejemplo, en unas doscientas cincuenta especies de aves se observa una de las costumbres más misteriosas del mundo animal: la actividad llamada *einemsen* o «baño de hormigas». Este bello concepto fue inventado en 1935 por el ornitólogo alemán Erwin Stresemann y expresa la acción de restregarse el plumaje con hormigas y otros insectos, así como con sustancias aromáticas como bolas antipolilla, cebolla cruda, agua jabonosa, trozos de manzana, vinagre, colillas de cigarrillos y frutos cítricos. Otros pájaros prefieren quedarse pasivos y que se lo hagan, mientras ellos permanecen con las alas desplegadas junto a un hormiguero o encima de él. Este frotamiento se ha observado también en jóvenes cuervos domesticados, sin que se lo hubieran enseñado antes otros cuervos mayores que ellos.

Pero no sólo las aves practican este deporte tonto, sino que también se ha observado en ardillas y especialmente en las ranas. El investigador especializado en ranas Martin Eisentraut describió en 1953 este comportamiento con todo detalle: «Si una rana entra en contacto con determinadas sustancias de sabor u olor característicos, sobre todo con aquellas que le resultan nuevas o a las que no está acostumbrada, se pone a lamerlas y, en algunos casos, a tomarlas con la boca y masticarlas con un vivo interés. Surge así un estado de excitación creciente y el animal segrega abundante saliva, que mediante los movimientos de masticación se

⁷ Este término es un invento ornitológico alemán que en inglés se traduce a veces por *anting*. En castellano sería 'baño de hormigas'. (N. de la t.)

convierte en una masa espumosa. Pasado cierto tiempo, la rana dobla el cuello hacia atrás mediante una contorsión especial y escupe la saliva sobre su piel rugosa o, mejor dicho, la lanza con la larga lengua que se dispara rápidamente hacia fuera». La excitación sólo empieza a disminuir cuando ya ha escupido varias veces. En los experimentos se ha conseguido que las ranas realicen este proceso utilizando en ellos pegamento, tabaco, sudor, jacintos, perfume, jabón, tinta de imprenta, tintura de valeriana, tejidos animales en putrefacción, piel de sapo, otra rana, humo de cigarro y olor de barniz. Incluso ranas muy jóvenes muestran este comportamiento, que en un experimento es difícil de provocar;

Eisentraut sospecha que para ello «la rana debe encontrarse en un determinado estado de ánimo».

Tanto las ranas como las aves, cuando se encuentran en esa situación, producen en el observador la impresión de estar extáticas, pero a veces pierden el equilibrio y se desmayan de alegría.

¿Y para qué todo esto? La hipótesis, una y otra vez formulada de que el ácido fórmico de las hormigas contribuye a mantener a raya los parásitos que viven en el plumaje (o en la piel rugosa o en el pelaje), fue refutada en 2004 por las ecólogas Hannah Revis y Deborah Waller: la concentración de ácido fórmico que hay en los cuerpos de las hormigas evidentemente no basta para frenar la proliferación de bacterias y hongos. Por lo tanto, mientras no surja otra cosa, hemos de suponer que los animales se hacen esas friegas, o sea se dedican a bañarse con hormigas, sólo para divertirse cuando da la casualidad de que hay cerca unos biólogos cámara en ristre. Pero, bueno, lo que seguramente los animales tampoco comprenden del todo es por qué nosotros nos desmayamos de entusiasmo.



Capítulo 13

Escritura del indo

Una vez que lo hubieron hecho, consumaron el sacrificio, según la costumbre, para no ser acusados de falta de devoción.

Traducción de inscripciones latinas de HENRY BEARD, *El latín completo para todas las ocasiones*

El valle del Indo está considerado como uno de los lugares clave de las primeras culturas escritas desde que en 1872 unos arqueólogos británicos descubrieron, en lo que ahora es la zona fronteriza entre Pakistán y la India, los primeros sellos de la cultura de Harappa, que tiene unos cinco mil años de antigüedad. Actualmente se conocen en total entre cuatro y cinco mil objetos con inscripciones, entre los que hay vasijas, fragmentos de barro, sellos de piedra y metal, amuletos, placas de cobre, armas y herramientas. Sin embargo, la escritura no se ha podido descifrar y el lenguaje utilizado sigue siendo desconocido. Se han publicado más de cien intentos de decodificación, y por cada teoría que es refutada, surge luego otra nueva.

En el supuesto de que se lograra descifrar la escritura del Indo, no se enriquecería mucho el mundo de la literatura, ya que la inscripción más larga que se conoce tiene sólo diecisiete signos, y por término medio los textos apenas tienen cinco signos, lo cual en el mejor de los casos alcanza para narraciones del tipo «chico encuentra chica». A pesar de esto, sería muy interesante saber en qué idioma se escribieron las inscripciones, si es que se trata de un idioma. Dada la brevedad de los textos, el historiador Steve Farmer, el indólogo Michael Witzel y el lingüista Richard Sprot defienden la hipótesis de que los símbolos de la cultura del Indo son más bien escudos de armas, marcas de propiedad o bonos, es decir, inscripciones que sólo servirían como identificación. Si se tratara de una auténtica escritura,

tendrían que verse más repeticiones de símbolos en las inscripciones, al menos esto es lo que se observa en escritos realizados con otros tipos de caracteres.

Por otra parte, la tesis tradicional de que se trata de una escritura está apoyada por el hecho de que los símbolos están ordenados en líneas, y no con una forma estética sin más, o colocados simplemente allí donde hay sitio para ello. Hacia el final de la línea se estrechan a veces, como si el autor no hubiera querido separar una palabra. En opinión de los defensores de la tesis de la escritura, los textos conservados son tan cortos porque los más largos se habrían escrito sobre algún material que no pudo sobrevivir cinco mil años. Farmer, Witzel y Sproat señalan, por contra, que todas las culturas antiguas conocidas que han tenido escritura han dejado textos más largos escritos en materiales duraderos, incluso aquellas de las que tenemos escasas noticias. (Si no salen pronto al mercado unos aparatos que graben el contenido de los discos duros en tablillas de barro, es posible que algún día pertenezcamos a las culturas de las que sólo quedarán algunos botones de chaqueta con una inscripción misteriosa).

Steve Farmer prometió en 2004 un premio de 10 000 dólares para quien encuentre la primera inscripción del Indo con más de 50 caracteres. Sin embargo, no es cuestión de que durante las próximas vacaciones en Pakistán saquemos de un montón de escombros una tablilla, y tampoco vale grabar los símbolos de puño y letra en una piedra, ya que la pieza de prueba debe proceder de una excavación oficial y ser reconocida como auténtica por personas expertas.

En el caso de que los escépticos tengan razón y los símbolos del Indo signifiquen en realidad nada más que «Estacionamiento prohibido» o «Denominación de origen Harappa», que no se desanimen los aficionados al deporte de pensar, porque aún quedan suficientes escrituras que nunca han sido descifradas: lineal A, meroítica, los ideogramas protoelámicos, unas veinticinco tablillas rongorongo de la Isla de Pascua, unas trece mil inscripciones etruscas, el «disco de Faistos» y, para los más avanzados será sumamente atractivo el «bloque de Cascajal», que está considerado como un imposible. Todo lo que se necesita es tener buenos conocimientos de algunas lenguas muertas y un poco de paciencia.



Capítulo 14

Estatura del ser humano

No puedes enseñar a ser más alto.

FRANK LAYDEN, Ex entrenador de baloncesto

Al nacer, los seres humanos son diminutos, pero al cabo de pocos años llegan a tener un tamaño que hace difícil imaginarse cómo en otro tiempo ese niño pudo caber todo él en el vientre de su madre. Después de unos veinte años, la mayoría de nosotros llega a tener una altura que varía entre ciento cincuenta centímetros y dos metros. El proceso de crecimiento que tiene lugar durante esos años se produce de manera furtiva y sólo se observa de forma indirecta en las rayas marcadas en la puerta de la cocina. Se desconoce gran parte de lo que sucede durante el crecimiento, por lo que este proceso no se comprende del todo. Sobre todo se plantea la pregunta relativa a qué es lo que determina finalmente la altura de una persona. ¿Por qué algunos seres humanos son más bajos y otros más altos?

Para la altura definitiva de una persona es importante su predisposición genética, que queda establecida de manera fija al crearse un nuevo ser humano. La herencia tiene importancia sobre todo al considerar casos individuales significativos: los padres bajos rara vez traen al mundo hijos altos, y los padres altos no suelen tener hijos bajos. No obstante, si se manejan las alturas medias de pueblos completos, desaparecen ampliamente esas diferencias individuales. La medida en que la estatura corporal media de un gran grupo de población es independiente de los genes, se puede constatar al comparar dos grupos tomados como muestra que, por término medio, presenten parecidas condiciones genéticas, pero hayan estado separados el uno del otro durante largo tiempo. Una separación de este tipo se mantuvo en Alemania durante más de cuarenta años, sin que mediara convenio alguno con la comunidad científica. El resultado de este «experimento» fue que los ciudadanos de la República Democrática Alemana eran al final, por término medio, un centímetro más bajos que los de la Alemania occidental. Las diferencias son aún más drásticas en el caso de Corea del Norte y Corea del Sur. Daniel Schwekendiek,

de la Universidad de Tubinga, descubrió que en el año 1997 los jóvenes norcoreanos de dieciséis años medían más de 15 centímetros menos que los surcoreanos de la misma edad. Dado que los norcoreanos y los surcoreanos no han podido tener en unas pocas décadas un desarrollo genético diferenciador, han de existir otros factores que sean responsables de la diferencia de estaturas. Actualmente se acepta mayoritariamente que la herencia sólo establece el límite superior de la altura. La medida en que un individuo se acerca a ese límite quedará determinará de alguna otra manera.

Como todos los seres vivos, las personas crecen sólo si se les suministran sustancias a partir de las cuales pueden construir nuevas células: proteínas, hidratos de carbono, grasas, grandes y agua, y además una larga lista de otras sustancias. Actualmente está fuera de toda duda que la composición y la cantidad de los alimentos influyen en la estatura media del cuerpo. Sin embargo, la alimentación no es con mucho el único factor. Las enfermedades, por ejemplo, frenan el crecimiento, porque el niño utiliza sus fuerzas en otra lucha y apenas le queda energía para crecer. Por lo tanto, es evidente que la prevención de las enfermedades desempeña un papel importante en el crecimiento, luego hay que suministrar vacunas a los niños, hay que hacerles reconocimientos preventivos y hay que vigilar regularmente su salud. Además, parece que la estatura de las personas depende también de factores difíciles de calibrar, como el cariño, la seguridad y la felicidad. Se cuenta que Maria Colwell, una niña inglesa nacida en la década de 1960, dejaba de crecer cuando la dejaban con sus padres, porque con ellos lo pasaba mal. Si la llevaban al hospital, empezaba a crecer inmediatamente. Otros niños, como Oskar Matzerath⁸, están convencidos de que no vale la pena crecer, y también Pipi Calzas Largas juraba a los dioses de los niños, unos dioses con forma de guisantes: «¡Querido pequeño Krummelus, nunca me haré mayor!». Ambos existen sólo en los libros, por lo que su opinión no tiene demasiado peso. Parece ser que la estatura corporal definitiva está determinada también por una serie de factores difíciles de precisar que en conjunto reciben el nombre de «nivel de vida biológico».

⁸ Protagonista de la novela de Günter Grass *El tambor de hojalata*, y de la película del mismo título dirigida por Volker Schlöndorff. (N. de la t.)

Recordemos lo anteriormente dicho: esto sólo tiene sentido cuando se trata de grandes grupos de población, por lo que en la mayoría de los casos no procede quejarse a los padres de que el suministro de alimentos ha sido insuficiente, sólo porque uno se ve demasiado pequeño.

Los especialistas en la relación entre la estatura humana y el nivel de vida, por ejemplo, John Komlos, de la Universidad de Munich, o el estadounidense Richard Steckel, se llaman «auxólogos». Su objetivo no es sólo descubrir las causas del crecimiento, sino, una vez que éstas se conocen, utilizar las estaturas corporales como indicador de la calidad de vida de las personas.

Se trata de una tarea importante, porque para incrementar la felicidad y el bienestar de la humanidad, es necesario primero encontrar un modo de determinar estas magnitudes difícilmente mensurables. Hasta bien entrado el siglo XX, conceptos tan ambiguos como «bienestar de la sociedad global» se determinaban sobre todo con la ayuda de datos procedentes de la investigación económica, a falta de otras alternativas. Se partía simplemente de la idea de que a las personas les iría mejor si aumentaba el producto nacional bruto. Por desgracia, muchos factores extraordinariamente importantes para la calidad de vida, como el índice de delincuencia, la seguridad del puesto de trabajo o la cercanía a la piscina, dependen sólo muy indirectamente de los meros datos económicos. Por eso, sería deseable que el bienestar de la humanidad se midiese con criterios mejores y más objetivos, por ejemplo, utilizando la estatura corporal. Esto tiene una ventaja importante: es fácil de determinar, incluso en el caso de personas que están muertas desde hace siglos.

En muchos casos es indiscutible la relación entre el nivel de vida y la estatura de las personas.

Basta con pensar que la estatura media en la mayoría de los países del llamado primer mundo ha aumentado continuamente desde que comenzó la industrialización. El europeo medio mide actualmente unos veinte centímetros más que hace ciento cincuenta años. Por otro lado la estatura media se reduce claramente en tiempos de guerra y destierro, como sucedió, por ejemplo, durante la Guerra Civil Americana. También las diferencias sociales se hacen notar: en la segunda mitad del siglo XVIII, los jóvenes de catorce años de las clases más

desfavorecidas eran más o menos una cabeza más bajos que los jóvenes de la misma edad pertenecientes a familias acomodadas.

Los africanos negros que viven en los prósperos Estados Unidos son bastante más altos que sus antiguos paisanos de África. Finalmente, el crecimiento deficiente de los niños de Corea del Norte se puede achacar evidentemente al hambre y a la quiebra económica. Hoy en día, las personas de estatura más elevada viven en Holanda y Escandinavia, es decir, en países con unos ingresos medios altos y una seguridad social amplia y completa.

Sin embargo, hay excepciones importantes que ponen de manifiesto la gran cantidad de factores diferentes que intervienen en el crecimiento físico de los seres humanos. Un ejemplo es la paradoja llamada «Antebellum Puzzle»: aunque en el siglo XIX, durante la primera fase de la industrialización, la renta per cápita subió claramente, las personas se encogieron, y lo hicieron en todos los países investigados hasta ahora. ¿Cómo se puede explicar esto? ¿Por qué se redujo la estatura? Se les daba más dinero, ferrocarriles y fábricas modernas, lo cual se entendería como mejor calidad de vida, ¿y ellos lo agradecen volviéndose más bajitos? No faltan en absoluto explicaciones plausibles. El experto en historia económica Jörg Baten investigó el fenómeno según los datos de la región de Münster. En esta región, la aparición del ferrocarril fue responsable del retroceso en estatura corporal, pero sólo en las zonas rurales situadas en el entorno cercano de la capital. La razón era que algunos productos, como la leche y los huevos, de repente podían venderse en el lucrativo mercado de la ciudad, con lo cual empeoró la situación alimentaria en el campo. Precisamente la proteína de la leche de vaca tiene una importancia especial en el proceso de crecimiento. En algunas zonas se observa una clara relación entre la estatura física y la cantidad de vacas por habitante. Aunque el ferrocarril y las vacas mejoran la calidad de vida, en determinadas circunstancias la acción conjunta de estos dos factores ocasiona una reducción de la estatura en las personas.

Como se ve en este ejemplo, las causas de la estatura que alcanzan finalmente los seres humanos pueden ser de lo más complejas. Por esta razón, los críticos de la auxología, por ejemplo la economista estadounidense Mary Eschelbach Hansen, piden mayor precaución para no interpretar precipitadamente las estaturas. La

relación entre calidad de vida y estatura no es clara y evidente en todos los casos. En ocasiones todo depende de cosas aparentemente triviales, como, por ejemplo, las vacas y los trenes.

He aquí uno de los mayores misterios de la auxología: los europeos han superado a los estadounidenses. A mediados del siglo XIX, los estadounidenses eran por término medio seis centímetros más altos que los europeos. El país era grande, las posibilidades ilimitadas, los alimentos estaban disponibles en abundancia y las enfermedades graves escaseaban. Incluso los esclavos negros llegaron a ser en aquella época manifiestamente más altos que los ciudadanos europeos. Sin embargo, entonces sucedió algo extraño. Mientras en Europa, desde hace doscientos años, han estado aumentando continuamente las estaturas medias, los estadounidenses en un momento dado simplemente dejaron de crecer. En la época de la Primera Guerra Mundial la diferencia era todavía de cinco centímetros, pero a mediados del siglo XX los europeos superaban en altura a los estadounidenses. Actualmente los europeos son por término medio varios centímetros más altos que los estadounidenses. Concretamente los europeos de mayor estatura, es decir, los holandeses y los escandinavos, superan a los estadounidenses en más de siete centímetros. Y eso a pesar de que los estadounidenses, desde hace alrededor de un siglo, son las personas más ricas del planeta.

¿Dónde queda entonces la prometida relación entre nivel de vida y estatura corporal? Si no se trata de dinero, ¿qué es entonces lo que favorece ese imparable crecimiento corporal en los europeos, pero no en los estadounidenses? John Komlos y Richard Steckel han comprobado numerosas posibilidades y han refutado la mayoría de ellas. La teoría favorita es actualmente la que atribuye el estancamiento de los estadounidenses a las siguientes causas: una mayor desigualdad social, una seguridad social más deficiente y una peor asistencia sanitaria. Holanda, por ejemplo, donde los hombres tienen una altura media de más de 1,80 metros, y donde además se bebe leche de vaca muy a gusto, tiene uno de los mejores sistemas sanitarios del mundo, especialmente por lo que respecta a las mujeres embarazadas y a los niños pequeños. En Europa, el acceso a los servicios de salud está organizado, en general, de una manera más justa que en Estados Unidos; la

proporción de personas que carecen de un seguro médico es claramente inferior, y la diferencia entre pobres y ricos es menor.

Pero si la desigualdad social fuera la explicación, ¿no tendría que haber en algún lugar de Estados Unidos un grupo de personas que no sigue la tendencia general, sino que continúa creciendo? Algún grupo de población ha de poder permitirse la atención en los mejores hospitales, con los mejores médicos, el consumo de los alimentos más caros y el empleo de la mayoría de las niñeras. Además, sería de esperar que las grandes diferencias entre pobres y ricos o a lo largo de los años. La pena es que todo esto no se puede demostrar. Parece como si todas las capas sociales de Estados Unidos, ricos y pobres, negros y blancos, cultos e incultos, hubieran dejado de crecer hace algunas décadas. Tiene que haber algún factor que ha impedido crecer a los norteamericanos, y que hasta ahora no se ha detectado.

Además no está claro qué altura pueden alcanzar los seres humanos. En algún punto están los límites genéticos, y posiblemente físicos, del crecimiento, porque un ser humano no es una jirafa.

Los primeros datos relativos a lo que sería la medida óptima de una altura corporal «sana» nos llegan de Noruega, que también es un país de gigantes modernos: la mortalidad de los noruegos se reduce a medida que aumenta la altura, alcanza un mínimo alrededor de 1,90 metros y después vuelve a aumentar. Por consiguiente, y desde un punto de vista estadístico, un ser humano vive más cuando es relativamente alto, pero sin ser un gigante. Las personas muy altas enferman, por ejemplo, de cáncer con mayor frecuencia, y se dan más golpes en la cabeza con los marcos de las puertas. Sin embargo, el límite superior sano no se ha alcanzado todavía como media en ningún lugar; incluso el holandés medio puede seguir creciendo sin preocuparse.

Para terminar, y simplificando el tema, supongamos que las personas que llevan una vida sana crecen más (seguimos hablando de un promedio). Incluso en este caso se plantea la duda de si ésa es una buena vida. Durante los siglos V y VI, las personas que vivían en la atrasada Baviera eran altas, sanas y longevas; sin embargo, muchas de ellas posiblemente habrían preferido vivir unos pocos siglos antes en Roma, donde todo era mucho más interesante, y tenían lugar unos combates entre gladiadores que, desde luego, eran muy entretenidos, aunque allí

las estaturas corporales hubieran dejado de aumentar. Por lo tanto, quizá habría que reflexionar un poco más sobre cómo se echa de menos la auténtica calidad de vida, y no «sólo» el buen nivel de vida biológico.



Capítulo 15

Estrella de Belén

Los de Star Trek eran todos judíos.

WILLIAM SHATNER

Cada año, aproximadamente dos semanas después de Navidad, las iglesias cristianas celebran la festividad de los Reyes Magos. El origen de esta tradición se basa en el Evangelio según san o nació Jesús, en Belén de Judea, bajo el reinado de Herodes, unos magos de Oriente se presentaron en Jerusalén y preguntaron:

“¿Dónde está el rey de los judíos que acaba de nacer? Porque vimos su estrella en el Este y hemos venido a adorarlo”».

Pocas veces un párrafo tan corto contiene tantos enigmas. No incluye ninguna información sobre los nombres de los astrólogos (lo que es seguro es que no se llamaban Melchor, Gaspar y Baltasar), sobre cuántos eran (es posible que varios, aunque se supone que fueron tres) ni de dónde venían (acaso de Persia o Babilonia). Además, en ninguna parte se dice de qué misteriosa estrella se trataba. La estrella de Belén, la estrella más influyente de la historia de las estrellas (a excepción del Sol), es un gran secreto.

La siguiente discusión parte de una serie de presuposiciones sin las cuales no tiene ningún sentido buscar la estrella de Belén. A partir de este punto daremos por sentado que Jesús de Nazaret fue una figura histórica y que el texto de los evangelistas del Nuevo Testamento se basa en testigos de la época, de tal modo que la historia de la estrella no se añadió a posteriori para darle más brillo al acontecimiento. Estas presuposiciones no están aceptadas al cien por cien; cada tanto tiempo, por ejemplo, resurge la teoría de que el Nuevo Testamento fue modificado por los romanos para fomentar la desunión entre los judíos. Sin embargo, estas tres suposiciones son mucho menos cuestionables que la explicación astronómica de la estrella de Belén. Además, hay que presuponer que la estrella de Belén no fue un acontecimiento mediático organizado por algún dios de la época, pues ese punto de vista descarta también la posibilidad de cualquier debate.

Si aceptamos esas hipótesis previas, a continuación resultaría de gran utilidad saber cuándo tuvo lugar realmente el nacimiento de Jesús. Por desgracia, la información al respecto es muy inexacta. De acuerdo con san Lucas y san Matías, el acontecimiento tuvo lugar durante el reinado de Herodes que, al oír la noticia del nacimiento del nuevo rey, mandó asesinar a todos los recién nacidos de la región de Belén para quitarse de encima a la competencia. Sin embargo, Herodes murió poco antes del inicio del sistema cronológico actual (que fue introducido a posteriori), por lo que nuestro calendario resulta algo inexacto para este propósito. La fuente más importante para determinar la muerte de Herodes es el historiador judío fariseo Flavio Josefo, que aproximadamente ochenta años más tarde escribió que Herodes había muerto poco después del eclipse lunar, pero antes de la Pascua Judía, cuya celebración tuvo lugar a continuación. Durante mucho tiempo se dio por sentado que Flavio Josefo se refería al eclipse lunar del mes de marzo del año 4 a. C. Sin embargo, entre éste y la Pascua judía transcurrieron tan sólo cuatro semanas durante las cuales tuvieron lugar numerosos acontecimientos históricamente probados y que requieren tanto tiempo como ejecuciones, sublevaciones o las larguísimas exequias de Herodes.

Por motivos que se desconocen, las copias del informe Flavio anteriores a 1552 huyen una fecha distinta para la muerte de Herodes: el eclipse lunar de enero del año 1 a. C., que sí dejaría tiempo suficiente hasta la Pascua. Por desgracia, hasta hace unos siglos la reproducción de documentos era bastante susceptible a error; de esa labor se encargaban los monjes, que debían copiar laboriosamente los documentos originales, pues en la Edad Media, y por motivos religiosos, a éstos les estaba prohibido el uso de fotocopiadoras. Tal vez hace quinientos años un monje cansado que trabajaba a la luz de una vela confundiera la fecha.

Una limitación ulterior aportará aún otra pista, ya que por ejemplo Lucas escribió en su Evangelio: «Todos iban a inscribirse en el censo, cada uno a su ciudad». Ese ominoso «censo» fue el que llevó a José y María a Belén. Esencialmente, existen tan sólo dos acontecimientos históricos que podrían haber motivado esa peregrinación masiva. Por un lado, y por motivos impositivos, todo ciudadano debía visitar su lugar de nacimiento cada veinte años, una exigencia a todas luces exagerada si la consideramos desde un prisma moderno. Uno de esos registros fiscales tuvo lugar

en el año 8 a. C., suficiente tiempo antes de la muerte de Herodes. Por otro lado, en el otoño del año 3 a. C., el emperador Augusto exigió que todos los ciudadanos del Imperio hicieran un juramento de lealtad al emperador, es decir, a él mismo. En particular los judíos debían prometer no intentar destronarle jamás de los jamases. Tanto eso como el registro fiscal justificarían un regreso a la propia ciudad de nacimiento. Si aceptamos el año 1 a. C. como el de la muerte de Herodes, sería posible que lo que llevó a María y a José a desplazarse a Belén fuera el acto de reafirmación imperial. La conclusión es que Jesús tuvo que nacer en algún momento entre el año 8 y el 1 a. C. Así pues, habrá que buscar el célebre fenómeno celeste en ese espacio temporal.

Llegados a este punto, es importante tener en cuenta no sólo el punto de vista astronómico, sino también el astrológico: la estrella no tiene que ser especialmente brillante, sino relevante.

Hace dos mil años, astronomía y la astrología eran aún la misma cosa. Las estrellas eran objeto de observación y cartografía, se estudiaban sus órbitas, se buscaban regularidades y todo ello servía para predecir el futuro. Con los siglos, sin embargo, la astrología fue perdiendo algo de prestigio, pues científicamente resultó cada vez más difícil justificar la relación entre nuestro futuro y unos lejanísimos globos gaseosos. Durante la época de Augusto la astrología era ciertamente popular, pero eso excluía a los judíos, que la consideraban una blasfemia. Así pues, no resulta tan sorprendente que a nadie en Israel se le ocurriera relacionar un fenómeno celeste con el nacimiento de un nuevo rey. Para los sabios, sin embargo, las señales celestes debieron de ser tan claras que exclamaron « ¡Hurra!», montaron en sus caballos y se pusieron en marcha hacia Occidente.

Algunos candidatos, largo tiempo discutidos, parecían muy apropiados, pero nuevos descubrimientos los han descartado casi por completo. Un buen ejemplo de ello es el afamado cometa Halley, que en numerosas imágenes aparece representado como la estrella de Belén. El Halley aparece regularmente cada setenta y cinco años, por lo que es de suponer que lo hizo en el año 12 a. C. Sin embargo, y tras todo lo dicho, se trataría de una aparición demasiado precipitada para anunciar el nacimiento de Jesús. A pesar de este percance, no tenemos motivos para compadecer al Halley, que se ha acabado haciendo igualmente famoso.

Teóricamente, cabría considerar también otros cometas que los atentos astrónomos chinos de la época se encargaron de documentar en el año 4 o tal vez incluso en el año 5 a. C. Sin embargo, se dan dos problemas: por un lado, los sabios (como su nombre indica) eran sabios y no tontos. Los cometas se distinguen perfectamente de las estrellas en el firmamento, pues se mueven de otra forma y a menudo dejan una estela a su paso, dos diferencias que por aquel entonces se conocían desde hacía ya tiempo.

Entonces ¿por qué iban los sabios a hablar de «estrella» si en realidad se referían a un cometa? Por otro lado, en el Imperio romano y en Persia los cometas solían anunciar desgracias, por lo que ante un descubrimiento así sería mucho más probable que los sabios se colocaran una bolsa de papel en la cabeza y corrieran a esconderse en el sótano a que anunciaran la gloriosa llegada de un rey.

La segunda posibilidad que baraja la astronomía moderna es la de la aparición en el firmamento de una supernova, una «nueva estrella». Actualmente, los fenómenos de ese tipo se observan regularmente con grandes telescopios y responden o bien a la agonía de una estrella especialmente grande o a la explosión de una estrella que ya ha muerto. Hoy en día, y sabiendo eso, a nadie se le ocurriría la idea de que una supernova pudiera ser un presagio de fortuna y gloria, pero en tiempos de Augusto no se sabía nada de los golpes del destino que rigen la vida de las estrellas. En cualquier caso, el efecto es impresionante: en un punto del cielo en el que antes había solo oscuridad, aparece de repente una estrella rutilante. Johannes Kepler, que en el año 1604 fue testigo de una de esas supernovas, fue el primero en relacionar la estrella de Belén con un fenómeno de esa naturaleza. En cualquier caso, es extraño que a los grandes astrónomos chinos de la época en cuestión no se les ocurriera también algo parecido. Además, una supernova no se mueve en relación con el resto de estrellas, sino que se mantiene siempre en la misma constelación, algo que contradice claramente el relato de san Mateo. Por ello, y aunque resultaría de lo más cómoda, la mayoría de expertos descartan también la teoría de la supernova.

Las teorías que actualmente se consideran plausibles no se centran en un único objeto celeste, sino en la interacción de varios planetas, tal vez con la ayuda también de la Luna. Para hacer honor a la verdad, a partir de ahora tendríamos que

eliminar del árbol de Navidad todas las estrellas con cola y sustituirlas por tres puntos de luz móviles. En la época en cuestión se produjeron diversas alineaciones planetarias singulares. En el año 7 a. C., Júpiter y Marte se encontraron tres veces en un período de apenas siete meses y, por si eso fuera poco, todas ellas en la constelación de Piscis, símbolo antiquísimo del judaísmo. Júpiter era considerada la estrella del rey y Saturno la protectora de los judíos, por lo que su alineación podía interpretarse perfectamente como el nacimiento del rey de los judíos, tal como apuntó el astrónomo austriaco Konradin Ferrari d'Occhieppo en la década de 1960. En el año 6 a. C. Júpiter, Saturno y en esta ocasión también Marte volvieron a encontrarse bajo el signo de Piscis que en aquella época, al parecer, era un punto de reunión popular entre los planetas. Ése es, teóricamente, otro fenómeno que debemos considerar.

Sin embargo, la sucesión de fenómenos que tuvieron lugar entre los años 3 y 2 a. C., simultáneamente con el solemne juramento de lealtad del emperador Augusto, resultan mucho más impresionantes. En mayo del año 3 a. C., Saturno y Mercurio se alinearon en un espacio estrechísimo. A continuación, Saturno siguió su órbita y en junio se encontró con Venus. Por si eso no era suficiente, Venus, siempre tan ávido de placeres, tuvo en agosto una cita con Júpiter y, a los pocos días, otra con Mercurio. Diez meses más tarde, en junio del año 2 a. C., Júpiter y Venus volvieron a encontrarse y en esta ocasión se juntaron tanto que para el ojo humano debieron de parecer un único cuerpo, extremadamente brillante. Además, el fenómeno tuvo lugar en la constelación del León, el rey del zodiaco. Así pues, Júpiter, el planeta real, se funde con Venus en una constelación también real, al mismo tiempo que había luna llena. ¿Existe alguna forma más refinada de anunciar la llegada de un nuevo rey? Unas pocas semanas más tarde, Júpiter, Venus y Mercurio se encontraron de nuevo bajo el signo del león; sólo Saturno faltó a la cita sin avisar. En la misma época, Júpiter realizó una breve gira por el cielo: primero, en el año 2 a. C., circundó la estrella de Régulo, el cuerpo más luminoso de la constelación de Leo, conocido como estrella del rey, antes que en diciembre del año 2 a. C. prácticamente se detuviera durante varios días ni más ni menos que en medio de la constelación de Virgo. Visto desde Jerusalén, durante aquellas noches Júpiter indicó el camino hacia Belén... donde aguardaban la Virgen con el rey (Júpiter) en sus

entrañas. Tampoco en esa ocasión se podía acusar al cielo de expresarse de forma equívoca.

Fue el historiador y meteorólogo Ernest L. Martin quien, en 1991, presentó esos dos espectaculares años de fenómenos planetarios constantes como explicación astronómica de la estrella de Belén. Sin embargo, para ello sería necesario que Herodes no hubiera muerto ya en el año 4 a. C., algo que aún debe demostrarse de forma concluyente.

Pero es que existe una teoría aún más reciente de otro astrónomo. Michael R. Molnar se dedica en su tiempo libre a coleccionar monedas antiguas. Analizando monedas romanas, llegó a la conclusión de que era un error dar por supuesto que la constelación de Piscis representaba a los judíos. En cambio, aseguró que existen pruebas que establecen un fuerte vínculo simbólico entre el judaísmo y la constelación de Aries. Eso, según Molnar, lo cambia todo. En abril del año 6 a. C., el Sol, Venus, Marte, Júpiter y durante un instante también la Luna (es decir, todos los personajes importantes del sistema solar) coincidieron en la constelación de Aries. Molnar está convencido de que ésa fue la «estrella» de Belén. O, tal como dijo su colega Brad Schaefer: «Caramba, a los astrólogos se les debió de caer el turbante». La teoría de Molnar ilustra el problema principal de las indagaciones sobre la estrella de Navidad: no basta con encontrar una estrella brillante, sino que también es necesario conocer el significado de la estrella en cuestión. El resultado final es una especie de astrología del absurdo, que no sirve para predecir el futuro sino para aclarar el pasado.

Como suele suceder con la mayoría de enigmas antiguos, es muy probable que la historia de la estrella de Belén no llegue a aclararse nunca. Por otro lado, la astronomía reciente ha realizado progresos importantes al respecto, por ejemplo descartando la teoría de la supernova o estableciendo la fecha exacta de la aparición del cometa Halley. Además, algunas de las teorías más plausibles han sido sugeridas en los últimos veinte años y eso resulta realmente esperanzador.

Sin embargo, hay algo que conviene no perder de vista: existe la posibilidad de que, al final, todos los intentos serios de explicar el fenómeno resulten ser vanos. De hecho, podría muy bien ser que un perro volador con un farolillo de papel en la boca hubiera inducido a los sabios a error.



Capítulo 16

Eyacuación femenina

¡Y, sin embargo, chorrea!

oekonews.de: «ÖKO-TEST Feste
Wandfarbe

Por un lado es sorprendente que no se sepa desde hace tiempo todo lo relativo a cuestiones tan elementales y comparativamente tan agradables de investigar como la eyacuación femenina y la zona de Gräfenberg (alias punto G). Por otro lado, la literatura médica no descubrió el clítoris hasta el siglo XVI, o sea, unos cien millones de años después de su introducción en el mercado.

Desde luego, se puede suponer que los profanos interesados en el tema lo habrían detectado ya alguna vez, pero, sin embargo, en el siglo XVII el anatomista danés Gaspar Bartholin criticaba a sus predecesores el hecho de que pretendieran adjudicarse ese supuesto descubrimiento, porque, según decía él, el clítoris ya era conocido en la Roma antigua, cosa que no parece del todo improbable.

También la eyacuación femenina se describió no pocas veces en textos como el Kamasutra, los escritos de Aristóteles y de otros griegos, y la literatura pornográfica. Hasta entrado el siglo XVIII, se sospechaba asimismo que sin el «semen femenino» no podía producirse en ningún caso la fecundación. Incluso en los textos científicos sobre sexología escritos a principios del siglo XX, como los de Richard Krafft-Ebing, Max Marcuse, Havelock Ellis y Magnus Hirschfeld, aparecen las llamadas «poluciones femeninas». Poco después, la eyacuación femenina se convirtió en un tema pasado de moda, al menos en la literatura médica, y durante algunas décadas se consideró como un mito acuñado por las ensoñaciones masculinas.

En general, después de un breve florecimiento durante las décadas de 1920 y 1930, la sexología ha avanzado con paso lento, lo cual podría deberse, entre otras cosas, a que, cuando en las universidades se ha pretendido investigar el orgasmo, siempre se han alzado voces airadas, tanto en Estados Unidos como en Europa. De todos modos, el contribuyente siempre ha sospechado que en la universidad se investiga

mucho sobre orgasmos y se trabaja poco. Así se podría explicar, quizá, que por ahora la mayoría de los médicos y estudiantes de medicina sepan sobre las partes y funciones del cuerpo femenino implicadas en el tema menos que el consumidor de pornografía medianamente avisado. Desde hace pocos años se considera en cierto modo indiscutible la existencia de la eyaculación femenina, pero los detalles siguen sin estar claros. Por el contrario, en el caso de los hombres se sabe con exactitud cómo, por qué y mediante qué órganos se produce una eyaculación.

El anatomista holandés Regnier de Graaf fue uno de los primeros que se dedicaron a investigar la cuestión de los órganos implicados. En 1672 escribió sobre una «próstata femenina», que al igual que la masculina estaría situada en torno a la uretra y cuya descarga «produce tanta voluptuosidad como la de la “próstata” masculina». Citaba al anatomista griego Herófilo de Calcedonia (300 a. C.) y al médico griego Galeno (siglo II), que asimismo informaron sobre la existencia de una próstata femenina, y sospechaba que la secreción de dicha próstata salía en parte por unos orificios situados en la uretra. Aconsejamos al lector que no se precipite a soltar una carcajada al leer esto, porque unos pocos párrafos más adelante podría arrepentirse.

En 1880, el ginecólogo escocés Alexander Skene describió las llamadas glándulas de Skene (también conocidas como glándulas parauretrales), situadas junto a la uretra femenina, pero no pudo averiguar cuál era la función de estas glándulas. En 1926, el ginecólogo holandés Theodoor Hendrik van de Velde en su libro *Het volkomen huwelijk*⁹, que fue un éxito de ventas, habló detalladamente de la posible existencia de una eyaculación femenina. «No hay duda de que se produce, al menos en parte de las mujeres», afirmaba. Sospechaba que el origen de ese flujo se encontraba en las glándulas de Bartholin, que también se encargan de la humidificación de la parte anterior de la vagina. Las glándulas de Skene serían, según Van de Velde, demasiado pequeñas «para hacer posible una acumulación de secreción que pueda eyacularse».

El ginecólogo alemán Ernst Gräfenberg describió finalmente en 1950 una «zona erógena en la pared vaginal anterior, a lo largo de la uretra», cuya existencia él mismo confirmaba mediante su propia «experiencia con numerosas mujeres». Del

⁹ El matrimonio ideal, Bruguera, Barcelona, 1968. (N. de la t.)

artículo se desprende claramente que Gräfenberg había obtenido sus datos en el ámbito privado. Desde entonces, rara vez se ha visto tanta franqueza en los textos sexológicos. Según Gräfenberg, algunas mujeres en el momento del orgasmo expulsan por la uretra gran cantidad de un fluido claro, que no es orina (cosa que él, sin embargo, no había investigado en el laboratorio). Como consecuencia de esa hipótesis formulada con tanta prudencia, probablemente se trataba de secreciones producidas por algunas glándulas ubicadas en el interior de la uretra, que estaría en relación con las zonas erógenas descritas. No es cuestión de pensar en una función como lubricante, ya que el fluido no se segregaría hasta producirse el orgasmo. En 1953 apareció este artículo de Gräfenberg reelaborado como capítulo de un libro de sexología. Ya no se trataba de una apreciación personal; el párrafo referente a la eyaculación femenina fue suprimido por razones desconocidas.

En un primer momento el artículo de Gräfenberg fue ampliamente ignorado. El sexólogo Alfred Kinsey y sus colaboradores, en su famosa obra de 1953 *El comportamiento sexual de la mujer* [*Sexual Behavior in the Human Female*], se limitaron a mencionar que «las contracciones musculares de la vagina que se producen como consecuencia del orgasmo [...] [pueden] expulsar algunas secreciones vaginales y, en unos pocos casos, lanzarlas hacia fuera con una cierta fuerza».

El hecho de que la «zona» de Gräfenberg no aparezca en el libro de Kinsey y que se presente el interior de la vagina como insensible se debe fundamentalmente a que Kinsey quería borrar de la faz de la Tierra la idea de un «orgasmo vaginal» decisivamente acuñada por Freud. Desde un punto de vista científico esto no era jugar limpio, pero muchas mujeres pudieron estar agradecidas a Kinsey: durante décadas se había esperado de ellas que en el transcurso de su «maduración psicosexual» aprendieran a renunciar al orgasmo clitórico en beneficio del orgasmo vaginal, que respondería a un comportamiento «más maduro».

Después de Kinsey, durante veinticinco años no hubo muchas novedades, salvo una o dos tímidas menciones del tema en la literatura especializada. También los sexólogos William Masters y Virginia Johnson, en su rompedor estudio *Human Sexual Response*¹⁰, de 1966, para el cual se utilizaron por primera vez datos sobre

¹⁰ *Respuesta sexual humana*, Intermédica, 1967. (N. de la t.)

el comportamiento sexual humano obtenidos en el laboratorio, calificaron la eyaculación femenina de «concepto erróneo, aunque muy difundido».

Más tarde, Masters y Johnson reconocieron que en algunas mujeres podía llegar a producirse una reacción sexual parecida a una eyaculación, pero explicaron el fenómeno como una incontinencia urinaria y recomendaron que en ese caso se consultara con un médico.

Hubo que esperar hasta finales de la década de 1970 para que la eyaculación femenina fuera redescubierta en el contexto del movimiento feminista, y en los diez años siguientes fuera incluida en algunos estudios y encuestas. En 1982, los psicólogos y orientadores sexuales Alice Kahn Ladas, Beverly Whipple y John D. Perry publicaron el libro *The G-Spot and Other Recent Discoveries About Human Sexuality*, que hizo popular la expresión «punto G», que todavía se usa, aunque puede inducir a error, para referirse a la zona descrita por Von Gräfenberg. Por primera vez se abrió un amplio debate sobre la zona de Gräfenberg en círculos especializados.

Hoy en día aún se menciona una y otra vez, aunque hasta ahora no se haya conseguido demostrar la existencia de abundantes terminaciones nerviosas u otras peculiaridades anatómicas en el lugar indicado de la pared vaginal. No era así como se describía en la tesis de Gräfenberg, según la cual se trataría de una zona erógena porque desde allí se podría estimular el sensible tejido glandular situado tras la pared de la vagina y en torno a la uretra.

Desde la década de 1980, el discutido fenómeno de la eyaculación femenina se investigó ocasionalmente en condiciones de laboratorio. Por desgracia no es fácil obtener ese fluido aislándolo de los otros que están presentes en la actividad sexual. Sin embargo, en el análisis se encontró en muchos casos, aunque no siempre, una elevada concentración (en comparación con la de la orina) de una sustancia llamada «fosfatasa ácida prostática» (PAP en inglés, FAP en castellano), así como fructosa, siendo ambas sustancias características de las secreciones masculinas de la próstata. En cambio, las concentraciones de urea y creatinina, que son importantes componentes de la orina, eran bajas en la mayoría de los casos. Más tarde se constató, sin embargo, la presencia de la FAP también en la secreción vaginal, con lo que surgió una pregunta relativa a si el fluido no viene, al menos en

parte, de la vejiga y simplemente sucede que se mezcla con secreciones glandulares en la uretra. Para poner las cosas aún más difíciles se constató el hecho de que las mujeres, tanto individualmente como según la fase del ciclo en que se encuentren, producen fluidos de distintas composición: unas veces la secreción recogida era blanquecina, otras veces transparente, en unas muestras se encontraban más similitudes con la orina, en otras menos, y también la cantidad descrita en las distintas publicaciones oscilaba entre 10 y 900 mililitros. En contra de la teoría de la orina se puede alegar que el característico olor a espárragos que, por causas genéticas, aparece en la orina en más o menos la mitad de las personas después de consumir esta hortaliza, no se presenta en el líquido eyaculado por las mujeres (y tampoco en el de los hombres). Un experimento privado, hasta ahora no reproducido, de una alumna del investigador canadiense Edwin Belzer dio como resultado que el fluido eyaculado no presentaba influencia alguna de un medicamento que tiñe la orina de un intenso color azul.

A finales de la década de 1980, dos amplios estudios realizados en Estados Unidos y Canadá dieron como resultado que el 39,5 por ciento de las mujeres encuestadas habían eyaculado un fluido una o varias veces. El 65,9 por ciento afirmaba tener una zona sensible en la vagina, y un 72,6 por ciento de estas mujeres podía llegar al orgasmo por estimulación de dicha zona, más de la mitad de ellas sin estimulación adicional del clítoris. En este subgrupo del 65,9 por ciento había un 82,3 por ciento que habían experimentado la eyaculación femenina.

Por lo pronto, se dedujo de todo esto que entre el 10 y el 40 por ciento de todas las mujeres habían eyaculado al menos algunas veces, y se mantuvo esta idea hasta que el sexólogo Francisco Cabello Santamaría en 1996 analizó orina femenina y halló el llamado antígeno prostático específico (PSA en inglés, APE en castellano), que, como su nombre indica, sólo se produce realmente en la próstata del hombre. Comprobó que en el 75 por ciento de las pruebas la concentración de APE era más elevada después del orgasmo que antes. Cabello Santamaría concluye a partir de esto que, aunque todas las mujeres son capaces de eyacular, en la mayoría de los casos el fluido producido va a parar a la vejiga; un fenómeno que también se da en los hombres y se conoce como «eyaculación retrógrada». En un experimento realizado en 2001, por el sexólogo Gary Schubach se puso de manifiesto que las

personas que tuvieron una eyaculación y cuya vejiga había sido vaciada antes con un catéter, produjeron durante el orgasmo entre 50 y 900 mililitros de fluido. Este fluido se extrajo asimismo mediante un catéter especial para la vejiga, con el fin de asegurar que provenía sólo de ésta. El contenido de urea y creatinina que se halló en este fluido era claramente inferior al que suele hallarse en la orina. Dado que el catéter introducido cerraba toda comunicación entre la vejiga y la uretra, el fluido recogido no podía provenir de las glándulas que desembocan en la uretra. Sin embargo, no se explicaba en esta investigación la causa de que en una vejiga recién vaciada se acumulara tanto líquido con unas características tan inusuales para la orina. La prudente conclusión de Schubach fue que la excitación sexual influía en la producción de fluido dentro de la vejiga.

Es poco lo que ha quedado realmente claro. En todo caso, ahora ya no se discute que, no sólo en algunas mujeres, sino en todas, existe en torno a la uretra femenina un tejido glandular que, tanto en su estructura como en su función, se parece a la próstata masculina y cuyos conductos desembocan en parte dentro de la uretra, y en parte a la salida de ésta. Este tejido glandular está activo en sus funciones y no es, como se creía todavía a finales de la década de 1980, simplemente un resto atrofiado.

Por ahora, la investigación no ha ido mucho más allá. Sigue sin estar claro si la «próstata femenina», es decir el tejido glandular situado en torno a la uretra, hace que la zona de Gräfenberg sea erógena, o si esas glándulas no son, al menos en algunas mujeres, más grandes o más productivas que lo que se ha supuesto hasta ahora, con lo cual se explicaría que las cantidades de fluido sean en algunos casos abundantes. También habría que averiguar si la eyaculación es una parte integrante de la respuesta sexual, y en caso afirmativo qué objeto podría tener, o si se trata más bien de un efecto secundario. En caso de que el fluido, en parte o en su totalidad, provenga de la vejiga, se plantea la pregunta de si el músculo que cierra ésta se puede abrir a causa de la estimulación, y cómo sucede esto. Como la mayoría de los hombres sabe por experiencia propia, la excitación sexual no trae consigo de modo alguno la apertura del músculo que cierra la vejiga, sino todo lo contrario. No hay razón para pensar que en las mujeres vaya a suceder de otra manera. En vez de investigarse primero el punto G con mayor precisión, se sacan al

mercado continuamente nuevos puntos, como el punto K (el clítoris, que ya no es tan nuevo), el punto U (en el meato urinario) y finalmente en 2003 el punto A, que tendría que estar entre el punto G y el cuello del útero. Quedan libres 22 puntos entre la B y la Z, por lo que hay un campo muy amplio para los investigadores que quieran hacerse famosos.

Si se llegara a aclarar finalmente cuál es exactamente la composición del fluido eyaculado y dónde se produce, éste sería un conocimiento útil, entre otras razones porque los resultados influirían en el trabajo de los censores del British Board of Film Classification: en Inglaterra están prohibidas todas las imágenes que tengan relación con juegos urinarios en las relaciones sexuales, y la eyaculación femenina está considerada por el BBFC simplemente como una expresión inventada para minimizar esas cochinadas ilegales.

Otro dato interesante es que en la literatura especializada del siglo XX, en la medida en que se basa en encuestas realizadas a mujeres, como por ejemplo el superventas de Shere Hite titulado *El informe Hite*, para el que se preguntó a apenas 2000 mujeres sobre sus experiencias sexuales, la eyaculación femenina no parece desempeñar papel alguno. En cambio, si actualmente alguien plantea una pregunta al respecto en los foros de Internet o en las listas de correo, son numerosas las mujeres que toman la palabra inmediatamente para mostrar que conocen el fenómeno y tienen opinión propia al respecto. Mientras la eyaculación femenina no existía oficialmente, lo que sucedía era que o bien no se mencionaba, o no se consideraba digna de mención, o se silenciaba por vergüenza a causa de la teoría que la relacionaba con la incontinencia urinaria. Así es como la ciencia, en ocasiones, hace que crezcan nuevos órganos.



Capítulo 17

Follaje otoñal

*¡Piensa en el castaño, que después de la
exuberancia del verano da paso a un
Nude Look provocadoramente mezquino!*
HILFSCHECKERBUNNY

La cuestión relativa al porqué del cambio de color que experimentan los árboles en otoño está siempre presente en cualquier recopilación de preguntas planteadas frecuentemente por los niños, y también por los adultos. La respuesta suele ser la siguiente: cuando la clorofila, es decir, la sustancia que tiñe de verde las hojas, se descompone, aparecen en primer plano los demás colorantes vegetales que estaban hasta entonces encubiertos. Por lo que respecta a los carotinoides, que son los causantes de los colores amarillos y anaranjados, esta explicación es correcta, pero los tintes rojos que aparecen en el follaje otoñal de muchos árboles (los antocianos) no se forman hasta el momento del cambio de color. Aquí es donde surge la pregunta:

¿Por qué se toma la naturaleza tantas molestias? Pues la naturaleza es perezosa y no mueve un dedo si no hay una buena razón para ello al contrario que los diligentes biólogos, a los que las numerosas preguntas abiertas dieron el motivo necesario para celebrar en 2001 un simposio sobre el tema «Por qué las hojas se vuelven rojas».

Comencemos por los hechos conocidos: en otoño las hojas de muchos árboles frondosos cambian de color dentro de una gama determinada. Cuando los días se hacen más cortos y las temperaturas descienden, los árboles comienzan a almacenar nutrientes que necesitarán de nuevo en primavera, y los que están en las hojas pasan a depositarse en capas profundas de la corteza y en las raíces. Aparecen en el follaje unos colores especialmente luminosos cuando hace frío y brilla el sol al mismo tiempo, por ejemplo a la mañana después de una noche despejada. Si el otoño se presenta lluvioso y con nieblas, puede que no se fabrique una cantidad suficiente de azúcar por falta de oportunidades para realizar la

fotosíntesis, y ese azúcar es necesario para la producción de antocianos. Los distintos tipos de árbol tienen preferencias diferentes por lo que respecta al cambio de color: los alisos y las hayas toman un color amarillo; los robles, un marrón rojizo; los arces, amarillo, naranja y rojo; pero a las coníferas no les afecta para nada este asunto, salvo alguna que otra excepción.

Los antocianos fueron descritos por primera vez en 1835 por el farmacéutico alemán Ludwig Clamor Marquart en su trabajo titulado «Los colores de las florescencias», que dice lo siguiente:

«El antociano es la sustancia colorante que está en los azules, violetas y rojos, y proporciona asimismo el color a todos los marrones y a muchas flores de color anaranjado». Al principio se pensó que los antocianos que aparecen en el follaje otoñal eran un producto de desecho de la descomposición de la clorofila, pero más tarde se comprobó que la producción de antocianos tiene lugar a menudo antes de que desaparezca la clorofila. A finales del siglo XIX, los botánicos observaron que la producción de antocianos aumenta tanto a temperaturas bajas, como cuando la radiación luminosa es más fuerte. En consecuencia, se supuso que los antocianos protegen las hojas contra la luz y el frío. A mediados del siglo XX se descubrió que también los rayos ultravioleta reactivan la producción de antocianos. Actualmente se sospecha que los antocianos protegen las plantas contra los daños que produce la luz ultravioleta. Por desgracia se observó en la década de 1980 que los antocianos a penas proporcionan protección en un espectro de radiación ultravioleta especialmente perjudicial y además se forman en el interior de las hojas, con lo cual la supuesta protección es algo tan absurdo como beberse la leche solar en vez de extenderla en la piel. En cualquier caso, durante la década de 1980 surgió la sospecha de que los árboles en otoño se apresuran a acumular las sustancias perjudiciales en las hojas, con el fin de librarse de ellas —una especie de recogida de basura.

En las últimas décadas, gracias a la mejora de los métodos de medición, se ha conseguido averiguar más cosas sobre el cambio de color de las hojas. La antigua teoría de la protección contra la luz volvió a cobrar vida en la década de 1990, cuando el experto en botánica tropical David Lee y el fisiólogo Kevin Gould pudieron justificar que las hojas con pigmento rojo se adaptan a una iluminación fuerte y

oscilante mejor que las verdes. La fotosíntesis, por ejemplo, alcanza su mejor nivel cuando hay una claridad homogénea y el aparato fotosintetizador puede adaptarse exactamente a esas condiciones de luz. Varios estudios demostraron durante los años siguientes que las hojas más viejas son más propensas que las jóvenes a una inhibición de la fotosíntesis producida por el exceso de luz. Quizá por esa razón necesitan en los últimos años de su vida una protección especial, que consiguen gracias a los antocianos.

Pero los antocianos pueden hacer aún más cosas: si se alimenta a los ratones con arándanos o se da a las personas vino tinto (ambos alimentos de color intenso y con alto contenido en antocianos), entonces sólo estas últimas se emborrachan, pero en la sangre de todos ellos, humanos y ratones, aumenta la cantidad de antioxidantes, que retienen los radicales libres. Estos han perdido uno de sus electrones, o a los que les apetecería tener uno más que antes, y por eso se ponen agresivos para conseguir apropiarse de ese electrón, que pueden tomar del ADN, de las membranas celulares o de cualquier proteína importante. Lo sensato es impedirselo, porque esos daños pueden, entre otras cosas, producir un cáncer. Para investigar si esta función favorece también a las hojas de plantas vivas, Kevin Gould y sus discípulos picaron unos agujeros en hojas rojas y verdes de una planta neozelandesa. Los radicales libres que surgían en las zonas agujereadas desaparecían en las hojas rojas mucho más rápidamente que en las verdes. Pero ¿cómo hacen los antocianos para proteger a la planta contra esos daños? «Es un fenómeno bastante misterioso», reconocen Lee y Gould, ya que los antocianos se fijan mayoritariamente en las vacuolas celulares (unas grandes burbujas llenas de líquido), mientras que los radicales libres realizan su trabajo en otros lugares de la hoja.

No obstante, hay varias funciones de protección de los antocianos que están ya bien demostradas, aunque no tan bien explicadas. Lo que, sin embargo, sigue siendo una incógnita es por qué los árboles invierten tanta energía en la protección de unas hojas que de todos modos se van a caer. ¿Para qué esforzarse en pintar un coche al que sólo le quedan tres días de ITV? Puede que los antocianos se encarguen de desmontar y almacenar el complejo laboratorio de la fotosíntesis. Quizá se trate también de recuperar el nitrógeno que va unido a los utensilios de la fotosíntesis y

que, en otro caso, no haría sino caer del árbol; las plantas no están dispuestas a desprenderse del nitrógeno que han conseguido tan fatigosamente —les pasa lo mismo que a los humanos con el dinero.

Otro modelo explicativo es el que propuso el biólogo estadounidense Frank Frey, quien en 2005 regó unas semillas de lechuga con extractos de hojas amarillas, verdes y rojas: las semillas tratadas con extracto de hojas rojas de arce brotaron y crecieron claramente peor. Según la hipótesis de Frey, los árboles que tiene un follaje otoñal especialmente rico en antocianos envenenan el suelo para otras especies cuando sus antocianos pasan del follaje a la tierra. Se sabe que los avellanos, los castaños, los manzanos y los pinos eliminan la competencia con unas técnicas igualmente innobles.

A partir de una idea del teórico de la evolución William D. Hamilton, en el año 2000 los biólogos Archetti y Brown empezaron a desarrollar la «hipótesis de las señales» del follaje otoñal, según la cual los árboles sanos y resistentes se adornan en el otoño con unos colores especialmente llamativos. De esa manera indican a los parásitos, sobre todo a los pulgones de las hojas, que se pueden permitir unos colores costosos y que, por lo tanto, tampoco van a ahorrar esfuerzos para defenderse, algo no muy distinto del caso de esos humanos que con su color de piel hacen saber que tienen dinero suficiente para pagarse un solarium. Entonces, los pulgones listos se instalan para pasar el invierno en árboles menos resistentes. Por ahora, la hipótesis de las señales se basa en reflexiones puramente teóricas, y en su contra se puede alegar una relación existente entre la concentración de antocianos y la de determinados anticuerpos, que ha demostrado el biólogo Martin Schaefer. Según esto, el árbol no tendría interés alguno por enviar mensajes a los pulgones de las hojas; los pulgones listos podrían quizá darse cuenta de que existe una relación entre el color y el veneno.

En el año 2004 unos biólogos israelíes que trabajaban con Simcha Lev-Yadun publicaron una teoría según la cual las distintas coloraciones del follaje servirían en general para no facilitar el camuflaje de los insectos. Así los insectos verdes que devoran las hojas estarían en otoño más expuestos a que sus enemigos los devoraran. Dado que el cambio de color que se produce en otoño dura poco, la presión de la selección para que los insectos se adapten no es muy fuerte; en

cualquier caso, ningún insecto verde ha sido hasta hora suficientemente refinado como para cambiar de color al mismo tiempo que lo hace el follaje. La psicóloga Linda Chalker-Scott desarrolló la teoría de que los antocianos sirven como protección contra las heladas: al contrario que la clorofila y muchos otros tintes, los antocianos se disuelven en el agua, y el agua en que están disueltas las sustancias se congela a una temperatura más baja que el agua normal. Pero, también se podría pensar que los antocianos frenan el crecimiento de ciertos hongos. Esta hipótesis surgió en la década de 1970, cuando se observó que las hormigas criadoras de hongos ponían sumo cuidado en no alimentarlos con hojas rojas. Quizá las hormigas tuvieran sus razones para actuar así, y en cualquier caso una razón mejor que el hecho de que no les gustara el color rojo, porque también un estudio de la Universidad de Friburgo demostró que los extractos de antocianos frenan el crecimiento de los hongos en los frutos de las plantas.

En conjunto, durante los últimos diez años se han registrado grandes avances en el tema del follaje otoñal, aunque quedan algunas preguntas abiertas. Por ejemplo, ¿qué función tiene el cambio al color rojo que se observa a veces en hojas jóvenes? ¿Por qué algunas plantas son rojas durante todo el año? ¿Por qué sucede a menudo que árboles muy próximos, de la misma especie, cambian a colores muy diferentes en el otoño? ¿O es que en realidad no cambian de color? Quizá se trata sólo de nuestros ojos, que se preparan para el otoño.



Capítulo 18

Goteo

¿Cuál es el verdadero aspecto de una lágrima? No muy romántico, la verdad. Primero es como una naranja ensartada en una aguja de punto y luego como una hamburguesa.

IAN STEWART, Matemático

Un grifo que gotea sería mucho menos molesto si conociéramos el porqué de esas gotas. En cambio, a uno no le queda más remedio que pasar las noches en vela y oír en cada gota un nuevo recordatorio de lo poco que sabe del mundo. La buena noticia es que, en un caso ideal, los matemáticos son capaces de explicar aproximadamente ese goteo. A cambio, hay dos malas noticias; la primera, por el momento nadie es capaz de predecir el goteo de la mayoría de grifos reales. ¿Cómo se va a predecir el tamaño de cada gota? ¿Cuánto tiempo pasará hasta que caiga la siguiente? ¿Y qué es lo que sucede exactamente dentro del grifo? La segunda mala noticia es que el goteo de un grifo plantea algunos de los problemas más complejos (y al mismo tiempo más importantes) de los que puede ocuparse la ciencia.

Cuando un líquido sale lentamente de un cilindro vertical vacío (una manguera o un tubo), se forman gotas. Primero se acumula un poco de líquido al final del tubo y luego se produce un estrechamiento entre las nuevas gotas y el borde del tubo. Este estrechamiento va adquiriendo longitud, con la gota colgando del final. Finalmente tiene lugar la descomposición de tan inestable formación. Además de las gotas grandes, en el resto del filo del tubo van formándose también gotas más pequeñas. La ciencia moderna ha descubierto que es sumamente difícil producir siempre gotas de un mismo tamaño. En realidad, las gotas que caen de un grifo comprenden un amplio espectro de tamaños.

La fuerza motriz detrás de la formación de gotas es la tensión superficial. Los líquidos se resisten a ampliar su superficie; por eso, si dejamos a un lado la

resistencia del aire y otros efectos similares, las gotas tienen forma esférica. Sin embargo, el paso de un líquido cohesionado a gotas es sumamente complejo.

Cualquier fabricante de impresoras de chorro de tinta o de pintura de pared desea comprender cómo se forman las gotas; los primeros desean que se formen unas gotas lo más regulares posibles y los segundos, si puede ser, que no se forme ninguna. La teoría básica del problema de las gotas, la llamada hidrodinámica de los líquidos en movimiento, también desempeña un papel crucial en la construcción de aviones y barcos, a la hora de comprender los procesos que tienen lugar en la atmósfera terrestre y en la Vía Láctea, o en la discusión sobre cuánto tiempo necesitará un queso viscoso para cubrir una maratón. Generalmente, para describir el movimiento de un líquido es necesario resolver las ecuaciones de Navier-Stokes, bautizadas en honor a dos matemáticos del siglo XIX que describieron los cambios de velocidad en relación con los cambios de presión en los líquidos. Hasta hoy, sin embargo, se desconoce la solución de dichas ecuaciones y, lo que es peor, nadie sabe si existe realmente una solución. Si alguien logra encontrarla, se hará rico, pues la solución de las ecuaciones de Navier-Stokes figuran, junto con la demostración de la hipótesis Riemann y del problema P/NP, entre los problemas matemáticos del milenio, para los que se ofrece una recompensa de un millón de dólares.

Las ecuaciones de Navier-Stokes sólo se pueden resolver en casos muy concretos. Se simplifican mucho, por ejemplo, cuando se trata de agua. El agua es lo que se llama un líquido newtoniano: independientemente de lo bruscos que seamos con ella, se comportará siempre como un líquido, su viscosidad (resistencia) depende sólo de la temperatura y la presión, y no de las fuerzas que la afecten. Por su parte, los líquidos no newtonianos, entre los que se encuentran el agua azucarada, la miel, la sangre, la mostaza, el pegamento, la pintura, el metal líquido y también la arena poseen una viscosidad que varía en función de qué hagamos con ellos. Así, por ejemplo, si mojamos fécula de maíz con agua, podremos caminar sobre la papilla resultante si lo hacemos con movimientos rápidos y enérgicos, pero nos hundiremos en cuando nos quedemos quietos.

Aplicando la fuerza necesaria, la mezcla se comporta como un flan, mientras que con una fuerza menor lo hace como un líquido. No obstante, no con todos los

líquidos no newtonianos sucede lo mismo; algunos, por ejemplo la sangre, reaccionan exactamente al contrario o simplemente se comportan de otra forma.

Un grifo que gotee ofrece aún otra posibilidad: la simplificación de las ecuaciones y, al mismo tiempo, el modelo de formación de las gotas. En un caso ideal, el grifo tiene una forma cilíndrica simétrica o, lo que es lo mismo, tiene el mismo aspecto vistas desde cualquier ángulo, y la fuerza de la gravedad empuja la gota de forma regular hacia abajo. Eso, sin embargo, varía si el grifo está torcido, o si está medio taponado por la cal, o si lo que gotea no es agua sino sangre (algo de lo más normal en las películas de terror). Por no hablar, naturalmente, de las complicaciones añadidas con que nos encontraremos si abrimos el grifo y la sangre ya no gotea, sino que sale a chorro limpio.

En cualquier caso, la solución más obvia del problema del goteo se conoce desde ya hace tiempo: uno debe mudarse a un lugar en el que no haya ni líquidos ni películas de terror, por ejemplo a la Luna. Allí no pasará nunca más la noche en vela por culpa de un grifo que gotea.



Capítulo 19

Hawai

Desciende, viajero audaz, por el cráter del Snaefellsjókull, que la sombra del Skartaris acaricia antes del 1 de julio, y alcanzarás el centro de la Tierra, tal y como yo lo he hecho ya.

Arne Saknussem

JULIO VERNE, *Viaje al centro de la Tierra*

Desde hace algunos años ha dejado de estar claro, una vez más, por qué existe Hawai. Aún peor: tampoco sabemos cuáles son los orígenes de Islandia, cuáles han sido los hechos que nos han obsequiado con unas islas como las Azores y por qué se elevan sobre el Pacífico Sur las islas Pukapuka.

Durante varios miles de años los brujos de Hawai manejaron la siguiente hipótesis de trabajo:

Pele, la diosa del fuego, que además era ninfómana, estaba huyendo de la furia de su hermana Namaka-o-kaha'i, a cuyo esposo había seducido. Llegó a una isla desierta y, con su bastón de cavar, se puso a horadar fatigosamente una cueva para alojarse en ella, pero justo entonces su hermana, que era de profesión diosa del mar, inundó el islote. Pele se trasladó a la isla siguiente, pero allí volvió a suceder el mismo drama familiar. En su camino hacia el sureste fue dejando tras de sí una cadena de islas con grandes agujeros. Pele consiguió por fin ponerse a salvo en lo alto de la montaña Mauna Loa, que era demasiado alta para las mareas de su hermana, y desde entonces se dedica a echar lava sobre Hawai. Esta teoría se considera actualmente errónea.

Desde la década de 1970, los geólogos prefieren la hipótesis de los penachos mantélicos (*mantle plumes*). *Plume* es una palabra inglesa que se podría traducir como penacho¹¹ y, en este caso, es el nombre que se da a un flujo de materia caliente que sale hacia arriba desde las profundidades de la Tierra. Muchas islas

¹¹ También se dice 'pluma'. (N. de la t.)

volcánicas, como Hawai, Islandia y las Azores surgen siguiendo este modelo de la manera siguiente: el penacho caliente está fijo dentro del manto terrestre y, desde abajo, proporciona calor a un determinado lugar de la superficie terrestre. Sin embargo, esta superficie es en general móvil, porque las placas de la corteza terrestre se desplazan de manera lenta, pero segura. A causa de esto, el punto caliente que produce el penacho se va desplazando también lentamente, durante mucho tiempo, por la superficie terrestre, empuja materiales calientes y fluidos hacia arriba, y los lanza al exterior, describiendo un arco a bastante altura. Al llegar arriba, las rocas se solidifican y forman en torno al volcán calentado por el penacho una isla con palmeras y seres humanos vestidos con unas chabacanas camisas de muchos colores. Después de unos pocos millones de años, el penacho se ha ido más allá y el volcán muere, pero la isla permanece. Así que la historia de Pele, la diosa del fuego, no era del todo falsa.

Para escenificar la aparición de Hawai a causa de un penacho mantélico basta con tomar un encendedor y sostenerlo bajo una hoja de papel que movemos lentamente sobre la llama. Si se hace bien, o sea tal como lo hace la Tierra, surge una cadena de volcanes de papel extinguidos y con bordes negros. Hawai se menciona a menudo como ejemplo típico de penacho mantélico. Hay discusión sobre la cantidad de penachos que pueden existir en total; las estimaciones de las últimas décadas oscilan entre un puñado y unos cinco mil. Los elegantes penachos explican de la mejor manera toda una serie de cosas; por ejemplo, muchas características de la cadena de islas hawaianas. Esta cadena está situada en posición transversal en el Pacífico y cuanto más nos alejemos de la novísima «isla grande» del sureste en dirección noroeste, más antiguas son las islas que vamos encontrando, hasta que al final de la cadena llegamos a una antigüedad de 50 millones de años. La posición y la antigüedad de las islas coinciden aproximadamente con el movimiento de la placa del Pacífico, a través de la cual sube ardiendo el supuesto penacho. Algo especialmente notable con respecto a los penachos mantélicos es que surgen en el manto terrestre, quizá incluso en el núcleo, por lo tanto a cientos o tal vez miles de kilómetros bajo la superficie.

Dado que esas regiones sólo se pueden visitar siendo un personaje de un libro de Julio Verne, estos penachos profundamente arraigados, si realmente existen, podrían contarnos cosas importantes sobre el corazón de las tinieblas.

Lo que pasa es que, por ahora, eso no está tan claro. Aunque la elegante hipótesis de los penachos se ha instalado durante décadas en la mayoría de los libros de texto (y, por lo tanto, en nuestras cabezas), siempre ha habido algún escéptico renitente, por ejemplo Don Anderson, un geólogo del Californian Institute for Technology. Sin embargo, durante la última década han aumentado considerablemente el número de críticos, su sonoridad y su cantidad de publicaciones.

Se pronuncian conferencias que versan exclusivamente sobre las razones por las que no pueden existir los penachos, sobre todo en lugares de origen volcánico. Finalmente, Gillian Foulger, de la Universidad de Durham, una de las más prominentes detractoras de los penachos, creó en Internet un portal especial que bombardea a los interesados con innumerables detalles sobre la actividad sísmica, las anomalías en la temperatura y la de laminación litosférica. Si a alguien todo esto le resulta demasiado, le presentamos aquí resumidos algunos argumentos a favor y en contra de la hipótesis de los penachos.

Por ejemplo, la producción de magma de los volcanes hawaianos no es en absoluto constante, sino muy variable. También dicen los críticos que podría no tratarse de una auténtica chimenea de penacho que bombeara lava de manera imperturbable. Sólo en los últimos cinco millones de años se ha multiplicado por diez la cantidad de lava expulsada: cada año salen en todo el mundo unos cien millones de metros cúbicos de roca procedentes del interior de la Tierra. A los defensores de la hipótesis de los penachos esto no les parece un problema y argumentan de otra manera: si no hubiera penachos mantélicos, tendríamos enormes dificultades para explicar el origen de esas grandes cantidades de magma.

El argumento que plantean a continuación los detractores de la hipótesis de los penachos es el de la «dorsal del Emperador», una cadena de antiguos volcanes extinguidos que se agrega por el noroeste a la cadena de las islas hawaianas y que tendrían que haber sido creados por el mismo penacho. Sin embargo, las islas Hawai y la dorsal del Emperador no se confunden las unas con la otra, sino que

forman un ángulo de unos 60 grados, lo cual configura claramente un codo en la cadena de volcanes. Si existiera un penacho fijo en la zona, que hubiera formado primero las islas Emperador y luego las islas Hawai, entonces la placa del Pacífico tendría que haberse girado brusca y repentinamente hace unos cincuenta millones de años. Pero las placas tectónicas, al igual que los trenes de mercancías y las delegaciones de Hacienda, tienen dificultades para cambiar de dirección abruptamente. No obstante, los defensores de la hipótesis responden que es muy posible que el penacho haya sido desviado por unas corrientes existentes en el manto terrestre y que, en consecuencia, se haya desplazado la «mancha caliente» de la superficie: se trataría de un penacho mantélico con los cabos sueltos.

Los defensores de la hipótesis de los penachos lo tendrían más fácil para imponerse a sus detractores si pudieran mostrar una imagen de un penacho mantélico. En teoría tendría que verse en esa imagen un penacho que desde las regiones inferiores del manto terrestre devorara unos tres mil kilómetros de materiales a través de la tierra, como un gusano dentro de una manzana. Sin embargo, por ahora se carece de pruebas que sean así de directas e irrefutables. El método ya está inventado: con numerosos aparatos de medición distribuidos por todo el planeta se detecta la propagación de los terremotos en el interior de la Tierra, y se extrapola a partir de los resultados el aspecto que tiene lo que hay allí abajo. No obstante, hasta ahora no se ha conseguido seguir la pista de un penacho mantélico con claridad hasta las profundidades del manto terrestre.

Entretanto se han propuesto muchas alternativas a la hipótesis de los penachos. Una de ellas da una explicación «superficial» del caso de Hawai, utilizando las propiedades de la corteza terrestre. El suelo que tenemos bajo los pies no es en absoluto tan estable como lo percibimos. A veces la corteza se rompe en pedazos, se deforma y produce hendiduras y grietas. Además, a la Tierra le salen granos y espinillas: la temperatura y el ensamblaje de las placas no son siempre y en todas partes iguales, sino que cambian continuamente a causa del movimiento de las placas. Se forman puntos de fractura por los que salen sustancias repugnantes, y en los que se desertizan o surgen en su totalidad los más maravillosos paraísos de los mares del sur. En este caso no se necesita ningún penacho mantélico que se abra paso desde el núcleo de la Tierra perforando hacia arriba. Todo lo que se

necesita para comprender la aparición de las islas Hawai tiene lugar en las capas superiores del manto terrestre. Gillian Foulger y otros detractores de la hipótesis de los penachos están convencidos de que esta teoría explica las características de Hawai y de otras zonas mucho mejor que la historia del penacho mantélico.

Si los debates entre defensores y detractores de los penachos, algunos de los cuales se desarrollan a veces en Internet, son tan acalorados, ello no se debe seguramente a que el tema trate de lava ardiente. Podría ser que fuera inminente un cambio de paradigma en la cuestión de Hawai, o podría ser que no. Si Julio Verne en su *Viaje al centro de la Tierra* no nos hubiera mentido tan desvergonzadamente como lo hizo, sabríamos más sobre el tema.



Capítulo 20

Hipótesis de Riemann

«Si me despertara después de dormir durante mil años, mi primera pregunta sería: ¿Se ha demostrado ya la hipótesis de Riemann?».

DAVID HILBERT

Hacia el penúltimo cambio de siglo, el ya entonces famoso matemático de Gotinga David Hilbert confeccionó una lista de los 23 problemas matemáticos más importantes no resueltos. En el octavo lugar de aquella (desordenada) lista figuraba «Distribución de los números primos y conjetura de Riemann». Cien años más tarde, el Clay-Institut de Estados Unidos intentó algo similar y, además, estableció un premio de un millón de dólares por la resolución de cada uno de los siete «problemas matemáticos del milenio». En el primer lugar de la lista aparece uno de los pocos problemas de Hilbert que ha seguido sin solución después de cien años de esfuerzos intensivos por parte de los matemáticos: la hipótesis de Riemann, o la eterna pregunta sobre la pauta que sigue la distribución de los números primos.

A todos aquellos que piensen que este millón de dólares es dinero fácil de ganar, es preciso indicarles que todos y cada uno de los siete problemas son huesos duros de roer y se necesita al menos haber conseguido un diploma universitario de matemáticas para poder entrar en materia.

Para la conjetura de Riemann no sólo se necesitan los números complejos y un cálculo diferencial avanzado, también se han de poder manejar series infinitas de cifras, una habilidad que en la vida cotidiana no aporta absolutamente nada. La matemática moderna es un conjunto de conceptos que parecen a primera vista terriblemente inútiles, lo cual no deja de ser una conclusión errónea. En realidad, determinados aspectos del modelo físico del mundo son tan abstractos que las matemáticas correspondientes aún no se han inventado.

El mundo es todavía más complicado que las matemáticas. No obstante, ha de ser posible expresar la idea básica de la hipótesis de Riemann de una manera lo

suficientemente sencilla como para que la gente no caiga en profundas depresiones y, sin embargo, lo bastante correcta como para que los matemáticos no se pongan a tirarle piedras.

Uno de los grandes misterios de este mundo son los números primos, es decir, aquellos números que sólo son divisibles por uno y por sí mismos, por ejemplo 2, 3, 5, 7, 11 etc. Hay muchas historias curiosas en torno a estos extraños números. Hace ya más de dos mil años, Euclides demostró que todo número natural mayor que uno, o bien es él mismo un número primo, o puede expresarse como producto de números primos. Por ejemplo, el número 260 no es un número primo, pero es el resultado de efectuar $2 \times 2 \times 5 \times 13$, todos ellos números primos. Los números primos son cada vez más escasos a medida que uno se acerca a números más grandes; entre los 10 primeros números hay cuatro números primos, pero entre los 100 primeros números hay 25 y entre los 1000 primeros hay sólo 168. No obstante, es infinita la cantidad de estos colegas indivisibles; también esto lo demostró Euclides. Pero ¿dónde están los números primos? ¿Se sitúan sin más allá donde les apetece? ¿O siguen un orden, aunque éste sea muy complicado?

Es a Carl Friedrich Gauss a quien hemos de agradecer una indicación importante con respecto a esta cuestión. En el año 1791, cuando sólo tenía catorce años de edad, Gauss sospechó que la «densidad de los números primos», es decir, la cantidad de estos números que hay entre el cero y un número determinado, puede predecirse mediante una sencilla fórmula. Dos ejemplos: entre 0 y 1000 tenemos 168 números primos, luego la densidad es un 16,8 por ciento. Con la fórmula de Gauss se obtiene el 14,4 por ciento, que no es un mal resultado, pero se desvía bastante de la realidad. Para el intervalo numérico comprendido entre 0 y un millón, la densidad de los números primos se queda en el 7,8 por ciento, y la fórmula predice el 7,2 por ciento, que es un resultado casi correcto. Cuanto más grandes sean los números, más se acerca la densidad real al valor fácilmente calculable mediante la fórmula. En sentido estricto, se puede decir que la densidad de números primos oscila incansablemente en torno a ese valor, pero la amplitud de las oscilaciones se reduce cada vez más a medida que crecen los números. Pasaron más de cien años hasta que el francés Jacques Hadamard y el belga Charles de la Vallée Poussin pudieron demostrar la genialidad del joven Gauss. Con ayuda de este

teorema se puede al menos calcular aproximadamente la probabilidad que tiene un número incómodo cualquiera, pongamos 3 608 152 892 447, de ser número primo (la palabra «aproximadamente» es, sin embargo, un poco molesta).

La conjetura de Riemann facilitaba el paso siguiente hacia una distribución más precisa de los números primos.

En este punto, vale la pena hurgar un poco en los abismos de la matemática moderna. Quien quiera renunciar a ello, haría mejor saltándose los dos párrafos siguientes. Ahora bien, si lo hace, se perderá uno de los logros más importantes de los tiempos modernos: los números complejos.

Desde la antigüedad se sabía que la representación usual de los números obstaculizaba seriamente el avance de los seres humanos hacia la comprensión del mundo. Por ejemplo, la longitud de las diagonales de un cuadrado sólo se puede determinar utilizando números que después de la coma, en su parte decimal, son infinitamente largos y presentan unas cifras que surgen a trompicones de manera caótica. Son los llamados números irracionales. Lo mismo se puede decir de la longitud de una circunferencia, que es un múltiplo de π , un número asimismo irracional cuyo valor es 3,141592654... (etc., hasta llenar el libro, y luego aún más). El concepto de número volvió a ampliarse cuando a un alma desventurada se le ocurrió la idea de sacar la raíz cuadrada de -1 , lo cual es imposible con los números conocidos hasta entonces. El resultado fue la introducción de los «números complejos»: a cada número real se le añade una parte llamada imaginaria y formada simplemente por un múltiplo de « i », que es el nombre que se da a la raíz cuadrada de -1 . Un número complejo de uso corriente sería, por ejemplo, $3 + 8i$. Este nuevo tipo de número resulta ser una ayuda extraordinariamente práctica para los físicos en sus tareas habituales. Concretamente una gran parte de nuestro modelo moderno del mundo se basa en unas matemáticas que trabajan con números complejos, aunque cuando vamos al supermercado no podemos comprar ni un solo producto con precio imaginario.

A continuación, necesitamos hacernos una idea de lo que es una función. Se puede decir que una función es la máquina omnipotente que poseen las matemáticas para hacer salchichas: la función toma un número (carne) y hace con él otro número (la salchicha), y lo hace aplicando unas instrucciones dadas, que podrían ser, por

ejemplo, «girar a la derecha al llegar a la curva» o «calcular la raíz cuadrada». Aplicada esta última al número nueve, produce el valor tres. Hay funciones de muchos colores, formas variadas y gustos diferentes; unas son muy sencillas y otras resultan extremadamente complicadas. También hay funciones para los números complejos, y funcionan igual que con otros: toman un número, hacen algo con él y al final producen otro número. Lo mismo se puede decir de la llamada función zeta de Riemann, sobre cuyo comportamiento hace la hipótesis de Riemann una importante predicción. Por desgracia, esta máquina especial de hacer salchichas es bastante complicada, y las instrucciones correspondientes son interminables, lo cual desaconseja que las enumeremos aquí. Se ha investigado mucho y bien sobre la función zeta de Riemann. Por ejemplo, se sabe lo que sucede cuando se aplica a números pares negativos, o sea a -2 , -4 , -6 , etc.: el resultado es cero. Por eso, los números pares negativos reciben el nombre de «ceros triviales» de la función zeta de Riemann. La hipótesis de Riemann dice: todos los demás ceros tienen una característica determinada, a saber, que su parte real es siempre exactamente $1/2$. Después de tanto ir y venir, esto suena terriblemente inútil, pero si se juega con ello un poco más, se obtiene una sorprendente predicción sobre el orden en que se ubican los números primos.

A partir de aquí, se puede hablar otra vez de una manera normal. Como se ha dicho, la densidad de los números primos tiene para los números grandes una oscilación en torno a un determinado valor fácil de calcular. Si la conjetura de Riemann es acertada, la densidad no oscila de una forma totalmente arbitraria, sino que lo hace según una probabilidad bien conocida y calculada. Cuando se lanza una moneda, aunque el resultado es de antemano completamente desconocido, se sabe que en la mitad de los casos sale cara. A partir de esto, se puede predecir la probabilidad de obtener un resultado concreto. Del mismo modo, mediante la conjetura de Riemann se puede predecir la probabilidad de una determinada densidad de números primos.

Gracias a esto, no nos sentimos tan desamparados cuando buscamos números primos: sin la conjetura de Riemann se puede decir más o menos qué probabilidad tiene un número concreto de ser número primo. Con la conjetura de Riemann se sabe además lo lejos que podríamos estar de encontrar ese número primo. La

hipótesis de Riemann nos pone en la mano algo así como una varilla de zahorí: nos indica el camino hacia las posiciones de los números primos. O tal como lo expresa el matemático Peter Sarnak: Sin la conjetura de Riemann tenemos que trabajar sólo con un destornillador en la jungla de los números primos. En cambio, la conjetura de Riemann es una niveladora.

Hasta aquí, todo esto suena muy académico. Nos podríamos preguntar: ¿Qué nos importan a nosotros los números primos? La respuesta es que, por ahora, dependemos de estas pequeñas bestias. En la época de las comunicaciones electrónicas nada funciona sin codificaciones. Cada vez que sacamos dinero de un cajero automático, cada vez que pagamos una factura en Internet, las informaciones que suministramos, ya sean códigos secretos o números de tarjetas de crédito, se transmiten de manera codificada. Por desgracia, las técnicas modernas de codificación han de ser forzosamente costosas y complicadas, porque los estafadores (y sus ordenadores) se han vuelto con el tiempo cada vez más listos. Los números primos son la base de la mayoría de las sofisticadas técnicas que utiliza la criptografía. La posibilidad anteriormente mencionada de escribir cada número como producto de factores primos desempeña un papel importante. La seguridad de la codificación se basa en el supuesto de que, en el caso de números muy grandes, esta descomposición en factores primos sólo se puede hallar invirtiendo en su cálculo una cantidad de tiempo tan enorme que resulta inviable, aunque se utilicen unos ordenadores tan rápidos como los que existen hoy en día. Pero, si se supiera más sobre la distribución de los números primos, esto podría cambiar.

Aquí entra en juego la hipótesis de Riemann. Existe el peligro de que, a través de una demostración correcta, salgan a la luz algunos conocimientos que desencadenen una terrorífica revolución en el tema de los números primos y simplifiquen la descomposición en factores primos. Muchos temen ese momento. Otros alimentan teorías conspirativas según las cuales la hipótesis de Riemann estaría demostrada hace ya tiempo, pero nadie estaría autorizado para conocer esa demostración. Por lo tanto, no se trata sólo de un premio de un millón de dólares, sino que es la protección de datos en todo el mundo lo que está en juego. Dejando a un lado estas consecuencias de largo alcance, existe, sin embargo, una

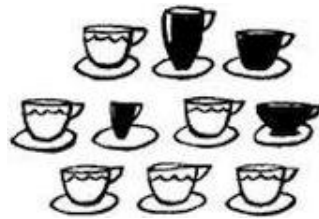
motivación mucho más importante para mantener esfuerzos continuos en la investigación de este tema. ¿Por qué desean los matemáticos demostrar la conjetura de Riemann? ¿Por qué quieren los seres humanos escalar el Everest? En palabras de George Mallory (que llegó a la cumbre de esta montaña en 1924): «Porque está ahí».

Hoy en día la mayoría de los matemáticos cree que la conjetura de Riemann es acertada. Al menos los primeros diez millones de ceros no triviales de la función zeta de Riemann se encuentran donde este matemático sospechaba que estaban. Por supuesto, esto no prueba nada; puede que los números siguientes no cumplan lo mismo, porque hay infinitos números, más que granos de arena en el mar. El matemático alemán Bernhard Riemann, que es quien nos ha dejado esta herencia, fue un hipocondríaco introvertido, que además solía estar realmente enfermo. Su trabajo titulado «Sobre la cantidad de números primos menores que un número dado» se publicó en 1859 y sorprendentemente sólo tiene ocho páginas. En comparación, uno de los últimos intentos de demostrar su hipótesis, publicado por Louis de Branges en 2004, es bastante más largo: consta de nada menos que 41 páginas de apretada escritura.

Además de De Branges, que en la última década ha intentado una y otra vez hacer esta demostración, hasta ahora sin conseguir nada que pueda considerarse un éxito rotundo, desde 1859 los mejores matemáticos de cada generación se han devanado los sesos para probar la hipótesis de Riemann. Durante mucho tiempo se tuvo la esperanza de que el propio Riemann pudiera haber dejado en algún sitio una indicación, y, en efecto, se encontró una nota que parecía hacer referencia a que la famosa conjetura no le había caído del cielo, sino que se derivaba de algo que este matemático no se atrevía a publicar. No se sabe qué podría haber sido aquello exactamente.

Entretanto muchos expertos consideran posible que la demostración de la hipótesis de Riemann, cuando llegue, no venga del campo de las matemáticas, sino de una rama vanguardista de la física teórica que se llama teoría del caos cuántico, porque es evidente que existen profundas vinculaciones entre el mundo de los números primos y el de los quanta o cuantos. Si algún día esto llega a funcionar, un físico

habrá ganado un millón de dólares, y el mundo habrá perdido un enigma bello y consistente. Entonces a alguien se le tendrá que ocurrir algo nuevo.



Capítulo 21

Lluvia roja

Algunos responden a la otra pregunta diciendo que no se trata de sangre auténtica, sino que sólo es agua espesa y sucia que, al ser hervida por el sol, adquiere ese color rojo. [...] Pero más bien dudo de que esas causas sean aceptables en todos los casos.

GOTTFRIED VOIGT, *Pasatiempos físicos*,
1670

Es fácil acostumbrarse a que caiga agua normal del cielo. En algunos lugares de la Tierra incluso se han resignado a ver de vez en cuando ranas y peces que caen del cielo, porque antes en algún otro lugar unos fuertes vientos los han arrastrado hacia arriba, lo cual sería un fenómeno completamente normal. En Kerala, una región de la India, hubo una lluvia extraterrestre. Al menos eso decía la explicación que dio el físico indio Godfrey Louis sobre la lluvia teñida de rojo que cayó allí en el verano de 2001. Otros expertos contemplan esta hipótesis con escepticismo.

Las conjeturas sobre el enigma de la lluvia roja empezaron cuando, entre finales de julio y finales de septiembre, en el sur de Kerala, se produjeron unas precipitaciones esporádicas durante las cuales cayó al suelo un líquido que, en cierto modo, tenía aspecto de sangre. Las precipitaciones de lluvia roja se limitaron a caer en pequeñas zonas de unos pocos kilómetros cuadrados de superficie, mientras que junto a ellas llovía al mismo tiempo de una forma completamente normal. Se informó además sobre unos ruidos que eran como explosiones y que precedieron a la primera lluvia roja. Hasta hoy no se ha aclarado a qué causa o situación especial deben agradecer los indios un fenómeno meteorológico tan extraordinario.

Pronto se descartó una explicación trivial: no se trata en ningún caso de polvo que haya traído el viento desde el desierto hasta Kerala. Desde luego, el polvo del desierto puede teñir las precipitaciones con colores curiosos (en Siberia, por

ejemplo, a principios del año 2007 cayó nieve teñida de amarillo a causa de una tormenta de arena), pero la lluvia de Kerala no contiene polvo de ninguna clase, como han demostrado las investigaciones pertinentes, sino que los componentes rojos parecen células orgánicas. Si la causa hubiera sido una tormenta de arena, sería de esperar que hubieran caído precipitaciones rojas en grandes zonas, y no unas lluvias tan estrictamente locales como las que se observaron. En cualquier caso, es aún menos probable que la lluvia hubiera sido teñida de rojo por un meteorito en tránsito que hubiera dejado en la atmósfera terrestre un reguero de polvo.

Otros sospecharon al momento lo siguiente: si la lluvia tiene aspecto de sangre, entonces es posible que se trate de sangre. Esta creativa historia hablaba de una bandada de murciélagos que volara a gran altura y hubiera sido alcanzada por un objeto duro, por ejemplo un meteorito. En consecuencia, habría sido quizá la sangre de los murciélagos la que tiñó la lluvia de rojo. La duda es dónde habrían quedado los otros restos de los murciélagos muertos, ya que estos animales no están formados exclusivamente por sangre. Además, también habría que explicar de dónde había venido una cantidad de murciélagos tan considerable que podía proporcionar lluvia roja durante meses.

En noviembre de 2001, unos científicos indios publicaron un informe, según el cual la lluvia roja contenía las esporas de un género de algas, en definitiva, unas células germinales a partir de las cuales podrían crecer nuevas algas. Incluso se consiguió obtener algas a partir de las células rojas contenidas en la lluvia. Esas mismas algas crecen de forma natural en la zona de la que proceden los informes sobre lluvia roja. Lo que no aparece en el informe de los científicos indios es de dónde podrían venir tantas toneladas de esporas, cómo habían ido a parar a las nubes, por qué se produjo una distribución extrañamente irregular de la lluvia roja, y si hay alguna relación con los sonidos de explosiones. A pesar de todo, se consideró durante un tiempo que con aquello prácticamente se había cerrado el caso.

En el año 2003, Godfrey Louis y su discípulo A. Santhosh Kumar explicaron que las células rojas contenidas en la lluvia no eran de este mundo. Es de suponer que su teoría no encontró aceptación de manera inmediata, ya que pasaron tres años hasta su publicación oficial en la revista especializada *Astrophysics & Space Science*. Louis

no había encontrado ADN en las partículas rojas y, como éste es un componente importante de cualquier célula existente en la Tierra, llegó a la conclusión de que las partículas «posiblemente no serían de origen terrestre».

Dicho de otro modo, proceden del universo exterior. Louis añadía que un meteorito habría explotado en la parte superior de la atmósfera terrestre con un gran estruendo y habría liberado así grandes cantidades de células biológicas extraterrestres, que luego caerían en forma de lluvia roja sobre la superficie de la Tierra. Si esto fuera cierto, tendríamos la primera demostración de la llamada hipótesis panspérmica, según la cual las células vivas están ampliamente diseminadas por el universo. Y sería también la primera prueba de que existe vida fuera de nuestro planeta. Eso hizo que Louis fuera de repente famoso.

Pero, como suele pasar con las teorías, quizá todo esto no sea en absoluto cierto. Por ejemplo, no queda claro en qué medida se puede aceptar la afirmación de que las células no contenían ADN. Investigaciones posteriores realizadas en centros de investigación británicos aportaron indicios de la existencia de ADN, si bien «aún no confirmada del todo». Carl Sagan precisó en una ocasión que «las afirmaciones insólitas requieren pruebas también insólitas» —por ejemplo, en este caso habría que presentar un extraterrestre adulto, no bastaría con unas cuantas células teñidas de rojo. Hasta que llegue el momento, podemos seguir ocupándonos de las algas.



Capítulo 22

Manuscrito Voynich

*Pada ata lane pad not ogo wart alan ther
tale feur fa rana lant tal told.*

Charles Kinbote, en *Fuego pálido*, de
VLADIMIR NABOKOV

El manuscrito Voynich fue redactado hace cuatrocientos años por un autor anónimo en un alfabeto desconocido y una lengua enigmática (no, no era francés). Su redescubrimiento se lo debemos al archivero Wilfrid Michael Voynich, que en 1912 se lo compró en secreto a los jesuitas italianos, que necesitaban dinero. Desde entonces, legiones de lingüistas, criptógrafos, matemáticos, medievalistas y expertos en literatura intentan desentrañar un texto al lado del cual resultan inteligibles incluso las obras de Niklas Luhmann.

Originalmente el manuscrito tenía 272 páginas de pergamino de distintos tamaños de los cuales se han conservado tan sólo 240. Está dividido en secciones que (a tenor de las numerosas ilustraciones) tratan probablemente de hierbas, astronomía, biología, cosmología y farmacia. Al final hay una última sección sin dibujos que suele llamarse «Recetas», pero que también podría tratarse de horarios de transporte o de una miscelánea de noticias. Las páginas fueron encuadernadas a posteriori; también la numeración de páginas y los colores de las imágenes son posteriores a la redacción del texto.

La primera parte del manuscrito Voynich incluye la detallada descripción de diversas plantas no identificadas. La sección sobre astronomía contiene al parecer descripciones de los signos zodiacales conocidos y de las estaciones del año, y por lo menos en este fragmento está claro que las ilustraciones representan los movimientos de las estrellas y los planetas. A partir, entre otras cosas, de la vestimenta y el corte de pelo de las figuras humanas que aparecen en las ilustraciones (en realidad básicamente a partir del corte de pelo, pues en la mayoría de los casos se trata de mujeres desnudas), se ha deducido que el manuscrito debió de ser redactado en algún país europeo entre 1450 y 1520. En cualquier caso, la

prueba más antigua de la existencia del texto es una carta de 1639 en la que el alquimista de Praga Georg Baresch le pide al jesuita Athanasius Kircher ayuda para descifrarlo. Esta carta, hecha pública por primera vez en la década de 1970, liberó a Voynich de la sospecha constante de embargo, de momento no se conocen más detalles sobre el origen del texto.

El manuscrito se ha analizado utilizando los medios de la lingüística computacional más moderna. El resultado: se ciñe a las reglas estadísticas de las lenguas naturales, que no se describieron científicamente hasta principios del siglo XX (es improbable que un hipotético falsificador del siglo XVI hubiera podido ser tan previsor). Por otro lado, en el texto hay muy pocas palabras que aparezcan regularmente en la misma estructura e incluye repeticiones de palabras poco habituales en las lenguas naturales. En general, el léxico del texto es sorprendentemente escaso, aunque tampoco eso debe ser motivo de desconfianza, pues en ese caso la mayoría de bestsellers actuales serían sospechosos de falsificación.

La colección de hipótesis sobre el significado que recogemos a continuación es tan sólo una pequeña selección e ilustra, por un lado, el desvalimiento y, por otro, la gran imaginación de los expertos en el manuscrito Voynich. William Romaine Newbold, profesor de filosofía de la Universidad de Pennsylvania, fue el primero en anunciar, en 1921, que había descifrado el manuscrito. Según éste, cada vocal contendría unas rayas diminutas que sólo se podrían ver con la ayuda de una lupa y que corresponderían a un viejo sistema taquigráfico griego. Según Newbold (y tal como el propio Voynich había intentado ya demostrar), el texto había salido de la pluma del filósofo inglés Roger Bacon y describía, entre otras cosas, la invención del microscopio. Sin embargo, pronto se comprobó que aquellos teóricos símbolos microscópicos no eran más que grietas naturales de la tinta.

Otra hipótesis creativa fue la que propuso en 1978 el filólogo amateur John Stojko, que afirmó que el texto había sido escrito en ucraniano (aunque sin vocales) y que trataba sobre una guerra civil. Por desgracia, su traducción no coincidía ni con las ilustraciones del texto ni con la historia ucraniana. En 1987, el físico Leo Levitov atribuyó el texto a los cátaros, una secta medieval francesa: el léxico contendría una mezcla de elementos procedentes del flamenco, del antiguo francés y del

antiguo alto alemán. En cambio, el escritor James Finn asegura en su libro *Pandora's Hope*, publicado en el 2004, que el texto está escrito en un hebreo ligeramente cifrado; como sucede con muchas otras interpretaciones, este sistema permite atribuirle al texto un número prácticamente ilimitado de significados. El lingüista Jacques Guy apuntó la posibilidad de que se tratara de una lengua asiática escrita en un alfabeto inventado. No sólo se trata de una hipótesis plausible, sino que también encaja con la estructura verbal del documento; por otro lado, sin embargo, las ilustraciones no parecen en absoluto asiáticas. A finales del 2003, el polaco Zbigniew Banasik aventuró que tal vez nos encontremos ante un texto manchú. Sin embargo, Banasik no habla manchú y, hasta la fecha, no se han encontrado hablantes competentes que puedan corroborar esa teoría.

En el año 2003 hubo el último intento de desentrañar el significado que provocó un cierto debate público. El psicólogo e informático británico Gordon Rugg demostró en su tiempo libre que se podía elaborar un texto de características parecidas a partir de una tabla de prefijos, radicales y sufijos inventados, que se podían combinar utilizando una pauta de papel como la que se utilizaba hasta mediados del siglo XVI para descifrar textos. La prensa publicó los resultados del estudio de Rugg y lo felicitó por haber resuelto el misterio del manuscrito Voynich. Sin embargo, la teoría de Rugg sólo demuestra que habría sido teóricamente posible elaborar un texto similar en poco tiempo sirviéndose de los medios disponibles en la época. Sea como fuere, hasta hoy se desconoce si ésa fue la verdadera génesis del texto. El propio Rugg manifiesta en su página web que, en su opinión, se trata de una falsificación. Económicamente habría sido una empresa lucrativa, pues alrededor del 1600 el emperador Rodolfo II, coleccionista apasionado de manuscritos alquímicos y otras curiosidades, adquirió el documento por 600 ducados de oro.

Siglos más tarde, en 1961, el anticuario Hans P. Kraus esperó sacar una tajada semejante y se hizo con el manuscrito por 25.000 dólares. Sin embargo, no encontró comprador y terminó donándolo a la Universidad de Yale. El texto está a disposición de los descifradores espontáneos, escaneado en la página web de la universidad o en la primera edición del manuscrito completo, publicado por Jean-Claude Gawsewitch en 2005 bajo el nombre *Le Code Voynich*.



Capítulo 23

Materia oscura

Un kilo de materia oscura pesa más de diez toneladas.

Profesor FARNSWORTH, *Futurama*

Sólo una pequeña fracción de la materia del universo es visible. El resto, y no nos referimos a las cosas que han desaparecido bajo la cama, se conoce como materia oscura. En total es más lo invisible que lo visible que hay en el universo: entre cinco y diez veces más. Lo que no está claro por ahora es a qué nos referimos al hablar de lo invisible.

Se sabe de la existencia de la materia invisible porque ésta se percibe de forma indirecta a causa de su masa: las masas se atraen entre sí, según afirma con razón la ley de la gravitación universal, y por eso la materia oscura influye a través de la fuerza gravitatoria en el movimiento de objetos visibles, como las estrellas, que por el contrario es observable. Una parte esencial del trabajo de los astrónomos es ocuparse de lo invisible. Cuando se observa con exactitud lo que sucede en el cielo, a menudo sucede que el movimiento de los cuerpos celestes sólo se puede explicar si se supone la presencia de otros cuerpos celestes que permanecen en la oscuridad, ya sea porque son verdaderamente invisibles (como los agujeros negros) o porque tienen un brillo demasiado débil para poder ser observados con los telescopios existentes. A medida que los telescopios se vuelven más potentes, son más los cuerpos «invisibles» que se vuelven de repente visibles. En 1844, Friedrich Wilhelm Bessel, a partir de los movimientos de la brillante estrella Sirio, dedujo que ésta tenía un acompañante invisible que giraba en torno a ella. Pasaron dieciséis años hasta que Alvan G. Clark, provisto de un telescopio de mayor potencia, pudo ver un acompañante de brillo extraordinariamente débil: Sirio B se hizo famosa enseguida, porque se trataba de un cadáver estelar caliente; pertenecía a una clase de objetos que posteriormente recibirían el nombre de «enanas blancas». Como en el caso de Sirio B, en los últimos diez años se han hallado más de cien planetas situados fuera de nuestro sistema solar, y estos hallazgos se han realizado de

manera indirecta, a través de la fuerza gravitatoria ejercida por estos cuerpos: es imposible verlos, pero atraen y arrastran a sus propios soles con tanta fuerza que los hacen agitarse un poco hacia aquí y hacia allá. Es esta agitación la que nos permite encontrar mundos desconocidos que con nuestras técnicas actuales son invisibles. Lo que realmente es misterioso en relación con la materia oscura no es su presencia, sino lo sorprendentemente grande que es la cantidad de esta materia. El primero que afirmó esto fue el astrónomo suizo Fritz Zwicky en el año 1933. Observó los movimientos de las galaxias en la constelación Coma Berenice, una zona del cielo que está plagada de estas agrupaciones de estrellas. Las fotografías de esta región del espacio muestran una apreciable cantidad de manchas de niebla, que al ser observadas más de cerca (con telescopios más potentes), se ven como galaxias, como muchos miles de vías lácteas, que se componen a su vez cada una de ellas de muchos millones de estrellas, una visión que pone de manifiesto que el universo está empeñado en hacer que nos sintamos poca cosa. Zwicky descubrió que las galaxias existentes en este hormiguero se movían con demasiada rapidez: la masa de la materia visible no es ni de lejos suficiente para mantener unidos esos montones de galaxias. En realidad tendrían que haberse disgregado hace miles de millones de años, con lo que ya no podríamos verlas actualmente. Tiene que haber una especie de «pegamento» adicional, la fuerza gravitatoria de la materia oscura, que evite que las galaxias se disgreguen. Aunque Zwicky lo formuló de una manera bastante más complicada, sus conocimientos fueron ampliamente ignorados. De nuevo hubo que esperar, esta vez casi cuarenta años, para que la existencia de la materia oscura se aceptara de forma generalizada, pero, una vez que se dio esta aceptación, son miles los astrónomos que se han ocupado de esta cuestión día y noche, sobre todo de noche.

La gran brecha que abrió el descubrimiento de la materia oscura tuvo su origen en la investigación de la rotación de las galaxias. Del mismo modo que los planetas giran en torno al Sol, las estrellas de una galaxia se mueven en torno al centro de la misma. El Sol, por ejemplo, lo hace con una velocidad atterradoramente alta de unos 250 km/s. Al mismo tiempo, por una parte, experimenta la atracción que ejerce sobre él el centro de la Vía Láctea mediante la fuerza gravitatoria. Por otra parte, la rotación en torno a ese centro genera la fuerza centrífuga, que es una fuerza

dirigida hacia fuera, cuya existencia podemos percibir sencillamente cuando vamos en coche y tomamos una curva a gran velocidad. En conjunto, la acción simultánea de la fuerza centrífuga y la fuerza gravitatoria hace que el Sol no caiga hacia el interior de la galaxia, ni vuele hacia fuera, sino que se mueva dócilmente alrededor del centro, con una velocidad que viene determinada únicamente por la distribución de la materia en la Vía Láctea. Por lo tanto, a partir de la velocidad de la materia visible pueden sacarse conclusiones sobre la cantidad de masa que hay dentro de la galaxia y sobre el lugar en que esta masa se encuentra. Gracias a este análisis, a principios de la década de 1970 se llegó a una conclusión deprimente: los objetos que están en las zonas más externas de las galaxias, y esto vale para todas (hay muchas, como ya hemos mencionado anteriormente), se mueven a una velocidad excesivamente grande en torno al centro, tan rápido que, como el coche en una curva, saldrían disparados hacia el exterior, si no fuera por la existencia de algo pesado, pero invisible, que se lo impide: la materia oscura.

Entretanto se ha «comprobado» la presencia de materia oscura en muchos lugares diferentes del universo. Se puede encontrar en nuestra Vía Láctea, en las galaxias elípticas, en las galaxias enanas, en los cúmulos de galaxias y en los supercúmulos, que son aún más grandes. En ningún sitio sucederían las cosas tal como deberían suceder según nuestras previsiones, si sólo existiera lo visible. Recientemente especulan algunos sobre la existencia de materia oscura también en nuestro entorno inmediato: las sondas espaciales *Pioneer* n.º 10 y n.º 11, cuya tarea principal consiste en espiar a los grandes planetas Júpiter y Saturno, se ven atraídas en dirección al Sol por una fuerza que no presagia nada bueno, y en consecuencia se desplazan cada vez con mayor lentitud: hasta ahora no se ha podido explicar este fenómeno; las posibles causas varían desde una fuga de combustible hasta la materia oscura, que tiraría con toda su potencia de estas pobres naves espaciales.

Ahora bien, ¿qué es esa materia oscura? ¿Es peligrosa esa extraña cosa? ¿Puede explotar, o es quizá comestible? Podríamos librarnos de este problema con suma elegancia si negáramos la existencia de la materia oscura y explicáramos los efectos anteriormente mencionados modificando sin vacilaciones la ley de la gravedad. Todos los fenómenos que apuntan a la existencia de la materia oscura, lo hacen sólo porque damos por supuesta la universalidad de la ley de la gravedad. Quizá

tenemos simplemente una idea errónea de lo que es la gravedad. Ésta es la idea básica en que se apoya la teoría de la «dinámica newtoniana modificada» [*Modified Newtonian Dynamics*, cuya abreviatura es MOND], y otras construcciones mentales de tipo similar. En el contexto de la MOND, propuesto por el cosmólogo Mordehai Milgrom en 1983, la gravedad ya no se comporta de esa forma tan intransigente que conocemos desde siempre, sino que cambia su manera de actuar cuando se aplica a objetos que están muy alejados, lo cual sucede muy a menudo en el universo. Así, la MOND puede, por ejemplo, explicar las curvas de rotación de las galaxias, que en cualquier otro caso requieren grandes cantidades de materia oscura. Sin embargo, no funciona ni mucho menos en todos los casos y genera una serie de problemas adicionales. Hasta ahora nadie ha inventado una teoría modificada y completa de la gravedad que se pueda aplicar sin dificultades tanto en el dormitorio como en la sala de estar, y también en la cocina del universo. Por eso continúa la búsqueda de la materia oscura.

Las teorías relativas a su naturaleza se dividen en dos clases: por una parte, podría tratarse de objetos pesados, pero sin brillo o poco brillantes, contruidos con los mismos materiales que todo lo que conocemos hasta ahora. A ésta se le llama materia oscura «bariónica», porque su masa se encuentra en gran medida en determinadas partículas elementales, concretamente en los protones y los neutrones, que se llaman también bariones. Unas buenas candidatas para esta categoría de materia oscura son las ya mencionadas enanas blancas, además de las llamadas enanas marrones, que ya volveremos a mencionar más adelante, y los agujeros negros. Estas oscuras sombras suelen englobarse en la sigla MACHO, correspondiente a *massive compact halo objects* «objetos de halo masivo compacto».

Por otra parte, la materia oscura podría estar formada por una gran (o incluso enorme) cantidad de partículas elementales que interaccionan muy débilmente con el resto del mundo y vuelan como autistas a través de las personas, la Tierra y el universo. Las primeras teorías de este tipo partían de unas partículas «calientes», o sea muy cargadas de energía. El mejor candidato para esto ha sido durante mucho tiempo el neutrino, una partícula fantasmal que se libera, por ejemplo, en las centrales nucleares o cuando se producen explosiones de estrellas y para cuya

constatación son necesarios al menos veinticinco años. Sin embargo, parece claro que la masa de los neutrinos es demasiado escasa para explicar los fenómenos asociados a la materia oscura. Hay también modelos que trabajan con materia oscura «fría» y son muy prometedores. Las partículas en cuestión reciben nombres curiosos, como neutralino, axión, gravitino o incluso Wimpzilla, y hasta ahora existen sólo en las mentes de algunos teóricos. Por el momento ninguno de ellos ha sido constatado de una manera que no deje lugar a dudas. En ocasiones, estos seres exóticos e hipotéticos reciben conjuntamente el nombre de WIMP, que significa *weakly interacting massive particles* [«partículas con masa que interaccionan débilmente»], y en última instancia lo que queda claro es que la investigación de la materia oscura es también una competición por el mejor acrónimo: ¿MOND, MACHO o WIMP?

Durante los últimos treinta años, los frentes de la investigación sobre la materia oscura han cambiado de posición muchas veces. En la década de 1970 se aceptaba principalmente la hipótesis de que se trataba de materia bariónica, con un tipo de objetos que posteriormente se denominarían MACHO. En la década de 1980 se pasó página y se hicieron populares los neutrinos, junto con los WIMP «fríos» y otras partículas elementales exóticas. A principios de la década de 1990 volvieron en primer lugar los MACHO, pero durante los años siguientes perdieron su ventaja a causa de nuevas observaciones. De vez en cuando se utilizaron también modelos híbridos: «El mundo necesita tanto MACHO como WIMP ¹²» afirmaba el astrofísico inglés Bernard Carr en 1994. Llenos de esperanza, los expertos se refieren a estas teorías llamándolas escenarios con «dos hadas de los dientes»: cuando un niño pierde un diente de leche, lo pone por la noche bajo la almohada y espera al hada de los dientes, que se lo cambiará por una moneda. Todavía está por ver si el problema de la materia oscura se resuelve con dos hadas de los dientes, es decir, con dos tipos diferentes de partículas.

Un buen ejemplo de lo que han supuesto las modas en la investigación de la materia oscura lo constituyen las llamadas enanas marrones. Al contrario que las estrellas, estos objetos no poseen en absoluto calentamiento interno. Mientras las estrellas, durante muchos millones de años, «queman» en su interior hidrógeno,

¹² Juego de palabras en inglés: wimp significa 'canijo' o 'alfeñique'. (N. de la t.)

generando así helio, las enanas marrones son demasiado pequeñas para producir las temperaturas que requiere este proceso. A eso se debe que su brillo sea débil y que sean, por consiguiente, difíciles de detectar. Si una estrella fuera una vela que arde de manera continua y segura, una enana marrón sería un pedazo de metal incandescente que se va enfriando poco a poco. Desde la década de 1960 se especula sobre la existencia y las características de las enanas marrones, pero hasta hace poco no se habían podido ver ni investigar, entre otras cosas porque los telescopios eran demasiado poco potentes. Como no se sabía nada sobre cuántas podía haber en la Vía Láctea, fueron durante casi veinte años las mejores candidatas para constituir la materia oscura. Incluso en 1994, un año antes del descubrimiento de la primera enana marrón, a saber, de un objeto que recibió el lamentable nombre de Gliese 229B, Bernard Carr se refirió a las enanas marrones diciendo que eran la explicación «más plausible» para la gran cantidad de cosas invisibles que había en el espacio. Sin embargo, en unos pocos años se derrumbaron las grandes esperanzas que se habían puesto en aquella rareza oscura y dudosa; se descubrieron numerosas enanas marrones, pero no fueron ni de lejos suficientes para dar ni siquiera una pista sobre lo que puede ser la materia oscura. Un destino parecido al de las enanas marrones fue el que sufrieron también el resto de los candidatos a ser MACHO, así como los neutrinos: estas cosas existen, por supuesto, pero si se contabilizan todas ellas juntas, se obtiene sólo una pequeña fracción del total de materia oscura.

Hoy en día, no queda prácticamente otra salida que creer en la existencia de materia oscura fría en forma de WIMP o algo parecido: unas partículas elementales que se relacionan con el resto del universo casi exclusivamente por la fuerza gravitatoria que ejercen. Por lo demás, hasta ahora nadie sabe qué son exactamente. Por eso se produjo un acontecimiento de lo más emocionante cuando, a finales de la década de 1990, un equipo de investigadores formado en torno a Rita Bernabei pretendió haber acreditado por primera vez la presencia de materia oscura en la Tierra.

Con ayuda de unos pesados cristales salinos, enterrados en las profundidades de los Apeninos italianos, para protegerlos de radiaciones que pudieran interferir, se descubrió una señal que se atribuyó a la incrustación de partículas hasta la fecha

invisibles dentro del cristal salino. ¿Sería que podían recogerse WIMP en el centro de Europa? Por desgracia esta noticia sensacional no sobrevivió a las comprobaciones críticas realizadas posteriormente. Por lo tanto, todo sigue como estaba, y el concepto de «materia oscura» sigue siendo, como reconoce el astrónomo estadounidense David B. Cline, una «expresión de nuestro desconocimiento» vacía de todo contenido.

Además, por lo que sabemos hasta ahora, el universo se compone de materia visible y materia oscura en una proporción que no va más allá de entre un cuarto y un tercio. El resto se ha calificado en su totalidad como «energía oscura», sólo por llamarlo de alguna manera, y con ello se alude a una fuerza misteriosa que acelera la expansión del universo. Quizá no haga falta decir que tampoco se sabe prácticamente nada sobre la naturaleza de esa energía oscura.



Capítulo 24

Miopía

La costura, vivir en espacios cerrados y leer por la noche libros escabrosos han producido como consecuencia esta degeneración de los ojos; se dificulta el uso de la vista para mirar a lo lejos y esto influye en todos los aspectos de la vida humana.

ROR WOLF, El gran manual universal de Raoul Tranchirer para todos los casos del mundo

Hacerse adulto es algo que en muchos aspectos está muy bien: uno puede hacer muecas sin que se le queden grabadas en la cara, puede ir a nadar con el estómago lleno e incluso, cuando la temperatura baja de cero, lamer el poste de la farola más cercana. Uno se hiela lo mismo, pero está ejerciendo el justo derecho de todo ciudadano adulto. Sin embargo, no está clara de una manera definitiva la importante cuestión de si uno puede dedicarse de modo ininterrumpido a los juegos de ordenador y leer con una linterna debajo de las sábanas, o si, por el contrario, tenían razón nuestros padres y abuelos en aquello de que la vista se echa a perder haciendo esas cosas.

Unos dicen una cosa, y otros otra, y, como adultos que somos, vamos a examinar ahora los argumentos de ambas partes. Por si acaso, quizá sería mejor no acercar el libro demasiado a los ojos. Lo que sucede realmente con los ojos, y en qué orden, cuando una persona se vuelve miope, es una cuestión complicada y que aún no se entiende del todo. Para lo que nos interesa aquí podría bastar con decir que el globo ocular pierde su forma, concretamente se alarga demasiado.

Entonces la imagen del mundo exterior no puede proyectarse como es debido en la retina, y se hace necesario el uso de gafas. El hecho de que, de vez en cuando, surja un ser humano cuya vista sea normal se puede considerar como un pequeño

milagro, porque para ello, durante el crecimiento, el cuerpo debe ajustar con precisión milimétrica la longitud del globo ocular, con independencia de distintos factores. La miopía puede aparecer en cualquier edad durante la vida de una persona. En general, cuanto antes aparezca, más dioptrías se llegará a tener a lo largo de la vida; sin embargo, el avance de la miopía puede detenerse temporal o definitivamente en cualquier momento. No se trata aquí de examinar más de cerca las distintas variantes de la miopía, sino de dedicarnos sin obstáculos a examinar la pregunta «¿Por qué y cómo llegan las personas a ser miopes?».

Para esto sería útil aclarar al principio cuántos y qué tipo de seres humanos se ven afectados por la miopía. Dado que los investigadores no pueden ir llamando a todas las puertas para hacer pruebas de visión a todo el mundo, lo que se hace es examinar a menudo lo que sucede entre los escolares, los universitarios, los soldados y otros grupos de población que están indefensos ante tales procedimientos. Mediante estos reconocimientos se llega a saber mucho sobre los escolares, los universitarios y los soldados, pero lamentablemente no se averigua gran cosa sobre la población en general. Éste fue el caso de un estudio realizado con pilotos israelíes: sus resultados «sólo eran aplicables a pilotos de Israel». Como en el estudio en cuestión se examinó a más de mil doscientos pilotos, lo más probable es que ya no quedara ningún otro piloto en Israel al que pudieran aplicarse los resultados de dicho estudio. Además, en la mayoría de los estudios y de los países se utilizan métodos de medición diferentes. Sin embargo, aunque todas las cifras se valoren con el debido escepticismo, hemos de fijarnos en dos cosas que aparecen con mucha claridad: hay una gran cantidad de miopes en el mundo (hasta 2.300 millones), y la probabilidad de ser uno de ellos depende en gran medida del país en que se vive. Parece haber relativamente pocos miopes en Sudamérica (por debajo del 10 por ciento), África y Australia (más o menos entre un 10 y un 20 por ciento, en ambos casos), y la India (en torno al 10 por ciento). Europa occidental y Estados Unidos están en un término medio con resultados que oscilan entre el 10 y el 30 por ciento, mientras en Japón, Taiwan, Corea del Sur y Singapur es probable que se burlen de uno en la escuela si no lleva gafas, porque allí el porcentaje de miopes está entre un 50 y un 80 por ciento. Pero no sólo el país constituye un factor importante: los estadounidenses que viven más cerca de la costa son más miopes

que los que viven en el interior; los estudiantes son más miopes que los que trabajan como peones; los blancos lo son más que los negros, y los que viven en las ciudades tienen más miopía que los que viven en el campo.

¿Ha sido esto siempre así? El hecho de que en las excavaciones de asentamientos de la edad de piedra no se haya encontrado gafas, lamentablemente no demuestra nada. Desde luego cabe pensar que los miopes, a falta de gafas, serían devorados por los tigres dientes de sable en menos de lo que se tarda en decir «test visual». Sin embargo, no se puede negar que hoy en día en la mayoría de las zonas del mundo viven más miopes que hace unas pocas décadas —sólo los australianos afirman que entre ellos la miopía está controlada—. Pero ¿qué se puede esperar de un país donde incluso los animales funcionan de una manera diferente a como lo hacen en el resto del planeta? Especialmente dramático es el aumento de la miopía en los países asiáticos; en Singapur en 2004 alrededor del 80 por ciento de los varones reclutados para el servicio militar eran miopes, en comparación con el 25 por ciento en 1974. Muchos investigadores sospechan que tampoco los países occidentales se van a librar de una evolución similar.

Dado que la miopía genera costes y, por el aumento simultáneo del riesgo de contraer determinadas enfermedades oculares, puede llevar a la ceguera, es grande el interés que suscita la investigación de este fenómeno. Consecuentemente son numerosos los modelos explicativos, que pueden clasificarse en tres grupos: la miopía puede estar condicionada genéticamente, o puede desencadenarse por las condiciones ambientales y los hábitos de comportamiento, o ambas cosas a la vez. Los resultados de las innumerables investigaciones sobre la miopía son, dicho con todas las precauciones, muy variados. Si se acepta la hipótesis de que todos los estudios se han realizado cuidadosamente y sin hacer chapuzas, esto podría significar que quizá haya un elemento común y necesario, hasta ahora desconocido, en la aparición de la miopía, o sencillamente que son muchos los caminos imprevisibles que conducen a la miopía.

Hasta muy entrado el siglo XX, este campo de estudio era más fácilmente abarcable: se pensaba que la miopía estaba condicionada genéticamente, que los padres miopes tenían hijos miopes, y los porcentajes de miopía que se observaban en determinados grupos de población estaban considerados como una mera

curiosidad estadística. De hecho parecen existir personas que sencillamente no se vuelven miopes, aunque pasen por alto todos los buenos consejos de los expertos en miopía. En contra de la teoría genética se da el caso de que los indios criados en la India son miopes con mucha menor frecuencia que los indios que viven en Singapur. Además, se da el caso de los niños tibetanos y sherpa del Nepal, que genéticamente no se distinguen mucho unos de otros, y, sin embargo, los porcentajes de miopía difieren ampliamente de un grupo al otro.

No pueden ignorarse las grandes diferencias en cuanto a la incidencia de la miopía en distintos grupos étnicos, pero faltan estudios que separen con claridad los factores genéticos y los ambientales. ¿Y si lo único que sucede es que los niños miopes adoptan las malas costumbres de sus padres?

El genetista británico Christopher Hammond publicó en 2004 los resultados de un estudio realizado con gemelos que apuntaba a la posibilidad de un defecto del gen PAX6, un gen importante para el desarrollo de los ojos. También otros estudios realizados con gemelos indican la tendencia a una fuerte influencia genética y una menor influencia del ambiente, y los altos índices de miopía se observan entre estudiantes, pero también se ha observado lo contrario. Entre otras cosas porque el rápido aumento de los índices de miopía en muchos países asiáticos pone en duda una causa genética, se enuncia actualmente la teoría de los genes con ciertas modificaciones: la predisposición a la miopía podría ser hereditaria, mientras que los factores ambientales determinan la aparición y la evolución de la enfermedad.

Pero ¿cuáles son los factores ambientales que inciden en esto? Según una teoría, el ojo, en su crecimiento longitudinal, se orienta a las exigencias de una precisa visión de lejos, y deja para la musculatura ocular la regulación de la visión de cerca. La escasez de ocasiones de mirar a lo lejos influiría en este ajuste preciso del crecimiento de los ojos. Todavía no hay acuerdo sobre si agudizar la vista muy a menudo (como antes se creía) o con demasiado poca frecuencia (que es más bien lo que se sospecha actualmente) propicia el desarrollo de este defecto. En todo caso, hay numerosos indicios de que trabajar mucho mirando algo de cerca lleva a la miopía: cuanta más formación escolar, más alto es también el porcentaje de miopía, y en los escolares la miopía avanza más rápido en épocas de aprendizaje intensivo, mientras que lo hace más despacio durante los períodos vacacionales. También aquí

resulta difícil aislar la influencia de la lectura o la escritura, ya que al mismo tiempo hemos de tener en cuenta que este trabajo se realiza en espacios interiores, donde puede haber mucha o poca luz, o iluminaciones de mala calidad. Además, los niños que leen mucho son a menudo introvertidos y permanecen en general más tiempo que otros niños dentro de sus casas. O quizá no sea así: un estudio del investigador estadounidense Donald Mutti muestra que el comportamiento en tiempo de ocio a penas se diferencia entre los niños miopes y los de visión normal. En cualquier caso, todo el mundo está de acuerdo en que los niños que hacen mucho deporte son menos miopes, lo cual podría deberse a que las distancias a las que miran son mayores o más variables, a que el riego sanguíneo es mejor en la retina, o a que en ésta hay más cambios de luminosidad. Los detalles se han de investigar todavía.

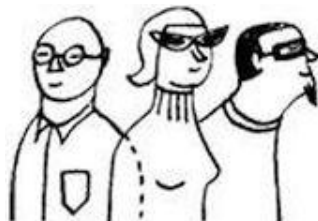
No se sospecha sólo de los comportamientos de trabajo o de ocio como factores que influyen en la miopía, sino también de los hábitos alimentarios, el estrés y muchos otros factores. Otra explicación de la ya mencionada diferencia entre los niños sherpa y los tibetanos señala que el sistema escolar es más duro y el desarrollo técnico está más avanzado en el Tíbet, lo cual conlleva cambios en la alimentación y un mayor estrés. El estrés agudo puede llevar a la miopía, como muestra un estudio realizado después de un terremoto, y también se sospecha que la alimentación puede ser en muchos casos el desencadenante. Se echa la culpa al exceso de pesticidas, al exceso de fluoruros, a la falta de cobre, cromo, manganeso, selenio y calcio, a la escasez de vitamina D que se produce por exponerse poco a la luz del sol, y a la falta de casi todas las demás vitaminas. Pero, por otra parte, está el hecho de que la miopía afecta sobre todo a personas de países desarrollados, en los que la escasez no supone en realidad un gran problema.

¿Podría ser que dependiera de la proporción de los nutrientes, y no de su cantidad en términos absolutos? ¿O son culpables, como pasa tantas veces, el consumo de azúcares y la falta de fibra? Los miopes tienen claramente una mayor frecuencia de caries dentales que las personas de visión normal, pero también para esto puede haber muchas otras causas. En sentido estricto apenas hay un factor que no caiga inmediatamente bajo la sospecha de desencadenar la miopía: los cambios en el ritmo de los días y las noches, la iluminación demasiado fuerte, la mala iluminación, las luces de noche en la habitación de los niños, el trabajo con ordenador. Con

respecto a esto último, puede suponerse que la sustitución de los monitores de tubo de rayos catódicos por pantallas de cristal líquido influirá en el desarrollo de la miopía.

¿Es posible frenar o detener el avance de la miopía? En muchos estudios antiguos se indicaba que las lentes de contacto duras eran buenas para ello, pero otros estudios más recientes no han podido demostrar ventaja alguna. Lo que se da como seguro es que las lentes de contacto blandas no sirven para nada en este sentido. Parece que hay una relación entre los trastornos visuales y el parpadeo poco frecuente, y los usuarios de lentes de contacto parpadean más. El hecho de que, cuando tenemos que concentrarnos mucho en el trabajo, parpadeamos menos, establece de nuevo una relación con el modo de vida. Es posible que las lentes de contacto favorezcan el riego sanguíneo de los ojos, y esto puede ser de ayuda. Los nuevos modelos de lentes de contacto, que en estos aspectos se diferencian de los antiguos, complican aún más la investigación. A partir del hecho de que los animales a los que se obliga a usar lentes para la miopía se vuelven miopes, se puede deducir como algo seguro que se estimula el crecimiento del globo ocular cuando la retina recibe imágenes de mala calidad. Del mismo modo tendría que ser posible, aunque por ahora no lo es, inventar unas gafas que frenaran o detuvieran el avance de la miopía, haciendo que en la retina se proyectara una imagen mejor. También hay fármacos que influyen en el crecimiento del globo ocular. Sin embargo, hasta ahora no se ha conseguido desarrollar un medicamento que sea beneficioso para esto y que no tenga fuertes efectos secundarios, pero la situación puede cambiar.

La investigación progresa de manera imparable y probablemente conseguirá algún día encontrar un procedimiento para prevenir la miopía. Esperemos que esto no implique un laborioso cambio de nuestros hábitos de vida, sino que sea quizá algo tan sencillo como un juego de ordenador o una gragea con sabor a frambuesa.



Capítulo 25

Olfato

El proceso de oler consiste en que el movimiento del objeto oloroso se capta, se mide y, a través del cerebro, se lleva al alma, para que ésta pueda percibir y reconocer las características de dicho objeto.

JOHANN HEINRICH ZEDLER, «Olfato»,
Gran enciclopedia universal completa de todas las ciencias y las artes, 1732-1754

Actualmente, para la mayoría de las personas, oler es sólo un hobby. En cualquier caso, rara vez es necesario para sobrevivir, porque preferimos confiar en nuestros ojos y, en menor medida, en nuestros oídos. Sin embargo, muchos animales se ríen de esa tendencia a usar la vista e insisten en reconocer su entorno con el viejo sentido del olfato, que les da muy buen resultado.

Seguramente no es una mala idea, ya que a menudo no hay luz eléctrica en los lugares donde viven.

He aquí *grosso modo* cómo se desarrolla el proceso de oler: el «olor» se compone de moléculas de sustancias olorosas. Éstas llegan a la mucosa olfativa, situada en la parte superior de la nariz, y son registradas por los receptores que allí se encuentran y que están especializados para cada molécula de olor primario. Al recibir la sustancia olorosa, los receptores emiten una señal eléctrica que es enviada al cerebro mediante las fibras del nervio olfatorio. Y, como sucede con los demás órganos de los sentidos, es en el cerebro donde se produce el análisis minucioso de las informaciones olfativas. En un proceso complicado y que todavía no se comprende del todo, se deducirá, a partir de los datos recibidos, todo aquello que para los seres humanos es importante: por ejemplo, si es una flor o una mofeta lo que tienen delante de la nariz.

Es mucho lo que se ha aprendido en los últimos años sobre los procesos que intervienen en el olfato. Se sabe que algunos mamíferos poseen alrededor de 1000 receptores diferentes (en el caso de los seres humanos sólo hay unos 350), con los que pueden distinguir más o menos 10.000 variedades de olores. El modo en que se configuran los receptores olfatorios está registrado en unos 1000 genes, que representan entre el uno y el cuatro por ciento de todo el genoma, dependiendo de la cantidad de genes que se atribuya en total al ser humano, lo cual es todavía objeto de discusión. En cualquier caso, está claro que el sentido del olfato es importante para el organismo. También lo es para el comité que otorga los premios Nobel, que en 2004 concedió el premio Nobel de medicina a Richard Axel y Linda B. Buck por las concienzudas investigaciones que realizaron durante más de una década en relación con el sistema olfativo, desde los receptores hasta el cerebro. No está claro por ahora cuál es el mecanismo que actúa al iniciarse el proceso olfativo, es decir, la interacción entre las moléculas olorosas, que son las «portadoras» del olor, y los receptores. ¿Qué sucede realmente cuando una molécula tropieza con un receptor? ¿En qué nota el receptor que una determinada sustancia ha entrado en la nariz? (Responder que «en el olor» sería demasiado simple). O también, considerándolo desde el otro lado: ¿Cuál es la característica de una sustancia que determina su olor? ¿Por qué algunas sustancias tienen un olor agradable y otras no? Según la opinión de la mayoría de los expertos, los receptores y las sustancias olorosas funcionan siguiendo el principio de la llave y la cerradura. Las moléculas del receptor constituyen la cerradura y poseen una forma determinada. Cuando llega al receptor una molécula que tiene justo la forma complementaria, es decir, que encaja como una llave en el receptor, la nariz se lleva una alegría y comunica el acontecimiento a los jefes, que están en el cerebro. Según este modelo «estereoquímico», propuesto inicialmente por el estadounidense John Amoore el año 1952, el olor de una sustancia viene dado por la forma y el tamaño de sus moléculas. Aunque el principio de la llave y la cerradura está ampliamente aceptado como base del mecanismo del olfato, dicho principio presenta algunas dificultades: como ya se ha dicho, la nariz humana sólo dispone de unos 350 tipos distintos de receptores. Si lo que de verdad importa fuera sólo la forma, únicamente podríamos distinguir en sentido estricto 350 olores diferentes;

sin embargo, está claro que distinguimos muchos más. Pues entonces, según dice, por ejemplo, Leslie B.

Vosshall, profesora de la Universidad Rockefeller de Nueva York, quedan las llaves un poco flojas en la cerradura. Lo que suena en principio como una chapuza, resulta ser una inteligente jugada de ajedrez: en realidad, de esta manera hay más moléculas olorosas que encajan en el mismo receptor (un poco mal, pero encajan), y hay distintos receptores que sirven para la misma molécula. Combinando las informaciones de distintos receptores, el cerebro podría percibir miles de olores diferentes.

Sin embargo, hay un serio problema con respecto a la teoría estereoquímica, y es el que plantean las moléculas que tienen formas similares, pero producen olores totalmente diferentes, o, al revés, que tienen aspectos totalmente distintos, pero huelen de manera similar. Por ejemplo, las moléculas de decaborano, una sustancia que, entre otras aplicaciones, tiene la de servir como combustible para cohetes, tienen un aspecto muy parecido a las de canfano (la única diferencia es que los átomos de boro se sustituyen por átomos de carbono). Mientras el canfano huele a alcanfor, una sustancia que forma parte de numerosos cosméticos y medicamentos, el decaborano huele claramente a azufre (un elemento que, para colmo, ni siquiera entra en su composición).

Muchas sustancias huelen a almendras amargas, aunque estén compuestas de una manera totalmente distinta al benzaldehído, el compuesto principal del aceite de almendras amargas. A causa de estas discrepancias, los investigadores están buscando ampliaciones del modelo estereoquímico, o modelos alternativos.

Una de estas alternativas va unida al nombre Luca Turin desde 1996, aunque la idea básica tiene casi sesenta años más y se debe a G. Malcolm Dyson. Este investigador pronosticó que lo decisivo no sería la forma de la molécula, sino las vibraciones que se produjeran en el interior de ella. Cuando se unen los átomos para formar una molécula, en ningún caso surge una estructura rígida e inmóvil. Hay que imaginarse los enlaces que hay dentro de la molécula más bien como plumas de las que penden unos pesos (los átomos) que vibran sin cesar de un lado para otro. No sólo vibran los átomos individualmente, sino que, si se trata de moléculas complicadas, lo hacen también grupos enteros de átomos. Las

vibraciones se producen con unas frecuencias determinadas, que dependen, entre otras cosas, del peso de los átomos implicados y de lo fuerte que sea el enlace. Cada molécula muestra un espectro de vibraciones característico que se puede utilizar, por ejemplo, para analizar la estructura de las moléculas. Turin afirma que la nariz hace exactamente lo mismo: funciona como un espectroscopio e identifica las sustancias olorosas según la frecuencia de vibración de las moléculas contenidas en dichas sustancias. Esto es, desde luego, más complicado que lo de la llave y la cerradura, y quizá por eso suena improbable. Pero el principio básico, es decir, la percepción de las vibraciones no es algo que extrañe al cuerpo: también el ojo y el oído perciben frecuencias, ya sea en forma de ondas electromagnéticas o acústicas. No obstante, en el caso de la nariz no está por ahora claro cómo pueden percibirse las vibraciones de las moléculas al nivel molecular. ¿Cómo «miden» los receptores el espectro de vibración de las sustancias olorosas? Una posible respuesta a esta pregunta es la que publicaron en 2006 Jennifer C. Brookes y sus colegas en Londres. El mecanismo que propusieron ya había sido mencionado por Turin en 1996 y se parece al de una tarjeta con banda magnética. Cuando una molécula con una frecuencia determinada entra en contacto con el receptor correspondiente, por decirlo de algún modo, se cierra un circuito: los electrones fluyen desde un donante a través de una molécula olorosa hasta el receptor, donde dejan la señal que será enviada al cerebro (es lo que dice la teoría). Futuros experimentos tendrán que aclarar si este mecanismo funciona también en la práctica y si está realmente instalado en la nariz.

La teoría de las vibraciones encontró una acogida sumamente escéptica entre los expertos, pero tuvo una marcha triunfal a través de los medios de comunicación. Turin escribió columnas sobre sus ideas olfativas para el *Neue Zürcher Zeitung*, la BBC hizo de él un retrato detallado, y el periodista estadounidense especializado en temas científicos Chandler Burr escribió todo un libro sobre Turin y su teoría. En 2006 se publicó por fin un libro escrito por el propio Turin, titulado *The Secret of Scent*. Sin embargo, la recién descubierta teoría no está en absoluto más libre de contradicciones que el principio de la llave y el cerrojo. Un problema son los enantiómeros, es decir, las moléculas enantiomórfas, que sólo se distinguen entre sí porque sus

estructuras aparecen en orden inverso, de tal modo que, si se colocan a ambos lados de un eje, se ven como en un espejo, algo así como la mano izquierda y la derecha. En esa simetría de espejo las frecuencias de vibración no cambian, por lo que las sustancias deberían oler igual. Sin embargo, no siempre sucede así: por ejemplo, un enantiómero de la molécula de carvona huele a comino, y el otro a menta.

Una prueba importante para toda teoría del olor son los experimentos con isótopos: se investigan moléculas en las que uno o más átomos han sido reemplazados por isótopos (el mismo átomo, pero con un número distinto de neutrones en el núcleo). El átomo de hidrógeno, que es el más sencillo que existe, no posee más que un protón en el núcleo. Si se le añade un neutrón, el resultado se llama deuterio. No obstante, sigue tratándose de hidrógeno, porque el neutrón tiene poca influencia en las propiedades químicas. Si en una gran molécula se cambian unos átomos de hidrógeno por otros de deuterio, apenas cambia la forma de la molécula, pero sí sus frecuencias de vibración, porque los átomos de deuterio son más pesados que los de hidrógeno normal. Si la forma fuera lo único que importa para el olor, esas moléculas «deuteradas» tendrían que oler igual, pero, si lo importante fueran las vibraciones, su olor tendría que ser distinto. Por lo tanto, en teoría se puede utilizar esas moléculas deuteradas para distinguir entre ambos modelos.

Un par de cucarachas consideran un deber hacernos saber que su sentido del olfato funciona más bien según el modelo de las vibraciones: si se deuteran unas moléculas que en las cucarachas producen un efecto afrodisíaco, cambian las reacciones de estos animales según la posición de los neutrones adicionales, tal como afirmaron en 1996 los químicos Barry A. Havens y Clifton E. Meloan, de la Universidad de Kansas. Además hallaron una relación entre el comportamiento vibratorio de las moléculas afrodisíacas y la actividad de las cucarachas, lo cual sería una alegría para Turin. Parece ser que también algunos peces pueden distinguir los isótopos por su olor, mientras que las moscas de la fruta hacen como si no supieran nada de neutrones. Pero ¿se puede confiar en estos bichos? Los experimentos con animales conllevan numerosas dificultades, en parte porque no se les puede preguntar detalles sobre el olor de las sustancias, como a uno le gustaría hacer. Hoy en día muchos creen que, en todo caso, para los seres humanos las

sustancias deuteradas y las que no lo están huelen igual. Así, unos experimentos que Vosshall y su colega Andreas Keller realizaron en 2004 dieron como resultado que las acetofenonas siempre huelen a almendras amargas (un olor frecuente en los laboratorios que experimentan con el olfato), con independencia del número de neutrones que posea el hidrógeno que contienen, tal como sería de esperar, si la forma de las moléculas determina el olor.

A fin de cuentas, la teoría vibracional de Turin sigue estando considerada hoy en día como una idea excéntrica. La mayoría de los especialistas piensa que la forma de las moléculas es el origen de los olores, aunque se admite que posiblemente podrían intervenir además otros factores. El propio Turin reconoció que su teoría era bastante «superficial». La cuestión decisiva para todos los modelos es en qué medida pueden predecirse los olores a partir de determinadas moléculas, antes de que alguien acerque la nariz. Así pues, estaría bien organizar una competición olfativa en la que los representantes de las distintas tendencias predijeran los olores y al final se compararan sus predicciones con la realidad. Quien tenga más aciertos, gana.



Capítulo 26

Parque nacional Los Padres

La tierra es seca, estéril y fría, si bien podría suceder asimismo que en ocasiones fuera mucho más caliente o, al menos, templada. A menudo, se ven por allí algunos saltamontes.

JOHANN HEINRICH ZEDLER, «California», en *Gran enciclopedia universal completa de todas las ciencias y las artes*, 1732-1754

El 21 de agosto de 2004 se declaró un incendio forestal en el Parque Nacional Los Padres, que se encuentra en California. La cosa en sí misma no tiene nada de particular y sucede allí tan a menudo, que el parque de bomberos daría la alarma si en un momento dado no hubiera nada que estuviera en llamas. Pero, como el suelo no se enfriaba varios días después de que se hubiera extinguido el incendio, los bomberos informaron preventivamente a Allen King, el geólogo del Parque Nacional. Mediante un vuelo de reconocimiento y la toma de fotografías sensibles al calor se descubrió que aquel insólito calor no lo había producido en absoluto el fuego, sino que éste se había declarado sobre una zona de unos 12 000 metros cuadrados provista de un sistema de calentamiento procedente del suelo. A una profundidad de apenas cuatro metros se midió en los lugares más calientes una temperatura de 307 °C. A una profundidad de sólo diez centímetros registraba el suelo una temperatura de 256 °C. Los puntos más calientes de la zona, según el resultado de mediciones posteriores más precisas, tenían unas dimensiones limitadas: alcanzaban menos de diez metros bajo tierra y medían apenas un metro cuadrado en la superficie.

Lamentablemente, durante los meses siguientes, o bien no se hicieron comprobaciones con la debida frecuencia, o los geólogos responsables tuvieron algo mejor que hacer que publicar con regularidad nuevos datos. En todo caso, según

una investigación posterior, realizada diez meses después, el suelo se había enfriado sólo ligeramente: en el punto de máximo calor se registraron entonces 296 °C. Son pocas las hipótesis que se han formulado para explicar esas insólitas temperaturas del suelo; no se conoce la existencia de grandes yacimientos de petróleo, gas o carbón en el entorno inmediato de la zona, y tampoco hay indicios de emisiones radiactivas, explosiones o actividad volcánica. En el Parque Nacional Los Padres hay fuentes termales, pero se encuentran en otras zonas.

Según el geólogo Allen King, a más o menos un kilómetro de distancia hay una falla de mayor tamaño, y numerosas fallas más pequeñas situadas en las proximidades de la zona caliente. De esas fallas podrían haberse deslizado gases combustibles, como el metano, que con anterioridad podrían haber permanecido almacenados en el subsuelo, y en un momento dado se habrían inflamado. Dado que la zona se encuentra en un lugar donde unos seis años antes se había producido un desprendimiento de tierras, se habla también de una reacción química entre el oxígeno del aire y los minerales de las rocas que se habían desmoronado y partido. King sospecha que los sulfuros contenidos en las rocas, concretamente la pirita y la marcasita, al contacto con el oxígeno se habrían calentado, y así se habría oxidado la materia orgánica contenida en las rocas, que antes había estado aislada del aire. En una expedición realizada en diciembre de 2005, no se encontró rastro alguno de pirita, pero sí muchos compuestos de hierro y oxígeno que podrían haberse generado al descomponerse la pirita. Si la pirita es tan escasa sobre el terreno, ¿su descomposición podría haber servido sólo como detonador para que se inflamase el gas natural que manaba del subsuelo? En cualquier caso, Scott Minor, del Servicio Geológico de Estados Unidos, considera que esto es posible. Como resultado de mediciones realizadas en la superficie, se pudo detectar monóxido y dióxido de carbono, lo cual respalda la idea de un proceso de combustión, pero se consigna la ausencia de una determinada variante del helio que es típica en los yacimientos de gas natural. Además, en el transcurso de esta misma investigación se puso de manifiesto que la temperatura, aunque había descendido en casi toda la zona, había aumentado de nuevo en dos de los puntos donde se llevaron a cabo las mediciones. Es muy considerado por parte del subsuelo californiano el tener la amabilidad de retrasar el proceso de enfriamiento hasta que los investigadores han comprobado

todos los modelos de explicación del fenómeno. Lamentablemente, por ahora sólo los especialistas conocen la situación exacta de la zona caliente. Por lo tanto, quien desee pasar sus vacaciones en el Parque Nacional Los Padres tendrá que llevar consigo su hornillo de camping y no confiar en que la bondadosa madre naturaleza le vaya a calentar la sopa de sobre.



Capítulo 27

Partículas elementales

No me creo nada. Los quarks son una bobada.

STEVEN WEINBERG premio Nobel de física

Desde hace mucho tiempo se sabe que el mundo está formado por gran cantidad de pequeñas partículas. En la historia del descubrimiento de estas partículas, los investigadores creyeron muchas veces haber descubierto por fin el ladrillo más pequeño con que estaba construida la materia, pero un par de años más tarde les llegaba la noticia de que no habían buscado con la precisión suficiente.

La moderna investigación sobre la materia comenzó hace unos dos mil quinientos años, cuando el filósofo griego Demócrito introdujo el concepto de «átomo», palabra que significa «indivisible». Según este filósofo, los átomos son partículas diminutas e indestructibles que están llenas de masa y tales que toda la materia está formada por ellas. A pesar de no estar demostrado, este postulado sobrevivió durante más de dos milenios, hasta que en el siglo XIX se hallaron los componentes negativos del átomo, llamados electrones, fuera de los átomos: sólo con someter un trozo de metal a una tensión eléctrica y calentarlo, se consigue que los electrones lo abandonen en manada. Se había terminado la indivisibilidad de lo indivisible.

En comparación con el desarrollo a paso de caracol que tuvieron las teorías sobre partículas elementales durante los pasados milenios, en el siglo XX sucedió todo de repente con gran rapidez, y paredes enteras de estanterías llenas de obras de premios Nobel quedaron obsoletas para el conocimiento de la estructura de la materia. Por de pronto, el físico inglés Ernest Rutherford reconoció que los átomos están en su mayor parte vacíos. El «modelo planetario» del átomo que se desarrolló posteriormente se basa en la idea de que en el núcleo del átomo se encuentran los protones, que tienen carga positiva y en torno a los cuales giran los electrones como los planetas alrededor del Sol. Casi al mismo tiempo, la física cuántica revolucionó nuestra manera de entender el mundo. En general, se descartó la idea de considerar las partículas elementales como cuerpos celestes, porque el universo

a escala atómica funciona con otras reglas. De Hendrik Kramers, uno de los descubridores de estas reglas, procede la afirmación de que existe una tendencia a aceptar con alegría durante unos meses la nueva mecánica cuántica, antes de romper a llorar. La física de partículas, que hasta entonces había sido una especie de juego de billar con unas bolas extraordinariamente pequeñas, se convirtió en un paisaje extraño que una mente sana difícilmente podía comprender.

Como todo país extraño, también el mundo de las partículas más pequeñas está poblado de animales exóticos. Pronto aparecieron en los laboratorios las primeras antipartículas: tienen exactamente el mismo aspecto que el ejemplar «normal» correspondiente, salvo que la carga es opuesta. Por ejemplo, la antipartícula del electrón se llama positrón, tiene carga positiva y se comprobó su existencia por primera vez en 1932. Aquel mismo año, el inglés James Chadwick descubrió el neutrón, que tiene tanta masa como el protón, pero no posee carga alguna. Unos años más tarde fueron descubiertos los muones, una especie de electrones con sobrepeso; diez años después surgieron los piones y, más recientemente, en la década de 1950, quedó el panorama bastante confuso cuando el kaón, el hiperón y diversos neutrinos entraron en el zoológico de las partículas elementales. La siguiente revolución llegó finalmente en 1968: se descubrió que los protones tienen de «elementales» tanto como los átomos. Una vez más, el mundo se fraccionó en partes aún más pequeñas.

Nuestros conocimientos sobre partículas elementales proceden sobre todo de experimentos en los que con un tipo de partículas se bombardea otro tipo de partículas. Por ejemplo, Rutherford disparó átomos de helio sobre átomos de oro, y sus proyectiles atravesaron en la mayoría de los casos el oro sin encontrar obstáculos. En el mismo principio se basa el experimento que llevó a descubrir los quarks: utilizando el acelerador de partículas de la Universidad de Stanford, se consiguió dar a unos electrones unas grandes velocidades y luego se les hizo entrar en colisión con protones. Por el modo en que los proyectiles eran desviados, se pudo deducir cómo es el interior de un protón. Está formado por tres «quarks», un nombre que se tomó de un poema de James Joyce, en el que se dice: «*Three quarks for Muster Mark!*».

Las cosas no podían seguir así. Para controlar la creciente diversidad de partículas y antipartículas, se creó a principios de la década de 1970 el llamado modelo estándar de la física de partículas: este modelo, no sólo aportó orden al zoológico de partículas, sino que también estableció finalmente unas reglas muy claras para la convivencia de aquellas pequeñas bestias.

Según el modelo estándar, el mundo está compuesto por doce «fermiones» diferentes (se trata del electrón, el muón, el tauón, tres clases de neutrinos y seis de quarks), sus doce antipartículas y diversos «bosones de Eich», que son los responsables de transmitir saludos y mensajes entre las partículas, la mayoría de las veces cosas como «Te encuentro muy atractiva», y de vez en cuando también «Me repeles». El fotón es una especie de bosón que trabaja de cartero y transmite fuerzas electromagnéticas, como la atracción entre partículas provistas de cargas opuestas. Otra partícula es el gluón, que adhiere los quarks al núcleo del átomo.

Hasta la fecha, el modelo estándar en muchos aspectos ha hecho honor a su pretencioso nombre. Predecía, por ejemplo, la existencia y las características de diversas partículas nuevas, antes de que fueran realmente descubiertas, lo cual era un gran avance, porque no se trataba ya de andar con la lengua afuera por detrás de la naturaleza, sino que se sabía de antemano cuál sería su siguiente jugada. Pero el contrario no se rendía tan fácilmente. Dentro del modelo estándar hay algunos problemas que siguen sin resolverse. Por ejemplo, a pesar de los grandes esfuerzos realizados, el bosón de Higgs no se pudo encontrar, siendo la última partícula del modelo estándar que queda por descubrir. El bosón de Higgs es el encargado de comunicar a otras partículas qué masa tienen (alguien tiene que hacerlo). Hay muchas características que la teoría no puede predecir, por ejemplo la masa de las partículas, ni tampoco la magnitud de las dimensiones del universo. Y, aunque existen en él tres fuerzas fundamentales, la cuarta inductora de cambios, la gravedad, cuya naturaleza básica se describe en la teoría general de la relatividad, se encuentra hasta ahora fuera, al otro lado de la puerta, y no puede entrar en el juego. El bosón de Eich de la gravitación, una partícula que comunica a la masa la presencia de otra masa, para que se comporte en consecuencia, y a la que precavidamente ya se ha bautizado con el nombre de gravitón, no se ha descubierto hasta ahora. En todos estos problemas se trabaja mucho actualmente.

Dos de las principales candidatas para una teoría polivalente, más allá del modelo estándar, que resuelva finalmente todos nuestros problemas, son las llamadas supersimetría y teoría de cuerdas. Para comprender en profundidad estas construcciones mentales extraordinariamente complejas, se han de realizar, por desgracia, ensayos difíciles de controlar, así como unos arduos trabajos con fórmulas que dan pánico. La supersimetría es un enfoque esperanzador, en el que a cada partícula del modelo estándar se le asigna una «superpareja», que sólo se diferencia del original en el impulso de rotación, que es exactamente de sentido contrario. Así, cuando una partícula rota hacia la derecha, entonces su superpareja rota hacia la izquierda. Mediante la duplicación de los tipos de partículas se puede resolver toda una serie de problemas, y se consiguen unas prometedoras partículas nuevas que quizá contribuyan a aclarar la cuestión de la materia oscura. Hasta la fecha no se ha encontrado ninguna de estas superparejas, y es un misterio la razón por la cual es tan rara su aparición en las partes del mundo físico que se han investigado hasta ahora.

Por otra parte, la teoría de cuerdas ha transformado el modelo estándar haciendo que las partículas elementales no sean consideradas ya como puntos, sino como «hilos»; así adquieren de repente una dimensión, mientras que antes no tenían dimensión alguna. Las distintas teorías de cuerdas tienen la ventajosa característica de poder predecir el número de dimensiones en el universo; como ya se ha dicho, el modelo estándar era incapaz de hacer esto. Según las variantes se alcanzan 10 u 11, o incluso 26 dimensiones; unas cifras tan confusamente distintas que es imposible valorar con exactitud si es mejor saber esto que no saber absolutamente nada. En todo caso estas dimensiones son en su mayoría demasiado pequeñas para que desempeñen un papel en la vida cotidiana; están «compactificadas» en el mundo de los cuantos. Hasta ahora las teorías de cuerdas se resisten tenazmente a cualquier confirmación experimental clara, por lo que algunos investigadores han empezado ya a preguntarse si no será una pérdida de tiempo andar a vueltas con tantos hilos.

Más allá de todos estos esfuerzos por hallar una teoría del mundo, sigue siendo posible que no hayamos encontrado todavía esas partículas elementales, es decir, los auténticos componentes más pequeños de la materia. Hay distintas teorías, sólo

un poco más recientes que el modelo estándar, según las cuales los electrones y los quarks están formados a su vez por partículas aún más pequeñas, que en general se llaman preones, aunque a veces también reciben los nombres de prequarks, rishones, tweedles o maones. Algunos físicos, como Haim Harari, que construyó el modelo del rishón, se preguntan: ¿Cómo es que nos ha tocado, precisamente a nosotros, ser la generación que ha chocado con la frontera fundamental, la de las partículas más pequeñas? (De inmediato podríamos preguntarnos también por qué las nuevas partículas elementales de Harari han de ser definitivamente las más pequeñas). Harari se preguntaba a continuación por qué ha de estar el mundo formado por tantas partículas elementales como prevén el modelo estándar y sus ampliaciones. En cualquier caso, sus rishones aparecen sólo con dos tipos diferentes y, según decía Harari al final del artículo en el que presentaba el modelo, dejan «muchas, muchas preguntas abiertas».

Quizá las preguntas abiertas sean realmente los componentes fundamentales de la materia.



Capítulo 28

Problema p/np

*Quien rehúse dedicarse a la aritmética
estará condenado a decir bobadas.*

JOHN MCCARTHY

Algunos problemas son tan desesperantemente difíciles, que vale la pena pararse a pensar si tienen en realidad alguna solución, si no existe solución alguna, o si buscar la solución es perder el tiempo. En el caso de los problemas llamados NP¹³, existe una solución obvia, pero le puede a uno crecer la barba hasta el suelo antes de llegar a encontrarla. Por desgracia es enorme la cantidad de personas que se interesan por los resultados de esas ecuaciones NP. Se sientan en sus despachos y se muerden las uñas llenos de impaciencia, mientras el ordenador no deja de hacer cálculos. El llamado problema P/NP es, grosso modo, la cuestión de si, en principio, hay soluciones rápidas para los problemas NP, y éste es uno de los siete enigmas matemáticos por cuyas soluciones el Clay-Institut de Estados Unidos ha ofrecido un millón de dólares. (Otro es la hipótesis de Riemann). La resolución del problema P/NP no supondría en absoluto la resolución rápida e inmediata de todos los problemas NP del mundo. Pero indicaría que se puede tener alguna esperanza. Sin embargo, como diría actualmente la mayoría de los expertos, el panorama es más bien oscuro.

Uno de los problemas NP más conocidos trata sobre un comerciante que recorre el país trabajosamente para vender sus mercancías. La cuestión es en qué orden ha de visitar las ciudades para acabar lo antes posible. El problema está resuelto hoy en día de la siguiente manera: lo que el buen hombre tiene que hacer es simplemente realizar sus ventas a través de eBay. Así se ganaría la libertad de poder viajar después para conocer el mundo. Se compra la guía *999 maravillas del universo que uno debe haber visto a lo largo de su vida* y emprende el viaje. Pero ¿por dónde empezar? ¿En qué orden hay que ir marcando esas atracciones turísticas, para visitarlas todas evitando que la hora de morir llegue antes de

¹³ Problemas P se refiere a problemas polinómicos, mientras que problemas NP significa problemas polinómicos en sentido no determinista. (*N. de la t.*)

terminar el viaje? ¿Se puede conseguir esto antes de cumplir los ochenta años? ¿No se convertirán las vacaciones en algo terriblemente agotador?

Si uno es moderado y sólo aspira a ver la torre Eiffel, el Big Ben, el foro romano y el parque ornitológico de Walsrode, el problema del turista desesperado es fácil de resolver: basta con buscarse un mapa de carreteras y calcular la longitud de los recorridos para todas las ordenaciones posibles de los lugares en cuestión. Así se averigua que, en vez de ir primero de Roma a Londres, luego a París y finalmente a Walsrode, hay que ordenar mejor los destinos. Si a estas atracciones turísticas se añaden unas pocas más, se dispara el trabajo de cálculo. Con cinco destinos hay ya 30 recorridos posibles; si son seis, las rutas posibles pasan a ser 180; con siete son ya 1260, y así sucesivamente. Si se desea visitar tan sólo diez castillos ruinosos en el sur de Francia, hay que calcular más de tres millones de trayectos para hallar la ruta más corta posible.

Para resolver el problema de las 999 atracciones turísticas, aunque se disponga de los ordenadores más rápidos, se requiere un tiempo más largo que el transcurrido desde la aparición del universo. Esto nos puede matar la ilusión que teníamos por hacer el viaje.

Aunque muchas generaciones de personas inteligentes se hayan quedado calvas discutiendo sobre este problema, y sobre otros planteados de forma parecida, por ahora no hay ningún método para determinar la mejor ruta que sea considerablemente más rápido que armarse de valor e ir calculando todas las posibilidades de una en una. Lo que tienen de especiales los problemas NP del tipo del dilema del turista desesperado, no es que exijan la realización de cálculos complicados, ya que, por el contrario, basta con sumar cifras. La dificultad está en la enorme cantidad de posibilidades que hay que contemplar. Los problemas NP son unos maratones, en comparación con el resto de los problemas matemáticos; no hay más que poder correr, pero muy, muy lejos.

Por el contrario, los llamados problemas P dan sólo unas cuantas vueltas al estadio y se pueden resolver haciendo todos los cálculos con un esfuerzo previsible. Un problema P sencillo es, por ejemplo, la suma de números. Si se trata de sumar cinco números, el cálculo se puede hacer en cinco pasos; si son mil números, serán necesarios mil pasos. En general, los problemas P tienen un comportamiento muy

formal: si se les añade una cantidad mayor de datos, no se ofenden, sino que optan por retirarse durante un montón de tiempo a trabajar. En cambio, los problemas NP, al crecer la cantidad de datos, aumentan exponencialmente la cantidad de pasos necesarios para los cálculos. Así pues, la pregunta decisiva es: ¿Son los problemas NP también problemas P? ¿Lo son quizá a escondidas, por la noche y bajo las sábanas, cuando nadie los ve? ¿Existe acaso una solución más rápida para los desesperados turistas y para muchos otros pesados y trabajosos enigmas NP? Ésta es una formulación simplificada del problema P/NP, que supone una tortura para mucha gente.

Y no se trata sólo de matemáticos y trotamundos. Muchos problemas importantes que agotan a los ordenadores en empresas, bancos y oficinas son problemas NP y están emparentados con el dilema de los turistas. Si se dispusiera de resoluciones relativamente rápidas, los serviciales ordenadores podrían hacer cálculos fabulosos. Si los problemas NP fueran al mismo tiempo problemas P, y se supiera con seguridad que existen resoluciones rápidas, se liberarían energías insospechadas para la búsqueda de estas resoluciones. Muchos expertos preferirían hacer lo contrario, si supieran con seguridad que los problemas NP no son problemas P. Por ejemplo, la decodificación de informaciones es asimismo en muchos casos un problema NP, por lo que también está en juego la protección de datos cuidadosamente custodiados. La cuestión relativa a si los problemas NP son problemas P, o no lo son, es por lo tanto mucho más importante que el fastidioso problema de los cordones, que siempre se sueltan en los momentos más inoportunos.

Para demostrar que los problemas NP son en definitiva problemas P camuflados, bastaría probar que uno de ellos tiene una resolución limpia y rápida. Ésta es la consecuencia de una característica extraordinariamente práctica que tienen muchos problemas NP importantes: si se encuentra una resolución rápida para uno de ellos, esto significa que existe lo mismo para todos los demás, aunque todavía no se conozca. Para demostrar con todas las garantías que los problemas NP no son al mismo tiempo problemas P, habría que proceder de otra manera: por ejemplo, se podría intentar demostrar que todas las resoluciones del problema de los turistas desesperados son largas y pesadas, y no sólo las que han sido comprobadas hasta

la fecha, sino también todas las que serán objeto de reflexión en el futuro. De nuevo bastaría con probar esto para un determinado problema NP. Si, por casualidad, hay un lector que sepa cómo se puede resolver en un santiamén el problema de los turistas desesperados, o si sabe con toda seguridad que no existe una resolución de este tipo, por favor que no se lo calle, porque su descubrimiento tendría consecuencias insospechadas para casi todos los problemas NP, llenaría de emoción a muchas personas que se dedican a la investigación o a la economía, y además le haría rico.

Y, por favor, que no le desanime el hecho de que la mayoría de los expertos consideren improbable que P sea igual a NP. Para los matemáticos la probabilidad no significa nada; sólo una demostración definitiva les deja dormir tranquilos. Además, nunca debe suponerse que algo no es posible, sólo porque todos afirmen que no lo es: si hace unos pocos miles de años le hubiéramos preguntado a un mono si sus sucesores iban a estar en condiciones de matar a un elefante, el primate en cuestión habría hecho el gesto del destornillador en la sien. Pero resulta que unos días después la caza del elefante no planteaba ya dificultad alguna. Lo mismo nos puede pasar a nosotros con los problemas NP, ante los cuales estamos ahora como los monos ante los elefantes.

Para dar respuesta a la cuestión del problema P/NP , no se necesitan posiblemente unas matemáticas que tengan una complicación de alto nivel; muchos afirman que bastaría con que a alguien se le ocurriera una idea realmente buena. Por eso, entre los que se dedican a resolver problemas por afición, es grande la cantidad de gente que propone demostraciones, un número mucho mayor que el de los que estudian otros enigmas matemáticos planteados hace milenios. En consecuencia, lo que necesitamos es más turistas inteligentes que se tomen realmente en serio la planificación de sus vacaciones.



Capítulo 29

Propina

... Si dejo una mancha de jarabe, mañana tendré que dejarle una propina a la camarera encima del escritorio. Aunque, ¿de verdad quiero hacerlo? Las propinas sólo se dan por miedo a una reacción adversa por parte de quien ofrece un servicio. Pero yo nunca sabré si la doncella limpió mi habitación de buen o de mal humor, pues cuando ella llegue yo ya me habré ido.

MAX GOLDT, «QQ»

Después de comer, millones de científicos aficionados se plantean regularmente la siguiente pregunta: ¿cuánto dinero hay que dejarle al camarero? En total, sólo en Estados Unidos, la cifra asciende hasta los veinte millones de dólares anuales. No sabemos por qué hay que pagar más de lo que dice la cuenta, por qué tiene que ser tanto dinero y de qué depende la cantidad que dejamos.

Los pocos científicos que se dedican a estudiar el fenómeno de las propinas no han dado aún con una respuesta clara sobre la acuciante cuestión de las propinas, pero pueden explicar varias cosas interesantes.

De entrada, no resulta sorprendente que la propina se incremente proporcionalmente con la cuenta, pues en cualquier libro sobre modales se puede leer que la costumbre es dejar de propina un tanto por ciento del importe de la cuenta. Por otro lado, una de cada cinco personas (por lo menos en Estados Unidos) no se ciñen, por motivos diversos, a esa convención, sino que dejan siempre la misma cantidad, independientemente de la cantidad que figure en la cuenta. Quienes pagan con tarjeta suelen ser más generosos, probablemente porque la presencia de la tarjeta despierta inconscientemente una fiebre consumista. Eso funciona también cuando (como sucede en Estados Unidos) la cuenta se presenta

sobre una bandejita donde aparecen las insignias de diversas tarjetas de crédito, aunque uno termine pagando en efectivo. Como restaurador, pues, no parecería tan mala idea realizar una inversión y tapizar todo el local con tarjetas American Express de tamaño sobrenatural.

Como era previsible, la propina depende también de la valoración que el cliente haga de la calidad del servicio: si se siente bien servido, paga un poco más, aunque no mucho más. En lugar de en ser eficientes, los camareros deberían concentrarse en ser amables o llamativos. Tal como descubrió el psicólogo Michael Lynn, las propinas aumentan considerablemente si el camarero «toca ligeramente el brazo, la mano o el hombro» del cliente, si «lo distrae o bromea con él» o si «dibuja una cara sonriente o algo parecido en la parte trasera de la cuenta». Si uno actúa así, probablemente pueda meter el dedo en la sopa y que al cliente le dé igual. Por lo demás, la propina depende entre otras cosas de las dimensiones de la ciudad, de la edad del cliente, de lo que cobra, de si el personal es atractivo, de si hace sol o si la previsión del tiempo dice que va a hacer sol el día siguiente, y de si la camarera lleva flores en el pelo. El cliente es un ser misterioso.

Pero ¿por qué pagamos propina si es voluntario y nos cuesta algo de dinero? La mayoría respondería: «Porque es lo que se hace». Por desgracia, ésa es una respuesta a todas luces insuficiente. He aquí una variante probablemente mejor: porque queremos asegurarnos de que en el futuro nos sirvan igual de bien o mejor. Sin embargo, si eso fuera cierto no dejaríamos propina cuando, previsiblemente, no vamos a volver a ver a la persona que nos ha servido; a un taxista, por ejemplo. En dicho caso sería mucho más razonable pagar antes de comer o antes de comenzara el trayecto para predisponer favorablemente al camarero o al conductor. La verdad, sin embargo, parece ir por otros derroteros.

Uno podría pensar también que la compasión desempeña cierto papel cuando se trata de dejar propina: hay que ayudar un poco a los pobres camareros. De hecho, algunos datos apuntan en este sentido; en un estudio realizado en Estados Unidos, el 30 por ciento de los encuestados declararon que dejaban propina porque tenían la sensación de que quienes la recibían, la necesitaban. Sin embargo, también en este caso se plantean numerosas contradicciones: si la compasión fuera tan importante, sería razonable pensar que las propinas aumentarían cuando la diferencia de salario

entre quien ofrece un servicio y quien lo recibe fuera mayor, por ejemplo en el caso de los limpiabotas. Tal vez tenga también su importancia la sensación de que los propietarios de los restaurantes atan corto a sus camareros, pero en ese caso los primeros no recibirían nunca una propina. Hasta la fecha, no se ha podido documentar ninguna de esas teorías.

También existe la posibilidad de que dejemos propina para que el mundo sepa cuán generosos somos, aunque tal vez eso no sea cierto. Pero, en ese caso, ¿por qué se dejan propinas también cuando el mundo no está presente, por ejemplo cuando uno está a solas con un taxista? El economista Robert Frank propone una posible respuesta: debemos demostrarnos constantemente a nosotros mismos lo generosos que somos para acallar la mala conciencia y mejorar el karma.

Esto, según el sociólogo Diego Gambetta, que por otro lado se ocupa también de la mafia, permitiría realizar un pronóstico: las personas realmente generosas dejarán menos dinero que los avaros, pues no lo necesitan. Aunque esa teoría aún no se ha puesto a prueba, no es cierta.

Además, tal como argumenta Michael Lynn de forma irrefutable, se trata de una teoría demasiado creativa para ser cierta.

Al final, es posible que no nos quede más explicación que la de «porque es lo que se hace».

Puede muy bien ser que cuando nos sentimos obligados a dejar propina estemos tan sólo sucumbiendo al deseo de no llamar la atención y comportarnos como lo hacen todos los demás porque es lo más fácil. Sin embargo, también los individualistas más feroces dejan propina y no menos que los demás. Por otro lado, no seguir la convención y no dejar propina tiene consecuencias desagradables por añadidura: los empleados que le han servido a uno le miran mal y uno se siente avergonzado y mala persona, más aún si dejar propinas es el único acto altruista que tiene para con el prójimo. Además, si bien el acto individual de dejar propina tiene poco sentido, como fenómeno de masa es de lo más positivo, pues supone un estímulo para la mejora del servicio. Sin embargo, ¿qué sucede si uno, saciado y satisfecho después de una buena cena en un restaurante, comienza a pensar en las consecuencias sociales de su acto?

Al parecer hay diversos motivos para no romper con la norma social y seguir pagando entre un 10 y un 15 por ciento de propina, aunque ni mucho menos se sepa por qué existe esa norma. Y también hay buenos motivos para seguir estudiando el fenómeno. Así descubre uno que las decisiones económicas de la gente no son tan racionales y egoístas como le gustaría a una economía exacta, que querría que calculásemos muy bien cuánto dinero dejamos y a cambio de qué. En lugar de eso, nos regimos por una intrincada mezcla de actitudes tradicionales, deseos diversos y algunos elementos racionales.

Por cierto: en muchos países, como en China, por ejemplo, no se deja propina. Una vez más, se desconoce el porqué. Difícilmente se podrá justificar aduciendo las décadas de educación comunista, pues lo mismo sucede en países capitalistas como Singapur o Australia. A veces, los taxistas asiáticos incluso redondean el precio a la baja, con lo que se genera una propina negativa o, dicho de otro modo, el taxista les da una propina a los clientes. Pero también en el paraíso de las propinas, Estados Unidos, se observan ciertas contradicciones. En el año 2006 el famoso restaurador y autor de libros de cocina Thomas Keller fue noticia tras anunciar que prohibía tajantemente las propinas en su restaurante de Nueva York. En el debate posterior, la costumbre de dejar propina fue calificada tanto de «típicamente americana» como de «muy poco americana», precisamente porque nadie sabe exactamente cómo funciona.

El estado actual de la ciencia queda perfectamente resumido en la escena inicial de la película de Tarantino *Reservoir Dogs*. Al final de un largo desayuno en un café y cuando llega el momento de pagar la cuenta, Steve Buscemi (alias «Señor Rosa») declara que él no cree en las propinas, y desencadena con ello una larga discusión. Utilizando muchas de las ideas expuestas en este capítulo, ofrece una explicación pormenorizada de por qué cree que las propinas son inútiles y no tienen sentido. Sin embargo, al final la sensación de tener que dejar propina acaba siendo más fuerte y el Señor Rosa paga lo que debe, agradecido por no tener que pagar la cuenta. Siempre encontramos un motivo u otro.



Capítulo 30

Rayos globulares

¡Y ahora un círculo luminoso que crepita!

Ésta es la prueba concluyente.

DANIEL DÜSENTRIEB

La humanidad lleva ya muchos siglos disfrutando con el misterio de los rayos globulares. Se han llenado varios miles de páginas con informes relativos a las observaciones de rayos globulares, y sus correspondientes teorías, experimentos y mitos. En el espectro de los intentos de explicación de este fenómeno se pasa con fluidez de uno a otro dominio, recorriendo la ciencia propiamente dicha, la ciencia ficción, el esoterismo y la chifladura, de tal modo que las teorías serias suelen exponer con una complejidad descorazonadora lo que otras teorías menos serias aprovechan desvergonzadamente. Además se insiste una y otra vez en que el rayo globular no es un fenómeno natural, sino simplemente una fantasía. Lo cierto es que hasta la fecha no existe una explicación ampliamente aceptada sobre el modo en que surgen y se comportan los rayos globulares.

La investigación de los rayos globulares se basa esencialmente en informes de testigos presenciales. Desde hace más de quinientos años están documentadas las observaciones de rayos globulares; un banco de datos ruso contiene por sí solo unos diez mil casos registrados durante las últimas décadas. Sin embargo, se trata de un fenómeno raro. Pocas son las personas que han podido vivir alguna vez esta cosa tan maravillosa. El periodista Graham K. Hubler describe un caso típico: «Estaba yo con mi novia en un parque de Nueva York, en concreto debajo de un templete. Llovía bastante copiosamente. Una bola de color blanco amarillento, más o menos del tamaño de una pelota de tenis, apareció a la izquierda de donde estábamos nosotros, aproximadamente a unos treinta metros de distancia. La bola planeaba a una altura de dos metros y medio sobre el suelo y se movía lentamente hacia el templete. Cuando llegó allí, cayó de repente al suelo y pasó más o menos a un metro de distancia de nuestras cabezas. Resbaló por el suelo, salió del templete,

ascendió a dos metros de altura, se alejó diez metros más, cayó de nuevo al suelo y desapareció sin que se oyera explosión alguna».

Muchos informes son como éste, o muy parecidos. De la enorme cantidad de descripciones, algunas muy precisas y detalladas, otras más bien confusas, surge una imagen de las características del rayo globular que es sólo parcialmente coherente y constituye la base de todos los intentos de dar una explicación: la mayoría de los rayos globulares se han observado mientras descarga alguna tormenta, pero no es así en todos los casos. A menudo, pero no siempre, van precedidos por un relámpago visible y normal. En muchas ocasiones la bola aparece a unos cuantos metros sobre el terreno, pero a veces también directamente por el suelo, y es menos frecuente que parezca caer del cielo, aunque también sucede así en algún caso. El diámetro de la bola oscila entre unos cuantos centímetros y un metro, siendo su tamaño típico más o menos el de un balón de fútbol. La mayoría de los rayos globulares tiene color amarillo, blanco o rojizo, y su luminosidad viene a ser la de una bombilla de poca potencia. Por su tamaño y su aspecto luminoso se parece a un farolillo para niños, pero sin cara pintada. Algunos rayos globulares planean, otros caen hacia abajo. Pueden emitir un silbido, apestar a azufre o hacer que el agua hierva. A veces botan sobre el suelo como una pelota de goma. La mayoría de las veces se desintegran al cabo de unos pocos segundos, pero a veces duran casi un minuto. Algunos son fríos al tacto, pero otros están muy calientes. Según ciertas observaciones parecen atravesar libremente muros o lunas de cristal, lo cual es una habilidad que plantea dificultades a la hora de formular cualquier teoría. Así entró un rayo globular en una iglesia del condado inglés de Devon en 1638, con el resultado de cuatro personas muertas y sesenta heridas. Rara vez se comportan los rayos globulares de una forma tan inconveniente; casi siempre se les ve tranquilos y mesurados, y desaparecen pacíficamente, aunque también se han observado algunos ejemplares extrovertidos que proyectan chispas y se acompañan de pequeñas explosiones. Por lo visto, suelen surgir casi siempre al aire libre, aunque a veces pueden aparecer en espacios cerrados. Incluso hay algunos que aparecen en submarinos y aviones. Una situación terrorífica: volamos en una noche de tormenta y del asiento contiguo al nuestro surge una bola luminosa que planea por el aire dentro del avión; a nadie se le puede reprochar que en una situación así

tire por la borda el sentido común y empiece a creer en cosas extravagantes. Por lo tanto, no es de extrañar que hasta hace unos cien años se pensara ante todo que los rayos globulares eran una manifestación sobrenatural.

Los intentos de explicar el fenómeno que se han ido recogiendo durante el último siglo pueden clasificarse en tres categorías: los rayos globulares son un fenómeno excepcional que se produce en la atmósfera terrestre, o son algo totalmente absurdo, o no existen en absoluto. Según esta última variante, la aparición de los rayos globulares se sitúa simplemente en los ojos o el cerebro del observador, por lo que el fenómeno se convierte en una ilusión óptica o, si vamos más lejos, en una alucinación. La posibilidad de zanjar así el problema es, por supuesto, completamente legítima: cuando no se puede comprender una percepción, quizá sea que ésta es en sí misma irreal.

Es cierto que «vemos» imágenes en movimiento en lo que llamamos películas, aunque en realidad allí no hay nada que se mueva, y se trata simplemente de varios fotogramas individuales donde no hay movimiento alguno.

Hasta la década de 1970 se desarrollaron serias construcciones teóricas según las cuales en el cerebro surgen imágenes luminosas, por ejemplo, como consecuencia de que la retina «siga ardiendo» después de haber visto relámpagos auténticos. Sin embargo, todos esos modelos son problemáticos; ninguno de ellos puede explicar de manera irrefutable las numerosas observaciones realizadas. Por ejemplo, la teoría de la retina que «sigue ardiendo» falla porque los observadores de rayos globulares no siempre habían estado mirando un rayo en un momento inmediatamente anterior. Además no está claro de dónde podrían proceder los interesantes ruidos y olores que algunos rayos globulares producen. Sin embargo, actualmente estas hipótesis no han llegado todavía a estar obsoletas, y por lo tanto no se considera del todo desatinado negar la existencia de estas misteriosas bolas de fuego. No obstante, la gran cantidad de observaciones documentadas así como las fotografías recopiladas hasta ahora, aunque estas últimas sean pocas, apuntan con bastante seguridad a la idea de que nos encontramos ante un fenómeno que se produce fuera de nuestras cabezas. Pero la existencia de los rayos globulares no estará realmente demostrada hasta que se consiga atrapar uno, domesticarlo y mostrarlo en congresos de expertos.

Actualmente son mayoría las explicaciones que se basan en la suposición de que la bola de fuego está formada por plasma. El plasma se forma cuando se calientan los gases durante tanto tiempo que las partículas individuales de dichos gases, a causa de las elevadas temperaturas, empiezan a desprenderse de electrones de una manera desesperada. El Sol, por ejemplo, está formado esencialmente por uno de estos gases que tienen carga eléctrica. Sin embargo, el plasma no explica por sí solo la aparición de rayos globulares, entre otras cosas porque las burbujas de aire caliente y con carga eléctrica tendrían que ascender, y eso rara vez lo hacen los rayos globulares. Además una bola de plasma tendría que deshacerse en cuestión de fracciones de segundo, y los rayos globulares se mantienen durante bastante más tiempo. Por lo tanto, la energía que consume el rayo globular no puede proceder sólo del plasma. Se necesita algún tipo de calentamiento adicional para producir bolas de fuego estables y de larga vida.

Estas reflexiones condujeron, a través de muchos pasos intermedios, a un modelo que hoy en día es muy popular y tiene relación con los llamados aerosoles, es decir, una acumulación de partículas en suspensión. En el año 2000, los neozelandeses John Abrahamson y James Dinniss plantearon una teoría según la cual los rayos globulares se forman cuando los rayos de la tormenta caen en la tierra. La elevada energía de los rayos calienta la tierra y forma remolinos de pequeñas partículas de polvo en el aire. Estas partículas, mediante procesos químicos, se unen en una complicada red que se convierte en una bola lanosa de tierra con la consistencia de una madeja de polvo. Esta malla flotante rodea y penetra la burbuja de aire caliente y forma así el rayo globular.

De esta manera, la energía del rayo contenida en la bola de polvo se almacena de forma química, algo parecido a lo que sucede en una batería. La burbuja arde lenta y regularmente, liberando así la energía del rayo.

El planteamiento de los aerosoles no es nuevo en la investigación de los rayos globulares, pero el modelo propuesto por Abrahamson y Dinniss puede explicar muchas observaciones, al menos en teoría. En la realidad las cosas son un poco diferentes. Las mejores imitaciones de rayos globulares, que se han realizado hasta ahora con base en la teoría de los aerosoles son las que presentó en enero de 2007 un equipo de investigadores brasileños que trabajaban dirigidos por Antonio Paváo y

Gerson Paiva. Colocaron placas de silicio entre dos electrodos, vaporizaron algunas zonas del silicio e hicieron finalmente que saltara un rayo artificial entre los electrodos.

El resultado fueron unos fenómenos luminosos que por su color y duración parecían auténticos rayos globulares, aunque algo pequeños y poco vistosos. Quedaría todavía por aclarar si eran realmente los primeros rayos globulares producidos artificialmente.

Otros expertos explican el fenómeno como la consecuencia de descargas eléctricas sobre superficies de agua. Al contrario que en la teoría de los aerosoles, en este caso el rayo no cae en la tierra, sino en un lago, en un recipiente con agua o incluso en un charco, como sucede a menudo cuando hay tormenta. En 2002 un grupo de investigadores de San Petersburgo consiguió producir de este modo una especie de rayo globular en el laboratorio. Cuatro años más tarde, unos científicos alemanes que trabajaban con el físico experto en plasma Gerd Fussmann lograron reproducir el mismo resultado. Aplicando unas descargas de alto voltaje sobre agua salada (un procedimiento bastante cercano a la realidad) produjeron unas bolas luminosas de entre 10 y 20 centímetros de diámetro, claramente mayores que las bolas de aerosoles obtenidas en Brasil. Los rayos globulares de Fussmann tampoco sobrevivían más allá de una fracción de segundo, un tiempo demasiado corto para poder competir con los ejemplares auténticos.

La tercera teoría importante que surgió en el marco de la investigación moderna sobre rayos globulares fue la que propuso en la década de 1950 el físico soviético Piotr Kapitza. Esta teoría se basa también en la hipótesis de que un rayo globular es una bola de plasma, sólo que Kapitza considera que sus cuerpos de plasma se calientan desde el exterior, mediante potentes microondas que se originarían durante la tormenta en el entorno de los rayos normales. En consecuencia, el rayo globular sería una bola de aire extremadamente caliente que, atravesada por las microondas, queda en suspensión y resplandece durante la tormenta. También en el laboratorio se pueden conseguir bellos efectos luminosos utilizando microondas: así, en el año 1991, los japoneses Ohtsuki y Ofuruton produjeron bolas de fuego artificiales que planeaban a través de las paredes y contra el viento, es decir,

exactamente lo mismo que describieron algunos observadores. Sin embargo, ni el tiempo de vida, ni los tamaños, concuerdan con los rayos globulares auténticos.

Algo parecido consiguieron en 2006 los investigadores israelíes Eli Jerby y Vladimir Dikhtyar: hicieron pasar microondas por una varilla metálica hasta una esfera de cristal, donde apareció una zona caliente de cristal fundido. Al apartar la varilla metálica, la zona caliente «se fue» de la esfera de cristal y produjo así unas bolas de plasma en suspensión y estables. Por desgracia, estos procesos experimentales están algo lejos de la realidad, porque rara vez nos encontramos una esfera de cristal en medio de una tormenta. De todos modos, el experimento realizado en Israel funciona con los aparatos de microondas disponibles habitualmente en el mercado y, puesto que se puede encontrar las instrucciones precisas en Internet, cualquiera tiene hoy en día la posibilidad de producir imitaciones de rayos globulares en la cocina de casa. En un futuro no muy lejano, seguramente se sabrá apreciar la utilidad de las bolas de plasma como componentes de los sistemas modernos de iluminación de las cocinas.

Aparte de los éxitos actuales con los aerosoles, los charcos de agua y las microondas, en el pasado no fueron pocos los intentos de producir artificialmente rayos globulares. Son casi legendarios los experimentos del genio de la física Nikola Tesla que, con ayuda de complicados circuitos de conmutación, produjo también algunas bolas de fuego, aunque seguramente no eran lo que llamamos rayos globulares. A partir de aquellos experimentos y poco antes de su muerte, acaecida en 1943, Tesla construyó al parecer una máquina de rayos mortíferos, que despertó un gran interés en la CIA. A muchos les sucedió lo mismo que a Tesla: produjeron fuegos y ruido, pero en última instancia no pudieron demostrar que aquello tuviera algo que ver con el fenómeno que buscaban. Trágico fue el final que encontró en San Petersburgo un aventurero llamado Richmann, que en 1753 atrajo un rayo auténtico a su laboratorio, para jugar con él. Un rayo globular lleno de furia y del tamaño de un puño saltó de los aparatos y chocó contra Richmann, que murió en el acto.

En conjunto se puede decir que todas y cada una de estas teorías plantean algún problema. Los aerosoles, el agua y las microondas pueden producir bellas luminosidades, pero ninguna de las teorías existentes hasta ahora puede explicar

las características observadas en los rayos globulares, especialmente la forma, el color, el tamaño y el tiempo de vida. También son problemáticos los rayos globulares que aparecen en edificios, o incluso en aviones, donde normalmente no hay tierra ni charcos de agua. Por suerte, para la explicación del fenómeno se han realizado además numerosos intentos alternativos que, aunque tampoco han conseguido solucionar el problema, destacan por su notable riqueza imaginativa.

Desde la década de 1990, el cosmólogo Mario Rabinowitz y sus colaboradores no han dejado de afirmar la posibilidad de que en el interior de los rayos globulares se oculten diminutos agujeros negros, es decir, versiones en miniatura de cuerpos celestes extremadamente pesados, que proceden de la agonía de estrellas gigantes. Sin embargo, la existencia de esos pequeños agujeros negros microscópicos no ha sido demostrada hasta ahora. Otros expertos explican los rayos globulares diciendo que son la consecuencia de la penetración de la antimateria en la atmósfera terrestre: la Tierra es bombardeada de manera continua por materia normal procedente del universo, que se presenta en forma de meteoritos pequeños, grandes y medianos, cuya incandescencia se puede ver en las noches claras en forma de estrellas fugaces. También se sabe que por cada partícula elemental, como el electrón o el protón, existe una antipartícula de antimateria. Si hubiera meteoritos formados totalmente por antimateria, estos invasores desaparecerían al entrar en contacto con la atmósfera, porque la antimateria y la materia no se llevan bien, y al chocar se destruyen mutuamente. En ese choque, según dice la teoría, podrían formarse unas luminosidades parecidas a los rayos globulares. Sin embargo, por lo que respecta a la existencia de rocas de antimateria, de momento no hay pruebas convincentes.

En 2003, el físico John J. Gilman propuso la teoría de que los rayos globulares están formados por átomos extremadamente energéticos, es decir, átomos locamente excitados y agitados, que a causa de ese bombeo de energía han llegado a inflarse hasta alcanzar un tamaño gigantesco de varios centímetros de diámetro. No obstante, quizá esto también sea absurdo, y los rayos globulares sean «nodos electromagnéticos» (Antonio F. Rañada y José L. Trueba, 1996), «ondas de choque que se generan por explosiones puntuales en la atmósfera» (Vladimir K. Ignatovich, 1992), o incluso «tornados ardientes» (Peter E. Coleman, 1993). No faltan

propuestas interesantes, pero, sin embargo, ninguna está tan bien pensada que haya convencido a la mayoría de los investigadores.

Tal vez se trate en definitiva de meros ovnis, como afirma una cantidad nada despreciable de contemporáneos. ¿Por qué no pensar que los extraterrestres exploran la Tierra desde unas bolas de fuego? Están en su derecho, y es posible que les gusten las noches tormentosas. Además, nos ahorran así el esfuerzo de explicar el fenómeno recurriendo a cosas tan complicadas como las microondas o los aerosoles. No obstante, hemos de tener en cuenta que la existencia de ovnis está mucho menos demostrada que la de rayos globulares. Explicar algo desconocido utilizando otra cosa desconocida se suele considerar en los círculos científicos como una falta de estilo. Los escépticos piensan que es justo al revés, y que muchos avistamientos inexplicables de ovnis corresponderían en realidad a rayos globulares. También los círculos del maíz pueden explicarse como consecuencias de apariciones de rayos globulares.

Entonces ¿son los ovnis rayos globulares, o es que los rayos globulares son ovnis? Sospecho que todavía leeremos más de una vez que el fenómeno de los rayos globulares ya está definitivamente explicado.



Capítulo 31

Resfriado

Aunque el cerebro de la persona esté en cierto modo limpio y sano, en ocasiones irrumpen en él remolinos de aire y de otros elementos y hacen que ciertos humores de distintos tipos fluyan entrando y saliendo, y produzcan en las vías nasales y en la garganta un vapor nebuloso, de tal modo que se concentra un pus dañino como vapor de agua enturbiada.

HILDEGARDA DE BINGEN, «Sobre el resfriado»

En comparación con, pongamos por caso, el ciempiés *Illacme plenipes*, del que se vieron un total de trece ejemplares a lo largo del último siglo, el resfriado es un fenómeno relativamente cómodo de investigar, porque no hace falta buscarlo durante mucho tiempo. Los adultos de todo el mundo contraen esta enfermedad por término medio entre dos y cinco veces al año; los escolares entre cinco y siete veces. Aunque es en proporción fuerte la presión de la investigación sobre esta enfermedad, que no es precisamente exótica, por ahora no sabemos cuándo y por qué se resfrían las personas. Y esto a pesar de que en la larga historia de la investigación se ha descubierto de vez en cuando alguna cosa.

Se sabe, por ejemplo, que en los adultos entre el 30 y el 50 por ciento de las afecciones son desencadenadas por los rinovirus, mientras que el resto se las reparten unos cuantos virus más, entre ellos el metapneumovirus, descubierto en 2001, así como unas cantidades considerables de agentes patógenos desconocidos hasta ahora. Los síntomas típicos (vías respiratorias congestionadas, fiebre leve, dolor de garganta) no tenemos que agradecerseles a los agentes patógenos, sino de forma indirecta a la reacción de nuestro sistema inmunológico. Una vez que se

supera el resfriado, se adquiere una inmunidad frente al desencadenante, por lo que se enferma sólo una vez a causa de un virus determinado. Sin embargo, dado que existen entre 100 y 200 agentes patógenos del resfriado, la variedad es suficientemente amplia como para infectarse con nuevos virus cada año durante toda la vida. Por un lado, no es fácil identificar el agente patógeno, y por otro lado hay muchos otros microbios que pueden desencadenar los síntomas de un resfriado, lo cual no hace precisamente más fácil el trabajo del investigador. Por esta razón, muchos estudios tienen una validez limitada, ya que los casos supuestamente investigados no se han diferenciado claramente de otras enfermedades parecidas como la fiebre del heno o la gripe.

El gran misterio es cómo se adquieren los agentes patógenos. Lo único que está demostrado sin lugar a dudas es que en los experimentos de laboratorio se infecta un 95 por ciento de las personas que reciben los rinovirus mediante un goteo directo en la nariz. Entonces ¿cómo nos llegan los agentes patógenos cuando no se utiliza una pipeta? En un experimento clásico se colocó a personas sanas junto a otras resfriadas durante dos horas en una habitación, separándolas mediante una cortina, y al cabo de una hora, por precaución, se administró a las personas resfriadas una dosis de polvo estornutatorio, con el fin de que no fueran tan discretas como para guardarse todos los virus del resfriado para sí mismas. El resultado fue que enfermó sólo alrededor del 10 por ciento de las personas sanas, lo cual no contradice totalmente la idea de transmisión por el aire como modo principal de contagio, pero tampoco la apoya con mucha fuerza.

Con el fin de averiguar si la presencia de personas infectadas era realmente suficiente para distribuir por la habitación una cantidad apreciable de agentes patógenos, en otro experimento se fijó en la nariz de un voluntario sano un fino tubo por el cual, como en un auténtico resfriado, empezó a gotear un líquido incoloro, que llevaba añadido un marcador fluorescente. Al cabo de unas pocas horas, durante las cuales el falso enfermo habló y jugó a las cartas con otros participantes en el experimento, toda la habitación, incluidas las cartas y los muebles, estaba embadurnada de un color fluorescente, incluidas las narices de las otras personas sometidas a la prueba. Por una parte, decimos que no se sabe con exactitud cómo se realiza la distribución de fluidos corporales, por ejemplo en los

vagones del metro, y por otra parte no nos parece tan importante esa presencia generalizada de agentes patógenos: las madres no se contagian inevitablemente cuando están con sus hijos resfriados, no sucede tan a menudo que los cónyuges se infectan mutuamente, e incluso el hecho de besar a una persona resfriada se considera inofensivo.

Está claro que el proceso de contagio es en sí mismo un procedimiento complicado.

Ni siquiera la afirmación básica que dice «Uno se resfría cuando se contagia a través de otras personas que tienen resfriado» es indiscutible, aunque haya algo que la apoye: cuando sólo se podía llegar por barco a los pueblos de Groenlandia y Spitzbergen, la población estaba libre de resfriados durante los meses de invierno, que era cuando vivían incomunicados con respecto al mundo exterior; sin embargo, tras la llegada del primer barco, en primavera, se desencadenaban infaliblemente las toses fuertes y los estornudos. Estos hechos están bien documentados, pero hasta ahora no se ha explicado suficientemente por qué en el ártico el resfriado se comporta tan claramente como una enfermedad infecciosa traída del exterior, mientras que en una buena cantidad de grandes estudios se muestra que las epidemias de resfriado del hemisferio norte no se extienden continuamente a través de la población, como sería de esperar en el caso de una auténtica enfermedad infecciosa. Las grandes oleadas de resfriados aparecen casi al mismo tiempo en todas las zonas investigadas, con máximos anuales en enero, septiembre y noviembre.

Algo típico en la investigación del resfriado: aquí también hay más de un estudio minuciosamente realizado de cuyos resultados se desprende justo lo contrario.

Se intenta explicar estas oleadas de resfriados simultáneas partiendo de la hipótesis de que llevamos en nosotros mismos los virus correspondientes, que tienen la condición de «organismos comensales», por lo que la infección siempre dormita dentro de nosotros. Sin embargo, hace falta un estímulo exterior para desencadenar la infección. No está claro cuál puede ser este estímulo: los factores climáticos, como el viento, la humedad, el calentamiento o el enfriamiento repentinos del medioambiente o del cuerpo se mencionan siempre desde la antigüedad y se reflejan en muchas de las palabras indoeuropeas que se utilizan para referirse al «resfriado». Dado que tanto en el hemisferio norte como en el hemisferio sur la

incidencia del resfriado es mayor en las estaciones frías del año, parece acertado buscar una relación entre las condiciones climáticas y el resfriado, pero hasta ahora no se ha encontrado. A menudo se habla de que la causa puede ser que en invierno nos apiñamos en espacios mal ventilados. Pero, de hecho, en verano también pasamos la mayor parte del tiempo en los mismos espacios, lo cual hace que esta tesis resulte un poco difícil de creer. En cualquier caso, no es muy probable que el aire de la calefacción central sea el causante, ya que a las mucosas de la nariz les va bien el aire seco. Lo mismo en los experimentos que en el desierto, se observa que estas mucosas funcionan también sin problemas cuando la humedad del aire es extremadamente baja. Además, en muchos países el período de encendido de las calefacciones empieza claramente después de que se hayan producido las oleadas de resfriados del otoño. ¿Influye en la infección el hecho de tener un «sistema inmunológico debilitado»? ¿Protegen de alguna manera la felicidad, el pensamiento positivo o hacer flexiones con las rodillas ante una ventana abierta? Nada de esto se ha demostrado hasta ahora, es decir, todos estos factores se han comprobado ya, pero los muchos estudios realizados dan resultados contradictorios, o ningún resultado en absoluto.

¿Puede ser, quizá, que la temperatura ambiente no influya tanto, y que lo importante sea el enfriamiento del cuerpo? Numerosos experimentos se han hecho en los que unos voluntarios dignos de compasión tenían que estar con bañadores mojados en pasillos donde había corriente, pero nunca ha habido resultados satisfactorios. Los mencionados estudios realizados en Groenlandia y Spitzbergen contradicen la posibilidad de tal relación, ya que allí las epidemias tienen que ver con la llegada del primer barco, y no con el influjo de las temperaturas.

Especialmente pérdidas serían las posibles variantes, difíciles de investigar, en las que la infección y el enfriamiento del cuerpo tendrían que suceder en momentos diferentes para que se produjera un resfriado. Además, las investigaciones realizadas a bordo de barcos, o especialmente en submarinos, dan como resultado que, en parte, las tripulaciones que se encuentran en el mar y aisladas del resto del mundo padecen resfriados incluso más a menudo que la media de la población, y eso sucede, por ejemplo en los submarinos, en un entorno sin variables climáticas, ni sol, ni cambios de temperatura, ni viento. Y, finalmente, se publicó en 2005 un

estudio del Common Cold Centre de Cardiff que por primera vez después de mucho tiempo podía probar de nuevo la existencia de una relación entre un enfriamiento agudo por baños fríos de pies y la consiguiente «aparición de síntomas de resfriado». La cuestión de si esos síntomas están relacionados también con una infección, estaría aún por demostrar, según los autores del estudio.

Sorprende bastante que los modelos habituales para explicar fenómenos misteriosos (extraterrestres, antimateria, radiaciones telúricas, Nikola Tesla tiene la culpa de todo) no se hayan utilizado hasta ahora en relación con los resfriados. Quizá convendría estimular a los investigadores mediante la publicación de hipótesis absurdas, para que progresen los conocimientos relativos a esta enfermedad. Una propuesta sería explorar más detenidamente las costumbres de proliferación de pañuelos para la nariz: al fin y al cabo, esas pequeñas criaturas blancas, que parecen tan inofensivas, son las que actúan en secreto para que el resfriado no se extinga.



Capítulo 32

Rey de las ratas

Junto con las otras cosas influyen probablemente otras más, sobre cuyo significado seguimos sabiendo todavía muy poco. En todo caso, está claro que aún hay muchos enigmas por resolver.
ROR WOLF, Teoría mundial y realista del reino de la carne, la tierra, el aire, el agua y los sentimientos de Raoul Tranchirer

La expresión «rey de las ratas» designa un grupo de ratas que están atadas unas con otras por las colas. Suena como un chiste macabro, pero, según todo lo que se ha sabido hasta ahora, son las propias ratas quienes se buscan este destino. Sin embargo, afortunadamente este fenómeno es muy raro y no ocasionará la extinción de las ratas.

Por desgracia, son escasas las investigaciones que se han llevado a cabo sobre los reyes de las ratas. En el siglo XVI aparecieron los primeros informes relativos a este fenómeno, aumentando su frecuencia durante los doscientos años posteriores, para desaparecer finalmente en el siglo XX, por carecer de importancia. La cifra de reyes de las ratas que se han conocido a lo largo de los últimos quinientos años oscila en total entre 30 y 60; por razones que no podemos explicar, la mayoría se ha encontrado en Alemania. Se descubrían en el interior de viejas chimeneas, en montones de heno, en sótanos y en barracas de feria. Fueron pocos los hallazgos que se registraron oficialmente y se describieron con tanto detalle como un rey de las ratas de Lindenau, junto a Leipzig, que fue localizado por un ayudante de molinero el 17 de enero de 1774: aquel monstruo de dieciséis cabezas saltó sobre el mozo del molino y, a causa de ello, fue «liquidado al momento». Unos diez museos centroeuropeos tienen la suerte de poder mostrar un rey de las ratas.

El mayor de los reyes expuestos, un manojo desordenado de 32 animales, fue descubierto en 1828 y se encuentra actualmente en el museo Mauritanum de

Altenburg, una pequeña ciudad de Turingia; además es el único ejemplar que se conserva momificado. Las últimas noticias sobre reyes de las ratas nuevecitos proceden de Holanda (1963), Francia (1986) y Estonia (2005).

Casi todos los reyes de las ratas están formados por ratas negras, en latín *Rattus*. En Europa estas ratas negras vivieron sus mejores tiempos en una época en que todo estaba todavía bastante desordenado y sucio, es decir, antes de que se introdujera la canalización de aguas residuales y la recogida regular de basuras. Durante siglos las ratas negras pudieron encargarse sobre todo de la propagación de enfermedades mortales, sin que tal ocupación les resultara aburrida ni pesada. Desde el siglo XVIII, la rata negra está siendo desplazada por la llamada rata común, la *Rattus norvegicus*, que es más robusta y se desenvuelve mucho mejor en las modernas metrópolis. Además, esta rata posee una cola que es más corta que la de la rata negra, lo que posiblemente le evita el peligro de acabar formando un rey de las ratas. En cambio, la cola de la rata negra es perfecta para que se produzcan anudamientos involuntarios (no sólo es muy larga, sino que también se utiliza para agarrarse y escalar, por lo que puede enroscarse arbitrariamente alrededor de otra cola). Otros roedores rara vez están en situación de anudarse trágicamente, aunque se sabe de los casos de un rey de las ratas de campo hallado en Java, un rey de los ratones de campo en Holstein, y parece que también ha existido un rey de las ardillas. Afortunadamente, el fenómeno es totalmente desconocido en el caso de las ballenas azules.

Aunque hay escasez de investigaciones serias sobre este fenómeno, son muchas las menciones que se encuentran en la literatura no científica. El zoólogo y experto en reyes de las ratas Albrecht Hase recopiló durante la década de 1940 más de mil citas sobre reyes de las ratas en textos que tenían ya varios cientos de años, pero pocas de estas referencias tenían algún valor zoológico. Se trata sobre todo de mitos, especulaciones y literatura de ficción. En estos textos, el rey de las ratas aparece como un prodigio de cien cabezas, como un asiento similar a un trono destinado a algún «rey» de las ratas (de ahí su nombre), como mal presagio, como anunciador de enfermedades y muerte, o incluso como el propio Satanás. El entusiasmo que suscita el rey de las ratas sigue siendo imparable: figura en el reparto de películas de terror y ha dejado sus huellas en algunas novelas de James

Clavell y Terry Pratchett. Desde luego, es difícil liberarlo de la maraña de supersticiones, mitos y fantasías. Sin embargo, el profesor Hase, tras largos años de investigación, llegó a la conclusión de que no se trata de un ser mítico, sino de unas rarezas zoológicas «sobre cuyas causas sólo es posible formular hipótesis».

Las publicaciones científicas más recientes relativas a este tema pueden contarse con los dedos de una mano. En total existen en esencia dos posibles formas de comprender este fenómeno.

Si se cree lo que dicen algunos, los reyes de las ratas son producto de la mano del hombre.

Seguramente esto habrá sucedido alguna vez, ya que con un rey de las ratas se podría lograr un gran éxito en una barraca de feria. Además, se adecúa perfectamente a lo que tiene que ser un fantasma terrorífico. Sin embargo, es muy dudoso que todos los hallazgos se deban a la actividad de cazadores de ratas con grandes habilidades comerciales. Los nudos realizados por la mano humana en las colas de las ratas tendrían un aspecto mejor y más limpio que lo que se suele ver en los reyes de las ratas encontrados hasta la fecha. Las radiografías de reyes de las ratas muestran complicados y caóticos anudamientos de colas con los consiguientes deterioros en las columnas vertebrales de los animales, lo cual sería inadmisibles en un hábil anudador de ratas. Además, a menudo se descubren los reyes de las ratas años después de que éstas hayan muerto, en lugares inaccesibles y no exclusivamente en barracas de feria. Por lo tanto, hay que reflexionar seriamente sobre la razón por la cual las ratas pueden atarse ellas solas de una manera tan firme.

Y no parece que esto tenga que ser tan difícil. Una publicación del año 1965 informa sobre unos experimentos de laboratorio realizados por holandeses sobre los anudamientos más estables posibles. Para ello se pegaron limpiamente unas a otras las colas de los animales, por algunas zonas, y se puso en una jaula a todas aquellas ratas que estaban unidas a la fuerza. Empezaron a arrastrarse desordenadamente unas encima de otras y, en un tiempo sorprendentemente breve, formaron con sus colas un montón de espagueti firmemente anudados, obteniéndose así el primer rey de las ratas producido en un laboratorio. Estos ensayos se realizaron con ratas comunes, que reaccionaron con furia y disgusto

ante su desesperante situación. Ésta podría ser otra de las razones por las que casi nunca se han encontrado reyes de las ratas con animales de esta especie.

Si pueden, se matan a mordiscos las unas a las otras antes que pedir ayuda chillando todas juntas.

Esto último es lo que hacían en muchos casos las ratas negras, por eso han sido encontradas tan a menudo vivas en sus reyes de las ratas. Así pues, para formar estas estructuras, los roedores necesitan unos fluidos que produzcan un efecto adherente, como, por ejemplo, la orina, la saliva o restos de alimentos, y un espacio vital reducido, de tal modo que la familia de ratas sólo puede colocar sus colas manteniéndolas juntas en un gran montón. Un posible argumento en contra: en realidad la rata tiende a mantener una limpieza escrupulosa y dedica mucho tiempo a una higiene corporal básica. Por este motivo, reaccionaría sintiéndose contrariada cuando se le obliga a estar pegada a otros colegas.

Quizá suceda lo mismo sin necesidad de que las ratas estén pegadas. Los nuevos conocimientos sobre este fenómeno vienen por una vía insospechada: el físico Jens Eggers y sus colegas de Bristol publicaron en 2006 un trabajo sobre la formación de lo que se denomina «ensalada de cables», un anudamiento involuntario de cables, con lo que realizaron lo que era seguramente el primer estudio realizado sobre este tema. En la publicación aparecen dos constataciones que son importantes para aclarar el fenómeno de los reyes de las ratas: por una parte, la ensalada de cables se forma sólo a partir de una determinada longitud de cable, que en el experimento de Eggers es de unos dieciséis centímetros, lo que explica de nuevo por qué los reyes de las ratas se componen casi exclusivamente de ratas negras, que son las que tienen colas largas.

Por otra parte, hay que revolver los cables fuertemente sólo durante unos pocos segundos o minutos para que se formen los nudos. Imaginémonos ahora unos cables que, cuando ya se ha formado un primer nudo, sigan agitándose continuamente, y así podremos comprender que para hacer un rey de las ratas no hace falta posiblemente nada de pegamento, sino sólo un embrollo más o menos breve, según las circunstancias.

¿Cómo se genera la confusión en un nido de ratas y entre sus colas correspondientes? Por ejemplo, organizando un alboroto inesperado, que espanta a

los animales y los incita a huir. La confusión resultante podría motivar fácilmente los primeros anudamientos, que después, en un caos cada vez mayor, acabarían convirtiéndose en un rey de las ratas. Otras posibilidades son las luchas de los machos que pretenden aparearse con las hembras, o un agrupamiento gregario para protegerse del frío. Muchos reyes de las ratas se forman durante las primeras semanas de vida, cuando los animalitos todavía no pueden sobrevivir sin sus madres, pero se mueven ya con independencia por la zona. En estos momentos, la rata es un ser insensato con cola larga, unos presupuestos ideales para llegar a una absurda situación de peligro. En cambio, lo que se descarta claramente es que el rey de las ratas pueda surgir en el vientre materno. Un rey de las ratas no es una camada de siameses, no está compuesto por fetos de rata que se hayan desarrollado juntos antes del nacimiento, sino un agrupamiento de individuos independientes que son víctimas de un encadenamiento de circunstancias desafortunadas. En cualquier caso, parece claro que los reyes de las ratas sobreviven a menudo durante cierto tiempo, seguramente con gran tormento, pero en un estado en que se encuentran, sorprendentemente, bien alimentados. No se sabe si es cierto, como se ha dicho muchas veces, que la compasión mueve a otros de su especie a alimentarlos.

Parece más plausible la idea de que los miembros del rey de las ratas, durante un tiempo, se alimentan de los restos dejados por los colegas que viven en el entorno. Hay discusiones sobre cuál es el tipo de comportamiento, seguramente imprudente, que motiva la formación de un rey de las ratas. Además, dada la escasez de los hallazgos, es difícil decir algo. De todos modos, ahora que tenemos las cosas lo suficientemente claras como para partir del hecho de que el rey de las ratas no es Satanás personificado, sería el momento de ir esclareciendo poco a poco este misterio que dura ya siglos. Probablemente, en los lugares en que surgen los reyes de las ratas se esconden otras apariciones aún más interesantes, que serían adecuadas para ilustrar mundos de fábula y películas de terror, y contribuir así a nuestra diversión. Se sabe demasiado poco sobre los infiernos.



Capítulo 33

Ronroneo

El caso más extraordinario de un silencio casi absoluto, cuando no absoluto, entre los animales terrestres es el de la jirafa. Por lo que yo sé, hasta el momento sólo se le ha oído emitir un débil mugido cuando se le ofrece comida.

FLANN O'BRIEN, *Consuelo y consejo*

Los gatos son unas criaturas sencillas que vienen en varios colores y patrones, todos ellos muy decorativos. Sin embargo, y aunque hace siglos que los observamos y nos rascamos la cabeza (unas veces la nuestra, otras la suya, es cierto), muchos detalles de su existencia, como por ejemplo su hábito de ronronear, aún se nos escapan. Es relativamente sencillo diseccionar un animal muerto y describir con precisión científica lo que albergan en su interior. Pero los gatos muertos no ronronean. Asimismo, no se puede tomar el ronroneo y estudiarlo bajo el microscopio.

He aquí dos inconvenientes técnicos que dificultan considerablemente la investigación del ronroneo.

Efectivamente, a pesar de algunos intentos, aún no se ha encontrado una explicación a cómo ronronean los gatos ni por qué lo hacen. También están por responder las preguntas de si el ronroneo es un rasgo común a todos los miembros de la familia de los felinos o sólo a la mayoría, y de si todos lo hacen del mismo modo y quieren decir lo mismo con ello, pues aparte del caso de los gatos domésticos, apenas existen datos sobre conductas y técnicas de ronroneo. En principio, parece que debemos partir del principio que el ronroneo tiene que ver con la garganta del gato... o tal vez no: «Presuponer que el ronroneo del gato procede de la garganta es tan gratuito como presuponer que los personajes de nuestra serie televisiva favorita existen», escribió el veterinario Walter R. McCuistion en la década de 1960.

McCuistion dejó de creer que el ronroneo se originase en la laringe después de tratar en su consulta a un gato al que un perro le había perforado la garganta de un mordisco. El animal pasó varias semanas respirando a través de un tubo de goma, incapaz de maullar, pero no dejó de ronronear ni por un día. A continuación, McCuistion se enzarzó en una serie de desagradables experimentos consistentes en realizar varios cortes en el cuerpo de un gato joven y, palpando el interior con el dedo, intentar localizar el origen del ronroneo. Su teoría era que éste se origina en el «fremitus», en la región del diafragma, donde se producen turbulencias sanguíneas bajo la vena cava que, tras dar diversos rodeos, terminan repercutiendo en vías respiratorias superiores a través de la tráquea.

En la década de 1970, los fisiólogos John E. Remmers y Herny Gautier conectaron varios gatos a aparatos de medición para, entre otras cosas, determinar si los ronroneos tenían un origen local o, por el contrario, se desencadenaba por impulsos cerebrales. Para ese fin perforaron varios nervios importantes, pero el test del ronroneo dio un resultado positivo. Ambos científicos observaron actividad sincrónica con los ronroneos tanto en los músculos de la laringe como en los del diafragma, y apuntaron a la apertura y el cierre de la glotis como causa del ronroneo. Sin embargo, se desconoce aún el funcionamiento concreto del proceso y su significado. «El ronroneo parece ser la expresión de un estado de excitación provocado por la interacción entre el gato y otros seres de naturaleza amistosa».

Actualmente se supone que el ronroneo tiene su origen en algún punto de la garganta; una y otra vez se cita el experimento del gato sin laringe de McCuistion, aunque por lo visto nadie lo ha repetido de momento. En la década de 1980 se demostró que el ronroneo puede provocarse mediante la estimulación de determinadas regiones cerebrales y que, por lo tanto, tiene un origen claramente central. Sin embargo, no está claro qué frecuencia utiliza para ello el cerebro. Parece que los científicos están de acuerdo en que, independientemente del tamaño del gato, el ronroneo tiene una frecuencia de entre 23 y 31 hercios, que al ronronear el gato respira más profundamente y más rápido, y que el latido del corazón se acelera. El ronroneo sale al mismo tiempo por la nariz y el hocico del gato, que es capaz de maullar simultáneamente.

Sin embargo, aunque supiéramos exactamente con qué parte del cuerpo ronronean los gatos, quedaría aún por responder la pregunta de por qué lo hacen. ¿Es la forma que tienen las crías de decirles a sus madres que todo va bien? A favor de esa teoría estaría el hecho de que los gatos ronronean al beber. Sin embargo, que también los gatos adultos ronroneen, y no sólo por «la interacción con seres de naturaleza amistosa», sino al experimentar dolor o incluso mientras mueren, parece descartar esa hipótesis. Entonces ¿será un recurso del gato para sosegar? ¿O acaso el ronroneo favorece la liberación de endorfinas, una sustancia analgésica endógena?

En una conferencia celebrada en 2006 bajo el inspirado título de «12ª Conferencia Internacional sobre Sonidos de Baja Frecuencia y Vibración y su Control», los investigadores de la acústica Elizabeth von Muggenthaler y Bill Wright presentaron su teoría más reciente: sirviéndose, entre otras cosas, de unos sensores del tamaño de una cerilla colocados sobre el cuerpo de los gatos, habían vuelto a medir con gran precisión el fenómeno del ronroneo y habían constatado que las frecuencias que predominan son las mismas que, en la década de 1990, se demostró que tienen un efecto estimulante sobre el crecimiento óseo. Además, esas mismas vibraciones alivian el dolor, relajan la musculatura y fomentan su crecimiento y flexibilidad.

Ambos científicos sospechan que el ronroneo podría ser una especie de mecanismo de curación espontánea. Esos mecanismos son mucho más marcados en los gatos que en los perros, por ejemplo, por lo que los veterinarios les gusta bromear diciendo que, si se colocan en una misma sala todas las partes de un gato descuartizado, éstas volverán a unirse. Esa mayor relevancia de los mecanismos de curación espontánea se debe a que los gatos sufrieron una domesticación más tardía que los perros y, por lo tanto, son unos animales menos amanerados, pero tampoco se debe descartar que el ronroneo desempeñe un papel importante. Si la hipótesis de Muggenthaler y Wright resulta ser cierta, a los astronautas les bastaría con aprender a ronronear para compensar la mengua de densidad ósea y la atrofia muscular provocadas por la ingravidez. Por desgracia, se trata de una teoría difícil de probar, pues para ello sería necesario reunir un grupo de control formado por gatos sanos que no ronronearan. Sin embargo, generalmente los gatos que no

ronronean no están sanos. Tal vez una solución sería atar un gato que ronronee a un perro y luego medir la densidad ósea de éste.

Sin embargo, lo más probable es que eso no suceda, pues la ciencia alega que tiene cosas mejores a las que dedicarse. Por ello, ignoramos aún muchas cosas sobre los gatos. ¿Por qué vomitan siempre en la alfombra y no sobre el parquet o las baldosas? ¿Por qué prefieren mordisquear los cables de unos auriculares caros que un cordón de zapato barato? ¿Y por qué quieren echarse siempre encima del periódico que uno está leyendo? Quien logre resolver todas esas preguntas podrá sacar al mercado una nueva variedad de gato que, al tiempo que ronronea, vomite sobre el periódico ya leído.



Capítulo 34

Rotación de las estrellas

Los astrónomos jamás parecen dispuestos a hacer algo sencillo.

PETER B. STETSON, Astrónomo

Las estrellas surgen a partir de unos grumos que se forman en nubes gigantescas de gas y polvo. El material del que se hacen estaba con anterioridad en la nube, repartido en un volumen considerablemente más grande, siendo su densidad al principio notablemente menor que al final.

Ahora bien, esas nubes giran, como todo en el universo. Cuando algo que gira se contrae, la consecuencia es que gira cada vez más rápido. Esto se puede observar en los que hacen patinaje artístico sobre hielo, que para girar haciendo piruetas pegan los brazos al cuerpo. (A quien quiera probarlo por su cuenta, no le hace falta más que una silla giratoria y tomar un poco de impulso).

Por lo tanto, las estrellas jóvenes tendrían que girar muy rápido y, según cálculos astronómicos relativamente fáciles, dar una vuelta en torno a su eje en bastante menos de una hora.

Sin embargo, esto no es así: si hacemos que una esfera gire cada vez más rápido, en algún momento las fuerzas centrífugas son en la superficie mayores que las fuerzas que mantienen la esfera cohesionada (en el caso de las estrellas se trata de la gravedad), de tal modo que el objeto en cuestión se rompe. Asimismo, la velocidad de rotación a la que las estrellas se romperían se puede calcular bastante bien, y está muy por debajo de la velocidad que deberían tener las estrellas cuando acaban de nacer. La conclusión más sencilla es que las estrellas no existen, porque durante su formación giran cada vez más rápido, hasta alcanzarla velocidad a la que se desgarran y se destruyen. Sin embargo, creemos que esta conclusión no encaja con la realidad, porque las estrellas existen. Estamos ante un dilema que los expertos, desde la década de 1970, denominan «problema del impulso de rotación» en el contexto del nacimiento de estrellas.

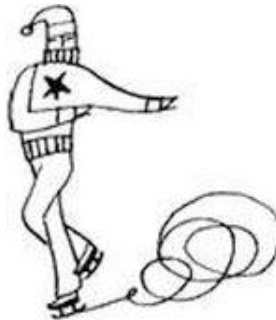
De alguna manera, pues, debe frenarse la velocidad de rotación de las estrellas. Por desgracia, no es fácil observar la «producción» de estrellas, porque éstas al principio están profundamente metidas dentro de la nube donde nacen. Hay que esperar más o menos un millón de años (esto se corresponde en la vida humana con la primera semana en el vientre materno) para que nos resulten claramente reconocibles, cuando la envoltura de gas y polvo se ha deshecho ya completamente. Queda la joven estrella con un disco que la rodea, compuesto por restos de la nube, y del que luego pueden surgir los planetas. Pero, en este momento, la rotación ya se ha frenado bastante. Así pues, mientras es un feto, la estrella lleva a cabo lo que es más interesante, en una fase de su vida en la que es difícil investigarla.

Desde hace varias décadas muchos expertos piensan que el campo magnético desempeña un importante papel en cuanto a salvar a la estrella de morir desgarrada. Según una teoría, que se denomina *Disk-Locking* (teoría del bloqueo de disco), la estrella y su disco se mantienen acoplados en el campo magnético de la primera: mientras la estrella gira, las líneas de su campo magnético recorren el material que la rodea como si fueran los surcos que hace un arado. Dado que el disco ofrece resistencia a este proceso, la rotación de la estrella se frena. Quien tenga a mano una silla giratoria puede sentarse en ella e intentar girar rápido, haciendo al mismo tiempo que sus brazos extendidos atraviesen una masa espesa de gas frío y polvo. Posiblemente a la joven estrella le sucede algo parecido, sólo que, en vez de los brazos, utiliza un campo magnético.

Esta idea del *Disk-Locking* parece correcta, al menos en teoría: las estrellas que poseen un disco giran realmente más despacio que las que no lo tienen. Esto se ha comprobado ya de una manera convincente en los lugares de nacimiento de numerosas estrellas. De algún modo, el disco ocasiona un efecto de frenado. Sin embargo, esta teoría plantea también muchos problemas: por una parte, no está claro si este mecanismo realmente funciona, ni si es suficiente para resolver el problema del impulso de rotación. En realidad, las líneas del campo magnético, al atravesar el disco, deberían curvarse, enmarañarse y deshacerse al cabo de poco tiempo, con lo que la buena relación entre la estrella y el disco se interrumpiría. Por otra parte, nadie sabe con exactitud en qué momento se dotan de un campo magnético las estrellas que están naciendo. Sin embargo, para el *Disk-Locking* se

necesita un campo magnético que sea en cierto modo estable y regular; si no es así, la teoría no tiene sentido.

Algunos expertos fruncen el ceño ante estos problemas; otros prefieren idear otras teorías completamente distintas que tienen que ver, por ejemplo, con tormentas iónicas, radiaciones y vientos propulsores. (Hay bastante desorden cuando una estrella viene al mundo). Para averiguar lo que sucede en realidad, quizá haya que esperar a conseguir unos testimonios más fiables sobre las primeras fases de la vida de una estrella. El nacimiento de las estrellas es imposible de ver observando la zona del espectro electromagnético al que tiene acceso el ojo humano. Pero los fetos de estrellas envían también otras señales: por ejemplo, radiación de infrarrojos o microondas. Recientemente se ha conseguido crear aparatos con los que esas señales pueden ser examinadas con gran detalle y precisión. Pronto tendremos fotos ultrasónicas de las estrellas recién nacidas y podremos ir por ahí enseñándoselas a todo el mundo.



Capítulo 35

Sonidos desagradables

Que yo, por miedo al ruido, dejara de tirar de la cadena del lavabo echó a perder nuestro matrimonio.

JOCHEN SCHMIDT, «Mi sensibilidad acústica», en *Mis principales funciones corporales*

En el mundo hay muchos ruidos desagradables; la mayoría de emisoras de radio no emiten otra cosa en todo el día. Y si sobre programas de radio hay opiniones para todos los gustos, existe unanimidad sobre arañar un plato de porcelana con un tenedor o sobre el rechinar de la tiza en la pizarra: determinados sonidos le ponen la piel de gallina a casi todo el mundo. El chirrido al frotar dos pedazos de porexpán o dos globos y el zumbir del taladro de dentista pertenecen también a esa categoría de sonidos. Pero ¿cuál es la causa? Los seres humanos son capaces de percibir sonidos de hasta veinte kilohercios y hasta hace poco se sospechaba que eran las frecuencias más altas las que producían esa repulsión: se trataría de una reacción de autodefensa porque esas altas frecuencias podrían dañar el oído de forma permanente. Tal como descubrieron Lynn Halpern, Randy Blake y Jim Hillenbrand en 1986, en uno de los pocos estudios sobre el tema, esos sonidos desagradables no se vuelven más soportables ni siquiera filtrando las altas frecuencias. De hecho, al parecer son las frecuencias bajas y medias, de entre tres y seis kilohercios, las que nos ponen la piel de gallina. El sonido más desagradable del experimento para todos fue el que produjo arañar una pizarra con una herramienta de jardín «Trae Value Pacemkaer». En 2006, y con veinte años de retraso, los tres investigadores recibieron el «Ig-Nobel Prize», conocido también como el premio anti-Nobel. Así pues, la reacción no sirve para proteger el oído. Halpern, Blake y Hillenbrand plantearon en sus estudios la pregunta de si esos ruidos podían recordarle al oído humano los gritos de alarma de los primates o de algún depredador y si, por lo

tanto, podía tratarse de una reacción innata. Esta hipótesis, sin embargo, no se ve sustentada por un estudio que el MIT realizó en 2004 con titíes (*Saguinus oedipus*), que reaccionaban con relativa indiferencia tanto si se les reproducía un sonido blanco como el de un arañazo en una pizarra. Blake defiende aún hoy la teoría de los primates, mientras que Hillenbrand, en cambio, se ha desmarcado de ella. En su opinión, lo que provoca aversión no es tanto el sonido en sí, como la visión de lo que genera dicho sonido. Hay varios experimentos que sustentan esa teoría, entre ellos los realizados en 1987 por el estudiante de psicología Philip Hodgson, de la Universidad de York. Hodgson había demostrado ya que las frecuencias próximas a los 2,8 kilohercios resultan especialmente desagradables para el oído humano. A continuación intentó resolver el enigma, por lo que tomó un grupo de ensayo formado por personas sordas de nacimiento y les preguntó hasta qué punto les resultaba desagradable ver cómo alguien arañaba una pizarra. El 83 por ciento de los sujetos declararon sentir cierto malestar y a la pregunta de en qué parte del cuerpo se localizaba dicho malestar, el 72 por ciento respondió que en los dientes. Sin embargo, Hodgson no fue capaz de encontrar una explicación al fenómeno.

El que quiera poner a prueba la propia sensibilidad, puede hacerlo en la página web «Bad Vibes» del profesor de acústica Trevor Cox (www.sound101.org), donde podrá escuchar y puntuar los 30 sonidos más desagradables. El propio Cox, sin embargo, asegura que a él todos esos sonidos lo dejan frío. Asimismo, no cree en la teoría de los primates, pero con sus datos recogidos a nivel internacional pretendía estudiar la posible existencia de divergencias regionales en la percepción de sonidos. Los primeros resultados publicados por Cox no incluyen datos separados por países, pero tras recibir 1,1 millones de votos, el número uno de los sonidos desagradables es el sonido de vómito, seguido por un micrófono que se acopla, unos gritos de bebé superpuestos y un chillido estridente. En la mayoría de casos, las mujeres se mostraron más sensibles que los hombres. Los resultados del estudio, según Cox, no se corresponden a una reacción de repugnancia pura y tampoco sirven para corroborar la teoría del grito de alarma. Por desgracia, no podemos esperar a que Cox publique más resultados, pues en lo venidero ha decidido dedicarse a buscar el sonido más agradable del mundo. Honestamente, no se le puede culpar de ello.



Capítulo 36

Suceso de Tunguska

—*Los cometas siempre provocan catástrofes* —dijo solemnemente el Snork.

— ¿Qué es una catástrofe? —quiso saber Schnüferl.

—Oh, cualquier suerte de calamidad —respondió el Snork—. Plagas de langostas, terremotos, tifones, huracanes y demás.

—Ruido, en otras palabras —dijo el Hemul—. No hay forma de que descanse uno.

TOVE JANSSON, *La llegada del cometa*

El 30 de junio de 1908, poco después de las siete de la mañana, hubo en Siberia una explosión «que podría reproducirse vagamente como un zabum» (Robert Gernhardt). O tal vez fueron más de una, y aquí es donde comienzan los problemas, pues según muchos testigos auriculares se oyeron hasta veinte detonaciones. Lo único que no admite discusión es que cerca de un afluente del Yeniséi que lleva el encantador nombre de Tunguska Pedregoso explotó algo, probablemente en el cielo. La explosión (según se ha podido calcular décadas más tarde a partir de los indicios disponibles) tuvo una fuerza explosiva de entre diez y veinte megatones de TNT, que corresponde a cinco o diez veces todas las bombas convencionales lanzadas durante la Segunda Guerra Mundial o, para convertirlo a una medida más al uso, a bastantes buscapíés. En el cielo se elevó una nube fungiforme, llovió barro y el temblor se registró en las estaciones sismográficas de Irkutsk, Taschkent, Tiflis e incluso en la de Jena, situada a más de cinco mil kilómetros de distancia. La onda expansiva fue detectada por diversos aparatos de medición en Inglaterra. El observatorio de Irkutsk, situado a 970 kilómetros de distancia, percibió alteraciones en el campo magnético terrestre similares a los que se producen tras la explosión de una bomba atómica.

Durante las 72 horas siguientes, se pudieron apreciar en toda Europa unas puestas de sol de colores inusitados y hubo noches de gran claridad; en la ciudad escocesa de Saint Andrews se podía jugar al golf a las dos y media de la madrugada. Por desgracia, esas hermosas puestas de sol habían llegado acompañadas de otros fenómenos, como nubes a gran altura, alteraciones atmosféricas y llamativos halos solares, visibles desde varios días antes de la explosión. Esos fenómenos fueron apreciables en una zona que abarcaba desde el Yeniséi en el este hasta la costa atlántica en el oeste, y que llegaba a la altura de Burdeos en el sur.

Durante las primeras décadas de investigación, todos los datos correspondían a las cinco expediciones realizadas entre 1921 y 1939 por el mineralogista ruso Leonid Kulik. La primera de esas expediciones fue una expedición genérica para la observación de meteoritos, al inicio de la cual, en la estación de tren de San Petersburgo, alguien le entregó a Kulik una hoja de calendario de 1910 en la que se informaba de un misterioso meteorito que había caído en Tomsk en 1908. La fecha resultó ser falsa, pero puso a Kulik sobre la pista del suceso de Tunguska. Esta primera expedición se quedó sin fondos antes de que Kulik lograra penetrar en la región de la explosión; sin embargo, y gracias a un cuestionario publicado en la prensa, logró tomar declaración a un gran número de testigos oculares.

La región de Tunguska no es precisamente un paraíso de vacaciones, sino una zona remota, infranqueable y plagada de mosquitos, demasiado calurosa en verano y demasiado fría en invierno. No es de extrañar que la densidad de población de la zona sea francamente baja. Por un lado se trata de una circunstancia afortunada, pues minimizó los daños personales: alguien se rompió un brazo, hubo varios cardenales y un anciano se murió del susto. Durante muchos años se intentaría encontrar una explicación plausible a la relación entre las dimensiones de la catástrofe y el reducido número de heridos. Por otro lado, hoy en día sabríamos mucho más sobre el llamado suceso de Tunguska si los testigos oculares no se hubieran tomado hace varias décadas; ni la policía de Berlín tarda tanto en llegar al lugar del suceso.

De las aproximadamente novecientas declaraciones de testigos oculares que aparecieron en la prensa rusa tras la catástrofe y que se fueron reuniendo durante los años siguientes mediante encuestas a la población, se desprende que en la

población más cercana al suceso, Vanavara, situada a 65 kilómetros, se divisó un brillo deslumbrante y las ventanas estallaron. Los encuestados mencionaron la sensación de calor sobre la piel, truenos y una onda expansiva. El sonido de repetidas explosiones se oyó en pueblos situados a más de 1200 kilómetros. En las primeras décadas de investigación, y partiendo siempre de las declaraciones de los testigos oculares, se dio por cosa cierta que la deflagración luminosa se había desplazado del sur al norte hasta que, en la década de 1960 y después de un largo tira y afloja, los expertos decidieron que el fenómeno lumínico había descrito una trayectoria este-suroeste-oestenoroeste. Sólo a principios de la década de 1980 apareció una nueva recopilación de testigos oculares de otras regiones que terminó de complicarlo todo: los habitantes de la región adyacente al río Angara, por un lado, y los de la Baja Tunguska, por el otro, describían el suceso y su desarrollo de formas tan distintas que no podía tratarse de un mismo acontecimiento. Además, la trayectoria reconstruida a partir de las descripciones de Angara no encajaba con el patrón de los árboles caídos. Finalmente, los testigos de Tunguska coincidían en afirmar que el suceso había tenido lugar por la tarde, mientras que los de Angara aseguraban que se había producido a primera hora de la mañana. No se trataba, pues, de pequeñas desavenencias, sino de dos grupos de testigos cuyas declaraciones no había forma de hacer encajar. (Sus motivos tienen los abogados para decir: «Es preferible no tener testigos a tener un testigo ocular»). Hasta hoy, la mayoría de investigadores se sirven de las declaraciones que se ajustan a sus teorías y ponen en duda la fiabilidad de las demás.

En 1927, tras una segunda expedición marcada por meses de penurias y escorbuto, Kulik llegó finalmente al lugar de la catástrofe. Diecinueve años después del suceso, en una superficie de más de 2000 kilómetros cuadrados había unos 60 millones de árboles sin ramas, sin corteza y arrasados como cerillas. Los árboles derribados indicaban dónde se encontraba el epicentro de la explosión, y en el centro aún había algunos de pie, pelados como postes de telégrafo. Los árboles de muchas zonas habían quedado carbonizados por culpa del incendio provocado por la explosión. Además, se podían apreciar ondas terrestres y numerosos agujeros en forma de cráter entre diez y cincuenta metros de diámetro. Por otro lado, y a pesar del tesón con que los buscó, Kulik no logró en sus siguientes expediciones encontrar ni el

cráter original ni los restos de ningún cuerpo celeste. En la década de 1960, los expertos concluyeron que la explosión debía de haber tenido lugar en el aire, encima del epicentro. A día de hoy, ése es uno de los pocos puntos compartidos por la mayoría de científicos, que no logran ponerse de acuerdo ni siquiera en el número de explosiones.

En un primer momento, el hipotético cuerpo celeste fue bautizado como el «meteorito Filimonovo» en honor al cruce ferroviario de Filimonovo, situado a 600 kilómetros del lugar del suceso y que aparecía mencionado en la hoja de calendario de Kulik. El astrofísico norteamericano Harlow Shapley fue el primero, en 1930, en sugerir que el episodio podía deberse a un cometa, es decir, no una roca sino un sucio pedazo de hielo (al que se le supone un diámetro de unos 40 metros) cubierto por una faja nebulosa. En 1934 dos astrónomos, el británico Francis Whipple y el ruso Igor Astapovich, asumieron también esa teoría, que a partir de aquel momento fue defendida y desarrollada fundamentalmente por los rusos, mientras que los científicos norteamericanos continuaban partiendo de la base de que el fenómeno había sido causado por un asteroide. Hay muchas versiones de asteroides; en relación con el suceso de Tunguska, se suele buscar un bloque de piedra de entre treinta y doscientos metros de diámetro. Hasta la década de 1990, estas dos religiones principales apenas tenían contacto entre sí, lo que a efectos prácticos quería decir que las publicaciones rusas no estaban disponibles en inglés y viceversa. Así, la teoría del asteroide ha logrado en los últimos años una cierta cuota de mercado, aunque ninguna de las dos ha logrado imponerse.

A favor de la teoría del cometa juega el hecho de que, hipotéticamente, éste se habría desintegrado en la atmósfera sin dejar rastro, pues a pesar de la minuciosa búsqueda aún no se ha hallado ni un solo fragmento de asteroide. Incluso meteoritos mucho más pequeños dejan normalmente restos, aunque sólo sea polvo. Asimismo, la cola que los testigos dijeron apreciar alejándose del Sol y que, por su contenido acuoso, podría ser la responsable de las inusitadas puestas de sol del año 1908 es algo exclusivo de los cometas. Un asteroide sería demasiado seco para formar las grandes nubes altas que desintegran la luz del Sol y provocan noches claras.

Sin embargo, lo que se conoce de la trayectoria del objeto parece reforzar la teoría del asteroide, pues es más propio de éstos. Así, en el año 2001 científicos italianos determinaron que de las 886 trayectorias posibles del cuerpo celeste, el 83 por ciento corresponden a trayectorias de asteroide y sólo un 17 por ciento a trayectorias de cometa. Además, desde la colisión del cometa Shoemaker-Levy con Júpiter (que se resolvió a favor de Júpiter) sabemos que para que se produzca una gran explosión, el cometa debe tener una masa de más de cien millones de toneladas.

A partir de su velocidad y de la altura de la explosión, se ha calculado que el objeto de Tunguska debía de tener una masa de aproximadamente cien mil toneladas. Contrariamente a lo que sucedería con un asteroide, un cometa de dimensiones tan reducidas no soportaría la presión que supone la entrada en la atmósfera. Si el objeto de Tunguska hubiera sido significativamente mayor, la explosión (según afirma el astrónomo estadounidense Zdenek Sekanina) habría oscurecido el sol y habría dejado un invierno nuclear tras de sí. Las consecuencias habrían sido tan dramáticas que hoy no habría ninguna discusión sobre el suceso de Tunguska porque no quedaría nadie vivo.

A eso hay que añadirle el argumento estadístico de que en el universo hay entre diez y cien veces más asteroides del tamaño adecuado que cometas. Finalmente, los cometas vuelan demasiado despacio para provocar una explosión de esa magnitud.

Naturalmente, las contradicciones de ambas teorías pueden resolverse proponiendo una tercera teoría. El astrofísico alemán Wolfgang Kundt lanzó en 1999 una nueva hipótesis según la cual la explosión de Tunguska fue provocada por diez millones de toneladas de metano que salieron de una grieta del suelo y ardieron. Está demostrado que eso mismo sucede de vez en cuando, sólo que a menor escala. Kundt aportó veinte argumentos para su teoría, los más importantes de los cuales son: Tunguska se encuentra en el punto de intersección de tres pliegues tectónicos y en el centro del antiguo cráter de un volcán. El patrón de los árboles derribados indica que se tuvieron que producir cinco o más explosiones a ras de suelo, algo que también encajaría en las declaraciones de testigos que hablan de varias deflagraciones. Las noches claras después del fenómeno quedan explicadas de

forma igualmente elegante por el hecho de que la mayor parte de los elementos de los gases volcánicos son lo bastante ligeros como para elevarse a 500 kilómetros de altura y desintegrar la luz, lo mismo que sucedió en 1883 con la erupción del Cracatoa. Además, en la región hay tantos yacimientos de gas natural como rocas de origen volcánico. El calor que los habitantes de Vanavara experimentaron sobre el rostro también se explica mejor si el cielo se llenó de gas ardiente que con otras teorías. El último argumento es de naturaleza puramente estadística: sólo el tres por ciento de los cráteres visibles en la Tierra tienen como origen un impacto de un cuerpo celeste. El 97 por ciento restante tienen un origen volcánico.

Contra la teoría de Kundt suele argumentarse que no existen casos similares, aunque en realidad eso no descarta nada. No sólo eso, sino que nadie sabe a ciencia cierta si un hipotético caso similar tendría de nuevo la delicadeza de devastar una zona casi totalmente inhabitada.

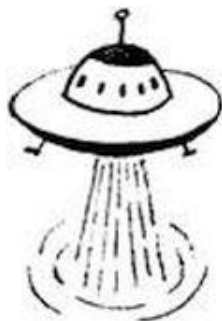
La teoría de Wolfgang Kundt parte de la misma base que el trabajo del científico ruso Andrey Oljovatov, que fue el primero en sugerir que la explosión podía ser de origen geológico. Su hipótesis se basa en la combinación de procesos aún por determinar en el suelo y en la atmósfera, una especie de rayo globular.

Científicos y aficionados de todo el mundo complementan estas tres grandes teorías con extravagantes adornos y modelos alternativos. Ya en 1941, el experto en meteoritos norteamericano Lincoln La Paz tuvo la idea de que se pudiera tratar del impacto de un fragmento de antimateria procedente del universo, y en 1965 fueron Willard Lobby, Clyde Cowan y C. R. Alturi quienes volvieron a sacar la misma teoría a la luz pública. Sin embargo, la antimateria se habría destruido al entrar en la atmósfera (y no poco antes de llegar al suelo), pues tiene una reacción sumamente alérgica al entrar en contacto con materia normal. El físico australiano Robert Foot, por su parte, esgrime la tesis de que éste como otros casos singulares se deben al impacto de una «materia espejo» procedente del espacio, una materia hipotética cargada con partículas elementales diametralmente opuestas (y no sólo distintas como en el caso de la antimateria) y cuyas propiedades físicas son, por lo tanto, distintas a las de la materia corriente. La hipótesis de que se trata del cráter que dejó una ballena caída desde el universo tiene un pero: que no se han encontrado aún restos de ballena en la zona. (Aunque sus defensores replican que tampoco los

ha buscado nadie). Los físicos teóricos A. A. Jackson IV y Michael P. Ryan Jr. sugirieron en 1973 que podría haberse tratado de un diminuto agujero negro que hubiera atravesado la Tierra y que hubiera vuelto a salir por el Atlántico Norte. Por lo que sabemos hasta hoy de los agujeros negros no sería completamente imposible, aunque falta encontrar un agujero de salida verosímil. Por otro lado, los agujeros negros tienen la gran ventaja de que son invisibles y está demostrado que la invisibilidad sirve para explicarlo todo, desde Dios hasta las ondas de radio: «¡Un perro invisible se me ha comido los deberes!». En la periferia del espectro se encuentran las teorías sobre ovnis y la conjetura, hermosa pero algo peregrina, de que el genial inventor Nikolai Tesla se equivocó de palanca durante un experimento de telecomunicación de energía.

Aunque aún hoy no resulta fácil llegar al lugar de la explosión (la estación de tren más cercana.

Se encuentra a 600 kilómetros), casi cada año se organiza alguna expedición a Tunguska. Aún hay la esperanza de que gracias a nuevas ideas o nuevas tecnologías se puedan obtener nuevos datos que permitan resolver la enigmática explosión sin lugar a dudas. Durante las últimas décadas, y tras una búsqueda minuciosa, se han encontrado nanopartículas de elementos raros (especialmente de iridio) en la resina de los árboles de Tunguska, en el suelo y en los pantanos de los alrededores, y también en las capas de hielo de la Antártida correspondientes a ese año, pero no ha sido posible establecer su relación con el suceso de Tunguska. De hecho, se puede encontrar polvo cósmico parecido en diferentes cantidades casi en cualquier parte de la tierra. Por si eso fuera poco, el iridio puede proceder tanto del universo exterior como del interior de la tierra, por lo que encajaría con cualquiera de las hipótesis. Pero es posible que aún haya indicios escondidos en alguna parte. Cuando finalmente alguien los encuentre y resuelva el enigma del suceso de Tunguska, se podrá preparar el factor desencadenante en un formato cómodo y portátil, que se pueda llevar siempre encima por si, un día, estando en Escocia, a alguien le entran ganas de jugar al golf a las dos y media de la madrugada.



Capítulo 37

Sueño

La conciencia, ese fatigoso estado entre siesta y siesta.

ANÓNIMO

Los mamíferos duermen, los pájaros duermen y duermen también los reptiles. También los anfibios y los peces están a veces menos atentos de lo habitual, y hace pocos años se descubrió que incluso los insectos duermen... aunque es cierto que viendo los mosquitos por la noche uno no lo diría nunca. El pequeño ratón de abazones duerme más de veinte horas al día, mientras que la jirafa sólo duerme dos. Muchos animales, como los gorilas, duermen muchas horas de un tirón, mientras que otros, como las vacas, lo hacen en pequeñas cabezaditas de varios minutos. Unos duermen de noche, otros lo hacen de día, y los animales crepusculares como los murciélagos tienen dos fases de vigilia.

En el ser humano, los hábitos del sueño se desarrollan progresivamente. Un bebé (por mucho que se quejen los padres primerizos) duerme siempre dieciséis horas repartidas a lo largo del día; los adultos se quedan en las ocho horas de media. La duración del sueño puede variar mucho de un individuo a otro; mientras que una persona necesita dormir cuatro horas, otra puede necesitar diez. Pero ¿qué es lo que lleva a humanos y animales a esta extraña conducta? ¿Por qué algunos animales son capaces de lograr lo que sea que se logra con el sueño en muchas menos horas que otros? ¿Cómo puede ser que la necesidad de dormir de todos los mamíferos, incluidos los humanos, disminuya con la edad? ¿A quién, aparte de a los fabricantes de camas, beneficia el sueño?

La investigación del sueño es una disciplina relativamente joven. Surgió a finales de la década de 1930 gracias al descubrimiento del electroencefalograma, que permite observar el cerebro durante el sueño. Pronto se descubrió que, contrariamente a lo que se creía, dormir no significa que se apaguen las luces del cerebro, sino que en éste tiene lugar un proceso que hasta el día de hoy aún no se ha acabado de comprender. Hubo que esperar hasta la década de 1950 para, con la ayuda de la

llamada polisomnografía, una combinación de diversos procedimientos de medición, poder identificar de forma fiable los diversos estadios y niveles del sueño. Durante el sueño, las neuronas comienzan a latir al compás, a un ritmo que (utilizando una gorra llena de cables no muy elegante) se puede medir; a partir del patrón resultante de esa medición, el sueño se ha dividido en cinco estadios. El estadio I corresponde a un ligero sueño inicial, el estadio II se extiende durante la mayor parte de la noche, mientras que en los estadios III y IV tiene lugar el sueño profundo. La quinta fase, la llamada fase REM, se distingue claramente de las otras cuatro: el cerebro está tan activo como durante la vigilia, pero la musculatura está completamente relajada.

Si durante una prueba se despierta al individuo durante la fase REM, éste casi siempre recordará qué estaba soñando. La fase REM se ha observado prácticamente en todos los mamíferos. Sin embargo, en el caso de los mamíferos los estadios del sueño están mucho más diferenciados que en el resto de animales; la teoría es que los cerebros que en el estado de vigilia deben analizar más cosas y hacerlo de forma más compleja, también necesitan un sueño más complejo. Para los pequeños animales, el ciclo de sueño completo es mucho más corto. Las musarañas, por ejemplo, despachan las cinco fases de sueño en tan sólo ocho minutos, mientras que un elefante necesita casi dos horas. Pero ni siquiera las musarañas saben por qué es así.

Determinar qué sucede durante el sueño no es nada fácil, por lo que una opción alternativa es estudiar qué sucede cuando un individuo no duerme. Para ello se coloca una rata en una plataforma rodeada de agua, tan pequeña que la rata se moja en cuanto se relaja para dormir. A las ratas les gusta tan poco mojarse que, en una situación como ésta, no logran dormirse. Pasadas entre dos y tres semanas de falta de sueño, la rata muere por motivos poco claros. No parece que la muerte pueda deberse a motivos evidentes, como una infección o un fallo cardíaco. Sin embargo, los críticos con este tipo de pruebas aducen que las consecuencias de la falta de sueño (ya sea para la simple supervivencia o para el correcto funcionamiento metabólico o cerebral) no quedan claramente diferenciadas de las consecuencias que el estrés provocado por esa situación excepcional conlleva para la rata. La falta de sueño, argumentan, no es el contrario que el sueño, sino un

estado anormal del que no se pueden extraer demasiadas cosas sobre la función del sueño.

El concepto más cercano a esta función del sueño es la llamada teoría del descanso y la reparación: cuando estamos agotados tenemos que dormir y, al despertar, nos sentimos menos cansados, lo que significa que durante ese tiempo pasa algo que revierte el desgaste del cuerpo.

Sin embargo, no puede ser tan sencillo. Por una parte, si esa hipótesis fuera cierta, tras su jornada de veintidós horas la jirafa tendría que dormir un buen rato y, sin embargo, no es el caso: cuanto más tiempo pasa despierto un animal, más breve es la fase de sueño que éste necesita, pues los animales (a diferencia de lo que, sin ir más lejos, hacen los programadores) se ciñen estrictamente a la jornada de veinticuatro horas. Por otra parte, apenas hay procesos en el cuerpo de los que sepamos con seguridad que se reviertan durante el sueño. En las fases III y IV del sueño se distribuyen las hormonas del crecimiento multiplicadas y hay indicios para pensar que existe una relación (aún no esclarecida) entre el sueño y la regulación del sistema inmunológico. Sin embargo, hasta hoy no se ha podido documentar ningún proceso de reparación.

El neuroendocrinólogo Jan Born apunta también que para descansar no tiene por qué ser necesario desconectar la conciencia. En primer lugar, este estado entraña para todo ser vivo el peligro de verse comido por un depredador, y en segundo lugar, al dormir (y especialmente en la fase REM) el cerebro no está ocioso sino, al contrario, muy activo. Born se adscribe a la teoría de la memoria según la cual durante el sueño se fijan los recuerdos. Existen numerosos experimentos en los que, tras un episodio de privación de sueño, los sujetos o animales sujetos a experimentación obtienen peores resultados en las pruebas de memoria. Esos experimentos permiten deducir claramente que la privación de sueño es perjudicial para las funciones de la memoria; sin embargo, eso no significa necesariamente que durante el sueño tengan lugar procesos importantes para la memoria. Muchos investigadores conjeturan que el conocimiento no se almacena directamente en la memoria a largo plazo, sino que primero pasa por una memoria intermedia y luego, mientras uno duerme, se copian en el disco duro, por así decirlo. Si para este proceso fuera importante bloquear la entrada de nueva información, tendría

realmente sentido impedir que el cuerpo pudiera observar, olfatear o caminar. Por desgracia, demostrar esa teoría no es tarea fácil. Para ello resultaría muy útil saber más cosas sobre cómo funciona realmente la memoria.

En 2004, los colegas de Born, Ulrich Wagner y Steffen Gais, lograron probar que el sueño favorece el proceso de memoria: sus sujetos de experimentación debían resolver un problema para el que había una solución laboriosa y otra sencilla. En comparación con los que se quedaron despiertos, más del doble de los sujetos que pudieron dormir entre dos pruebas llegaron a la solución sencilla. Así, es posible que dormir en lugar de trabajar suponga incluso un ahorro de tiempo. Es una lástima que la investigación del sueño no dedique más esfuerzos a una tarea tan importante como la de surtirnos de justificaciones para poder dormir en la oficina.

Entretanto, los fundamentos de la hipótesis de la memoria han sido confirmados por diversos laboratorios y mediante diferentes métodos, pero aun así sigue siendo objeto de discusión. Sus principales detractores, Jerome Siegel y Robert Vertes, aseguran que una hipótesis que ha sido objeto de estudios tan frecuentes debería estar corroborada por pruebas más concluyentes. Para explicar los resultados contradictorios se le ha echado agua al vino a la tesis original de la inutilidad del sueño, considerando sólo determinadas formas de memoria (la memoria del movimiento, por ejemplo) en función del resultado de cada experimento y excluyendo el resto.

El caso (afortunadamente infrecuente) de algunas personas que, debido a una determinada lesión cerebral, no experimentan la fase REM y, sin embargo, son capaces de llevar una vida normal demuestra que esa fase del sueño no es indispensable para recordar las cosas. Asimismo, hay numerosos pacientes que, como efecto secundario a la toma de medicamentos contra la depresión, se ven privados total o parcialmente de la fase REM y, sin embargo, no suelen sufrir problemas de memoria dignos de mención. Hace unos años se suponía que los sueños desempeñaban una función importante y que sólo tenían lugar durante la fase REM. Sin embargo, aún hoy nadie sabe cuál es la finalidad los sueños. La tesis de Freud según la cual en los sueños se experimentan emociones y deseos reprimidos ha pasado tanto de moda como la suposición de que los sueños son apenas subproductos insignificantes de la actividad cerebral durante el sueño.

A grandes rasgos, las teorías más recientes suponen que los sueños deben de cumplir alguna función, aunque se desconoce cuál puede ser. Tal vez sirvan tan sólo para que no nos aburramos mientras dormimos, como las películas en los vuelos de larga distancia.

Pero volvamos a la función del sueño: Siebel la compara con la de la hibernación y señala que ésta no la cuestiona casi nadie. La hibernación sirve para apartar de la circulación durante un tiempo a unos animales que, de todos modos, tampoco podrían hacer nada porque en el exterior está todo cubierto de nieve. (La hibernación, por cierto, no sustituye el sueño normal. La mayoría de animales invernantes, por extraño que parezca, deben despertar, calentarse un poco y gozar de un sueño «normal» de vez en cuando). Por especies, los animales carnívoros son los que duermen más y los herbívoros los que menos, mientras que los omnívoros (los humanos entre ellos) se encuentran a medio camino. A un animal que debe pasar el día pastando y protegiéndose de los depredadores no le queda demasiado tiempo para dormir, mientras que, después de devorar un antílope, un león puede permitirse el lujo de pasar el resto del día dormitando. En nuestro caso, como necesitamos menos de veinticuatro horas para realizar las funciones imprescindibles, es razonable que dejemos nuestro cuerpo aparcado durante la parte del día en que, por otra parte, experimentaría más pérdidas que ganancias. A los pequeños animales, la superficie de cuyo cuerpo es relativamente grande en comparación con su peso, yacer en un lugar caliente del nido debe de suponerles probablemente un ahorro de energía. Esta tesis se ve reforzada por el hecho de que, en el caso de los mamíferos marinos, la duración del sueño no disminuye, sino que aumenta a lo largo de la vida: en el mar no hay ni rincones resguardados donde los animales puedan dormir libres de peligros mientras son jóvenes, ni abismos en los que tropiecen por culpa de la oscuridad.

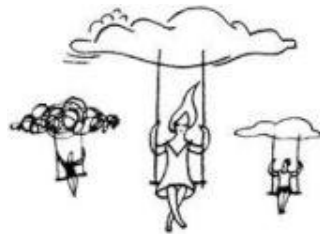
Existe una hipótesis parecida según la cual la duración del sueño está determinada genéticamente para garantizar el equilibrio ecológico. Así, los animales de rapiña duermen más que sus presas para evitar que se agote su coto de caza. También en este caso, pues, la función principal del sueño sería prevenir comportamientos contraproducentes. No resulta difícil imaginar cómo el departamento de programación de la evolución decidió incorporar esa característica en lugar de

introducir una función tan compleja como la razón: «Apaguemos el animal un rato y así, por lo menos, no hará tonterías».

Sin embargo, el motivo actual del sueño no tiene por qué ser necesariamente el mismo por el que éste se desarrolló de buen principio. Tal vez inicialmente el sueño cumplía con un propósito determinado y posteriormente, con el curso de la evolución, se le han ido asignando otras funciones aprovechando que, de todos modos, el cuerpo pasaba un buen rato inactivo. Sin embargo, existen muchos indicios que apuntan a que tiene que haber un buen motivo para el sueño.

No en vano, éste ocupa mucho tiempo de vida, es sorprendentemente similar en todas las especies y las ratas, por lo menos, mueren cuando se ven privadas de él. Si hubiera que darle un nombre a este argumento, sería el del investigador del sueño James Krueger, candidato bastante claro al premio Nobel.

Algunos investigadores objetan a todas estas hipótesis que la pregunta «¿Por qué dormimos?» está mal formulada: lo que deberíamos preguntarnos es por qué de vez en cuando nos despertamos. El sueño, aseguran, es el estado natural que compartimos con muchos animales simples y con las células de nuestro propio cuerpo. De vez en cuando lo interrumpimos para tomar algo de comida de la nevera y para asegurar la conservación de la especie. A efectos prácticos, es mucho más fácil responder a la pregunta de por qué nos despertamos que a la de por qué dormimos: en general es porque suena el despertador. Es una pena que eso no nos valga un premio Nobel.



Capítulo 38

Tamaño de los animales

Le dijo el ternero a la vaca: yo soy apenas la mitad de grande que tú. Te llevo a las ubres y poco más.

ARNOLD HAU

Algunos animales tienen cuernos y barbas, otros tienen aletas, alas, tentáculos, las patas peludas o no tienen patas, pero hay algo que todos los animales tienen: una estatura. Y como, aun cuando se trate de animales muertos o incluso fosilizados, ésta es fácil de medir y comparar con la de otros animales, se trata de un tema que se estudia a menudo. Hasta el momento se han hecho ya algunos descubrimientos: por ejemplo, los animales pequeños mueren antes, producen antes descendencia suficiente y comen (en relación con su tamaño) más que los animales grandes.

También hay muchas más especies de animales pequeños que grandes: si consideramos los mamíferos, por ejemplo, el 75 por ciento pesan menos de un kilogramo.

Una de las principales dificultades radica en el hecho de que, hasta la fecha, el cotejo y medición de la fauna ha ignorado muchas especies animales. Actualmente se conocen y se han descrito aproximadamente 1,5 millones de animales y plantas, pero la cartografía completa podría ascender hasta los cinco o, según la opinión de muchos biólogos, hasta los 50 millones. Y estas entre 3,5 y 48,5 millones de especies no descubiertas no son todas necesariamente pequeñas: cada año, para perplejidad de los biólogos, asoman la cabeza un puñado de nuevos grandes mamíferos.

El saola o buey del Vietnam, animal fácilmente distinguible a simple vista, fue descrito por primera vez en 1993. Por algún motivo, se tarda más tiempo en descubrir especies pequeñas que especies grandes: de entrada, anteriormente había un mayor interés en encontrar animales grandes y vistosos que, una vez disecados, provocarían la envidia de los vecinos. En segundo lugar, los animales grandes ocupan más espacio, de modo que uno se tropieza más fácilmente con

ellos. En tercer lugar, para no confundir un animal pequeño con alguno de sus parientes se necesitan técnicas mejores y más modernas, pues en los animales pequeños las diferencias no son tan marcadas. (Una especie, según la definición al uso, es aquella cuyos individuos que no se reproducen mejor con los de una especie vecina). Todo eso significa que es muy probable que haya no sólo *más* sino *muchas más* especies pequeñas que grandes.

Durante varios siglos de medición del reino animal se han observado y descrito determinadas regularidades. Al parecer, y si se observa durante un largo período de tiempo, normalmente el tamaño medio de los animales de una especie tiende a aumentar. Esta circunstancia se conoce como la «Ley de Cope» en honor al paleontólogo Edward Drinker Cope. Aunque discutida, la ley ha ganado adeptos en los últimos años. Sin embargo, no existe una explicación para ese patrón de crecimiento, si bien es probable que se deba a una suma de factores diversos; así, por ejemplo, las hembras mayores ponen más huevos, los machos de mayor tamaño tienen más posibilidades de éxito durante la reproducción y, en definitiva, ser grande suele ser una buena estrategia para asegurarse la supervivencia.

Pero si un gran tamaño tiene tantas ventajas, ¿por qué muchas especies animales no presentan una tendencia continua al crecimiento? «¿Por qué la ley de Cope es tan poco regular?», se preguntan los científicos, perplejos. ¿Y por qué algunos de los grandes animales del pasado, como el wombat gigante de Australia, con sus tres toneladas de peso, o el castor gigante de América, de tres metros y largo, se han extinguido? Si la ley de Cope es cierta, tiene que haber por lo menos un segundo mecanismo que dé sentido al aumento del tamaño de las especies. Los paleobiólogos Blaire van Valkenburgh, Xiaoming Wang y John Damuth propusieron en 2004 un modelo aplicable, por lo menos, a muchos ámbitos parciales: en los últimos 50 millones de años han surgido varias grandes especies carnívoras que se han extinguido sin que nadie sepa exactamente por qué. Según Van Valkenburgh, Wang y Damuth, los animales con un mayor crecimiento se especializaron en una alimentación exclusivamente carnívora, tal como se puede deducir a partir de sus dientes fosilizados. Sin embargo, esa especialización fue su perdición cada vez que las condiciones ambientales cambiaban, mientras que los omnívoros se mostraron mucho más capaces de adaptarse a la nueva situación. También resulta fácil

imaginar que para cada animal individual sea una ventaja sobresalir por encima de su congéneres, y que especies enteras pueden beneficiarse de un cierto crecimiento, pero que a largo plazo las especies medianas tienen más posibilidades. Además, el que se queda pequeño se beneficia también de que todo va más rápido: crecer requiere tiempo, por lo que aumenta la probabilidad de morir víctima de los parásitos o los depredadores antes de alcanzar la edad reproductiva.

Sea como fuere, los animales no pueden crecer a discreción ni siquiera en época de vacas gordas, valga la contradicción. Para los insectos, que respiran a través de tráqueas, el porcentaje de oxígeno del aire es un factor limitador, pues a partir de determinado tamaño el cuerpo no recibe suficiente oxígeno; por ello hoy en día las libélulas, por ejemplo, no miden setenta centímetros como en el Carbonífero Superior, pues entonces la atmósfera terrestre contenía más oxígeno que ahora. Sin embargo, los animales sin tráqueas tienen otros problemas, pues un aumento de la masa corporal implica también un aumento de sus costes de vida. Pasar de presas pequeñas a presas grandes implica un crecimiento potencial del gasto energético por parte de los depredadores, pues ya no basta con acechar o recoger tranquilamente las presas, sino que hay cazarlas y acabar con ellas. Por ello, los zoólogos Chris Carbone, Amber Teacher y Marcus Rowcliffe calculan que los depredadores pueden tener un peso corporal máximo de 1100 kilos (por ejemplo, el mayor depredador de la tierra, el oso polar, pesa actualmente unos 500 kilos). La existencia de depredadores mucho mayores, conjeturan, sólo fue posible gracias a un metabolismo que economizaba enormemente el uso de energía. La eficiencia estimada del metabolismo de los grandes depredadores, que llegaron a pesar hasta nueve toneladas, era equiparable a la de un mamífero de una tonelada de peso. De ahí se desprende la pregunta de si existe un tamaño óptimo al que terminarían adaptándose todos los animales si no se vieran amenazados por depredadores y dispusieran siempre de alimento. Muchos expertos dan por supuesto que, de vivir en condiciones idílicas, los animales irían replegándose con el tiempo hasta alcanzar una masa corporal de un kilogramo o menos. Otros, en cambio, no creen que exista tal cosa como un tamaño óptimo.

Por irritante que resulte, el tamaño de un animal depende también del tamaño de la masa terrenal en la que viva. Eso no significa que en Luxemburgo vivan unas

ardillas diminutas, sino que las especies grandes suelen encogerse si se trasladan a una isla o si un istmo desaparece y se convierten de pronto en animales insulares. Así, hasta hace unos dos mil quinientos años había en muchas islas del Mediterráneo elefantes que medían apenas un metro y medio de alto. También se sabe que el mamut de Wrangel se originó tras la desaparición del brazo de tierra que unía Siberia y la isla de Wrangel: tras apenas quinientas generaciones, el animal se encogió hasta el práctico tamaño (por lo menos si tenemos en cuenta las dimensiones de un mamut convencional) de 1,80 metros. Un espacio vital pequeño se traduce en animales pequeños; bien pensado, podría ser que fuera así de sencillo. Sin embargo, como en la naturaleza no hay casi nada sencillo, muchas especies que se mudan a una isla tienden a aumentar de tamaño en comparación con lo que medían en tierra firme. Estos dos fenómenos quedan descritos en la Ley Foster, bautizada en honor al zoólogo J. Bristol Foster. Según observó Foster en la década de 1960, en una isla los animales pequeños aumentan de tamaño, mientras que los grandes se vuelven pequeños. También en este caso se trata de una afirmación que se cumple sólo a medias: los roedores, murciélagos, artiodáctilos, elefantes, zorros, mapaches, serpientes y lagartos suelen ser más pequeños en las islas que en tierra firme. En cambio, con los animales cavadores, las iguanas, las tortugas y los osos sucede lo contrario. Y el mayor reptil de la Tierra, el dragón de Komodo, ha logrado alcanzar un tamaño considerable viviendo en una isla. De momento no se sabe por qué eso es así.

Es probable que los animales herbívoros sufran influencias evolutivas distintas que los depredadores y que los animales que ya existen se encojan debido a la competencia y a la falta de alimento, mientras que las especies que llegan a un lugar se encuentran con un medio libre de competencia y por eso crecen. Existen versiones de la regla de la isla que se aplican a otros ámbitos vitales, como el fondo del mar. Una de esas versiones dice que los peces que viven en riachuelos son más pequeños que los que viven en masas de agua más grandes. El hecho de que en los arroyos no haya tiburones parece corroborar esa tesis, aunque no la prueba.

El hecho de que los animales, criaturas veleidosas donde las haya, varíen de tamaño debido a factores de lo más diverso no facilita en absoluto el trabajo de los biólogos. Los paleobiólogos Gene Hunt y Kaustuv Roy publicaron en 2006 un estudio

sobre una especie de molusco que, a lo largo de los últimos cuarenta millones de años, aumentó de tamaño cada vez que descendió la temperatura de su entorno, y se mantuvo inalterable cuando la temperatura se mantuvo estable. La razón que motivaba ese cambio de tamaño en los moluscos es probablemente la misma que hoy hace que los animales que viven en regiones frías sean mayores que sus parientes cercanos de regiones más cálidas. La Ley Bergmann, formulada en 1847 por el anatomista y fisiólogo Carl Bergmann, describe este fenómeno basado en la relación entre superficie corporal y volumen: a un animal grande y pesado le es más fácil mantener el calor que a uno de pequeño. Por ese mismo motivo, en las regiones extremadamente frías no hay animales pequeños de sangre caliente.

Muchos animales no se ciñen a la Ley Bergmann, algunos animales de sangre fría que no tendrían por qué seguirla, como las tortugas, sí lo hacen y otros, como los lagartos y las serpientes, disminuyen de tamaño con los descensos de temperatura. Por si eso fuera poco, el biólogo Wayne A. van Voorhies publicó en 1996 un estudio según el cual las células de uno de los animales preferidos por los biólogos de laboratorio, el nematodo (*Caenorhabditis elegans*), aumentaban un 33 por ciento de tamaño a una temperatura de diez grados en comparación con su tamaño a veinticinco grados. En 2005, investigadores checos demostraron que en diversas especies de gecónidos existía una relación directa entre el tamaño de los glóbulos rojos y el tamaño de la especie en cuestión. Así, muchos animales son grandes no porque tengan más células que los animales más pequeños, sino porque estas células son mayores.

En general no resulta sencillo observar los hábitos de crecimiento de los animales. Es bastante probable que éstos se vean afectados por diversas fuerzas simultáneas que afectan su tamaño.

También en el mundo de las cosas inanimadas se observan fenómenos similares: los barcos son cada vez más grandes, mientras que los teléfonos se vuelven cada vez más pequeños. ¿Por qué las cosas son así y no al revés? Seguramente la ciencia lo descubra pronto.



Capítulo 39

Tectónica de placas

Por favor, no intente detener la deriva continental usted solo.

Lema de camisetas

A la pregunta sobre cuáles han sido los éxitos de la ciencia en el siglo XX, se responde a menudo mencionando la mecánica cuántica, la teoría de la relatividad o los viajes espaciales, pero rara vez se habla de la tectónica de placas. Sin embargo, la tectónica de placas es la historia de la deriva de los continentes, la gran teoría unificadora que sirve para comprender la Tierra.

Explica de dónde vienen las cordilleras, los océanos y la mayoría de los volcanes, cómo es que los terremotos sólo se producen en determinadas zonas, por qué algunas especies animales con muy cercano parentesco viven en distintos continentes, por qué ciertas rocas están donde están, cómo es que la costa este de Sudamérica encaja perfectamente con la costa oeste de África, y mucho más. Ninguna otra teoría nos da de golpe tantas explicaciones convincentes para tantos fenómenos curiosos. Pero la felicidad que nos transmite no es completa: al mismo tiempo plantea algunas incógnitas nuevas, que están todavía muy lejos de quedar resueltas.

Actualmente se reconoce, casi de forma unánime, que el «inventor» de la tectónica de placas fue Alfred Wegener, quien a principios del siglo XX aportó gran cantidad de argumentos para apoyar su teoría, aunque otros hombres inteligentes, como Francis Bacon o Benjamin Franklin, habían especulado ya, varios siglos antes que Wegener, sobre la posibilidad de una deriva continental. Sin embargo, pasó bastante tiempo hasta que la tectónica de placas fue imponiéndose progresivamente, durante la segunda mitad del siglo pasado, por lo que se puede hablar de una historia de éxitos basada en la tenacidad. Por desgracia, el movimiento de las partes de la Tierra es más lento que el crecimiento de la hierba (unos pocos centímetros al año); no es cuestión de que alguien se siente a contemplar estos fenómenos. En vez de hacerlo así, se argumenta de forma

indirecta: si una roca, cuyas propiedades apuntan a que su origen estuvo en la zona ecuatorial, aparece ahora en Siberia, hay que pensar que esta región tuvo que estar antes en otro lugar. No importa que actualmente haga como si no se hubiera movido.

Pero ¿cómo funciona exactamente la tectónica de placas? Aunque la idea parezca sencilla a primera vista, en realidad es tremendamente compleja. La superficie terrestre se compone de «placas» de unos cien kilómetros de espesor y en cierto modo sólidas, que se deslizan sobre una capa viscosa. En casi todas estas placas hay un continente y además algunas partes del suelo marino. En las fronteras que discurren entre las placas se producen hechos dramáticos: en el caso más simple sucede que dos placas se limitan a rozar la una con la otra. Esto puede llevar a la formación de complejas fallas, y a que se produzcan terremotos y actividad volcánica, como en California, donde entran en contacto la placa norteamericana y la del Pacífico. En otros casos se trata de dos placas que se mueven alejándose la una de la otra; el espacio que va surgiendo entre ellas se llena de un flujo de rocas fundidas, magma, que procede de capas más profundas de la Tierra. Se forma una nueva corteza terrestre, y el mar crece, en un proceso que actualmente tiene lugar entre las placas africana y sudamericana, en medio del Atlántico. Pero, si los océanos crecen y las placas se separan, ¿adónde van éstas?

Una primera explicación del fenómeno fue la siguiente: la superficie terrestre se infla cada vez más, como un globo lleno de aire. Esta hipótesis carece de pruebas creíbles; por ejemplo, nadie ha encontrado el lugar donde se sopla para inflar el globo. Por el contrario, se supo enseguida que, al mismo tiempo que se produce corteza nueva, en otros lugares se deshacen las placas: se destruyen mutuamente en gigantescas colisiones por alcance. Concretamente, una placa se desliza bajo la otra, proceso que se denomina subducción y que conlleva unas consecuencias espectaculares para las placas implicadas. Por ejemplo, el Himalaya existe sólo porque la placa de la India se introduce brutalmente bajo la eurasiática y la levanta, dándole un impulso desde abajo.

En este punto, si planteamos la típica pregunta del científico, o sea «¿Por qué?», habremos llegado a lo desconocido. ¿Por qué se mueven las placas? Todo el mundo está de acuerdo en que el «motor» de la tectónica de placas se encuentra a gran

profundidad en el interior de la Tierra, donde se libera energía a partir de ciertos procesos radiactivos: se trata de una central nuclear antediluviana. El calor producido es transportado hacia el exterior por la llamada convección, un mecanismo que conocemos por los pucheros de cocina: el agua caliente asciende en unos lugares, mientras en otros la fría desciende. En el manto terrestre sucede exactamente lo mismo. Para que todo suceda muy lentamente y el planeta no se desborde al hervir, la Tierra utiliza rocas en vez de agua. Hay muchas cosas relativas a la convección en el manto terrestre que no se entienden; por ejemplo, no sabemos con exactitud dónde ascienden realmente los materiales y dónde se hunden.

Sin embargo, parece seguro que los amplios movimientos de convección en cierto modo desempeñan un papel en cuanto a impulsar la deriva continental.

Por una parte, las placas pueden «cabalgar» directamente sobre las corrientes de convección.

Si los continentes tienen sus raíces unidas con suficiente firmeza a la masa viscosa que se encuentra debajo de ellos, no les queda más remedio que seguir el movimiento de dicha masa, aunque sea de mala gana.

Actualmente es posible que América del Norte se esté moviendo de esta manera hacia Asia.

Hay otros dos procesos que posiblemente están en marcha en zonas de subducción: por una parte, las corrientes de convección pueden «succionar» hacia abajo las placas, que están tan indefensas como un nadador arrastrado hacia el fondo por un tiburón. Por otra parte, las placas pueden simplemente sumergirse, es decir, hundirse por un extremo a causa de su propio peso en los estratos que están debajo de ellas. El resto de la placa se vería arrastrado, sin que sus espectaculares protestas, en forma de terremotos y tsunamis, pudieran cambiar en nada la situación. Ambos procesos, la «succión» y el «hundimiento», producen el mismo efecto final (el borde de la placa desaparece gradualmente de la superficie terrestre), pero la fuerza que los desencadena es en un caso la convección y en el otro la gravedad. Lo que sucede después con los trozos de placa que son arrastrados al interior de la Tierra tampoco está claro. ¿Se deshace la placa inmediatamente, o se mantiene durante un tiempo y se desliza recorriendo un total

de cientos de kilómetros hacia el interior del planeta? Lo cierto es que allí el entorno será oscuro y poco confortable.

Sospechamos que todos los procesos mencionados desempeñan un papel en el desplazamiento global de las placas. Las fuerzas que se ejercen en las zonas de subducción son evidentemente las más fuertes y pueden generar unas velocidades vertiginosas de 15 centímetros por año. Sin embargo, otros mecanismos pueden estar también implicados, porque no todas las placas limitan con zonas de subducción. En cualquier caso, no se excluye la posibilidad de que la atracción contribuya al movimiento de las placas.

Por lo que respecta a la forma en que se realiza la deriva de las placas, hay también una buena cantidad de preguntas abiertas. Sólo la historia «más joven» del planeta, la de los últimos quinientos años, puede reconstruirse de una forma en cierto modo segura: hace unos doscientos cincuenta millones de años todos los continentes conocidos actualmente formaban una única masa terrestre gigantesca, un supercontinente llamado Pangea. Todo el que tuviera buenos pies podía caminar desde Alaska hasta Australia, sin necesidad de inventar el bote de remos durante el trayecto. Un poco más atrás en el tiempo, hace aproximadamente mil millones de años, las partes de la Tierra se encontraban reunidas también en un supercontinente al que llamamos Rodinia. A diferencia de lo que sucede con Pangea, sólo hay unos mapas muy imprecisos de Rodinia: por ejemplo, hay una discusión sobre el lugar en que se encontraba Siberia en aquella época. Es una suerte que los viajes alrededor del mundo no fueran especialmente populares en aquellos tiempos, porque el viajero hubiera ido a parar a la caldera del mismísimo diablo. Pero, mil millones de años es menos de un cuarto de la historia de la Tierra, y es muy difícil explicar cómo estaban los continentes durante los primeros 3500 millones de años. Se sospecha que surgían y desaparecían continentes a intervalos de tiempo desiguales. Los indicios más antiguos de la existencia de un supercontinente, al que llamamos Vaalbará, se encuentran en algunas rocas que pueden estar desplazándose sobre la Tierra desde hace más de tres mil millones de años.

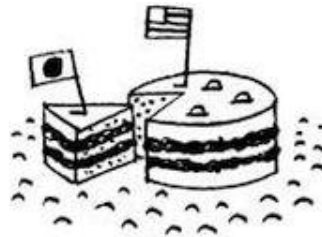
Sin embargo, por lo que respecta a los primeros tres mil millones de años de la historia de nuestro planeta, no está claro si existía algún tipo de tectónica de placas y, en caso de haberla, tampoco se sabe si funcionaba igual que ahora. ¿Qué

enormes fuerzas hicieron que entonces, en la noche de los tiempos, unos trozos de la Tierra se pusieran en movimiento? La Tierra de aquella primera época era en muchos aspectos diferente de lo que hoy en día conocemos; por ejemplo, su zona interna estaba más caliente. Ésa pudo ser la razón de que el impulso de la tectónica de placas, que, como hemos visto con anterioridad, tiene mucho que ver con las circunstancias del interior del planeta, haya sufrido grandes cambios con el paso del tiempo. El modo en que esto se produjo es objeto de discusión. La revista científica *Nature* informaba en julio de 2006 sobre un congreso, centrado en la historia de las primitivas placas, en el que los participantes, todos ellos expertos en la cuestión, llegaron a un acuerdo sobre el momento en que empezó a funcionar la tectónica de placas en la Tierra. En cualquier caso, la mayoría sospecha que el auténtico inicio de los desplazamientos se produjo antes, en algún momento entre hace tres mil y cuatro mil millones de años, pero hasta ahora no se ha podido dar una fecha más precisa. Otros creen que los comienzos fueron muy posteriores, y otros más no excluyen la posibilidad de que los continentes pudieran estar quietos algunas veces durante largo tiempo, probablemente porque tuvieran necesidad de tomarse un respiro.

La tectónica de placas no es en absoluto una cosa que tenga que existir obviamente en todos los planetas. No se sabe de otros planetas del sistema solar que hayan tenido este comportamiento tan inquieto, al menos en tiempos recientes. De todas maneras, Marte podría haber tenido alguna vez, en otras épocas, «continentes» a la deriva: ciertos datos recogidos por la sonda espacial *Mars Global Surveyor*, que orbita en torno a este planeta desde 1999, permiten pensar en una tectónica de placas tal como la conocemos en la Tierra. Sin embargo, esto tuvo que suceder hace miles de millones de años. Venus posee una superficie que en muchos aspectos recuerda a la de la Tierra, con cordilleras, volcanes y quebradas profundas, como el Gran Cañón; este tipo de formaciones surgen en la Tierra por la acción de la tectónica de placas. No obstante, Venus consigue tenerlas sin que se produzca esta acción (si esto es cierto y cómo sucede, es todavía objeto de discusión). Ganimedes, la gran luna de Júpiter (una de las cuatro que posee), descubierta por Galileo Galilei, muestra unas estupendas fallas del tipo de las producidas por la tectónica de placas y, por lo tanto, podría darnos unas valiosas indicaciones sobre

las causas y los efectos de la deriva continental. En cualquier caso, si nos dedicamos únicamente a la observación del planeta en que vivimos, en algún momento llegaremos a no conocer nuestro propio negocio.

El futuro de las placas es tan incierto como su pasado, aunque se intenta predecirlo a partir de los movimientos observados y medidos hasta ahora. Por ejemplo, actualmente se produce la colisión de África con Europa, un proceso que ha generado ya los Alpes y los Pirineos, y que puede llevar a que dentro de cincuenta millones de años el Mediterráneo se llene de arena procedente del Sahara. Además, en sólo doscientos cincuenta millones de años puede suceder, si se dan determinadas circunstancias, que América choque con un continente formado por la unión de Europa, Asia y África. No sabemos con exactitud si el ataque tendrá lugar en África Occidental o en el este asiático. No se puede decir si Nueva York en un futuro geológico cercano estará al lado de Namibia, o San Francisco junto al Japón. Pero, de esta manera, surgiría prácticamente de un día para otro un nuevo supercontinente que, según como resultara finalmente, podría llamarse Amasia o, de una manera menos creativa, Pangea Final. Por suerte, todavía queda tiempo suficiente para pensar un nombre más original.



Capítulo 40

Tendencias sexuales

RANDY MARSH: ¿Sabes, Token? Cuando un hombre y una mujer se quieren mucho, él le mete a ella el pene en la vagina. A eso se le llama «hacer el amor» y es de lo más normal.

TOKEN: Y si una mujer tiene cuatro penes en la vagina y luego se mea de pie sobre los hombres, ¿eso también es hacer el amor? Y si cinco enanos riegan a un hombre con salsa de cóctel y le dan una tunda, ¿también están haciendo el amor?
South Park

En los últimos años se ha puesto de manifiesto que la vida sexual de los animales no es una actividad tan regulada y temerosa de Dios como se creía antes: en muchas especies se han observado conductas homosexuales, los cisnes se enamoran perdidamente de patines acuáticos y un 60 por ciento de las truchas fingen el orgasmo (no, no nos lo estamos inventando). Con todo, el hombre es el animal que más lo ha complicado, hasta el punto de que no hay forma de abarcar la desconcertante diversidad de actitudes sexuales que uno puede observar en Internet. Este desarrollo (como el que ha llevado algunos omnívoros a ser críticos de restauración) es probablemente un efecto secundario de la diferenciación creciente de nuestro cerebro. Sin embargo, aunque a casi nadie le preocupa por qué la sopa de guisantes no le gusta tanto como a las demás personas, sí hay muchas personas que se interesan por el origen de sus inclinaciones sexuales. Hasta hoy no disponemos de respuestas convincentes.

Los conceptos mismos plantean ya una complicación: ¿hay que hablar de tendencias sexuales, de orientaciones o de identidad sexual? Las tres opciones conllevan sus propios problemas. Así, la homosexualidad y la heterosexualidad se definen

generalmente como orientaciones sexuales, al tiempo que la bisexualidad no suele merecer ese calificativo y la atracción por los pies o por las prácticas sadomasoquistas se suelen considerar tendencias sexuales independientes de la orientación. No obstante, esa clasificación no aporta pistas claras sobre el origen, la frecuencia o la invariabilidad de las tendencias sexuales, sino que tiene una naturaleza más bien histórica.

Resumiendo, se puede decir que aquellas prácticas sexuales que cuentan con el apoyo de un *lobby* suelen describirse como «orientación sexual» y gozan en algunos países de protección legal contra la discriminación.

Hasta el siglo XIX, las desviaciones de la norma sexual, en el caso de ser reconocidas siquiera, eran consideradas vicios. Durante el siglo XIX se pasó poco a poco del «las desviaciones sexuales son causa de enajenación mental» al «la enajenación mental es la causa de las desviaciones sexuales». Durante la primera mitad del siglo XX, los científicos sexuales más progresistas apuntaron que la homosexualidad se debía a una falta de testosterona y que podía curarse mediante el trasplante de testículos «heterosexuales». En esa misma época, Freud y sus sucesores desarrollaron la tesis que las relaciones familiares fuera de lo normal se traducían en conductas sexuales fuera de lo normal, pero que éstas podían curarse por medio del psicoanálisis.

Las desviaciones de las conductas sexuales eran consideradas un síntoma de «infantilismo psicosexual» o, lo que es lo mismo, de que un adulto se había quedado encallado en alguna de las fases de desarrollo normal de los niños. En la década de 1930, el doctor Theo Lang conjeturó que los homosexuales eran «seres humanos metamórficos» que, en realidad, pertenecían genéticamente al otro sexo. Veinte años más tarde, esa teoría quedó descartada con el descubrimiento de los cromosomas sexuales.

En cuanto a las teorías psicoanalíticas, en la década de 1950 se añadieron también las behavioristas: las tendencias sexuales fuera de lo normal podían venir condicionadas por algún episodio traumático experimentado durante la infancia. Dicho condicionamiento se veía reforzado más tarde mediante las prácticas sexuales. Uno de los inconvenientes de esa teoría es que su veracidad difícilmente puede comprobarse en los humanos. El hecho de que sea posible volver fetichista a

un animal de laboratorio tampoco significa demasiado. Por un lado, los animales en observación suelen tender de todos modos a desarrollar conductas sexuales fuera de lo corriente y, por otro, la mayoría de animales son fetichistas zoofílicos por naturaleza. El científico sexual Brian Mustanski lo expresa de la siguiente forma: «Las conductas propias de cada especie (por ejemplo saltar en el caso de las ratas) no sirven para obtener una visión general de las orientaciones sexuales humanas».

A partir de la década de 1970 se ha ido extendiendo un vacío argumental: las hipótesis anteriormente en curso han sido borradas del mapa, por lo menos en lo tocante a la homosexualidad. Las primeras en desaparecer fueron las teorías de la seducción y del contagio, que solían ser la excusa utilizada para la enérgica intervención de los legisladores. Ya nadie defiende seriamente la teoría de que la homosexualidad sea una orientación condicionada, inducida por la homosexualidad de uno de los padres o por algún trauma infantil. Esas teorías aún se publican en relación con otras conductas sexuales, pero han desaparecido por completo del debate sobre la homosexualidad. Hay que encontrar un sustitutivo, pero ¿dónde?

Desde principios de la década de 1990, la medicina y la psicología se centran cada vez más en la investigación «biológica», que no se fija tanto en las influencias sociales, sino en los efectos de genes, hormonas e infecciones. Ese desarrollo tiene que ver en parte con los métodos de estudio disponibles actualmente y, por otra, con la influencia cada vez menor del psicoanálisis. En el marco de estas nuevas tendencias, vuelve a ser objeto de investigación una observación que, en la década de 1930, inspiró a Theo Lang su teoría de los «seres humanos metamórficos»: cuantos más hermanos mayores tiene un hombre, mayor es la probabilidad de que sea homosexual. Esta circunstancia, por absurda que pueda parecer, se ha visto confirmada ya por más de veinte estudios. La existencia de hermanas mayores, en cambio, no ejerce ninguna influencia en este sentido y tampoco la homosexualidad femenina se ve condicionada por ese factor. Seguramente, Freud habría dicho que la presencia de hermanos mayores influye en las dinámicas familiares, pero lo cierto es que no es necesario que esos hermanos mayores estén presentes durante la etapa de crecimiento del niño en cuestión. Asimismo, los hermanos no carnales que están presentes en la familia no tienen ninguna influencia: sólo cuentan los hijos de la misma madre. Todo ello apunta pues a un factor que ejerce su influencia ya en el

útero materno y no en el cajón de arena. Aunque no está claro cómo funciona exactamente el proceso, lo que sucede es que el sistema inmunológico de la madre reacciona de algún modo contra las proteínas «masculinas». Y como la naturaleza tampoco quiere ponérselo fácil a los investigadores, la teoría sólo se cumple en el caso de los diestros.

Otra tesis de la investigación de la orientación biológica sostiene que la presencia de hormonas masculinas en el útero materno tiene efectos tanto en la futura orientación sexual del hijo como en la relación de tamaño entre los dedos índice y anular, circunstancia por otro lado mucho más fácil de medir. Sin embargo, los resultados de esos estudios se revelaron sumamente contradictorios, si bien eso se debe también al hecho de que otros factores, como por ejemplo el origen étnico, tienen un efecto directo en la relación de tamaño de los dedos. Los estudios en gemelos apuntan a una cierta (si bien es cierto que no demasiado marcada) influencia de la genética que, en cualquier caso, es mucho más acusada en los hombres que en las mujeres. Muchos investigadores sugieren que la homosexualidad masculina se encuentra ubicada en el cromosoma X, pues ésta es siempre más frecuente entre los familiares maternos que en los paternos. Otros objetan que la herencia por línea paterna es más difícil, pues los homosexuales no suelen tener hijos. En suma, parece que entre otras formas (y sin saber aún cómo se desarrolla) existe una homosexualidad de naturaleza biológica y que en el caso de las mujeres ésta surge de forma distinta que en el caso de los hombres.

De momento se desconoce por falta de trabajos de investigación si otras tendencias sexuales dependen también de condicionantes similares a los de la homosexualidad. Existen relatos anecdóticos de personas que a consecuencia de lesiones cerebrales o de la toma de medicamentos desarrollan de repente extrañas inclinaciones sexuales o se deshacen de ellas, pero no hay de momento estudios que analicen el comportamiento de fetichistas, sadomasoquistas o zoofílicos más allá de casos individuales. En ese sentido, se sabe poco sobre los hombres y aún menos sobre las mujeres. Muchos sexólogos sostienen que el impacto en el caso de las mujeres se limita a casos excepcionales. Tampoco parece que a corto plazo vayan a producirse grandes cambios en este ámbito tan poco agradecido de la ciencia. En el mundo, tan sólo un puñado de sexólogos se dedican a investigar las causas de las

tendencias sexuales, algo debido en parte al hecho de que tampoco es que haya demasiados sexólogos en el mundo. Médicos y psicólogos no se dedican actualmente a estos temas, pues hace falta contar con un *lobby* sólido, consciente y obstinado a combatir la discriminación para conseguir ayudas para la investigación y plazas en las universidades sin que los medios de comunicación te bauticen como «el científico que se chupaba el dedo». De momento, estas condiciones sólo se dan (y aún a medias) para la investigación de la homosexualidad.

De vez en cuando siempre hay alguien que, estudiando otros asuntos, descubre por error algo sobre la sexualidad humana: partiendo de sus estudios sobre dolores fantasma, el neurólogo americano Vilaynur S. Ramachandran ha lanzado una hipótesis que explicaría el fetichismo de pies por la razón de que la información procedente de los pies se analiza en el cerebro junto a la procedente de los genitales. Un paciente de Ramachandran explicó que, tras la amputación de un pie, experimenta los orgasmos en su pierna fantasma y que estos son mucho más satisfactorios que antes. En cualquier caso, esa teoría ayudaría a entender por qué a muchas personas les resulta agradable que les laman los dedos de los pies, pero aún no hay una explicación clara para el deseo de los fetichistas de lamer los dedos de los pies de otros. Ramachandran lo atribuye a las «neuronas espejo», que cuentan cada vez con más popularidad entre los neurólogos. Las neuronas espejo son células nerviosas que, al observar una actividad, le activan a uno el área del cerebro que funcionaría como si realizara esa misma actividad. Actualmente ésta es la gallina de los huevos de oro de la neurología, pues se puede relacionar con casi todo. Según Ramachandran, el deseo secreto de los fetichistas de pies sería que cada uno se dedicase a sus propios pies, algo que no se puede descartar pero que es poco probable. En cualquier caso, esa tesis es ya un paso adelante respecto a la suposición de los psicoanalistas Alfred Adler y Wilhelm Stekel, según la cual los fetichistas de pies son personas que, de bebés, tenían la costumbre de chuparse el dedo gordo del pie.

En general, la investigación del fetichismo (si es que podemos dar ese nombre a los escasos y dispersos intentos de encontrarle explicaciones) presupone que las partes del cuerpo a las que actualmente se les concede una significación sexual, a saber, boca, pechos, trasero y genitales, no se consideran fetiche, aunque su papel en la

procreación sea limitado. Las únicas partes del cuerpo que tradicionalmente se han considerado objeto de fetichismo son el pelo y los pies, algo que puede achacarse tanto a la historia de las ciencias como a convenciones sociales. Lo cierto es que la mayoría de partes visibles del cuerpo pueden convertirse en fetiches sexuales, especialmente si van ocultas en la vida cotidiana. Sin embargo, no se ha investigado aún la frecuencia de las tendencias fetichistas con objetos o partes del cuerpo y hasta qué punto éstas dependen de modas y condicionantes sociales. En general, disponemos de pocos datos de verdadera utilidad para el estudio de las tendencias sexuales que permitan, por ejemplo, comparar la situación en diversos países y detectar posibles influencias culturales.

Las psicólogas canadienses Patricia Cross y Kim Matheson revisaron en 2006 la teoría que sostiene que la sexualidad sadomasoquista es inducida por conflictos de personalidad. Ese enfoque no permitió demostrar ninguna de las teorías: los masoquistas investigados no sufrían sentimientos de culpa de índole sexual, tal como suponía el psicoanálisis, como tampoco tenían una tendencia a los problemas físicos o a la inestabilidad superior a la media. Se constató que los sádicos investigados no tenían un carácter autoritario en comparación con el grupo de control y no se observó en ellos ningún síntoma de trastornos de personalidad antisocial. En cuanto a los valores y roles sexuales, todos los sadomasoquistas estudiados se movían en un ámbito relativamente pro feminista. Tampoco se pudo demostrar la tesis del psicólogo Roy Baumesiter según la cual las prácticas masoquistas son la forma que tienen muchas personas de combatir la molesta conciencia moderna del yo.

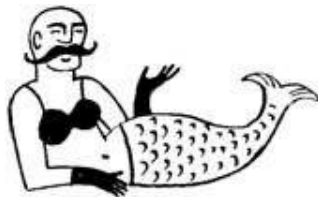
Cada tantos años se recogen datos relativos a la conducta sexual que muestran determinados grupos de población. Esos estudios muestran claramente que las desviaciones de la norma sexual casi nunca aparecen solas. Eso puede tener muchas explicaciones: ¿Acaso después de salir del armario a uno le da ya todo igual y puede declararse inmediatamente también fetichista del látex? ¿Es posible que las personas sexualmente más abiertas y con unos intereses más amplios sean más proclives a responder a una encuesta telefónica anónima sobre sus hábitos sexuales que a colgar el teléfono indignadas? ¿Existe una predisposición, más o menos

intensa, a desarrollar tendencias sexuales poco comunes que durante el proceso de evolución sexual se estructura en diversos

campos temáticos mediante influencias externas aún por determinar? Muchos encuestados declaran que sus tendencias sexuales adultas se manifestaron claramente ya antes de la pubertad.

Los científicos, sin embargo, no logran ponerse de acuerdo sobre si hay que dar credibilidad a ese tipo de afirmaciones o si se trata tan sólo de justificaciones a posteriori: «No puedo evitarlo, siempre fui así». Por el momento no está claro si a lo largo de la vida las preferencias sexuales cambian o pueden cambiar con las terapias adecuadas o si, por el contrario, éstas quedan fijadas para siempre como muy tarde después la pubertad. Muchos estudios apuntan a esta última opción, pero tanto los sectores más religiosos y conservadores como determinados ámbitos de la subcultura están tan interesados en que esa cuestión se resuelva de acuerdo con sus intereses, que las respuestas, sean en el sentido que sean, deben tomarse con escepticismo.

Desde luego, no parece que las complejas conductas que determinan la sexualidad humana pueden reducirse a causas sencillas. Probablemente las tendencias sexuales obedecen a causas múltiples y diversas, y seguramente diferentes personas desarrollan las mismas conductas sexuales por motivos diferentes. Tal vez sería mejor centrarse primero en descubrir por qué a algunos les gusta la sopa de guisantes y a otros no.



Capítulo 41

Túnel del metro norte-sur

Construir es luchar contra el agua.

BERND HILLEMEIER, Profesor de
materiales de construcción

Desde luego, las guerras mundiales constituyen una fuente fiable para averiguar algo sobre esas cosas de las que no se sabe nada con exactitud. Después de las guerras, aunque la mayoría de las veces nadie reconoce haberlo hecho, se da siempre la circunstancia de que unos documentos importantes se han quemado, se han extraviado durante la huida, o se los ha llevado el vencedor a casa como recuerdo. Tomemos un ejemplo: en los últimos días de la segunda guerra mundial, fue volada en Berlín la cubierta del túnel del metro que pasaba bajo el Landwehrkanal, con lo que gran parte del sistema de túneles se llenó de agua. Para qué podía servir esto, si fueron culpables los alemanes o los rusos, cuántas personas perdieron la vida a consecuencia de ello, e incluso la fecha y la hora de la explosión son cuestiones que actualmente siguen siendo objeto de discusión.

Sin embargo, se sabría aún menos sobre este dramático suceso, si el museo Kreuzberg de Berlín no se hubiera encargado de conseguir una documentación básica a principios de la década de 1990. En 1989, la junta de representantes del distrito de Kreuzberg había decidido colocar una placa en recuerdo de las víctimas y encargó al museo las investigaciones correspondientes. Sin embargo, como no se pudo averiguar qué era realmente lo que se debía poner en la placa, al día de hoy ésta todavía no existe. Si hubiera en Berlín un recorrido guiado de la ciudad que se centrara en las atracciones turísticas invisibles o inexistentes, se podría contemplar la placa ausente a medio camino entre las estaciones de metro de Möckerbrücke y Gleisdreieck. Por allí, hoy en día, se cruzan de nuevo el metro y el Landwehrkanal en dos niveles distintos, como es debido.

Uno de los pocos aspectos que no se discuten es el hecho de que el túnel fue destruido desde dentro hacia fuera (es decir, no mediante un disparo de artillería desde el exterior) y con una enorme fuerza explosiva. La cubierta de hormigón

armado, que en algunas zonas tenía más de un metro de espesor, se abrió en una longitud de varios metros. Según una empresa berlinesa de explosiones controladas, se tuvo que utilizar una cantidad de material explosivo del orden de varias toneladas y fueron necesarios unos conocimientos muy precisos de la zona. El agua que entró desde el Landwehrkanal fluyó hasta la estación de Friedrichstrasse, desde allí invadió lo que es ahora la línea U6, pasó por la estación de metro Stadtmitte para penetrar en la actual U2 y a la altura de la Alexanderplatz llenó las líneas U8 y U5. En consecuencia, la mayor parte de las vías de comunicación subterráneas de Berlín quedaron bajo el agua.

En muchos informes relativos a esta inundación se habla de una orden de voladura, que hasta la fecha no ha podido ser verificada. La especialista en historia de la civilización Karen Meyer alega en su informe realizado para el museo de Kreuzberg que los alemanes no podían estar muy interesados en esa voladura, ya que los túneles del metro y del metro les servían como último bastión. Por lo que respecta a los rusos, en el supuesto momento de la voladura les era seguramente más fácil entrar por la superficie, por lo que con la inundación del túnel prácticamente no se les impedía progresar en su avance. Por otra parte, si el Ejército Rojo hubiera tenido algún interés en provocar la inundación, porque así hubiera podido «limpiar» los últimos nidos de la resistencia alemana, hay que decir que en el lado ruso probablemente no disponían de los planos detallados del subsuelo berlinés que les hubieran sido necesarios para llevar a cabo la acción.

Algunos informes, cuyas fuentes no se han podido comprobar, hablan del 26 de abril como fecha de la voladura. *Die Befreiung Berlins 1945*, obra de referencia escrita en Berlín Oriental, menciona el 27 de abril, pero cita como fuente los documentos del archivo de la red ferroviaria del Reich (*Reichsbahnarchiv*), donde, sin embargo, sólo aparece el 2 de mayo. También se cita el 28 de abril en varias fuentes, sin justificación alguna, y un único informe, que no se puede verificar, habla de una explosión de polvo que habría tenido lugar el 29 o el 30 de abril en la estructura del metro junto al puente de Móckern, como consecuencia de la cual se habría dañado la cubierta de hormigón y habría entrado agua en el túnel. No obstante, una explosión de polvo no basta para producir los daños descritos.

Una alusión indirecta a la fecha de la voladura nos viene dada por la evacuación del búnker de la estación Anhalter, que existe aún actualmente. Entre 4000 y 5000 mujeres, niños y ancianos de las viviendas cercanas, que habían buscado refugio allí, fueron evacuados el 1 de mayo de 1945 por las SS a través del túnel del metro, aunque aquí la palabra «evacuados» ha de entenderse con el extraño significado de «expulsados de un lugar relativamente seguro». Los civiles pasaron por la estación del metro de la Potsdamer Platz hacia la estación de Friedrichstrasse, y de allí a la actual estación de metro de Zinnowitzer Strasse. En aquel momento había agua en algunos lugares del túnel, pero seguramente procedía de tuberías rotas por los disparos de la artillería y no cubría más que hasta las rodillas o la cadera. Este dato de la evacuación del búnker está bien documentado. Si la cubierta del túnel se hubiera volado ya en esa fecha, habría sido imposible avanzar por dicho túnel y, por otra parte, entre los muchos miles de evacuados se habría encontrado necesariamente cierto número de testigos que al menos hubieran oído la explosión. En el interior del túnel la explosión tuvo que ser mucho más sonora que los disparos simultáneos de la artillería; además, la onda expansiva tuvo que sentirse en las estaciones del entorno.

La mayoría de los informes habla del 2 de mayo como fecha de la voladura, sin mencionar las fuentes de las que se ha tomado el dato. En un informe interno realizado para la dirección de la red ferroviaria del Reich en Berlín, se dice lo siguiente: «El 2 de mayo, a las 7.45 de la mañana, una fuerte detonación sacudió la zona del cruce del Landwehrkanal con el túnel del metro norte sur...

». El autor, Rudolf Kerger, que como director del departamento de obras de la red ferroviaria del Reich era responsable de los trabajos de reconstrucción del túnel, por desgracia tampoco menciona fuente alguna, por lo que no está claro si otras dataciones parecidas han de atribuirse a Kerger, a sus fuentes o a otros documentos. En la obra publicada por el Berliner Landesarchiv con el título *Berlin. Kampf um Freiheit und Selbstverwaltung 1945-1946* se mencionan como pruebas para fijar la datación en la «mañana del 2 de mayo», por una parte, el informe novelado y no demasiado fiel a la realidad que aparece en el libro *Ten Days to Die*, del estadounidense Michael A. Musmanno, y asimismo dos fuentes que formaban parte de los fondos del propio archivo. Estas últimas no han podido encontrarse en

el curso de las investigaciones realizadas por el museo de Kreuzberg, que se quejó de que, mientras investigaba el archivo, las adversidades «habían superado lo que se podría considerar normal». Entretanto las citadas fuentes de los fondos del Landesarchiv han aparecido de nuevo, pero no contienen material alguno que pueda contribuir al esclarecimiento de las dudas planteadas.

Tras el anuncio que publicó el museo de Kreuzberg en la prensa berlinesa en 1991, se pusieron en contacto con el museo muchos lectores, de los cuales diez se acordaban de haber estado en el túnel durante la noche del 1 al 2 de mayo. Sin embargo, un artículo del año 1950 sitúa la inundación en la noche del 3 al 4 de mayo; su autor, Gerhard Krienitz, en una conversación con el museo de Kreuzberg, afirmó que en la mañana del 2 de mayo el túnel estaba todavía lleno de gente, de tal modo que, en su opinión, si se hubiera producido una inundación en aquel momento, el número de víctimas mortales tendría que haber sido mucho mayor.

Pero ¿cuántas víctimas mortales hubo? En agosto de 1945, la oficina de enterramientos de Kreuzberg solicitó al alcalde los datos correspondientes a las capacidades de los camiones disponibles, para rescatar «entre 1000 y 2000 cadáveres del túnel del metro». Dado que se supone, en principio, que el túnel se había inundado durante la evacuación, y por el hecho de que, curiosamente, en el verano de 1945 el entusiasmo de la prensa por noticias llenas de montañas de cadáveres no parece haber sido menor que el que demuestra la prensa hoy en día, algunas fuentes hablan de muchos miles de muertos. Sin embargo, durante los trabajos de desescombro del túnel se rescataron sólo unas cien víctimas que posiblemente habían fallecido ya antes de la irrupción del agua. Previamente en las estaciones del metro habían sido rescatadas del agua algunas víctimas, por lo que el museo de Kreuzberg considera que en una estimación realista se podría hablar de una cifra de entre cien y doscientos muertos.

Las preguntas que quedan aún sin respuesta son las siguientes: ¿Cuándo se produjo la voladura? ¿Por qué no hay alguien que recuerde la detonación o la onda expansiva? ¿La voladura se hizo cumpliendo una orden? Si es así, ¿quién dio esa orden y cuál sería su propósito? Las fuentes de información disponibles pueden considerarse agotadas, y los testigos son cada vez menos, pero quizá haya todavía materiales no utilizados que estén escondidos en el archivo de la red ferroviaria del

Reich o en los archivos soviéticos. Hasta que se dé respuesta a las preguntas mencionadas, no estará de más que, cuando recorramos el trayecto entre Yorckstrasse y Anhalter Bahnhof, ensalcemos brevemente el hecho de que el agua del Landwehrkanal fluya de nuevo por el lugar que le corresponde, y que en Berlín muchas cosas estén actualmente mejor que en mayo de 1945.



Capítulo 42

Vida

*Y Dios le dijo a un bloque de barro:
«¡Levántate!».*

*Yo, el barro, me levanté y vi la belleza
con la que Dios había dispuesto todo. El
único modo de sentirme un poquito
importante es pensar en todo el barro que
no puede levantarse y mirar a su
alrededor. ¡Recibí yo tanto, y la mayoría
del barro tan poco!*

KURT VONNEGUT, *Cuna de gato*

Incluso hoy en día, se desconoce absolutamente todo lo relativo a cómo surgió la vida sobre la Tierra. Aunque legiones de astrónomos, geólogos, químicos y biólogos trabajan el tema, «no estamos más cerca que los antiguos griegos» de encontrar una respuesta, según dicen los astrofísicos Eric Gaidos y Franck Selsis al hacer un resumen de la situación. Como sólo conocemos un caso de vida en el universo, precisamente la que se da en la Tierra, la investigación ha de limitarse forzosamente a este caso especial. Esto conlleva cierto riesgo, ya que la consideración de un caso especial puede desembocar en conclusiones totalmente erróneas. Pero no se puede hacer otra cosa. Serán las generaciones futuras las que averigüen si somos una forma típica de vida, o más bien el resultado de una evolución exótica.

No hace mucho, se partía de la idea de que la vida había surgido del agua o del barro, y lo había hecho de una manera en cierto modo espontánea. Al fin y al cabo había pruebas consistentes, porque cuando se depositaban residuos y basuras, y no se retiraban durante un tiempo, salían gusanos y ratas, aparentemente por sí mismos. El invento del microscopio en 1590 fue el principio del fin de esta bonita y sencilla teoría. Louis Pasteur le dio la puntilla en 1864: si se observaba con más detalle la materia supuestamente inanimada, se ponía de manifiesto que la vida era omnipresente, incluso en la basura.

La búsqueda del origen de la vida en la Tierra plantea una doble dificultad. En primer lugar hay que averiguar qué aspecto tenía el planeta poco después de su nacimiento, lo cual no es precisamente fácil. Y, cuando ya se sabe esto, hay que fabricar la vida a partir de lo que había en aquella Tierra inanimada, empezando por generar las piezas constituyentes más simples, especialmente los aminoácidos, que son los componentes de las proteínas, y con todo esto hay que crear formas de vida primitivas. El modo en que, a partir de estas formas, surgieron posteriormente los paramecios y los perros pastores alemanes es otro tema diferente, y no vamos a explicarlo ahora.

Actualmente tenemos una gran certeza de que la Tierra se formó hace 4600 millones de años (una cifra bastante exacta) a partir de un montón de escombros que giraba alrededor del Sol, siendo estos escombros los restos de los materiales que con anterioridad habían formado dicha estrella. Pero, por desgracia nadie ha documentado razonablemente la evolución del joven planeta Tierra. Todo lo que nos ha quedado de los primeros quinientos años es un puñado de piedras viejas. Y lo que es peor, la primera vida no nos ha dejado nada que podamos utilizar. Los primeros vestigios de vida son diversos organismos minúsculos petrificados, que podrían tener entre 3500 y 3800 millones de años (unas cifras muy discutidas, por supuesto). Estos seres vivos, y todos los que vinieron después, podrían proceder de un «último antepasado universal común» (*last universal common ancestor*, llamado también LUCA), del que se creyó durante mucho tiempo que entró en escena hace unos 3900 millones de años. Actualmente parece más probable que esto sucediera antes, en un planeta altamente inhóspito, sometido a bombardeos de cometas y meteoritos, y en el cual las erupciones volcánicas estaban al orden del día. Eran los tiempos oscuros y desolados de la historia de la Tierra, a los que ni siquiera se puede enviar un cámara de *National Geographic* para que tome unas imágenes en color.

Una cosa se sabe con certeza sobre aquellos primeros tiempos de una Tierra inhóspita: no había aún oxígeno, porque está claro que fueron los primeros seres vivos quienes lo produjeron para nosotros. Aparte de esto, hay discrepancias sobre la composición de la atmósfera primitiva.

Desde la década de 1920 se ha difundido siempre la idea de que se componía esencialmente de metano y amoniaco. El amoniaco tiene un olor bastante desagradable, lo cual no contribuyó precisamente a convertir aquella Tierra primitiva en un lugar más grato. Sin embargo, una atmósfera como aquélla tenía una ventaja decisiva: ofrecía unas condiciones favorables para la producción de aminoácidos. En 1953, Stanley Miller demostró esto por primera vez en el laboratorio. Miller, que entonces era todavía estudiante y trabajaba bajo la dirección de su profesor Harold Urey, llenó un depósito con la supuesta atmósfera primitiva y otro con agua, e hizo que ambos actuaran recíprocamente, sometiéndolos a descargas eléctricas que simulaban rayos. El resultado de esta sencilla receta de cocina fue realmente una sopa de aminoácidos, sobre cuyo sabor no se nos ha dicho nada. Charles Darwin había hecho ya en 1871 experimentos parecidos a los de Miller y Urey, aunque sólo mentalmente. Especuló sobre la posibilidad de una «charca caliente con todas las sales posibles de amonio y fósforo, con luz, calor y electricidad», en la que se formaron las piezas fundamentales de la vida. Por lo tanto, para que empiece la vida se necesita una atmósfera maloliente, agua suficiente y fuertes tormentas.

Naturalmente esto sólo puede ser cierto en el caso de que hubiera metano y amoniaco en la atmósfera primitiva, un supuesto sobre el que todavía no hay acuerdo. Pero, incluso si la atmósfera era así, como imaginaron Miller y Urey, ¿de dónde vino el agua? La existencia de agua es una condición necesaria para la aparición y el desarrollo de la vida; sin agua no habría un océano primitivo, y sin éste no habría una sopa primitiva. Sin embargo, el modo en que el agua pudo llegar a la Tierra sigue siendo una incógnita. En algunas teorías se supone que unos cuerpos celestes ricos en agua, por ejemplo meteoritos o cristales de hielo, suministraron agua a la Tierra, pero relativamente tarde. Otros investigadores dicen que nuestro planeta pudo formarse por la unión de varios pequeños embriones de planeta, algunos de los cuales tenían agua. Todas las teorías son problemáticas. En algunos casos aparece el agua, pero vuelve a desaparecer rápidamente; en otros, también aparece el agua, pero sólo si se tiene mucha suerte, porque a veces las cosas no van bien. Además, dado que se discrepa sobre el modo y el momento en que el agua llegó a la Tierra, también es cuestionable si funcionó la producción de

aminoácidos en los primeros tiempos del planeta, aunque se tuviera una atmósfera adecuada.

Pero ¿no pudieron ser las cosas mucho más fáciles? Esta pregunta se plantea una y otra vez desde que en las proximidades de Murchison, una pequeña ciudad de Australia, cayó en 1969 un meteorito, un pedazo de roca extraterrestre en cuyo camino se cruzó la Tierra (en principio, nada menos que una estrella fugaz especialmente grande). Lo más asombroso fue que el meteorito de Murchison contenía cierta cantidad de aminoácidos, exactamente lo que se intentaba producir químicamente en la primitiva Tierra. Se puede pensar que otros meteoritos parecidos suministraron al planeta en sus primeros tiempos los componentes de la vida. Desde luego, lo que esto no aclara es cómo llegaron los aminoácidos a aquel pedazo de roca procedente del universo.

Si aceptamos que de algún modo los aminoácidos y otros componentes básicos de la vida llegaron a la Tierra desde el exterior, nos tropezamos inmediatamente con otro misterio. Los aminoácidos son buenos y útiles, pero están todavía lejos de ser la vida. Al final tienen que salir colibríes y árboles del caucho, o por lo menos, para empezar, unas formas de vida tan sencillas como las bacterias. El desarrollo posterior desde las moléculas orgánicas, como los aminoácidos, hasta las primeras formas de vida no tiene aún explicación. Es casi imposible resumir todos los argumentos que se han esgrimido en pro y en contra de las más variadas teorías sobre la llamada evolución química. Desde la década de 1980 se han hecho populares los modelos según los cuales los procesos necesarios para la aparición de la vida tienen lugar cerca de manantiales calientes en aguas abisales, más o menos en lugares donde la lava volcánica fluye al mar. Otras teorías prefieren temperaturas normales en un entorno cambiante que pasa de húmedo a seco, por ejemplo en aguas costeras bajas. Tampoco está claro qué fue lo que se produjo a partir de los aminoácidos en el primer paso de la evolución química. El problema fundamental se puede esbozar de la siguiente manera. El modo en que se forman unas piezas básicas más complejas a partir de los aminoácidos está detallado en el banco de datos genético del ser vivo que se va a construir, o sea en el ADN. Pero éste no existe todavía. Es una situación muy embrollada: para configurar las instrucciones de uso, se necesitan las instrucciones de uso. Una variante moderna

que permite salir de este círculo vicioso es la llamada hipótesis del «mundo de ARN», en el que una pieza química fundamental y polifacética, a la que denominamos ARN, asume al mismo tiempo las funciones de arquitecto y albañil, de una manera primitiva, pero algo es algo. Sean cuales sean el lugar y el modo en que esto se produce, al final de la evolución química surge el LUCA, el antepasado común del ser humano, el mosquito y el microbio.

El LUCA tuvo que ser una forma de vida muy especial, ya que todos los seres vivos que son sus sucesores funcionan igual que él, es decir, basándose en las proteínas y el ADN. Sin embargo, es un misterio si la evolución química produjo únicamente el LUCA o también toda una serie de seres vivos primitivos diferentes de él. La primera posibilidad significaría que el proceso de creación de formas de vida no sería especialmente consistente y, por lo tanto, la vida sería un fenómeno más bien raro en el universo. Esto no se puede ni confirmar, ni refutar, porque todavía no nos hemos encontrado con nadie por ahí fuera. Pero, si al principio hubieran surgido más formas de vida diferentes, ¿acaso habría exterminado el LUCA a todos sus convecinos? Si así hubiera sido, ¿por qué habría hecho algo tan monstruoso? Ni el LUCA, ni sus hipotéticos competidores, dejaron escritos unos diarios que hubieran podido explicarnos los oscuros manejos que pudieron tener lugar durante los primeros seiscientos millones de años de la historia de la Tierra.

Actualmente tiene preferencia una teoría según la cual hubo en la primitiva Tierra distintos planteamientos de vida, algunos de los cuales pudieron haber funcionado de una manera totalmente diferente a todo lo que existe hoy en día en el mundo. Pero sólo uno de ellos, nuestro antepasado LUCA, habría sobrevivido a una terrible extinción masiva que se produjo hace 3900 millones de años, por ejemplo porque tuviera la posibilidad de refugiarse en las aguas marinas abisales. Una posible causa de dicha extinción masiva sería un bombardeo mortífero procedente del universo exterior, a saber, cierto número de grandes meteoritos que habrían caído sobre la superficie terrestre. Al menos en la Luna se ha podido comprobar un suceso de este tipo. Según esta hipótesis, el bombardeo esterilizó nuestro planeta y creó así en el desarrollo de la vida un «cuello de botella» que sólo dejó pasar una única forma de vida, que posteriormente se difundió sin impedimento alguno.

Sin embargo, existen teorías completamente diferentes, ya que la historia de los primeros tiempos de la Tierra se presta como ninguna otra época a suscitar interesantes hipótesis en los más renombrados centros de investigación. Por ejemplo, bien pudo ser que el LUCA hubiera emigrado temporalmente, y así habría escapado de la extinción: podría haberse encapsulado, para agarrarse después a un trozo de roca y volar a otro lugar a través del sistema solar. Esta teoría del autoestopista, aunque suene un poco absurda, no es una broma.

Nos podemos imaginar, por ejemplo, que un asteroide rozó la Tierra y, tras su impacto con algunos bloques de roca, arrastró consigo algunos LUCA, lanzándolos hacia el universo. De esta forma pudo ser que los LUCA anduvieran yendo y viniendo entre los planetas Tierra, Venus y Marte.

Si se acepta la idea de que existió un transporte interplanetario de componentes de la vida, entonces no hay por qué suponer que ésta surgió necesariamente por primera vez en nuestro planeta. Podría proceder asimismo de otros lugares, por ejemplo de Marte, lo cual explicaría también por qué no se encuentra en la Tierra vestigio alguno de esa forma primitiva de vida.

¿Somos nosotros, en definitiva, esos marcianos que se han buscado durante tanto tiempo, y que hace tan sólo unos pocos miles de millones de años llegaron de excursión al planeta azul? Esto haría que los científicos se pusieran casi histéricos, si se encontraran restos de vida en Marte.

Dado que el «planeta rojo» posee una superficie estable y tiene además un clima relativamente frío, podrían haberse conservado hasta nuestros días algunos vestigios de vida de tiempos muy remotos. En 1996 surgió algo sensacional: en una roca llamada ALH84001, que procedía de Marte y se encontró en la Antártida, se descubrieron huellas microscópicas de minúsculos seres vivos que, según se dijo en un principio, sólo podían proceder de Marte. Por desgracia, la alegría fue demasiado precipitada: las comprobaciones posteriores hacen suponer un origen más bien terrestre; por lo tanto, tampoco hay nada de vida extraterrestre en el ALH84001. Sin embargo, la teoría sigue siendo muy prometedora. Quizá Marte no sea el origen de la vida, pero pueda ser el archivo de su génesis.

Por último, también se puede pensar que la vida nació antes de que existieran los planetas.

Una primera teoría que apuntaba en esta dirección fue la hipótesis panespérmica, según la cual por todo el universo hay bacterias encapsuladas en esporas que están omnipresentes y vegetan por su cuenta en las nubes de gas y polvo de la Vía Láctea. Hasta la fecha no hay pruebas directas que confirmen esta hipótesis, si se deja a un lado el hecho de que algunos consideran el fenómeno de la lluvia roja en la India como una confirmación de la teoría panespérmica. Por otra parte, una variante parecida está adquiriendo una importancia cada vez mayor: la vida podría haber surgido sobre algunos asteroides, cientos de miles de pequeñas rocas insignificantes que vuelan de un lado para otro por el sistema solar. Con esta forma de asteroides hay en el sistema solar una multitud de pequeños mundos que muestran diferentes composiciones químicas, estructuras y temperaturas, un precioso cajón de arena gigantesco donde podrían jugar distintas formas de génesis de la vida, de tal modo que la probabilidad de que en algunos de esos cuerpos se haya desarrollado una bacteria es mucho mayor que la probabilidad de que haya sucedido lo mismo en algún planeta.

Como de costumbre, tampoco faltan en este caso buenos argumentos en contra. Por ejemplo, se plantea la pregunta sobre cuál puede ser la razón de que nunca se hayan encontrado formas de vida en algún asteroide, a pesar de que son muchos los asteroides que han sido examinados minuciosamente. Quizá esto se deba a que el único asteroide que transportaba vida aterrizó en la tierra hace cuatro mil millones de años, y fue el primero y único que posibilitó el desarrollo de formas de vida superiores. En este sentido, nuestra propia existencia sería en última instancia la prueba de la teoría de los asteroides, y también la de cualquier otra teoría.



BIBLIOGRAFÍA

La mención de todas las fuentes consultadas mientras se investigaba para escribir este libro llenaría un segundo tomo. Por lo tanto, a continuación se ofrecerá un máximo de cinco referencias por cada tema tratado. Por una parte, figuran los trabajos que resultaron de más ayuda para la investigación y, por otra parte, algunas referencias que ofrecen la posibilidad de que el lector obtenga más información sobre lo ya aportado en el libro. También se da la dirección de las publicaciones *on line* que tienen cierta probabilidad de existir de forma duradera en la red.

Muchas de las publicaciones citadas están disponibles de manera gratuita en Internet; el resto se pueden encargar a través del servicio de entrega de las bibliotecas en www.subito-doc.de.

SOBRE LO IGNORADO EN GENERAL

JOHN BROCKMAN (editor), *What We Believe But Cannot Prove*, Edge Foundation, Inc., The Free Press, 2005 (Científicos de primera fila especulan sobre posibles respuestas a las grandes cuestiones aún no resueltas por la investigación.) RONALD DUNCAN, MIRANDA WESTON-SMITH, *The Encyclopaedia of Ignorance - Everything You Ever Wanted To Know About the Unknown*, Pergamon Press 1977. (Una selección de lo que era desconocido en los años setenta.) JAY INGRAM, *The Science of Everyday Life, The Velocity of Honey, The Barmaid's Brain*. (A diferencia de otros autores que han escrito libros similares sobre la ciencia en la vida cotidiana, Jay Ingram se refiere a resultados contradictorios obtenidos en la investigación y a lo que queda sin explicar.) DONALD KENNEDY et al., «What Don't We Know? » *Science*, 2005, n.º 309, pp. 75 y ss.; también en: www.sciencemag.org/cgi/content/summary/309/5731/75.

JOHN MALONE, *Unsolved Mysteries of Science*, John Wiley & Sons, 2001.

JENS SÓNTGEN, «Forscher im Nebel» *duz Magazin* 02/2006; también en www.wzu.uniaugsburg.de/Projekte/Nichtwissenskultur/pdfsForscher_im_Nebel.pdf.

AGUA

MONWHEA JENG, «Hot Water Can Freeze Faster Than Cold?!?», 2005, [arxiv.org/abs/physics0512262](http://arxiv.org/abs/physics/0512262).

J. I. KATZ, «When Hot Water Freezes Before Cold», 2006, [arxiv.org/abs/physics0604224](http://arxiv.org/abs/physics/0604224).

KENNETH G. LIBBRECHT, «The Physics of Snow Crystals», *Reports of Progress in Physics*, 2005, vol. 68, pp. 855-895.

ROBERT ROSENBERG, «Why is Ice Slippery?», *Physics Today*, 2005, vol. 12, pp. 50-55.

ALUCINÓGENOS

DAVID E. NICHOLS, «Hallucinogens», *Pharmacology & Therapeutics*, 2004, vol. 101, pp. 131-181.

DAVID E. NICHOLS, «The Medicinal Chemistry of Phenethylamine Psychedelics», *The Heffter Review of Psychedelic Research*, 1998, vol. 1, pp. 40-45.

ROLAND R. GRIFFITHS et al., «Psilocybin can Occasion Mystical-type Experiences Having Substantial and Sustained Personal Meaning and Spiritual Significance», *Psychopharmacology* (edición On line), 2006.
www.hopkinsmedicine.org/Press_releases_2006/GriffithsPsilocybin.pdf.

AMERICANOS

JAMES ADOVASIO, JAKE PAGE, *The First Americans: In Pursuit of Archaeology's Greatest Mystery*, Modern Library, 2003.

MICHAEL D. LEMONICK, ANDREA DORFMAN, «Who Were the First Americans?», *TIME Magazine*, 2006, vol. 167, N° 11.

ANESTESIA

B.W. URBAN, «Current Assessment of Targets and Theories of Anaesthesia», *British Journal of Anaesthesia*, 2002, vol. 89(1), pp. 167-183.

ANGUILA

V. J. T. VAN GINNEKEN, G.E. MAES, «The European eel (*Anguilla anguilla*, Linnaeus), its Lifecycle, Evolution and Reproduction: a Literature Review», *Reviews in Fish Biology and Fisheries*, 2005, N° 15(4), pp. 367-398.

VINCENT VAN GINNEKEN, ERIK ANTONISSEN et al., «Eel Migration to the Sargasso: Remarkably High Swimming Efficiency and Low Energy Costs», *Journal of Experimental Biology*, 2005, vol. 208, pp. 1329-1335.

KATSUMI TSUKAMATO, IZUMI NAKAI, W. V. TESCH, «Do all Fresh-Water eels Migrate? », *Nature*, 1998, vol. 396, pp. 635-636.

DIETMAR BARTZ: «Der Analytiker und sein Phallusfisch», *taz Magazin* del 03-06-2006, pp.752-754.

BOSTEZO

R. R. PROVINE, B. C. TATE, L.L. GELDMACHER, «Yawning: No Effect of 3-5% CO₂, 100% O₂, and Exercise», *Behavioral and Neural Biology*, 1987, vol. 48(3), pp. 382-393.

S. M. PLATEK, S. R. CRITTON et al., «Contagious Yawning: the Role of Self Awareness and Mental State Attribution», *Cognitive Brain Research*, 2003, vol. 17(2), pp. 223-227.

S. M. PLATEK, F. B. MOHAMED, G. G. GALLUP JR., «Contagious Yawning and the Brain», *Cognitive Brain Research*, 2005, vol. 23, pp. 448-452.

M. SCHURMANN, M.D. HESSE et al., «Yearning to Yawn: the Neural Basis of Contagious Yawning», *Neuroimage*, 2005, vol. 24(4), pp. 1260-1264.

J. R. ANDERSON, M. MYOWA-YAMAKOSHI, T. MATSUZAWA, «Contagious Yawning in Chimpanzees», *Proceedings of the Royal Society of London, Biology*, 2004, vol. 271, Anexo 6, pp. 468-470.

CIEMPIÉS

L. A. H. LINDGREN, «Notes on the Mass Occurrence of *Cylindroiulus Teutonicus* Pocock in Sweden», *Entomologisk Tidskrift*, 1952, pp. 38-40.

HUGH SCOTT, «Migrant Millipedes and Centipedes in Houses, 1953-1957», *Entomologist's Monthly Magazine*, 1958, vol. 94, pp. 73-77.

HUGH SCOTT, «Migrant Millipedes Entering Houses 1958», *Entomologist's Monthly Magazine*, 1958, vol. 94, pp. 252-256.

KARIN VOIGTLÄNDER, «Mass Occurrence and Swarming Behaviour of Millipedes (*Diplopoda: Julidae*) in Eastern Germany», *Peckiana*, 2005, vol. 4, pp. 181-187.

K SAMSINAK, «Über einige in Häusern lästige Arthropodenarten», *Anzeiger für Schädlingkunde, Pflanzenschutz, Umweltschutz*, 1981, vol. 54, pp. 120-122.

CINTA AUTOADHESIVA

THE STRAIGHT DOPE SCIENCE ADVISORY BOARD, «How Does Glue Work?», 2006, www.straightdope.com.

ROBERT KUNZIG, «Why Does It Stick?», *Discover*, julio 1999, pp. 27-29.

ROLAND WENGENMAYR: «Das Geheimnis der Haftzeher», *Frankfurter Allgemeine Sonntagszeitung*, N° 4, 2006, pp. 64-65.

CÚMULOS GLOBULARES

JEAN P. BRODIE, JAY STRADER, «Extragalactic Globular Clusters and Galaxy Formation», *Annual Reviews of Astronomy and Astrophysics*, 2006, vol. 44, pp. 193-267.

MICHAEL J. WEST, PATRICK COTE et al., «Reconstructing Galaxy Histories from Globular Clusters», *Nature*, 2004, vol. 427, pp. 31-35.

KEITH M. ASHMAN, STEPHEN E. ZEPF, «The Formation of Globular Clusters in Merging and *Astrophysical Journal*, 1992, vol. 384, pp. 50-61.

DUNCAN A. FORBES, JEAN P. BRODIE, CARL J. GRILLMAIR, «On the Origin of Globular Clusters in Elliptical and cD galaxies», *Astronomical Journal*, 1997, vol. 113(5), pp. 1652-1665.

CURVA DE LAFFER

ARTHUR B. LAFFER, «The Laffer Curve: Past, Present, and Future», *Backgrounders*, 2004, n.º 1765, publicado por medio de «The Heritage Foundation» (www.heritage.org).

N. GREGORY MANKIW, MATTHEW WEINZIERL, «Dynamic Scoring: A Back-of-the-Envelope

Guide», *Harvard Institute of Economic Research*, 2005, Discussion Paper N° 2057.

PAUL PECORINO, «Tax rates and tax revenues in a model of growth through human capital accumulation», *Journal of Monetary Economics*, 1995, vol. 26, pp. 527-539.

PETER N. IRELAND: «Supply-side economics and endogenous growth», *Journal of Monetary Economics*, 1994, vol. 33, pp. 559-571.

DINERO

HELMUT CREUTZ, *Die 29 Irrtümer rund ums Geld*, Signum Wirtschaftsverlag 2005.

EINEMSEN MARTIN EISENTRAUT, «Vergleichende Beobachtungen über das Sichbespucken bei Igeln», *Zeitschrift für Tierpsychologie*, 1953, n.º 10, pp. 50-55.

ESCRITURA DEL INDO

STEVE FARMER, RICHARD SPROAT, MICHAEL WITZEL, «The Collapse of the Indus-Script Thesis: The Myth of a Literate Harappan Civilization», *Electronic Journal of Vedic Studies*, 2004, vol. 11(2), pp. 19-57.

THE STRAIGHT DOPE SCIENCE ADVISORY BOARD, «How Come We Can't Decipher the Indus Script?», 2005, straightdope.com.

ESTATURA DEL SER HUMANO

ROD USHER, «A Tall Story for our Time», *TIME Magazine*, 1996, vol. 148, pp. 92-98.

JOHN KOMLOS, MARIELOUISE BAUR, «From the Tallest to (One of) the Fattest: The Enigmatic Fate of the American Population in the 20th Century», *Münchener Wirtschaftswissenschaftliche Beiträge (VWL)*, 2003, p. 19.

RICHARD STECKEL, «A History of the Standard of Living in the United States», *EH.Net Encyclopedia*, ed. Robert Whaples, 2002, eh.net/encyclopedia.

JÓRG BATEN, GEORG FERTIG, «After the Railway Came: Was the Health of Your Children Declining? A Hierarchical Mixed Models Analysis of German Heights», 2000, Trabajo publicado para ES SHC Amsterdam.

ESTRELLA DE BELÉN

DIETER B. HERRMANN, *Rätsel um Sirius - Astronomische Bilder und Deutungen*, Verlag Der Morgen, Berlín, 1985.

SUSAN S. CARROLL, «The Star of Bethlehem: An Astronomical and Historical Perspective», sciaastro.net/portia/articles/the_star.htm. Bibliografía de ROBERT H. VAN GENT, www.phys.uu.nl/~vgentf/stella_magorum/stellamagorum.htm.

EYACULACIÓN FEMENINA

SABINE ZUR HIEDEN, *Weibliche Ejakulation*, Psychosozial-Verlag, 2004 (Beiträge zur Sexualforschung, vol. 84).

GARY SCHUBACH, «Urethral Expulsions During Sensual Arousal and Bladder Catheterization in Seven Human Females», *Electronic Journal of Human Sexuality*, 2001, www.ejhs.org/volume4/Schubach/abstract.html.

GARY SCHUBACH, «The Human Female Prostate and Its Relationship to the Popularized Term, GSpot », 2002, doctorg.com.

FOLLAJE OTOÑAL

DAVID W. LEE, KEVIN S. GOULD, «Why Leaves Turn Red», *American Scientist*, 2002, vol. 90, pp. 524-531.

MARCO ARCHETTI, SAM P. BROWN, «The coevolution theory of autumn colours», *Proceedings of the Royal Society of London, B*, 2004, vol. 271, pp. 1219-1223.

GOTEO

JENS EGGERS, «Drop Formation - an Overview», *Zeitschrift für Angewandte Mathematik und Mechanik*, 2005, vol. 85, pp. 400-410.

HAWAI

Bibliografía en www.mantleplumes.org

A. D. SAUNDERS, «Mantle Plumes: an Alternative to the Alternative», *Geoscientis*, agosto, 2003, www.geolsoc.org.uk.

HIPÓTESIS DE RIEMANN

KEITH DEVLIN, *The Millennium Problems*, Basic Books, 2002. Bibliografía de MATTHEW R. WATKINS: secamlocal.ex.ac.uk/~mwatkins/zeta/riemannhyp.htm.

LLUVIA ROJA

S. SAMPATH, T. K. ABRAHAM, V. SASI KUMAR, C. N. MOHANAN, «Coloured Rain: A Report on the Phenomenon», informe oficial publicado por el Centre for Earth Science Studies» y por el Tropical Botanical Garden and Research Institute, 2001, CESSPR 114-2001.

GODFREY Louis, A. SANTHOSH KUMAR, «The Red Rain Phenomenon of Kerala and its Possible Extraterrestrial Origin», *Astrophysics & Space Science*, 2006, vol. 302, pp. 175-187, también en: arxiv.org/abs/astroph/0601022.

MANUSCRITO VOYNICH

The Voynich Manuscript, www.voyrich.nu. *Voynich Manuscript Mailing List HQ*, www.voyrich.net.

MATERIA OSCURA

VIRGINIA TRIMBLE, «Existence and Nature of Dark Matter in the Universe», *Annual Reviews of Astronomy & Astrophysics*, 1987, vol. 25, pp. 425-472.

BERNARD CARR, «Baryonic Dark Matter», *Annual Reviews of Astronomy & Astrophysics*, 1994, vol. 32, pp. 531-590.

JEREMIAH OSTRIKER, «Astronomical Tests of the Cold Dark Matter Scenari», *Annual Reviews of Astronomy & Astrophysics*, 1993, vol. 31, pp. 689-716.

ROBERT H. SANDERS, STACY S. MCGAUGH, «Modified Newtonian Dynamics as an Alternative to Dark Matter», *Annual Reviews of Astronomy & Astrophysics*, 2002, vol. 40, pp. 263-317.

VARUN SAHNI, «Dark Matter and Dark Energy», *Lecture Notes in Physics*, 2004, 653, p. 141, y también en: arxiv.org/abs/astroph/0403324.

MIOPÍA

IAN G. MORGAN, «The Biological Basis of Myopic Refractive Error», *Clinical and Experimental Optometry*, 2003, vol. 86(5), pp. 276-288.

FRANK SCHAEFFEL, «Das Rätsel der Myopie», *Der Ophthalmologe*, 2002, n.º 2, pp.120-141.

KLAUS SCHMID, *Myopia Manual*, 2006, www.myopia-manual.de.

OLFATO

ANDREAS KELLER, LESLIE B. VOSSHALL, «A Psychophysical Test of the Vibration Theory of Olfaction», *Nature Neuroscience*, 2004, vol. 7, pp. 337-338.

RICHARD AXEL, «The Molecular Logic of Smell», *Scientific American*, 1995, vol. 273, pp. 154-159.

LUCA TURIN, «A Spectroscopic Mechanism for Primary Olfactory Sensation», *Chemical Senses*, 1996, vol. 21(6), pp. 773-791.

JENNIFER C. BROOKES, FILIO HARTOUTSIOU et al., «Could Humans Recognize Odor by Phonon Assisted Tunneling?», *Physical Review Letters*, 2007, 98, 038101.

BARRY R. HAVENS, CLIFTON E. MELOAN, «The Application of Deuterated Sex Pheromone Mimics of the American Cockroach to the Study of Wright's Vibrational Theory of Olfaction», en: *Food Flavours: Generation, Analysis, Process Influence*, ed. G. Charalambous, Elsevier Science 1996, pp. 497-524.

PARTÍCULAS ELEMENTALES

D. P. ROY, «Basic Constituents of Matter and their Interactions - A Progress Report», conferencia para el «3rd International Symposium on Frontiers of Fundamental Physics», Hyderabad, 17-21 de diciembre de 1999, arxiv.org/abs/hep-ph/9912523.

HAIM HARARI, «A Schematic Model of Quarks and Leptons», *Physics Letters B*, 1979, vol. 86(1), pp. 83-86.

ROGER PENROSE, *The Road to Reality. A Complete Guide to the Laws of the Universe*, Vintage, Londres, 2006 [traducción al castellano: *El camino a la realidad: una guía completa de las leyes del universo*, Debate, Barcelona, 2006.]

PROBLEMA P / NP

KEITH DEVLIN, «The Millennium Problems», Basic Books, 2002. MICHAEL SIPSER, «The History and Status of the P vs. NP Question», *Proceedings of the 24th Annual ACM Symposium on the Theory of Computing*, 1992, pp. 603-618.

PROPINA

DIEGO GAMBETTA, «What Makes People Tip? Motivations and Predictions», *Oxford Sociology Working Papers*, 2006, Paper N° 09.

MICHAEL LYNN, «Tipping in Restaurants and Around the Globe: An Interdisciplinary Review», en *Handbook of Behavioral Economics*, Morris Altman (ed.), Armonk, 2006.

RAYOS GLOBULARES

JOHN ABRAHAMSON, JAMES DINNISS, «Ball Lightning Caused by Oxidation of Nanoparticle Networks from Normal Lightning Strikes on Soil», *Nature*, 2000, vol. 403, pp. 519-521.

GRAHAM K. HUBLER, «Fluff Balls of Fire», *Nature*, 2000, vol. 403, pp. 487-488.

STANLEY SINGER, «Great Balls of Fire», *Nature*, 1991, vol. 350, pp. 108-109.

ANTONIO F. RAÑADA, JOSE L. TRUEBA, «Ball Lightning an Electromagnetic Knot? », *Nature*, 1996, vol. 383, pp. 32-33.

RESFRIADO

J. BARNARD GILMORE, *In Cold Pursuit*, Stoddard Publishing Co., Toronto, 1998

REY DE LAS RATAS

G. WIERTZ, «Experiment zur Bildung eines Rattenkönigs», 1966, *Zeitschrift für Säugetierkunde*, 1966, vol. 31, pp. 20-22. «Von Rattenkönig», *Orion*, 1952, vol. 3, pp. 118-120.

ALFRED EDMUND BREHM, «Die Nagetiere (Nager)», en *Brehms Tierleben*, 1864-1869 [traducción al castellano: *Vida animal*, Plaza & Janés Editores, Barcelona, 1997].

J. HICKFORD, R. JONES, S. COURRECH DU PONT, J. EGGERS, «Knotting probability of a Shaken Ball-Chain», *Physical Reviews*, 2006, E 74, 052101.

RONRONEO

W. R. MCCUISTION, «Feline Purring and its Dynamics», *Veterinary Medicine / Small Animal Clinician*, junio 1966, pp. 562-566.

J. E. REMMERS, H. GAUTIER, «Neural and Mechanical Mechanisms of Feline Purring», *Respiration Physiology*, 1972, vol. 16, pp. 351-361.

DAWN E. FRAZER SISSOM, D.A. RICE, G. PETERS, «How Cats Purr», *Journal of the Zoological Society of London*, 1991, vol. 223, pp. 67-78.

G PETERS, «Purring and Similar Vocalizations in Mammals», *Mammal Review*, 2002, vol. 32(4), pp. 245-271.

ROTACIÓN DE LAS ESTRELLAS

WILLIAM HERBST, JOCHEN EISLÖFFEL, REINHARD MUNDT, ALEKS SCHOLZ, «The Rotation of Young Low-Mass Stars and Brown Dwarfs», en: *Protostars & Planets V*, ed. B. Reipurth, D. Jewitt, K. Keil, University of Arizona Press, Tucson, 2007, pp. 297-311, también en: arxiv.org/abs/astro.-ph/0603673.

SEAN MATT, RALPH E. PUDRITZ, «Understanding the Spins of Young Stars», ponencia presentada en el congreso «Cool Stars, Stellar Systems, and the Sun», Pasadena, 6-10 de noviembre de 2006, 2007, arxiv.org/abs/astro-ph/0701648.

PETER BODENHEIMER, «Angular Momentum Evolution of Young Stars and Disks», *Annual Review of Astronomy and Astrophysics*, 1995, vol. 33, pp. 199-238.

SONIDOS DESAGRADABLES

D. LYNN HALPERN, RANDOLPH BLAKE, JAMES HILLENBRAND, «Psychoacoustics of a Chilling Sound», *Perception & Psychophysics*, 1986, vol. 39(2), pp. 77-80.

SUCESO DE TUNGUSKA

SURENDRA VERMA, *The Mystery of the Tunguska Fireball*, Icon Books, 2005.

WOLFGANG KUNDT, «The 1908 Tunguska catastrophe: An alternative explanation», *Current Science*, 2001, vol. 81(4), pp. 399-407. Página web de ANDREJ OLCHowATow: olkhov.narod.ru.

SUEÑO

JEROME M. SIEGEL, «Clues to the Functions of Mammalian Sleep», *Nature*, 2005, vol. 437(27), pp. 1264-1271.

JEROME M SIEGEL, «The REM Sleep-Memory Consolidation Hypothesis», *Science*, 2001, vol. 294, pp. 1058-1063.

ULLRICH WAGNER, STEFFEN GAIS et al., «Sleep Inspires Insight», *Nature*, 2004, vol. 427, pp. 352- 355.

JIM HORNE, «The Phenomena of Human Sleep», *The Karger Gazette*, 1997, n.º 61, pp. 1-5.

TAMAÑO DE LOS ANIMALES

CHRIS CARBONE, AMBER TEACHER, J. MARCUS ROWCLIFFE, «The Costs of Carnivory», *PloS Biology*, 2007, vol. 5(2), www.plosbiology.org.

GARY P BURNES, JARED DIAMOND, TIMOTHY FLANNERY, «Dinosaurs, Dragons, and Dwarfs: The Evolution of Maximal Body Size», *Proceedings of the National Academy of Sciences of the United States of America*, 2001, www.pnas.org.

DAVID W. E. HONE, MICHAEL J. BENTON, «The Evolution of Large Size: How Does Cope's Rule Work?», *Trends in Ecology and Evolution*, 2003, vol. 20(1), pp. 4-6.

C. A. ALLEN, A. S. GARMESTANI et al., «Patterns in Body Mass Distributions: Sifting Among Alternative Hypotheses», *Ecology Letters*, 2006, vol. 9, pp. 630-643.

TECTÓNICA DE PLACAS

E. RONALD OXBURGH, DONALD L. TURCOTTE, «Mechanisms of Continental Drift», *Reports on Progress in Physics*, 1978, vol. 41, pp. 1249-1312.

ALEXANDRA WITZE, «The Start of the World as We Know It», *Nature*, 2006, vol. 442, pp. 128-131.

PAUL J. TACKLEY, «Mantle Convection and Plate Tectonics: Toward an Integrated Physical and Chemical Theory», *Science*, 2000, vol. 288, pp. 2002-2007.

GÖTZ BOKELMANN, «Which Forces Drive North America?», *Geology*, 2006, vol. 30(11), pp. 1027- 1030.

TENDENCIAS SEXUALES

BRIAN S. MUSTANSKI, M. CHIVERS, J. M. BAILEY, «A critical Review of Recent Biological Research on Human Sexual Orientation», *Annual Review of Sex Research*, 2002, vol. 12, pp. 89-140.

RAY BLANCHARD, JAMES M. CANTOR et al., «Interaction of Fraternal Birth Order and Handedness in the Development of Male Homosexuality», *Hormones and Behavior*, 2005, vol. 49, pp. 405-414.

HARRY OOSTERHUIS, *Stepchildren of Nature: Krafft-Ebing, Psychiatry, and the Making of Sexual Identity*, The University of Chicago Press, 2000.

TÚNEL DEL METRO NORTE-SUR

KAREN MEYER, «Die Flutung des Berliner S-Bahn-Tunnels in den letzten Kriegstagen», en: *Nord-Süd-Bahn: Vom Geistertunnel zur City-S-Bahn*, ed. Berliner S-Bahn-Museum, Berlín 1999, pp. 47-85.

VIDA

ERIC GAIDOS, FRANCK SELSIS, «From Protoplanets to Protolife: The Emergence and Maintenance of Life», en: *Protostars & Planets V*, edición de B. Reipurth, D. Jewitt, K. Keil, University of Arizona Press, Tucson, 2007, pp. 929-944, también en arxiv.org/abstastro-ph/0602008.

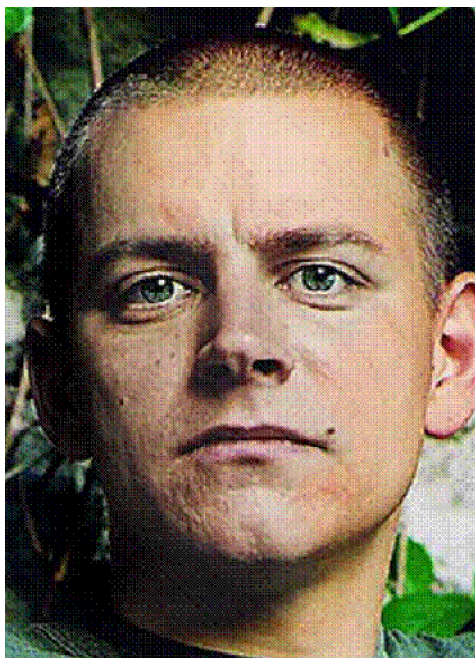
AGRADECIMIENTOS

Nuestro más sincero agradecimiento a los principales correctores de esta enciclopedia, Bernd Klóckener y Christian Y. Schmidt, a los expertos Helmut Baaske, Dietmar Bartz, Jörg Baten, Wolf Blanckenhorn, Gótz Bokelmann, Anton Deitmar, Steffen Gais, Gerald Hártlein, Hanns Hatt, Gerda Horneck, Wolfgang Kundt, Gustav Peters, Martin Schaefer, Bertram Schefold, Ralph Scheicher, Kai Schreiber, Martin Schürmann, Daniel Schwekendiek, Kathrin Worschech y Klaus Zimmermann, a los lectores Michael Brake, Holm Friebe, Johannes Jander, Angela Leinen, Ruben Schneider y Volker Scholz,

así como a Bettina Andrae, Christoph Danz, Uwe Heldt, Wolfgang Herrndorf, Ray Jayawardhana, Dieter y Gertrud Passig, Ulrike Richter, Tex Rubinowitz, Martin Rudolph, Jens Soentgen, Jochen Schmidt, Henning Scholz, Kai Schreiber, Ira Strübel, Stese Wagner y Klaus César Zehrer.



KATHRIN PASSIG (1970). Dirige la agencia de comunicación ZIA, ha traducido a George W. Bush y a Bob Dylan, y colabora en varias revistas alemanas como *GEO*, *FAZ* y *Spiegel Online*. Es asimismo redactora y programadora del weblog «Riesenmaschine», galardonado con el Grimme Online Award 2006. En 2006 mereció el premio Bachmann Preis y el Publikumpreis por el relato *Se encuentra usted aquí*.



ALEKS SCHOLZ (1975). Es astrónomo y se ha especializado en el estudio de la formación, el desarrollo y la estructura de las estrellas y los planetas. Trabaja en la Universidad de Saint Andrews, en Escocia, y es también redactor del weblog «Riesenmaschine». Dentro de ese proyecto, ha coeditado el libro *Riesenmaschine. Das Beste aus dem brandneuen Universum* (2007).