

TEMA 1. CIENCIA, TECNOLOGÍA E INNOVACIÓN

Libro 100. El.01.CienciaTecnologíaeInnovación
Esteban Fernández Sánchez
Borrador actualizado: Septiembre 2019

1. EL CONOCIMIENTO

- 1.1. Datos e información
- 1.2. El conocimiento: características básicas
- 1.3. Conocimiento explícito
- 1.4. Conocimiento tácito

2. LA CIENCIA

- 2.1. El paradigma científico
- 2.2. La investigación básica

3. LA TECNOLOGÍA

- 3.1. El paradigma tecnológico
- 3.2. La investigación aplicada y el desarrollo experimental (I+D)

4. RELACIÓN ENTRE CIENCIA Y TECNOLOGÍA

5. EL INVENTO Y LA INNOVACIÓN

6. TIPOS DE INNOVACIONES

- 6.1. Innovación radical *versus* incremental
- 6.2. Innovación original *versus* adaptada
- 6.3. Innovación tecnológica *versus* mercado
- 6.4. Innovación modular
- 6.5. Naturaleza funcional de las innovaciones
- 6.6. Innovación en la base de la pirámide
- 6.7. Innovación inversa
- 6.8. Innovación en el modelo de negocio

7. LA DIFUSIÓN DE LA INNOVACIÓN

- 7.1. Características de la innovación que afectan a la difusión
- 7.2. Dimensión temporal
- 7.3. Sistema social (mercado potencial)

8. IMITACIÓN DE LA INNOVACIÓN

9. APROPIACIÓN DE LOS BENEFICIOS DE LA INNOVACIÓN

1. EL CONOCIMIENTO

El conocimiento se deriva de la información y la información se extrae de los datos (figura 1). La cantidad de datos que hay en el mundo está creciendo desafortunadamente, desbordando no solo la capacidad actual de las máquinas, sino también la propia imaginación de las personas.

1.1. Datos e información

Los datos y la información, aunque no son conceptos intercambiables, en la práctica no resulta sencillo separarlos. Los datos se refieren habitualmente a observaciones en bruto sobre hechos, acontecimientos y entidades que las personas realizan por sí mismas o utilizando los instrumentos adecuados. Un hecho objetivo sería aquel que fuese verdadero para cualquier persona en todo momento, con independencia de su perspectiva particular sobre la realidad. Ahora bien, la evidencia del observador no cambia la verdad de las cosas, pero tiñe con tonalidades propias la forma en la que es percibido el hecho.

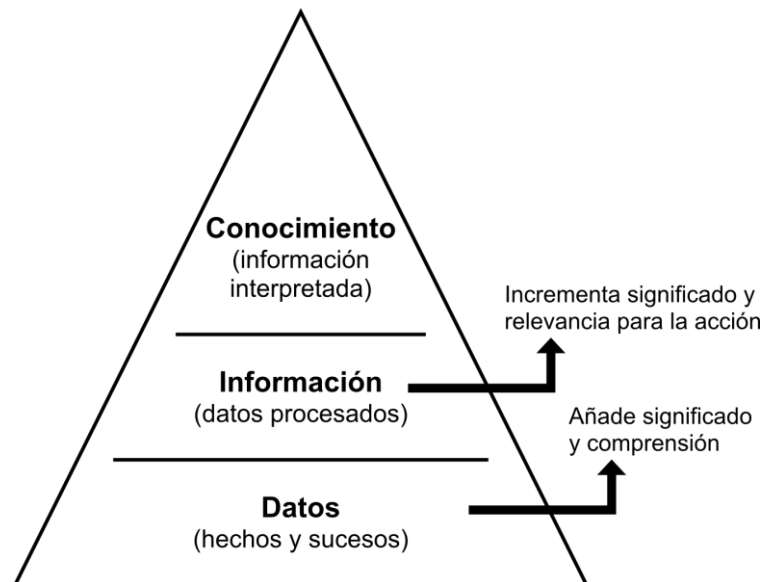


Figura 1: Datos, información y conocimiento (Jasimuddin, 2012)

Los datos pueden ser útiles sólo a nivel elemental y se representan de cuatro formas diferentes: números, letras, sonidos e imágenes. Los datos incluyen tanto descripciones cualitativas como medidas cuantitativas. Una vez registrados, es fácil almacenarlos y transferirlos a otros lugares. Carecen de un significado inherente, es decir, su importancia intrínseca es escasa, ya que son cualidades objetivas de las cosas. En consecuencia, no se utilizan en bruto. Los datos no se obtienen accidentalmente, sino que hay que encontrarlos, y no son fijos, sino que pueden cambiar.

La información son datos procesados (organizados) dentro de un determinado contexto que resulta útil para ciertos usuarios. Al transformarse en información, los datos adquieren significado para un individuo y, en consecuencia, alteran o refuerzan su comprensión de un hecho. Así pues, la información son datos dotados de relevancia y propósito (Drucker, 1988).

No existe un acuerdo general sobre la forma de procesar los datos o, lo que es lo mismo, sobre la forma de transformarlos en información. En general, un individuo procesa los datos de acuerdo con sus creencias previas, el contexto donde se realiza y su percepción de la realidad que estudia. Por tanto, dicho procesamiento es subjetivo, pues las personas pueden interpretar las observaciones de distinta manera. Por ejemplo, Lavoisier encontró oxígeno donde Priestley había visto aire deflogistizado y donde otros investigadores no vieron nada en absoluto (Kuhn, 1962). Por otra parte, la información adquiere significado dentro de las prácticas sociales de la comunidad que la utiliza (lectura 1). En suma, una característica de la información es su subjetividad, lo cual genera discrepancias acerca de su validez, utilidad e importancia.

LECTURA 1: EXPLICACIÓN DE UN HECHO

Un estudio comparó la cobertura mediática de los periódicos estadounidenses y japoneses de escándalos financieros, como el del “comerciante deshonesto” Nick Leeson, cuyos negocios ilegales acabaron creando una deuda de 1.400 millones de dólares que provocó la quiebra del Barings Bank en 1995, o de Toshihide Iguchi, cuyos propios negocios ilegales costaron al Daiwa Bank 1.100 millones de dólares ese mismo año. Los investigadores observaron que los periódicos estadounidenses tenían más tendencia a explicar los escándalos haciendo referencia a la acción individual de los estafadores, mientras que la prensa japonesa hacía más énfasis en factores institucionales, como una mala supervisión de la alta dirección. Ya sea que consideremos los resultados como merecedores de escarnio o de alabanza, los miembros de sociedades individualistas asignan la responsabilidad a un individuo, mientras que las sociedades colectivistas ven los resultados como indisolublemente vinculados a los sistemas y contextos.

Fuente: Menon, T.; Morris, M. W.; Chiu, C. y Hong, Y. (1999): “Culture and the construal of agency: Attribution to individual versus group dispositions”, *Journal of Personality and Social Psychology*, vol. 76, pp. 701–717.

Los procesos a los que se someten los datos para transformarlos en información se pueden clasificar en cinco grupos (Davenport y Prusak, 1998): a) contextualización o el propósito con el que se recopilan los datos; b) categorización, esto es, delimitar las unidades de análisis o componentes clave de los datos; c) cálculo, es decir, análisis matemático o estadístico de los datos; d) corrección o eliminación de los errores de los datos y e) condensación o resumen de los datos de alguna forma más concisa.

Para que la información sea útil debe ser de calidad, oportuna, pertinente, completa y valiosa. La calidad de la información está determinada por la exactitud y la fiabilidad en la descripción real de un hecho o situación. En otro caso, es probable que induzca a error o se ignore. Una información oportuna es aquella que está al día en los asuntos de interés y disponible cuando se necesita, por lo que resulta básica para la toma de decisiones. Decisiones relativas a la compra de materias primas basadas en la cotización del mercado, podrían ser muy costosas para la empresa si la información que recibe está desfasada. Una información pertinente está circunscrita a los hechos adecuados. Una empresa que comercializa sus productos en el extranjero necesita información específica de cada país, que ha de contemplar, entre otros aspectos, las diferencias culturales. Una información completa significa que debe abarcar todo lo relevante, y no debe presentar en forma selectiva la evidencia que apoya solamente un punto de vista en detrimento de otros. Finalmente, una información es valiosa si quien la posee obtiene ganancias mayores que en ausencia de ella. La información es un bien económico, ya que no puede obtenerse gratuitamente, ni siquiera está al alcance de todos, lo que ocasiona asimetrías de información (Arrow, 1962a): unos agentes poseen información de la que

otros carecen. En suma, la información posee ciertas características que la convierten en imprescindible en la toma de decisiones. De hecho, es su principal *input*.

La información puede ser digitalizada –codificada como un conjunto de bits– (Shapiro y Varian, 1999), por lo que se presta a ser copiada con exactitud y ser transmitida instantáneamente por todo el mundo. En la actualidad, la tecnología de la información y de las comunicaciones ha potenciado el acceso gratuito (o con un mínimo coste) a una ingente cantidad de información, que, por otra parte, ninguna persona es capaz de procesar en su totalidad. Como hoy en día la información está disponible de una manera rápida, universal y barata, no puede sorprendernos que todo el mundo se queje de la saturación informativa. Herbert Simon lo expresa con claridad al afirmar que «la riqueza de información provoca una pobreza de atención». El problema actual no es de acceso a la información, sino de sobrecarga de información. Por ello, los usuarios tienden a utilizar sólo la información que les filtran las personas y organizaciones especializadas que les asesoran o aquella que, de forma rutinaria, están acostumbrados a manejar (Simon, 1997). Esta circunstancia da lugar a una curiosa paradoja: el actual acceso a una enorme cantidad de información une al individuo en mayor medida a la subjetividad de los filtros (Sveiby, 1997), sean personas o máquinas programadas para ello. Es lo mismo que decir: hoy en día, en términos informativos, ¡más es menos!

1.2. El conocimiento: características básicas

Una persona puede aseverar que conoce un hecho, independientemente de si la conoce realmente o no. En realidad, tiene una opinión en la que cree firmemente. Trivializamos el conocimiento cada vez que pasamos por alto diferencias como las que rigen entre el conocimiento y la opinión; entre saber y aseverar (lectura 2).

LECTURA 2: LA OPINIÓN

El ámbito de la opinión es el de la pluralidad y la diversidad. Arraigada en la experiencia individual o social es un parecer que trasluce una visión del mundo. La opinión es fruto de cierta elaboración personal, o es recibida de una tradición reconocida y aceptada. Constituye, pues, una aprehensión empírica del mundo que nos rodea. A este respecto, puede ser reivindicada: “es mi opinión”. Nos remite a la experiencia y reflexión de aquel que la forma y se adhiere a ella.

Las opiniones representan una forma de conocimiento empírico. Tienen, por tanto, una función práctica indudable; muchos conocimientos técnicos, por ejemplo, se obtienen a base de tanteos y aproximaciones, por ensayo y error, no siendo por ello menos eficaces o útiles.

Tal es el caso de las fórmulas geométricas gracias a las cuales los egipcios recalculaban la superficie de campo tras la cosecha del Nilo. Obtenidas por intuición y avaladas *a posteriori* a través de su uso, no están fundadas, por tanto, sobre una demostración racional enteramente *a priori*, como sí lo están las fórmulas griegas. Las técnicas humanas, antes del nacimiento de las ciencias aplicadas, representan un tesoro de conocimientos empíricos indispensables para la vida y la supervivencia de las sociedades humanas. Hay sociedades sin ciencias, pero no sociedades sin técnicas.

Si el mundo fuese enteramente transparente a la ciencia, no habría lugar alguno para la opinión. La opinión ocupa un lugar decisivo allí donde es imposible determinar algo con exactitud en una proposición universal y necesaria. La opinión es plural y llama al debate y a la argumentación.

A pesar de su innegable utilidad, la opinión adolece de una falta de fundamento racional y no puede ser considerada como conocimiento seguro. Cambiante y en movimiento es fácilmente manipulable.

Para los filósofos, la opinión representa una fuente de insatisfacción debido a la insuficiente fundamentación de sus afirmaciones. Pero es fecunda en cuanto fuente de interrogaciones y problemas.

La opinión se sitúa en el ámbito localizado entre la ignorancia y la ciencia. Y aunque todas valen lo mismo en su insuficiente fundamento racional, hay que reconocer, no obstante, que hay opiniones verdaderas y opiniones falsas.

La opinión verdadera, inconsciente de sus razones, es, en este sentido, inferior a la ciencia; pero, en tanto que verdadera, es, en su aplicación, tan útil como esta.

La opinión es una conjetura que puede tener un valor, pero que se caracteriza por su saber insuficiente. Es por ello que, a menudo, toma la forma de una creencia más o menos consciente de esta insuficiencia.

Fuente: Raffin, F. (2011): *Pequeña Introducción a la Filosofía*, Alianza Editorial, Madrid.

Para algunos, el conocimiento es información combinada con la experiencia, el contexto, la interpretación y la reflexión (Davenport *et al.*, 1998). Según una escuela de pensamiento, que se remonta a los diálogos platónicos, el *Menón* y el *Teeteto*, y que ha gozado desde entonces de gran predicamento, el conocimiento es una creencia verdadera justificada. Esta concepción, aunque sencilla y de fácil comprensión, presenta algunas incómodas dificultades.

Creencia. Las creencias son las comprensiones del individuo de la relación existente entre una acción y un resultado. Por tanto, determinan en gran medida lo que el individuo percibe, absorbe y concluye de sus observaciones. Y permiten al sujeto tener una opinión formada.

La percepción es el proceso de adquirir datos por medio de nuestros sentidos, que son el primer canal hacia la comprensión del mundo, y nos brinda una justificación de nuestras creencias. Ahora bien, tenemos un sesgo confirmativo, que se manifiesta en la tendencia a ver el mundo a través de lentes que han sido coloreadas por nuestras creencias anteriores (Stiglitz y Greenwald, 2014). El mundo confirma nuestros conceptos previos. Por tanto, las personas con creencias distintas «ven» cosas diferentes en la misma situación y organizan su conocimiento sobre la base de sus valores (lectura 3). De esta forma, el conocimiento verdadero de una persona puede ser una fuente de escepticismo para otras, como ocurre con la teoría sobre el origen de las especies. ¿Y qué hubo más cierto, durante milenios, que el movimiento del Sol alrededor de una Tierra inmóvil?

LECTURA 3: DIFERENTES PERCEPCIONES DE UN SUCESO

Los científicos ven los hechos de forma diferente acorde con sus creencias y conocimientos. Un buen ejemplo nos lo proporciona el descubrimiento de Urano por Sir William Herschel. Al menos en diecisiete ocasiones diferentes, entre 1690 y 1781, una serie de astrónomos, incluyendo a varios de los observadores más eminentes de Europa, vieron una estrella en posiciones que suponemos actualmente que debía ocupar entonces Urano. Uno de los mejores observadores de dicho grupo vio realmente la estrella durante cuatro noches sucesivas, en 1769, sin notar el movimiento que podía haber sugerido otra identificación. Herschel, cuando observó por primera vez el mismo objeto, doce años más tarde, lo hizo con un telescopio perfeccionado, de su propia fabricación. Como resultado de ello, pudo notar un tamaño aparente del disco que era cuando menos muy poco usual para las estrellas. Halló en ello algo raro y, por consiguiente, aplazó la identificación hasta llevar a cabo un examen más detenido. Ese examen mostró el movimiento de Urano entre las estrellas y, como consecuencia, Herschel anunció que había visto un nuevo cometa. Sólo al cabo de varios meses, después de varias tentativas infructuosas para ajustar el movimiento observado a una órbita de cometa, Lexell sugirió que la órbita era probablemente planetaria. Cuando se aceptó esa sugerencia hubo varias estrellas menos y un planeta más en el mundo de los astrónomos profesionales. Un cuerpo celeste que había sido observado varias veces, durante casi un siglo, se percibe de modo diferente a partir de 1781.

Fuente: Kuhn, T. S. (1962): *The Structure of Scientific Revolutions*, University of Chicago Press, Chicago.

En suma, la creencia impregna el conocimiento de un cierto carácter subjetivo y provoca que los individuos desarrollen pautas de actuación diferenciadas ante el mismo hecho. En el límite, este carácter subjetivo del conocimiento negaría su existencia en el

sentido objetivo, al cuestionar la certeza absoluta: sólo hay que recordar la evolución de la investigación en la física y el cambio de sus principios y creencias.

Verdadera. La teoría de la correspondencia sostiene que cuando una verdad es pensada o expresada, hay un hecho del mundo que sustenta que esa creencia o aseveración sea verdadera. De ahí que no pueda haber verdades sin hechos. Así pues, un hecho es todo lo que puede ser acertadamente representado o descrito por las verdades. Es decir, las verdades son representaciones o descripciones acertadas de un hecho. Estoy sentado y pienso: estoy sentado. De acuerdo con la teoría de la correspondencia, el pensamiento es verdadero porque corresponde al hecho de estar sentado (Hetherington, 2007). Ahora bien, con demasiada frecuencia es difícil visualizar los hechos representados por las verdades. Por ejemplo, ¿cómo representamos las partículas fundamentales (o quarks)?.

LECTURA 4: FUENTES DEL CONOCIMIENTO

Históricamente, la filosofía ha tomado en cuenta dos fuentes del conocimiento: la observación (a través de los sentidos e incluso utilizando medios creados para ello) y la razón (inferir, pensar a través de implicaciones, intuir de forma puramente racional). Una tercera fuente es el pragmatismo

Los empiristas consideran que en última instancia todo conocimiento es de tipo observacional; pues, a fin de cuentas, toda buena sustentación de una opinión es siempre de tipo observacional. Por tanto, nuestro saber llega con la información que nuestros sentidos recopilan. Se puede ser un empirista selectivo. Por ejemplo, ser un empirista sobre la moralidad, pero no con respecto a las matemáticas.

Los racionalistas consideran que nuestro intelecto posee la capacidad de proporcionarnos conocimiento por sí solo. Por tanto, aseveran que por lo menos una parte del conocimiento es no observacional.

Los racionalistas están de acuerdo con los empiristas en que la gente obtiene conocimiento cuando procede racionalmente, es decir, cuando utiliza su razón. Pero los racionalistas, a diferencia de los empiristas, piensan que existe cierto conocimiento que es puramente racional. Esto es, alegan que algunas opiniones están bien sustentadas de una forma que es completamente intelectual o razonada, o proviene de un tipo no observacional de apercepción (*insight*, también intuición). Para hacer que una opinión se transforme en conocimiento de este modo, bastaría con el intelecto; no harían falta los sentidos.

Los filósofos han llamado al conocimiento no observacional conocimiento *a priori*. Se trata de un conocimiento sustentado de manera no observacional, supuestamente por vía de la razón pura. Su sustentación tiene lugar incluso con anterioridad a cualquier observación que pudiera apoyarlo. Y a este tipo de conocimiento se lo contrasta tradicionalmente con el conocimiento *a posteriori*, cuyo elemento sustentador es, en gran medida, la naturaleza observacional. Una opinión es conocimiento *a posteriori* sólo después de que un número suficiente de observaciones pertinentes pasan a formar parte de su sustentación. Emmanuel Kant (1724–1804), en su *Crítica de la Razón Pura*, decía que la razón y la observación deben aliarse para que una persona posea conocimiento. El pensamiento sin la sensación es vacío. La sensación sin el pensamiento es ciega. Tanto uno como otro son necesarios.

La validez del conocimiento pragmático se basa en la experimentación. Si el experimento puede explicarse, si los hallazgos del experimento concuerdan con la situación o alcanzan los resultados deseados, es un conocimiento válido. El pragmatismo subraya también la naturaleza experimental de ciertas formas de experiencia: los individuos prueban algo y descubren que funciona o que falla. La naturaleza de la experiencia cotidiana y del conocimiento cotidiano es que aprendemos mediante el experimento.

Fuente: Hetherington, S. (2007): *¡Filosofía! Una Breve Introducción a la Metafísica y a la Epistemología*, Alianza Editorial, Madrid. Scheffler, I. (1965): *Conditions of Knowledge*, University of Chicago Press, Chicago.

Justificada (bien sustentada). Nuestra creencia está bien justificada si y solo si podemos producir buenas razones en su apoyo; y nuestras razones son realmente buenas (en las más estrictas normas filosóficas) si y solo si podemos producir un argumento concluyente o formalmente válido, que relaciona la creencia con un punto de partida que no es cuestionado (Toulmin, 1976). Tener buenas razones para una creencia significa que se ha obtenido esa creencia por un proceso seguro. Empleamos bien la

razón cuando, como dijera Hume, adecuamos la creencia a las pruebas y a los argumentos. De acuerdo con algunos filósofos, una persona con una opinión bien sustentada está pensando o hablando como debiera, lo que se manifiesta no solo por el hecho de preocuparse por la verdad, sino también por el de comprobar cuidadosamente la evidencia de que se dispone (Hetherington, 2007). Tiene una enorme dificultad justificar la validez del conocimiento. El concepto de buena sustentación es objetivo, pero no absoluto. Una creencia puede descansar sobre una sustentación objetivamente buena, aunque esa sustentación haya podido ser mejor, o, si es el caso, peor. La justificación de una creencia tiene que poder responder con buenos argumentos a todos los críticos que piensen que la creencia no es justificada.

En filosofía se ha trazado una distinción entre razonamiento inductivo y razonamiento deductivo. Gran parte del razonamiento real es una mezcla de los dos, si bien el de la lógica y el de las matemáticas es casi siempre deductivo, mientras que en la ciencia y en la vida ordinaria es más habitual el inductivo. La deducción consiste en sacar las conclusiones que se siguen a ciencia cierta de sus premisas. Se dice que es “preservadora de la verdad” porque la verdad de la conclusión de una inferencia deductiva lógica está garantizada siempre que sus premisas sean verdad.

La inducción consiste en reunir hechos concretos y luego generalizar a partir de ellos; es decir, está basado en la experiencia, en la observación y en el experimento (lectura 4). El razonamiento inductivo como tal no es preservador de la verdad, porque las generalizaciones que hace son abiertas o universales, es decir, se refieren al futuro, al pasado y al presente. Por supuesto, las generalizaciones basadas en la observación de hechos a menudo son ciertas, que sepamos. La cuestión es que la verdad de un número finito de afirmaciones acerca de cosas concretas es compatible con la posible falsedad de una generalización universal basada en ellos (Teichman y Evans, 2010).

Hoy día los filósofos dudan cuando se trata de inferir, a partir del hecho de que una creencia es cierta y está bien sustentada, que se trata de conocimiento. Esa duda se debe, en parte, a un artículo de Gettier (1963), donde propuso varios ejemplos en los que intuitivamente nos damos cuenta de que alguien no sabe algo en realidad, aunque la creencia de esa persona sea verdadera y justificada. Por ejemplo, he quedado con mi amigo Pedro en su casa y, cuando llego, veo a través de la ventana que está sentado en la cocina. Lo cierto es que no es Pedro a quien veo, sino a su hermano gemelo; Pedro está en otra habitación. Mi creencia de que Pedro está en casa es verdadera; estoy seguro de haberlo visto, así que tengo un buen motivo para creerlo, pero es erróneo decir que yo sé que está en casa: no lo sé, porque la persona a la que he visto no es Pedro.

Una buena sustentación no prueba que una creencia sea verdadera. Solamente una sustentación contundente (perfecta) lo podría lograr, pero raramente se consigue. La buena sustentación suele dejar abierta la posibilidad de que aquello que sustenta como verdadero sea, a pesar de todo, falso; esa es precisamente la razón de que la sustentación sea sólo buena y no contundente. A veces, dicha posibilidad se materializa, y la opinión bien sustentada se revela como falsa. A pesar de todo, los filósofos consideran que la sustentación no tiene que ser contundente, ya que eso sería poco realista y carecería de sentido. Aceptan, en lugar de ello, alguna forma de falibilismo. Esto es, aceptan la posibilidad de disponer de una buena sustentación para la verdad de una creencia. Lo meramente bueno también es bueno: sólo que no es aún más bueno. No alcanza a ser

una sustentación contundente (que sería la única capaz de proporcionar certeza objetiva) y excluye de plano toda posibilidad de que la opinión sustentada sea falsa (Hetherington, 2007). El no poder probarse que una creencia sea verdadera conduce a una concepción gradacionista y falible del conocimiento.

Graduado. Un hecho puede ser conocido de mejor o peor manera. En general, conocer un hecho de manera normal consiste en conocerlo de forma aceptablemente sólida y segura, lo cual puede implicar que la creencia respectiva esté bien comprobada y ha sido obtenida de un modo razonablemente fiable. De esta forma, la creencia de alguien será conocimiento mejor o peor de un hecho particular dependiendo de la calidad de la sustentación que tenga para ella. Esta es una concepción gradacionista del conocimiento porque afirma que puede haber grados de conocimiento, e incluso grados de conocimiento de un mismo hecho. No entra en contradicción con la objetividad del conocimiento. El conocimiento sigue siendo conocimiento de los hechos. Es más, cualquier sustentación que vaya aparejada al conocimiento puede ser objetivamente mejor o peor. La concepción gradacionista niega, en cambio, que el conocimiento sea absoluto.

El gradacionismo defiende la idea de que existe el conocimiento falible. En efecto, el conocimiento falible de un hecho no es el mejor conocimiento posible de ese hecho. Es un conocimiento en el cual la sustentación no es perfecta ni irrefutable: es mejor de lo que hubiera podido ser, pero no tan buena como quisiéramos. El conocimiento infalible requeriría un tipo de sustentación incapaz de confundirnos, a saber, que nos condujera inevitablemente a formarnos una opinión correcta basándonos en ella. Se trata de un conocimiento que pocas veces, si alguna, nos es dado alcanzar. En suma, puede haber distintas formas de conocer un hecho particular; y nada evita que haya diferencias en la calidad –en los grados– del conocimiento de ese hecho. El gradacionismo es tolerante al respecto.

Falible. El principio de verificación se propuso en la década de los cincuenta para justificar lo verdadero. Si una afirmación podía ser verificada, ello la convertía en conocimiento en virtud de la exclusión de conjeturas sobre opiniones o preferencias. Pero este principio se encontró, entre otros, con el problema lógico de la inducción: aunque las pruebas de verificación podían aplicarse a cualquier generalización que se deseara, el estatus de esta resultaba siempre incierto, pues cualquier observación posterior podía contradecirla. Así pues, la verificación garantizaba muy poco, ya que, en principio, cualquier observación futura podía acabar con una generalización (lectura 5).

Imagine a dos personas que hubiesen vivido en Inglaterra hace 300 años, uno de los cuales hubiese visto a cien cisnes, y la otra a ciento diez cisnes, y donde todos los cisnes avistados hubiesen sido blancos. Ambas personas habrían tenido evidencia que proporcionaba lo que era, objetivamente hablando, una buena sustentación para la verdad de la creencia de que todos los cisnes son blancos. Una de ellas habría poseído una sustentación ligeramente mejor que la otra. Sin embargo, esa creencia se demostró que era falsa, ya que en 1697 Willem de Vlamingh en una expedición a Australia descubrió que los cisnes autóctonos de ese continente son negros (Hetherington, 2007). Observo cien cisnes y todos son blancos, infiero que todos los cisnes son blancos. Aunque no es posible demostrar que la conclusión sea cierta, es posible demostrar que es falsa. Para esto es suficiente un contraejemplo.

LECTURA 5: EL PAVO INDUCTIVISTA VERSUS EL CISNE NEGRO

Un ejemplo del problema de la inducción muy interesante, aunque bastante truculento, lo constituye la explicación de la historia del pavo inductivista, debido a Bertrand Russell. Este pavo descubrió que, en su primera mañana en la granja avícola, comía a las nueve de la mañana. Sin embargo, siendo como era un buen inductivista no sacó conclusiones precipitadas. Esperó hasta que recogió una gran variedad de circunstancias, en miércoles y en jueves, en días fríos y calurosos, en días lluviosos y en días soleados. Cada día añadía un nuevo enunciado observacional a su lista. Por último, su conciencia inductivista se sintió satisfecha y efectuó una inferencia inductiva: “siempre como a las nueve de la mañana”. Pero ¡ay! Se demostró de manera indudable que esa conclusión era falsa cuando, la víspera de navidad, en vez de darle la comida, le cortaron el cuello. Una inferencia inductiva con premisas verdaderas ha llevado a una conclusión falsa.

Demos un paso más y pensemos en el aspecto más inquietante de la inducción: el retroaprendizaje. Pensemos que la experiencia del pavo, más que no tener ningún valor, puede tener un valor negativo. El animal aprendió de la observación, como a todos se nos dice que hagamos (al fin y al cabo, se cree que es el método científico). Su confianza aumentaba a medida que se repetían las acciones alimentarias, y cada vez se sentía más seguro, pese a que el sacrificio era cada vez más inminente. Consideremos que el sentimiento de seguridad alcanzó el punto máximo cuando el riesgo era mayor. Pero el problema es incluso más general que todo eso, sacude la naturaleza del propio conocimiento empírico. Algo ha funcionado en el pasado, hasta que ... pues, inesperadamente, deja de funcionar, y lo que hemos aprendido del pasado resulta ser, en el mejor de los casos, irrelevante o falso y, en el peor, brutalmente engañoso.

El problema filosófico de la inducción es el nombre del Cisne Negro. Todos los cisnes son blancos. Esta afirmación fue irrefutable durante siglos. Esto fue así hasta 1697, cuando Willem de Vlamingh vio por primera vez un cisne negro durante una expedición a Australia. Los cisnes negros se han convertido en el símbolo de lo improbable. Un Cisne Negro es un suceso con los tres atributos siguientes. Primero, es una rareza, pues habita fuera del reino de las expectativas normales, porque nada del pasado puede apuntar de forma convincente a su posibilidad. Segundo, produce un impacto tremendo. Tercero, pese a su condición de rareza, la naturaleza humana hace que inventemos explicaciones de su existencia «después» del hecho, con lo que se hace explicable y predecible. Los cisnes negros surgen cada vez con más frecuencia y tienden a ser más relevantes.

El problema del cisne Negro: ¿cómo podemos conocer el futuro teniendo en cuenta nuestro conocimiento del pasado; o de forma más general, ¿cómo podemos entender las propiedades de lo desconocido (infinito) basándonos en lo conocido (finito)?

Popper expuso el mecanismo de las conjeturas y las refutaciones, que funcionan como sigue: se formula una conjetura (osada) y se empieza a buscar la observación que demostraría que estamos en un error. Si no logramos demostrar que es falsa no la convierte en verdadera, pero permite tener una buena sustentación de la verdad. Esta es la alternativa a nuestra búsqueda de casos confirmatorios.

Taleb, N. N. (2008): *El Cisne Negro*, Paidós, Barcelona.

Karl Popper (1902–1994), en su obra traducida al inglés en 1958 como *The Logic of Scientific Discovery*, preguntó cómo se distinguen las opiniones científicas de las que no lo son. Su respuesta fue elegante: las opiniones científicas son comprobables (*testable*), falsables (*falsifiable*). Una hipótesis (o conjetura) es científica si se somete a una prueba que permita falsarla (refutarla). Por tanto, una teoría puede ser comprobada experimental y observacionalmente, pero no existe una sustentación definitiva para confirmarla. Lo mejor que podemos hacer, según Popper, para comprobar (*to test*) una teoría es tratar de demostrar que es falsa; si no logramos demostrar que es falsa, podremos reconocer que la teoría ha cumplido con lo máximo que podíamos exigirle hasta el presente. Esto quizás no sea suficiente para convertir la teoría en verdadera (de acuerdo con la teoría de la verdad como correspondencia, que fue la aceptada por Popper). Pero si se hará que esté más corroborada hasta ese momento. Seguirá estando viva, lo cual constituye hasta cierto punto una buena sustentación de la verdad. El que no sea falsa no la convierte en verdadera. Pero sí le permite obtener una especie de sustentación real (y no solo aparente) en caso de que haya superado hasta ahora las comprobaciones a las que ha sido sometida (Hetherington, 2007).

No hay una verdad científica. Sólo existen conocimientos científicos, siempre parciales, relativos y provisionales. Pero esto, lejos de debilitar las ciencias, las fortalece. Las ciencias no progresan de «certeza en certeza», como pretendía Descartes, sino por «conjeturas y refutaciones», como propuso Popper, y así progresan mejor. Poner el acento en el carácter falible de la humanidad no equivale a renunciar al progreso del conocimiento, sino más bien a concederse los medios para comprenderlo y proseguirlo –sacando, a falta de certeza positiva, enseñanzas de nuestros errores– (Comte–Sponville, 2012).

Así pues, las ciencias no progresan pasando del hecho a la ley (es decir, de una verdad particular a una verdad más general: inducción), sino deduciendo la falsedad de una teoría a partir de enunciados empíricos singulares (falsación). Mil hechos no probarán nunca la verdad de una teoría (porque el mil uno puede invalidarla); y un solo hecho basta para probar su falsedad. Esta asimetría permite que el conocimiento avance en la dirección inductiva (hacia teorías cada vez más potentes), pero mediante inferencias deductivas (no inductivas). Por eso no tenemos certeza, científicamente hablando, más que de los errores que se han refutado. Es lo que impide erigir cualquier teoría científica como un absoluto, sin permitir a pesar de ello dudar de su progreso científico (Comte–Sponville, 2012). De esta forma, el problema lógico asociado a la verificación queda substituido por la promesa de una generalización cada vez más fiable (aunque, en última instancia, nunca cierta) que resistiera más y más intentos de falsación.

Conocimiento versus información. Conocimiento es lo que la gente sabe, información es lo que la gente comunica. El conocimiento es personal; para que el conocimiento de una determinada persona sea útil para otra debe ser comunicado de tal manera que pueda ser interpretable y accesible para ese otro individuo. No todo lo que se sabe se puede comunicar.

Los economistas normalmente han creído que la «reproducción» del conocimiento es sinónimo de comunicación, es decir, codificación, transmisión y recepción de información, por lo que no establecen diferencias entre información y conocimiento. Tener información es tener conocimiento, y aquel que tiene conocimiento será capaz de expresarlo como una información transmisible que, una vez recibida por otro, reproducirá el conocimiento original. Si fuera simétrico el proceso inverso de codificación del conocimiento, es decir, si la «descodificación» fuera aparentemente tan clara como la codificación, sería apropiado ignorar las distinciones entre información y conocimiento.

No obstante, aunque el conocimiento codificado pueda ser transmitido sin pérdida de integridad una vez que las reglas sintácticas (códigos de interpretación) requeridas para descifrarlo sean conocidas (Kogut y Zander, 1992), hay que tener en cuenta que los individuos de la organización se involucran en un proceso lingüístico propio (von Krogh *et al.*, 1994) que dificulta la comprensión (incluso, la valoración) del conocimiento codificado, no sólo por parte de agentes externos, sino también por el personal de la empresa alejado del sitio donde se generó el conocimiento. Esto significa que el lenguaje interno de comunicación en una organización nunca es «lenguaje de calle» (Nelson y Winter, 1982), sino que suele ser específico al sitio. Ello se debe a que el receptor y el transmisor no comparten el mismo contexto, por lo que surgen dificultades en su traslado a lugares alejados de la fuente. Los individuos no reciben el conocimiento pasivamente, sino que lo interpretan de forma activa para ajustarlo a sus

propias situaciones y perspectivas. Así, lo que tiene sentido en un lugar, puede adquirir otro significado, o incluso perderlo, cuando se comunica a personas situadas en un entorno diferente (Nonaka y Takeuchi, 1995). Por ejemplo, Fruin (1997), al estudiar el funcionamiento de la fábrica Yanagicho de Toshiba, observó que la manera de detallar los proyectos o dibujos de ingeniería entre las fábricas de Toshiba es diferente. Los dibujos son similares, pero no explican lo que se supone que son los términos comunes en cada fábrica. De esta forma, los dibujos técnicos de Yanagicho no podían interpretarse con éxito en otras fábricas de Toshiba en Francia, Estados Unidos, o incluso Japón. Dicho de manera más simple, no es adecuado traducir la documentación y las reglas de una fábrica sin el material de fondo que informa sobre ello (a menudo conocimiento no codificado de naturaleza tácita).

Desafortunadamente, el receptor de conocimiento codificado suele necesitar un conocimiento sustancial para reconstruir esta información y convertirla en conocimiento útil. Los defectos en los conocimientos y la experiencia del receptor, la incapacidad para encontrar representaciones adecuadas para el conocimiento y la inevitabilidad de los errores de transcripción garantiza que hasta los esfuerzos más sencillos para reproducir el conocimiento quedarán cortos de su objetivo.

La codificación de conocimiento funciona, pero imperfectamente (Steinmueller, 2002). Ello se debe a que todo conocimiento tiene una naturaleza explícita (saber qué) y otra tácita (saber cómo). Las habilidades prácticas son de naturaleza tácita y resultan difíciles de expresar: es el «saber cómo» que, en gran medida, significa lo mismo que poder hacer: saber cómo montar en bicicleta, es un buen ejemplo. «Saber cómo» es distinto de «saber qué». «Saber cómo» es tener una habilidad, «saber qué» es poseer una información. Poseer información significa normalmente creer esa información. Pero tener una habilidad no significa que uno crea en la habilidad, porque no tiene sentido creer en ella. Por tanto, si se acepta que las habilidades prácticas son un tipo de conocimiento, se sigue que no todo conocimiento implica una creencia (Teichman y Evans, 2010).

Conviene destacar que es difícil separar la naturaleza tácita de la explícita cuando trata de ponerse en práctica el conocimiento. Por ello, Polanyi (1966) no percibe los conocimientos tácito y explícito como dos tipos diferentes de conocimiento, sino que interpreta que el conocimiento tácito rodea todo tipo de conocimiento explícito.

A diferencia de la información que se acumula en las cosas, el conocimiento se acumula en los humanos y en redes de humanos. Por tanto, existe un límite a la capacidad de conocimiento que puede acumular un individuo. En el plano personal, acumular conocimiento es difícil porque aprender requiere experiencia. El aprendizaje es práctico y social. Las personas aprenden de las personas y de la experiencia. La acumulación de conocimiento tiene, por tanto, un sesgo geográfico.

1.3. Conocimiento explícito

El conocimiento explícito se concibe como información, básicamente cuantitativa, articulada que permite a los individuos expresar sus opiniones y creencias y confrontarlas. El conocimiento explícito normalmente está codificado. Los códigos son sistemas en los que se organizan los signos y que determinan como estos signos se relacionan los unos con los otros. La codificación es el proceso de conversión del

conocimiento en mensajes que pueden ser procesados como información. Esta codificación conlleva la pérdida de conocimiento, pues durante el proceso siempre suele quedar, residualmente, algo del mismo en el lugar de origen. Además, solo se puede codificar el conocimiento explícito. Una vez codificado, el conocimiento se vuelve independiente del individuo o grupo que lo posee, por lo que puede almacenarse, reproducirse y transferirse entre diversos individuos, grupos y organizaciones sin incurrir en nuevas pérdidas. El conocimiento explícito codificado se almacena en diferentes soportes, tales como libros, revistas, *software*, manuales de procedimiento o en los derechos de propiedad industrial.

Para construir una ventaja competitiva es esencial que el conocimiento pueda compartirse dentro de la organización. El resultado es la «paradoja de reproducción». Con el fin de utilizar el conocimiento para construir una ventaja competitiva tenemos que reproducirlo y la reproducción es mucho más fácil si el conocimiento está codificado. El potencial de reproducción es máximo si el conocimiento puede trasladarse (en planos de construcción o en disquete de ordenador), interpretarse (susceptible de reducción a ecuaciones, símbolos o términos comúnmente aceptados) o asimilarse (independiente de cualquier contexto cultural) con relativa facilidad (Hamel *et al.*, 1989). La desventaja del conocimiento codificado es que resulta difícil impedir que llegue a mano de los competidores, incluso por medios no siempre honestos, como el espionaje industrial. Promover la reproducción interna al tiempo que se limita la externa supone un desafío clave para las empresas (Kogut y Zander, 1992).

No conviene olvidar que una vez codificado el conocimiento continúa permaneciendo en los individuos que lo desarrollaron o que lo utilizaron y, por tanto, se traslada con ellos cuando abandonan la empresa.

Por otra parte, el conocimiento explícito también está insertado en objetos físicos, como los productos y las máquinas. Los competidores tienen fácil acceso a los productos, ya que pueden adquirirlos sin dificultad en el mercado y, por tanto, pueden obtener el conocimiento que los sustenta aplicando la ingeniería inversa. Sin embargo, las máquinas desarrolladas por la empresa para su propio uso están normalmente fuera de su alcance. No solamente los mensajes sino también la mayoría de los objetos están compuestos de información.

Propiedades del conocimiento explícito (información). Por conocimiento entendemos una «unidad o pieza de conocimiento», es decir, el conocimiento total necesario para lograr un resultado. Por tanto, el conocimiento es un bien indivisible (un activo fijo), ya que se posee o no se posee, no tiene valor ni utilidad conocer una parte de la unidad (o pieza) de conocimiento. Entre sus características más significativas destacamos las siguientes.

El proceso de creación del conocimiento es un coste irrecuperable, pues, si fracasa o no resulta útil, la empresa sólo puede recuperar una mínima parte (o ninguna) de las inversiones realizadas en su generación.

El conocimiento explícito es un bien público, ya que es a) no excluible y b) no rival. Samuelson (1954) puso mucho énfasis en la exclusión: la posesión de información por parte de un individuo no impide a otros utilizarla. Los bienes de cuyo uso se podían excluir a los individuos se consideraban bienes privados. Los economistas han

reconocido que es necesario proporcionar algún incentivo a los factores económicos si se trata de producir conocimiento y hacerlo disponible para un intercambio. Un segundo atributo de los bienes es la rivalidad (denominada a veces sustrabilidad) que indica que el uso por parte de una persona sustraerá los bienes disponibles para otros (Ostrom y Ostrom, 1977). Por tanto, «no rival» significa que su consumo por una persona no reduce la cantidad disponible para otros, es decir, el uso del conocimiento no merma el *stock* disponible. De esta forma, tiene un coste de oportunidad nulo.

El conocimiento explícito es asimétrico. Cada individuo sabe cosas que los demás no saben. Desde la perspectiva del individuo, un avance en su conocimiento consiste en saber algo que antes no sabía (Stiglitz y Greenwald, 2014). Por ello, en muchos casos, el conocimiento que un individuo obtiene, ya es conocido por alguien más.

El conocimiento que ha sido producido dentro de un contexto determinado para un propósito particular no tiene por qué tener un valor general y, por tanto, es idiosincrásico al sitio. Por otra parte, la mera existencia del conocimiento no garantiza su disponibilidad para la gente que lo necesita en el lugar apropiado y en el momento oportuno.

El intercambio de conocimiento contribuye a aumentar la base de conocimiento de las entidades participantes. Esta circunstancia la expresa magníficamente George Bernard Shaw, premio Nobel de Literatura en 1925, en una de sus célebres frases: “Si tú tienes una manzana y yo tengo una manzana e intercambiamos las manzanas, entonces tanto tú como yo seguiremos teniendo una manzana. Pero si tú tienes una idea y yo tengo una idea e intercambiamos las ideas, entonces ambos tendremos dos ideas”. De hecho, cuanto más gente comparta conocimiento útil, mayor es el bien común que poseen.

El conocimiento es un bien duradero, que se puede usar una y otra vez sin que se deteriore. Es decir, su uso no lo minusvalora para posteriores o concurrentes aplicaciones. Al contrario, su uso proporciona nuevas ideas y líneas de investigación, así como mejoras (aprendizaje por el uso) y nuevas aplicaciones, lo que aumenta su valor. El fármaco Gleevec de Novartis, se desarrolló al principio para combatir una leucemia infrecuente que afecta a unos pocos pacientes en todo el mundo. Después de que se desarrollara y se introdujera en el mercado, se descubrió que también era eficaz para combatir otras seis enfermedades; al final alcanzó unas ventas de 3.700 millones de dólares (Dávila y Epstein, 2014). A su vez, cuanto más se use el conocimiento mayor confianza genera a la par que mejora su comprensión, lo que elimina muchas de sus incertidumbres y facilita su adopción por las personas aversas al riesgo. La depreciación del conocimiento se debe al olvido y a la obsolescencia y no al desgaste por el uso.

El conocimiento es acumulativo, ya que la principal fuente del conocimiento, las ideas, tienen un potencial de crecimiento ilimitado. Una idea conduce a otra, y esta última a otra más y así hasta que alguien combina todas en algo nuevo. Parafraseando a Isaac Newton: “Si he logrado ver más lejos, ha sido porque me subí a hombros de gigantes”. El conocimiento se puede combinar con otro conocimiento (incluso con conocimiento fallido) para producir nuevo conocimiento.

La transmisión de conocimiento a terceros proporciona un valor añadido al transmisor, derivado de los efectos del aprendizaje por la enseñanza y potenciado por retroalimentaciones, ampliaciones y modificaciones de su base cognitiva. En suma, el

conocimiento presenta rendimientos marginales crecientes, al contrario que los recursos físicos. De esta forma, la utilización de conocimiento (escalabilidad) sólo está limitada por el tamaño del mercado, ya que su uso contribuye a aumentar el valor del mismo.

Los productos del conocimiento difieren mucho de los productos convencionales. El conocimiento tiene un coste fijo elevado y un coste marginal bajo (a menudo, insignificante). En otras palabras, a pesar de que resulta muy costoso producir el primer bien de conocimiento, una vez obtenido, su reproducción supone un coste muy bajo (Arrow, 1996). No hay límites naturales a la producción de copias adicionales. De esta forma, sin una protección especial, cualquiera que tenga acceso al conocimiento puede reproducirlo a un coste reducido.

La apropiación de las ganancias que el conocimiento proporciona exige excluir de su uso a los no propietarios. La manera más directa de evitar la reproducción consiste en asignar derechos de propiedad al conocimiento. Al crear «propietarios legítimos» del conocimiento se establecen las condiciones iniciales para el funcionamiento de un mercado. Sin embargo, las prácticas de no rivalidad y los bajos costes marginales de reproducción dificultan la tarea de velar por los derechos de propiedad. Ninguna protección legal puede convertir algo tan intangible como el conocimiento en un bien completamente privado. Por tanto, rara vez puede excluirse totalmente a los no propietarios del reparto de las ganancias que genera. Esta exclusión parcial, o no exclusión, da lugar a la existencia de beneficios para los no propietarios y a la ausencia de control, en el estricto sentido legal del término, sobre la propiedad del conocimiento (Lev, 2001). En suma, la producción del conocimiento conlleva externalidades positivas. Por tanto, genera un «efecto derrame» (*spillover*), pues se filtra a terceras partes a través de diferentes mecanismos, tales como la movilidad de los trabajadores, las publicaciones técnicas, la ingeniería inversa, los contactos informales, el espionaje o la piratería, entre otros. De hecho, el uso mismo del conocimiento en cualquier forma productiva necesariamente lo revela a terceros, por lo menos en parte (Arrow, 1962a). Tampoco puede retrotraerse su posesión a quien haya tenido acceso al mismo, como mucho se podrá impedir su uso. Además, la apropiación indebida (espionaje industrial) del conocimiento por terceros resulta difícil de detectar, al ser un bien invisible (intangible) que, por otra parte, no aparece en el balance de situación de la empresa.

El conocimiento es un «bien de experiencia», ya que los consumidores tienen que pasar por la experiencia de probarlo para poder evaluarlo (Nelson, 1970). Una persona que posee conocimiento, pero que teme su expropiación subrepticia, no está dispuesto a proporcionar dicho conocimiento para que el cliente pueda estudiarlo antes de un acuerdo de compra. Esto plantea al comprador dificultades para evaluar la calidad del conocimiento (Stiglitz y Greenwald, 2014). Por tanto, está sujeto a la «paradoja de la información» (Arrow, 1962a): el precio del conocimiento es difícil de determinar en la medida en que el comprador tiene, en principio, que efectuar su adquisición sin una valoración adecuada del mismo, ya que, si el conocimiento le fuera revelado con antelación, entonces lo habría adquirido sin coste alguno.

1.4. Conocimiento tácito

Gran parte del conocimiento no puede ser enseñado formalmente, o comunicado por medio del mensaje. Es el resultado del aprendizaje, pero del aprendizaje en forma de experiencia personal. La experiencia nunca puede transmitirse; produce una

transformación en los individuos –con frecuencia una transformación sutil– y no puede separarse de ellos (Penrose, 1959). El conocimiento tácito es un conocimiento específico para un contexto determinado en el que se genera y en el que es posteriormente aplicado. Está orientado hacia el cumplimiento de tareas específicas y se deriva de un aprendizaje práctico. Tácito significa no articulado y no codificado. Se aprende por el mecanismo «aprender haciendo» y no requiere el uso del lenguaje. El conocimiento tácito es un conocimiento personal por naturaleza, pertenece al individuo, emana de sus vivencias y experiencias. Al no estar articulado, este conocimiento no se puede codificar. De acuerdo con Polanyi (1966), podemos conocer mucho más de lo que podemos expresar. El conocimiento tácito individual surge, pues, de una íntima familiarización laboral alimentada con años de esfuerzo y reside en las relaciones del individuo con determinada tarea y herramientas dentro de un contexto, no en un procedimiento ni en cualquier conjunto de reglas. Se trata de un núcleo de conocimientos muy importante, aunque desorganizado, informal y relativamente inaccesible, que no resulta eficaz para la instrucción directa (Wagner y Sternberg, 1985). Tampoco puede recibir el apelativo de científico, por no contener reglas universales de aplicación general, pues se refiere al conocimiento de las circunstancias particulares de tiempo y lugar.

El conocimiento tácito es un conocimiento difícil de transportar, memorizar, recombinar y aprender. Puede ser individual y colectivo. Consiste en la capacidad de llevar a cabo acciones, aun sin poder explicarlas.

	Tácito individual	Explícito
Contenido	No codificado	Codificado
Articulación	Difícil	Fácil
Localización	Cerebro personas	Documentos, artefactos
Comunicación a terceros	Difícil	Fácil
Medios de transmisión	Contacto directo	Tecnologías información
Almacenamiento	Difícil	Fácil
Propiedad	Persona	Organización
Protección	Política de RRHH	Propiedad industrial
Imitación	Difícil	Fácil

Figura 2: Conocimiento explícito versus tácito

Conocimiento tácito individual (know-how). El *know-how* es la habilidad práctica acumulada que permite a alguien hacer algo sin dificultad y con eficiencia (von Hippel, 1988). Se trata de un conocimiento práctico subjetivo que adquiere una persona por haber estado inmersa en una actividad durante un largo período de tiempo (figura 2). Este conocimiento está, pues, profundamente enraizado en la acción y en el cometido personal dentro de un determinado contexto (Nonaka, 1991). Su carácter específico (o idiosincrásico) hace que pierda valor (incluso todo su valor) cuando se utiliza en otras aplicaciones y relaciones distintas a la de su uso en la tarea donde se desarrolló (Liebeskind, 1997). Por ejemplo, el conocimiento tácito de Miguel Indurain para mantenerse en equilibrio en la bicicleta no tiene ningún valor en su trabajo como articulista de un periódico deportivo. La transferencia de conocimiento tácito entre personas es lenta, costosa e incierta.

En consecuencia, el conocimiento tácito individual se revela mediante la aplicación y se adquiere a través de la práctica. Es decir, se crea por un proceso de *learning-by-doing* (Arrow, 1962b), que significa que el conocimiento acerca de cómo realizar eficientemente una tarea se adquiere y aprende con la experiencia acumulada mediante la repetición de la actividad. Aunque se describa muy detalladamente en un documento la forma de montar en bicicleta, esta habilidad, por mucho que se estudie y se memorice el documento, no se adquirirá hasta después de haberla practicado durante un cierto tiempo. Debe resaltarse que es la práctica la que hace la perfección, no el transcurso del tiempo. El paso del tiempo sin la realización de la actividad no conlleva ningún aumento del *know-how*. Más bien lo contrario, contribuye a deteriorarlo por su falta de uso y la correspondiente pérdida de habilidad.

El *know-how* se genera a través de largos períodos de experimentación y realización de una tarea, durante los cuales el individuo desarrolla una habilidad y una aptitud para hacer juicios intuitivos sobre cómo ejecutarla de forma satisfactoria (Choo, 1998). Para poder valorarlo adecuadamente, sólo debe recordarse lo mucho que hay que aprender sobre cualquier ocupación después de haber terminado la formación teórica, la gran parte de tiempo de la jornada laboral que se dedica a asimilar determinadas rutinas y lo valioso que es en todas las profesiones el conocimiento de la gente, de las condiciones locales y de las circunstancias especiales (Hayek, 1945). Desde esta perspectiva, el conocimiento tácito resulta difícil de transferir dentro de la organización y, sobre todo, de una organización a otra, lo que, por otra parte, supone una ventaja: dificulta su imitación por los competidores. A mayor diferencia en el contexto de las organizaciones entre las que se quiere transferir, más difícil resulta la transferencia y menos valor tiene (Garud, 1997).

Conocimiento tácito colectivo. El conocimiento tácito colectivo es aquél que el individuo comparte con otros, pasando, de este modo, a formar parte del sistema social en el que se desarrolla. El conocimiento tácito colectivo depende de las experiencias compartidas y de las interacciones entre los individuos (Penrose, 1959). Acumular conocimiento colectivo requiere tiempo para ensayar, confianza y un lenguaje común. Este conocimiento existe entre los individuos y no dentro de ellos, está siempre vinculado a la acción, proviene de ella y conserva el potencial para que otros lo utilicen (Dixon, 2000). Puede estar presente en un grupo de trabajo (conocimiento de equipo) o formar parte del conocimiento más general de la empresa (rutina organizativa). El conocimiento tácito colectivo desaparece cuando se rompe (se disuelve) el sistema social que lo sustentan.

Conocimiento de equipo. Los equipos de trabajo poseen dos clases de conocimiento. El primero, en mayor proporción de naturaleza explícita, se refiere al conocimiento que el equipo en conjunto posee en relación con la comprensión de la tarea que tiene asignada. El segundo, de naturaleza esencialmente tácita, está relacionado con la manera en que los miembros del equipo se relacionan y trabajan juntos en la ejecución de una tarea.

El conocimiento de equipo de naturaleza fundamentalmente explícita se refiere a los conocimientos que cada uno de sus miembros, y el equipo en general, poseen respecto a la tarea encomendada. De este modo, los miembros de un equipo como un todo pueden proporcionar una explicación más completa de por qué un proceso (o un producto) se comporta de la forma en que lo hace.

En el caso de un equipo de trabajo encargado de realizar una tarea compleja, cada individuo puede estar especializado en una o varias actividades que se acoplan con el total de actividades que realizan el resto de miembros para completar la tarea. En la realización de la tarea, el equipo necesita el conocimiento y participación de todos los miembros.

Un equipo de trabajo, a lo largo del tiempo, acaba desarrollando un conocimiento tácito sostenido por un lenguaje común y un determinado comportamiento. Todo ello es producto de unas interacciones únicas y recurrentes dentro de un contexto específico. Este conocimiento tácito se refiere al conocimiento de los miembros del equipo acerca de cómo trabajar juntos y está incrustado en las relaciones que desarrollan, siendo, por tanto, difícil de imitar. Como la mayoría del trabajo es una operación colectiva y cooperativa, también una parte del conocimiento resultante es básicamente colectivo, retenido por los individuos que comparten el equipo de trabajo (Brown y Duguid, 1991). Se trata de un conocimiento muy etéreo y amorfo, que contribuye a crear la química apropiada del equipo, permitiéndole alcanzar una conjunción armoniosa apoyada en las experiencias compartidas en el pasado, que han establecido un sistema de comunicación, sumamente detallado y específico, que subyace al comportamiento de los individuos (Nelson y Winter, 1982). Dicha armonía permite la lectura sutil de signos entre los trabajadores compenetrados, que no es traducible a medidas concretas y verificables, pero que facilita la realización eficiente de las tareas (Ouchi, 1984). Este conocimiento lo desarrolla un equipo deportivo cuando decimos: ¡qué bien juega! No nos referimos a la habilidad de cada uno de los integrantes del equipo, sino al acoplamiento entre sus miembros, que les permite actuar de forma coordinada y sistémica, es decir, como un todo. En suma, el conocimiento del equipo es mayor que la suma del conocimiento de los individuos que lo componen, lo que, a su vez, permite que la productividad del equipo sobrepase la suma de las productividades individuales de sus miembros (Alchian y Demsetz, 1972).

Rutinas organizativas. La rutina organizativa representa una pauta regular y predecible de actividad, integrada por una secuencia de decisiones coordinadas, que se ponen en funcionamiento ante un problema o estímulo específico (Nelson y Winter, 1982). La coordinación que se despliega en la realización de rutinas organizativas es, al igual que la desplegada en el ejercicio de las habilidades individuales, el fruto de la práctica. Se trata de un comportamiento aprendido, repetitivo o cuasi-repetitivo, fundado en parte en el conocimiento tácito. Las rutinas son la base de la memoria organizativa, constituyen el principal sistema de almacenamiento de conocimiento operativo de la organización y definen en cada momento lo que puede y no puede hacerse. La información sobre decisiones tomadas y problemas resueltos a lo largo del tiempo conforma el núcleo de la memoria organizativa. De esta forma, la empresa puede conocer, recordar y saber cosas que ninguno de los individuos o equipos que operan dentro de ella saben en su totalidad (Ouchi, 1984). Memoria organizativa es todo aquello que existe en una organización y que puede ser recuperado de alguna manera. La naturaleza tácita de las rutinas hace que resulten difíciles de imitar. De hecho, son un conocimiento específico de la empresa y el resultado del aprendizaje colectivo de ésta, en especial sobre cómo coordinar diversas tecnologías y habilidades (Hamel y Prahalad, 1994). Una empresa es en sí misma una compleja jerarquía de rutinas organizativas que definen los conocimientos de nivel inferior, cómo se coordinan y los procedimientos de decisión en el nivel superior de la organización, que sirven para seleccionar lo que se hará en los niveles inferiores (Nelson, 1991).

Las empresas desarrollan rutinas porque, de este modo, se economizan recursos dedicados a la toma de decisiones y se facilita la coordinación en un contexto caracterizado por la incertidumbre y la racionalidad limitada de los decisores. Comportamiento rutinario no significa comportamiento inflexible o pasivo. Las empresas modifican sus rutinas a fin de mejorarlas y adaptarlas a las circunstancias cambiantes del entorno.

2. LA CIENCIA

El creciente cuerpo de ideas que nos ayuda a enriquecer el mundo para hacerlo más confortable se conoce como ciencia. El conocimiento científico (ciencia) en la Grecia clásica se conoce como *episteme* (conocimiento especulativo, teórico o abstracto). No obstante, el término ciencia procede del verbo *scire*, que significa saber. Luego, etimológicamente, la ciencia equivale al saber. Aunque no todo el que conoce algo de un dominio del saber posee ciencia de él, sino sólo quien ha penetrado sistemáticamente en él y que, además de los detalles, conoce las conexiones entre los contenidos¹. Esto significa que hay saberes que no son ciencia. Por ejemplo, Platón² distingue entre saber y opinión, y advierte que ésta no es simple no saber, sino algo situado entre la perfecta ciencia y la absoluta ignorancia. La ciencia es un sistema de ideas establecidas provisionalmente (conocimiento científico) y una actividad productora de nuevas ideas (investigación científica). Así pues, la ciencia puede caracterizarse como conocimiento racional, sistemático, objetivo, verificable y, por consiguiente, falible. Por medio de la investigación científica, el hombre ha alcanzado una reconstrucción conceptual del mundo que es cada vez más amplia, profunda y exacta (Bunge, 2013).

Las ciencias se dividen en formales (o ideales) y fácticas (o materiales) (Bunge, 2013). Las formales no proporcionan un conocimiento objetivo. Así, la lógica y la matemática son racionales, sistemáticas y verificables, pero no son conocimientos objetivos. No nos dan informaciones acerca de la realidad, ya que no se ocupan de los hechos. La materia prima que emplean los lógicos y los matemáticos no es fáctica sino ideal: inventan entes formales, que sólo existen en la mente humana, si bien a menudo son una abstracción de fenómenos reales (como las figuras geométricas), y establecen relaciones entre ellos. Ahora bien, la física, la química, la psicología, la economía y demás ciencias fácticas recurren a la matemática y la lógica, empleándolas como herramientas para representar con precisión las complejas relaciones que se encuentran entre los hechos, y entre los diversos aspectos de los hechos.

¹ La comunidad científica describe la ciencia como aquellos conocimientos explícitos en los que todos están de acuerdo, al menos en lo relacionado con determinadas leyes y teorías.

² Sócrates pone estas palabras en boca de Diotime, mujer de Mantinea muy versada en todo lo concerniente al amor y a muchas otras cosas, con la que sostuvo la siguiente conversación:

Diotime: ¿Y crees que no se puede carecer de ciencia sin ser absolutamente ignorante? ¿No has observado que existe un término medio entre la ciencia y la ignorancia?

Sócrates: ¿Cuál es?

Diotime: Tener formada una opinión verdadera sin poder dar la razón de ella; ¿no sabes que eso no es ni ser sabio, porque la ciencia debe fundarse en razones, ni ser ignorante, puesto que lo que participa de la verdad no se puede llamar ignorancia? La opinión verdadera ocupa, pues, el justo término entre la ciencia y la ignorancia.

Sócrates: Confesé a Diotime que decía verdad.

Fuente: Platón: *Diálogo El Banquete*

Por otra parte, mientras que los enunciados formales consisten en relaciones entre signos, los enunciados de las ciencias fácticas se refieren en su mayoría a sucesos y procesos. En cuanto al método para probar los enunciados verificables: mientras las ciencias formales utilizan la lógica para demostrar sus teoremas, las ciencias fácticas confirman sus conjeturas mediante la observación y/o el experimento. La demostración de los teoremas es una deducción confinada a la esfera teórica. La matemática y la lógica son, en suma, ciencias deductivas. La verdad matemática no es absoluta sino relativa al sistema de ideas admitido previamente, en el sentido de que una proposición que es válida en una teoría puede dejar de ser lógicamente verdadera en otra teoría (por ejemplo, en el sistema aritmético que utilizamos para contar las horas del día, vale la proposición $24 + 1 = 1$). El proceso constructivo se limita a la formación de los puntos de partida (axiomas). Tan sólo las conclusiones (teoremas) tienen que ser verdaderas: los axiomas pueden elegirse a voluntad. Las ciencias fácticas emplean símbolos interpretados, necesitan la racionalidad –es decir, la coherencia con un sistema de ideas aceptado previamente– y exigen de los enunciados que sean verificables en la experiencia, sea indirectamente (en el caso de hipótesis generales), sea directamente (en el caso de las consecuencias singulares de las hipótesis). Por esto el conocimiento fáctico verificable se le llama a menudo ciencia empírica (Bunge, 2013).

Las ciencias formales demuestran o prueban; las ciencias fácticas verifican (confirman o desaprueban) hipótesis que en su mayoría son provisionales. La demostración es completa y final; la verificación es incompleta, y en consecuencia temporal. Por ello, mientras las teorías formales pueden ser llevadas a un estado de perfección (o estancamiento), los sistemas teóricos relativos a los hechos son esencialmente defectuosos; cumplen, pues, la condición necesaria para ser perfectibles (Bunge, 2013). Por ser una ficción seria, rigurosa y a menudo útil, pero ficción al fin y al cabo, la ciencia formal requiere un tratamiento especial. En lo que sigue nos concentraremos en la ciencia fáctica.

La ciencia, en su concepción clásica, es el conocimiento seguro, fiable y verdadero. No hay ciencia si no hay verdad. Ciertamente que ahora sabemos que la verdad absoluta no es accesible, ni siquiera posible, pero eso no debe hacernos olvidar que, sin aspirar a la verdad, aunque sea parcial y tentativa, no es posible la ciencia.

La ciencia asume que hay un mundo real, independiente de nuestras teorías, y que podemos conocerlo. La tierra giraba alrededor del sol, al margen de que la gente apoyara la postura de Ptolomeo o la de Copérnico.

En la actualidad, la importancia de la ciencia (fáctica) parece derivarse de dos factores. En primer lugar, las teorías científicas se basan en la cuidadosa observación y el experimento controlado, no se trata de especulaciones infundadas. En segundo lugar, las teorías científicas suelen ser bastante abstractas: usan conceptos no evidentes para los sentidos a fin de explicar en función de factores que no se pueden ver sucesos que son familiares. Esta combinación de observación y teoría no tiene precedente en el pensamiento humano.

La ciencia (fáctica) se caracteriza por orientarse a la búsqueda de conocimientos sobre un determinado objeto o realidad dentro de un proceso de pensamiento. En este sentido, Hodgson (1984) define la ciencia como un cuerpo integrado de conocimientos, principalmente cuantitativos, construido por los esfuerzos dinámicos del hombre para comprender su entorno y a sí mismo de manera sistemática y comunicable. Por ello, una

característica del conocimiento científico es la de su vínculo entre la hipótesis y la prueba. Una hipótesis científica nueva surge cuando la visión científica al uso empieza a dar problemas. A continuación, se proponen conjeturas –variaciones o alternativas a la visión estándar–, que se contrastan mediante la experimentación y la acumulación de datos. El resultado es que algunas conjeturas sobreviven y acaban aceptándose como nuevas teorías. Por tanto, otra característica básica de la ciencia es la de la formulación precisa, generalmente cuantitativa de sus teorías. Si las teorías se formulan con precisión, tendrán consecuencias concretas. Pequeñas discrepancias con los datos observados aumentarán la posibilidad de que las teorías falsas puedan identificarse y refutarse.

En la ciencia, los conceptos están relacionados entre sí de forma definida y coherente, constituyendo una estructura, por lo que la comprensión del todo requiere el conocimiento previo de las partes y sus relaciones. De igual forma, no es posible comprender una parte cualquiera hasta que no se conozca cómo está relacionada con las otras partes y con el todo.

La ciencia es objetiva y permanente. La evidencia objetiva no depende del observador ni de sus circunstancias y sentimientos particulares. Ser universal quiere decir que no depende del espacio ni del tiempo.

La ciencia es el resultado de la comprensión intelectual de una realidad que existe independientemente del acto de conocer y del sujeto cognoscente. Esto no quiere decir que sea fácil determinar cuándo un conocimiento tiene el carácter de científico. De hecho, existe históricamente una fuerte discrepancia entre los filósofos de la ciencia acerca del método científico. A su vez, la ciencia ha de ser comunicable y sistemática, porque, de otro modo, no constituiría un cuerpo coherente de conocimiento (Hodgson, 1984). También se afana continuamente por ampliar sus fronteras, extendiéndose a nuevos campos y profundizando en su comprensión. Este planteamiento lo sintetiza Watkins (1982) al considerar que la ciencia: a) consta de verdades profundas con mucho rigor explicativo, ya que, para resolver los secretos más íntimos de la naturaleza, debe ir hasta el fondo de las cosas alcanzando sus explicaciones últimas; b) proporciona verdades exactas en todos los niveles: si se hace más vaga a medida que es más profunda, el resultado sería, en gran medida, decepcionante; c) es verdadera y deja constancia de que lo es y d) conecta sus afirmaciones de manera lógica y no de forma imprecisa y descuidada. Por otra parte, la ciencia es más fructífera cuando genera la capacidad de pronosticar hechos relativos a fenómenos sin experimentación ni observación, o antes de ellos. Consideramos científico un método basado en el análisis profundo de los hechos, las teorías y las opiniones, presuponiendo discusiones y conclusiones abiertas sin prejuicios ni temores.

La ciencia es un trabajo en progreso y, por eso mismo, abierto a rectificaciones. La imperfección de la ciencia es la clave de su progreso. Sin lugar para nuevas ideas, la ciencia de hoy sería la misma que la ciencia de ayer. Muchos descubrimientos comienzan con alguna anomalía: un fenómeno o dato que no parece estar de acuerdo con lo esperado. La ciencia es conservadora, y no acepta sin oposición ideas o resultados contrarios al consenso, actitud prudente que la protege de cambios prematuros.

El colectivo de científicos son los jueces de la ciencia, un juicio que tiene dos aspectos: la coherencia de lo nuevo con lo ya conocido (en todas las disciplinas y entre disciplinas) y los resultados de la confrontación empírica. De acuerdo con estos criterios, la comunidad determina la validez de su contenido, y establece un cuerpo de conocimientos aceptados

por consenso. El consenso logrado por la ciencia se refiere a los resultados de la ciencia, pero no a la dirección en la cual se encamina la investigación científica ni a sus aplicaciones. La ciencia es una actividad social estructurada de tal forma que luego de un proceso de aceptación por consenso, sus resultados se incorporan al conocimiento público.

La comunidad científica tiene dos mecanismos (primordialmente la revisión por pares y la replicación de los resultados) para descubrir errores y fraudes. Pero, los revisores, por muy expertos que sean, no pueden delimitar si hubo errores de procedimiento, si algún instrumento no funcionó correctamente o si hubo manipulación de los datos. Esto se descubre en el momento en el cual otros investigadores independientes tratan de repetir el estudio. El experimento repetible es la prueba necesaria para convencer al que duda o para demostrar que se cometió un error o fraude. La ciencia ha de ser reproducible, que quiere decir que es posible obtener la evidencia por diferentes observadores, con instrumentos distintos y en diferentes sitios.

La ciencia se concibe como conocimiento explícito, que utiliza un lenguaje y un contexto claros, lo que facilita su difusión social, sobre todo a través de su publicación como artículos en las revistas científicas especializadas. El soporte final de la actividad científica suele ser una formulación escrita, el artículo científico, que informa de un hallazgo experimental o de una nueva posición teórica. Para publicar un artículo es necesario un informe favorable de *referees* anónimos: especialistas en la materia tratada, propuestos por el editor de la revista. Con ello, se pretende que los trabajos que se publiquen alcancen un elevado nivel de calidad técnica, argumentación lógica y exposición coherente. La crítica es el arma más poderosa de una metodología y de unos resultados científicos: es la única seguridad que tiene el investigador de que no persistirá en el error (lectura 6). Posteriormente los libros de texto recogen los resultados más importantes para su divulgación a través de las aulas.

LECTURA 6: FRAUDES CIENTÍFICOS

Se supone que la ciencia es una metodología precisa y objetiva, cuya noble misión es descubrir toda la verdad sobre el mundo que nos rodea. En las sociedades modernas, los expertos científicos gozan de un enorme prestigio social, como demuestra el hecho de que políticos, periodistas y ciudadanos de a pie acuden a ellos para comprender los aspectos más complejos de la realidad. Sin embargo, en la ciencia –como en cualquier otra esfera de la sociedad humana– también existe la mentira, el engaño y la picaresca.

El Hombre de Piltdown. Uno de los fraudes más escandalosos de la historia de la ciencia es el llamado *Hombre de Piltdown*, un fósil que se presentó a bombo y platillo como el eslabón perdido de la evolución humana, y que finalmente resultó ser una enorme patraña. *Eoanthropus dawsoni* o el *Hombre de Piltdown* fue hallado en 1912 por Charles Dawson en una gravera en Piltdown, Sussex. Durante cuarenta años estuvo expuesto con su enorme cráneo humano y mandíbula de simio en el edificio del actual Museo de Historia Natural de Londres, como un ejemplo del famoso eslabón perdido entre los seres humanos y sus ancestros los primates.

Sin embargo, el 21 de noviembre de 1953, los científicos declararon que el hallazgo, es decir, la unión de un moderno cráneo humano y una mandíbula de orangután, era una burda falsificación. Llegaron a la conclusión de que alguien había colocado adrede el conjunto de fragmentos fósiles de Piltdown, que contenía incluso una ridícula pala de cricket prehistórica.

El mundo de la paleontología se moría de vergüenza y los partidarios de la teoría de la conspiración se volvieron locos. No faltaban candidatos para el escarnio y durante las siguientes cinco décadas no dejaron de nombrarlos. En el reparto de probables bromistas potenciales en esta novela de antropología policíaca figuran aficionados entusiastas, profesionales apasionados y guasones desinteresados.

Los teóricos han llegado a señalar con el dedo a un sacerdote jesuita, el padre Teilhard Chardin, que póstumamente se convirtió en un gurú del *New Age*, y al mismo creador de Sherlock Holmes, Sir Arthur Conan Doyle, quien escribió su propio thriller paleontológico, titulado *El Mundo Perdido*. De lo que no hay duda es de que todo lo que se encontró en la cueva había sido colocado de manera fraudulenta por alguna mano experta.

La vergüenza de Bell. Jan Henrik Schön, un joven investigador de los Laboratorios Bell de Nueva Jersey, llevaba publicados cinco artículos sobre cuestiones avanzadas de electrónica en *Nature* y siete en la revista *Science* entre 1998 y 2001. Los hallazgos eran inauditos hasta el punto que Schön estaba considerado por sus colegas como una estrella fulgurante.

En el año 2002 un comité averiguó que el joven había maquillado sus resultados al menos en dieciséis ocasiones, avergonzando públicamente a sus colegas, a su jefe y a los equipos editoriales de las dos publicaciones que aceptaron sus resultados. Schön, que entonces apenas tenía 32 años, dijo: “Tengo que admitir que en mi trabajo científico he cometido algunos errores de los que me arrepiento profundamente”. *Nature* también publicó una declaración posterior suya que decía: “Realmente creo que los resultados científicos aportados son auténticos, estimulantes y merece la pena trabajar por ellos”. Desde entonces no se ha vuelto a saber nada de él.

Fuente: Radford, T. (2003): “Los grandes fraudes de la ciencia”, *El Mundo*, 17 de noviembre, pp. 37.

La ciencia con su publicación en revistas científicas, libros e Internet y presentación en congresos pasa a ser un bien público: es difícil impedir que otros se beneficien de ella (no excluible) y el coste marginal de su utilización por una persona adicional es cero (no rival). Por eso, al igual que la mayoría de los bienes públicos, no suelen suministrarla los mercados privados, sino el Estado u organizaciones sin ánimo de lucro.

2.1. El paradigma científico

En su libro *La Estructura de las Revoluciones Científicas*, Kuhn (1962) considera que no se puede pasar de lo viejo a lo nuevo mediante una simple adición a lo que ya era conocido. Defiende dos formas de progreso científico: el cambio normal y el cambio revolucionario. Este último es el que provoca grandes avances de la ciencia, al romper con todo lo anterior y lograr explicar nuevos e importantes fenómenos de la naturaleza. Los cambios revolucionarios son en cierto sentido holistas, ya que no pueden hacerse poco a poco. Por otra parte, el cambio normal de la ciencia tiene como resultado el crecimiento, aumento o adición acumulativa de lo conocido con anterioridad.

El cambio revolucionario proporciona descubrimientos que no pueden acomodarse dentro de los conceptos generalmente aceptados hasta ese momento. Una revolución supone, pues, abandonar una estructura teórica y reemplazarla por otra incompatible, por lo que contrasta abiertamente con los cambios acumulativos. Para explicar los cambios revolucionarios Kuhn acuña el término «paradigma», que, muy a su pesar, llegó a significar casi todo para todas las personas, lo que, de alguna forma, contribuye a explicar las razones del éxito del libro y del uso del término.

Un paradigma está constituido por los supuestos teóricos generales, las leyes y las técnicas para su aplicación, que adoptan los miembros de una determinada comunidad científica (Kuhn, 1962). En cierto sentido, el paradigma entraña un determinado marco conceptual a través del cual se ve el mundo y también se le describe, así como un conjunto de técnicas experimentales y teóricas que permiten compaginar el paradigma con la naturaleza. Por tanto, un paradigma es, en primer lugar, un logro o realización científica fundamental, que incluye a la par una teoría y algunas aplicaciones ejemplares a los resultados del experimento y la observación. De igual forma, es una realización cuyo término queda abierto, ya que permite toda suerte de investigaciones. Y, finalmente, es una realización ampliamente aceptada por la comunidad científica, que no intenta rivalizar con ella ni crearle alternativas. Por el contrario, dichos miembros intentan extenderla y explotarla de muy diversos modos (Kuhn, 1963). En suma, un paradigma establece las normas necesarias para legitimar el trabajo dentro de la ciencia

que rige, al igual que coordina y dirige la actividad de «resolución de problemas» que efectúan los científicos que trabajan dentro de él. De esta forma, no es necesario expresar de manera explícita las reglas y suposiciones en las que se asienta la investigación.

Los científicos que desarrollan su labor dentro de un paradigma practican lo que Kuhn llama «ciencia normal», que se caracteriza por producir los ladrillos que la investigación científica está continuamente añadiendo al creciente edificio del conocimiento. De esta forma, el paradigma presenta a los investigadores un conjunto de problemas definidos, junto con unos métodos que ellos confían en que serán los adecuados para solucionarlos. Por tanto, la ciencia normal es una actividad de resolver problemas gobernada por las reglas del paradigma que lo sostiene. Los problemas serán de naturaleza teórica y experimental.

La ciencia normal es el ámbito propio de la actividad científica y está basada en algún logro científico significativo del pasado. Este logro histórico marca los límites de la agenda de investigación del trabajo científico; agenda que no incluirá una investigación crítica del paradigma establecido. Al contrario, es el paradigma el que define los términos de referencia de la ciencia normal, así como el alcance y los límites de la agenda de investigación. Se sigue de ello que la ciencia normal es una ocupación tenaz en la que los científicos concentran sus esfuerzos en resolver enigmas (problemas). Así pues, un científico no debe criticar el paradigma en el que trabaja. Sólo de esa manera es capaz de concentrar sus esfuerzos en la detallada articulación del paradigma y efectuar el trabajo necesario para explorar la naturaleza en profundidad.

Los períodos de ciencia normal proporcionan a los científicos la oportunidad de desarrollar los detalles esotéricos de una teoría. Trabajando dentro de un paradigma cuyos fundamentos se dan por sentados, son capaces de llevar a cabo el duro trabajo teórico y experimental necesario para que el paradigma se compagine con la naturaleza en un grado cada vez mayor. Gracias a su confianza en la adecuación de un paradigma, los científicos pueden dedicar sus energías a intentar resolver los detallados problemas que se les presentan en vez de enzarzarse en disputas sobre la licitud de sus supuestos y métodos fundamentales. Por tanto, es necesario que la ciencia normal sea, en gran medida, acrítica. Si todos los científicos cuestionaran el marco conceptual en el que trabajan, no se llevaría a cabo ninguna investigación detallada (Chalmers, 1984).

La ciencia normal se desarrolla sin perturbaciones mientras la aplicación del paradigma explique satisfactoriamente los fenómenos a los que se aplica. Ningún paradigma resuelve completamente sus problemas y, la no resolución de un problema, se atribuye a un fracaso del investigador más que una insuficiencia del paradigma. Los problemas que se resisten a ser solucionados se consideran anomalías. Una anomalía es particularmente grave si afecta a los propios fundamentos del paradigma. La presencia de una o dos anomalías no es suficiente para producir el abandono de un paradigma. Si las anomalías están fuera de control, entonces se desarrolla un estado de crisis. La crisis se resuelve cuando surge un paradigma completamente nuevo que se gana la adhesión de un número de científicos cada vez mayor hasta que finalmente todos abandonan el paradigma original, acosado por problemas. El establecimiento del nuevo paradigma cumple de una sola vez tres importantes exigencias: a) ofrece una solución a la crisis científica, b) proporciona una nueva «visión científica del mundo» y c) ofrece una agenda de investigación alternativa para que los científicos puedan trabajar. En lo

sucesivo la revolución amaina y la ciencia normal reasume su papel de resolver enigmas de un nuevo paradigma (figura 3). Hay que destacar que los paradigmas competidores son inconmensurables³ (aunque no totalmente). Reflejan orientaciones conceptuales diferentes. Quiénes proponen paradigmas alternativos observan de manera diferente cierto tipo de fenómenos (Losee, 1981).

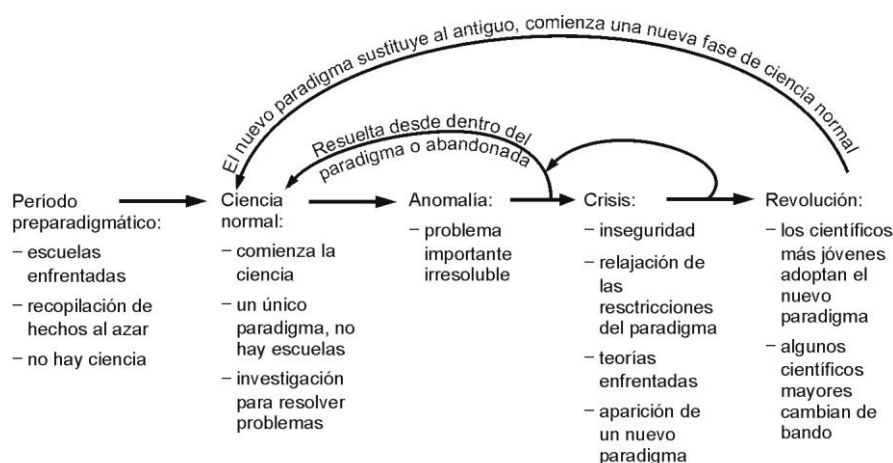


Figura 3: El carácter revolucionario de los paradigmas (Leahey, 2004)

El nuevo paradigma, lleno de promesas y no abrumado por dificultades en apariencia insuperables, guía entonces la nueva actividad científica normal hasta que choca con serios problemas y surge una nueva crisis seguida de una nueva revolución (Chalmers, 1984). La tradición científica normal que emerge de una revolución científica es no sólo incompatible, sino también inconmensurable, con la que existía con anterioridad. En suma, los paradigmas obtienen su estatus como tales, debido a que tienen más éxito que sus competidores para resolver unos cuantos problemas que el grupo de investigadores ha llegado a reconocer como agudos (Kuhn, 1962). Los paradigmas no se sustituyen unos a otros repentinamente y, en cualquier caso, los paradigmas nuevos no surgen y se asientan de repente, sino que obtienen la victoria después de un largo proceso de competencia intelectual.

La existencia de anomalías es lo que sirve de estímulo para la invención de paradigmas alternativos. De esta forma, sólo la aparición de un paradigma alternativo que pueda resolver o alejar (al menos) las anomalías más importantes proporcionan la condición suficiente para el rechazo del paradigma dominante existente. Kuhn sostuvo que, en último término, la ciencia normal tiende únicamente al reconocimiento de anomalías y crisis, que se resuelven no por deliberación e interpretación, sino por un acontecimiento relativamente repentino, discontinuo y no estructurado, denominado «revolución científica».

Una revolución científica corresponde al abandono de un paradigma y a la adopción de otro nuevo e incompatible, no por parte de un científico aislado, sino por parte de la comunidad científica en su totalidad. Casi siempre, los científicos que realizan las

³ Inconmensurable en el sentido de que no hay una unidad de longitud contenida un número entero de veces sin resto en cada miembro del par. Por ejemplo, la hipotenusa de un triángulo isósceles es inconmensurable con su lado.

investigaciones fundamentales de un nuevo paradigma han de ser muy jóvenes o muy noveles en el campo cuyo paradigma cambian. Evidentemente se trata de científicos que, al no estar comprometidos con las reglas tradicionales de la ciencia normal, debido a que acumulan poca práctica anterior, tienen muchas probabilidades de darse cuenta que esas reglas no definen un juego que pueda continuar adelante y de concebir otro conjunto que pueda reemplazarlas (Kuhn, 1962). A medida que, por diversas razones, se convierten más científicos al paradigma hay un creciente cambio en la distribución de las adhesiones profesionales. Al principio, los científicos se resisten a cambiar de paradigma. La fuente de la resistencia reside en la confianza de que el paradigma de mayor antigüedad finalmente resolverá todos sus problemas y de que la naturaleza puede compendiarse dentro de las marcas proporcionadas por el paradigma.

Kuhn vincula el cambio de la adhesión por parte de los científicos de un paradigma a otro alternativo e incompatible con un «cambio de *gestalt*» o una conversión religiosa. No existe ningún argumento puramente lógico que demuestre la superioridad de un paradigma sobre otro y que, por tanto, impulse a cambiar de paradigma a un científico racional. Una razón de que no sea posible esa demostración estriba en el hecho de que en el juicio de un científico sobre los méritos de una teoría intervienen muchos factores. La decisión del científico de cambiar de paradigma dependerá de la prioridad que les otorgue. Los factores incluirán cuestiones tales como la simplicidad, la conexión con alguna necesidad social urgente, la capacidad de resolver algún determinado tipo de problema, etc. (Chalmers, 1984).

El proceso de las revoluciones científicas descrito por Kuhn es, como mínimo, la línea más importante de distinción entre ciencias «maduras» e «inmaduras». Las ciencias maduras están dominadas por un paradigma, mientras que las ciencias inmaduras no, ya que, al encontrarse aún en el período de su prehistoria científica, están infectadas de paradigmas o escuelas en competencia. Así, mientras la ciencia madura avanza dentro de los límites de un paradigma ampliamente aceptado por la comunidad científica, la ciencia inmadura se encuentra en un estado precientífico caracterizado por el total desacuerdo y el constante debate de lo fundamental, lo que da lugar a la existencia de varios paradigmas en competencia, resultando imposible abordar el trabajo detallado.

Al aprender un paradigma, el científico adquiere al mismo tiempo teoría, métodos y normas, casi siempre en una mezcla inseparable. A falta de un paradigma o de algún candidato a paradigma, todos los hechos que pudieran ser pertinentes para el desarrollo de una ciencia dada tienen probabilidades de parecer igualmente importantes, por lo que resulta difícil avanzar.

En el *postscriptum* de la segunda edición de su libro, Kuhn considera que si se hiciera una revisión del mismo, el término paradigma tendría dos usos muy diferentes: uno amplio y otro restringido. Denomina al primero «matriz disciplinar» y al segundo «ejemplos compartidos». La matriz disciplinar es un conjunto de ideas básicas, un sistema de nociones o marco conceptual (constitutivos del paradigma) que abarcan cierto número de teorías, hipótesis y conceptos (correspondientes a las partes del paradigma o paradigmáticos). De esta forma, la matriz disciplinar es una completa constelación de creencias, valores, técnicas, etc., compartidas por los miembros de una comunidad científica. Los investigadores practicantes pueden compartir un compromiso con la existencia de entidades teóricas (espacio absoluto, átomo, campos, genes, etc.). Además, pueden estar de acuerdo acerca de qué tipos de investigación y explicación son

importantes (estudios in vivo frente a estudios in vitro, acción por contacto frente a interpretaciones de campos, explicaciones deterministas frente a explicaciones probabilísticas, etc.). Esos compromisos y creencias forman parte de un paradigma en sentido amplio. Una matriz disciplinar incluye también uno o más paradigmas en sentido restringido (Losee, 1981).

En sentido restringido, un paradigma es un «ejemplo compartido», una representación influyente de una teoría científica. Por lo general, los ejemplares se consignan, desarrollan y revisan en los libros de texto, los cuales contienen ilustraciones y ampliaciones estereotipadas de una teoría (Losee, 1981). Quizá la función más importante de los ejemplos compartidos es su contribución tácita, aunque fundamental, para el desarrollo del lenguaje científico común y, lo más importante, una forma común de acercarse a problemas de un tipo dado, pero todavía no encontrado por el científico individual.

2.2. La investigación básica

La ciencia se obtiene a través de la investigación básica, realizada de forma desinteresada, sin intencionalidad operativa y por mera curiosidad intelectual, una de las características que distingue al individuo de los demás animales (Bunge, 1982).

Imaginemos, como propone Cooper (2007), que hace un siglo alguien hubiera puesto en marcha un programa de investigación aplicada y desarrollo tecnológico cuyo objetivo fuera conseguir una tecnología capaz de permitirnos escuchar en el salón de nuestra casa un concierto de música con el mismo nivel de calidad de audición que si estuviéramos en una sala de conciertos. Es fácil imaginar qué tipo de proyecto habría salido de allí: una especie de fonógrafo gigante con toda clase de artilugios mecánicos para producir y transmitir música. Imagínese a Edison concursando a esa convocatoria de proyectos. Podría haber salido cualquier cosa salvo, con toda seguridad, lo que realmente se hizo con posterioridad: utilizar la microelectrónica, el láser, la teoría de la comunicación y los algoritmos más avanzados del cálculo numérico para conseguir los sistemas de reproducción, almacenaje y transmisión de música que tenemos hoy. Ningún proyecto de desarrollo tecnológico podría haber previsto eso porque, sencillamente, es el resultado de descubrimientos posteriores de la investigación científica básica, que por definición eran impredecibles hace más de un siglo.

La investigación básica (científica, pura o fundamental) consiste en la búsqueda del conocimiento científico (ciencia) sin ninguna finalidad determinada (lectura 7). Sus resultados están sometidos a contrastes cuantitativos. En otras palabras, persigue el descubrimiento *per se* y se concreta en un nuevo hecho, principio, hipótesis, ley o teoría relacionada con los sucesos que se observan directamente o a través de sus efectos. Esta tarea la realizan investigadores que generalmente desempeñan su labor en universidades o instituciones sin ánimo de lucro. Sus incentivos para divulgar el conocimiento en libros, revistas científicas y congresos, entre otros medios de difusión, son fuertes, debido al prestigio científico, la cooperación con los colegas y la promoción académica.

LECTURA 7: FORMULANDO UNA TEORÍA

Empecé con una observación sencilla: las hormigas sacan a sus muertos del nido. Las de algunas especies simplemente se deshacen de los cadáveres al azar, fuera del hormiguero, mientras que las de otras especies los colocan en montones de desperdicios que bien pudieran calificarse de «cementeros». El

problema que yo vi en este comportamiento era sencillo pero interesante: ¿cuándo sabe una hormiga que otra está muerta? Para mí era evidente que el reconocimiento no se hacía mediante la vista. Las hormigas reconocen un cadáver incluso en la oscuridad completa de las cámaras de nidificación subterráneas. Además, cuando el cuerpo está fresco y se halla en una zona iluminada, e incluso cuando se halla panza arriba con las patas al aire, las otras hormigas lo ignoran. Solo después de un día o dos de descomposición, el cuerpo se convierte en un cadáver para otra hormiga. Barrunté (planteé una hipótesis) que las hormigas funerarias utilizaban el olor de la descomposición para reconocer la muerte. Pensé, además (segunda hipótesis), que su respuesta se desencadenaba solo a partir de unas pocas de las sustancias que exuda el cuerpo del cadáver. La inspiración para la segunda hipótesis era un principio evolutivo establecido: los animales de cerebro pequeño, que son la inmensa mayoría de los animales de la Tierra, tienden a usar el conjunto más simple de pistas disponibles para que los guíen a través de la vida. Un cuerpo muerto ofrece docenas o cientos de pistas químicas entre las que elegir. Los seres humanos pueden seleccionar estos componentes. Pero las hormigas, con un cerebro cuyo tamaño es una millonésima parte del nuestro, no pueden.

De modo que si las hipótesis son ciertas, ¿cuál de estas sustancias podría desencadenar la respuesta funeraria: todas ellas, unas pocas o ninguna? De proveedores químicos obtuve muestras sintéticas puras de varias sustancias de descomposición, entre ellas el escatol, la esencia de las heces; la trimetilamina, el olor dominante del pescado en descomposición; y varios ácidos grasos y sus ésteres de un tipo que se encuentran en los insectos muertos. Durante un tiempo, mi laboratorio olía como una combinación de osario y cloaca. Puse cantidades minúsculas de estas sustancias en cadáveres de hormigas de mentirijillas, hechos de papel, y los inserté en colonias de hormigas. Después de una maloliente serie de pruebas y errores, encontré que el ácido oleico y uno de sus oleatos desencadenan la respuesta. Las otras sustancias o bien fueron ignoradas, o bien causaron alarma.

Para repetir el experimento de otra manera (y, debo admitirlo, para diversión mía y de otros), embadurné con pequeñas cantidades de ácido oleico el cuerpo de hormigas obreras vivas. ¿Se convertirían acaso en muertos vivientes? Ciertamente, se transformaron en zombis, al menos si estos se definen de una manera amplia. Fueron agarradas por compañeras del nido y, mientras agitaban sus patas, fueron transportadas al cementerio y depositadas allí. Después de haberse limpiado durante un rato, se les permitió volver a la colonia.

Después se me ocurrió otra idea: los insectos de todas las especies que se alimentan a base de carroña, como las moscardas y los escarabajos estercoleros, encuentran su camino hasta los animales muertos o los excrementos siguiendo el olor. Y lo hacen utilizando un número muy reducido de las sustancias químicas de la descomposición que hay presentes. Una generalización de este tipo, aplicada ampliamente, con al menos algunos hechos aquí y allá y algo de razonamiento lógico detrás, es una teoría. Harán falta muchos más experimentos, aplicados a otras especies, para convertirla en lo que puede llamarse con seguridad un hecho.

Wilson, E. (2014): *Cartas a un Joven Científico*, Debate, Barcelona.

La investigación básica consiste en trabajos experimentales o teóricos que se emprenden fundamentalmente para obtener nuevos conocimientos acerca de los fundamentos de fenómenos y hechos observables, sin pensar en darles ninguna aplicación o utilización determinada. Con frecuencia se afirma que el progreso intelectual tiene lugar en las épocas en que los científicos realizan experimentos y observan críticamente sus resultados, mientras que habrá un estancamiento intelectual en períodos en que los experimentos son realizados acríticamente o no son realizados en absoluto. Sin embargo, sería simplista llegar a la conclusión de que el progreso científico depende por completo de la disposición o no para llevar a cabo experimentos. Sería poco realista insistir en que los científicos deberían volver a realizar todos los experimentos que ya han sido llevados a cabo anteriormente, pues esto significaría que no habría tiempo para nuevos experimentos (Teichman y Evans, 2010).

La investigación básica analiza propiedades, estructuras y relaciones con objeto de formular y contrastar hipótesis, teorías o leyes. La referencia a “sin pensar en darle ninguna aplicación o utilización determinada” en la definición de investigación básica es crucial, ya

que el ejecutor puede no conocer aplicaciones reales cuando lleve a cabo la investigación. Los resultados de la investigación básica no se ponen normalmente a la venta, sino que generalmente se publican en revistas científicas o se difunden directamente a colegas interesados. En ocasiones, la difusión de los resultados de la investigación básica puede ser considerada «confidencial» por razones de seguridad nacional. La investigación básica se lleva a cabo normalmente por científicos, quienes tienen libertad para fijarse sus propios objetivos. Esta investigación normalmente se efectúa en el sector de la enseñanza superior, pero también, en cierta medida en el sector administración pública. La investigación básica puede estar orientada o dirigida hacia grandes áreas de interés general, con el objetivo explícito de un amplio abanico de aplicaciones en el futuro (OCDE, 2003).

El científico debe ejercer su juicio personal en la elección de los temas de investigación, en la formulación de las preguntas que ha de contestar o en los métodos y procedimientos a emplear. También tiene libertad para cambiar los planes de investigación cuando lo estime oportuno a medida que enfrenta nuevos retos y surgen nuevas posibilidades. Algunos de los adelantos científicos más notables han resultado de proyectos de investigación iniciados con fines muy distintos a los propuestos (Nelson, 1959). A pesar de su libertad, en la práctica las decisiones relacionadas con la orientación de la investigación están muy restringidas por el estado actual del conocimiento, por la relación del científico con sus colegas y por los recursos técnicos y económicos que tiene a su disposición (Ziman, 1980). La investigación básica es, pues, un proceso, más bien muchos procesos, cuyo fin es generar conocimiento científico formado por acciones dirigidas a través del tiempo, que no pueden almacenarse, aunque sí el conocimiento resultante.

El número de científicos dedicados a investigar un campo específico evoluciona con el tiempo en función de los descubrimientos realizados y de las perspectivas de investigación futuras. En la figura 4 la curva P representa el número de personas que participan en la investigación básica, con objeto de lograr un avance en un campo concreto de la ciencia. Como se observa, se incrementa con el tiempo hasta alcanzar un máximo. Posteriormente, disminuye incluso hasta alcanzar el valor cero (cuando ningún científico se dedica a investigar en ese campo). La cantidad de ignorancia (I) desciende exponencialmente, ya que está determinada no sólo por el número de investigadores, sino también por las comunicaciones que mantienen, lo que mejora la eficiencia de la búsqueda individual y colectiva del saber (Holton, 1985). La curva A representa la aplicación de los conocimientos. Tales aplicaciones incluyen su uso en el campo principal y en otros campos relacionados.

El campo de investigación, a medida que decrecen las ideas interesantes que quedan por descubrir, se irá haciendo menos atractivo. La curva P si bien indica directamente el volumen de la profesión en cualquier momento, indirectamente representa la intensidad del interés o atractivo del campo. Se observa que el número de investigadores en activo alcanzará un máximo cuando una gran parte del total de ideas interesantes haya sido descubierta. Muchos investigadores se introducen en el campo al tener conocimiento de la potencialidad del mismo a través de las publicaciones de sus colegas. Lo que ocurre es que las publicaciones comienzan a aflorar cuando las investigaciones han permitido lograr muchos de los descubrimientos posibles.

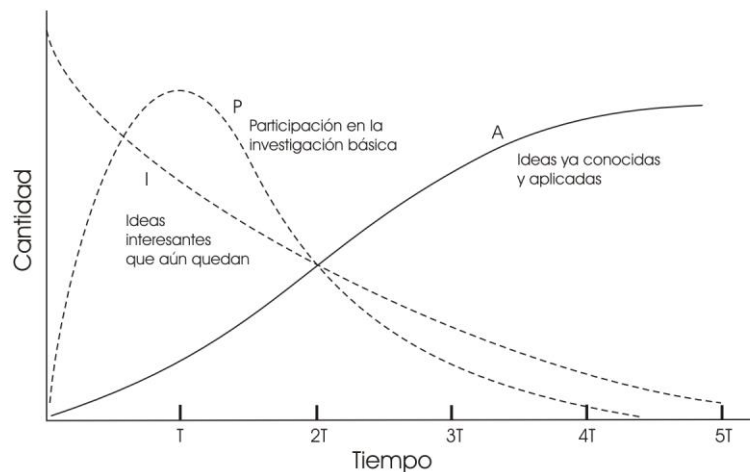


Figura 4: Participación de investigadores e ideas descubiertas en un campo de la ciencia (Holton, 1985)

En la figura 5a la curva D es la imagen refleja de la curva I, es decir, mientras que I representa la disminución de la ignorancia, D representa el aumento de los descubrimientos totales básicos realizados en el conjunto finito de ideas interesantes. Ahora bien, como se observa en la figura 5b, cada línea de investigación principal representada por la línea D en realidad es parte de una serie D_1 , D_2 , D_3 , etc. En este sentido, se considera que el crecimiento de la investigación científica se produce mediante el aumento a saltos del conocimiento —o quizás de nuevas áreas de ignorancia— y no por simple acumulación (Holton, 1985).

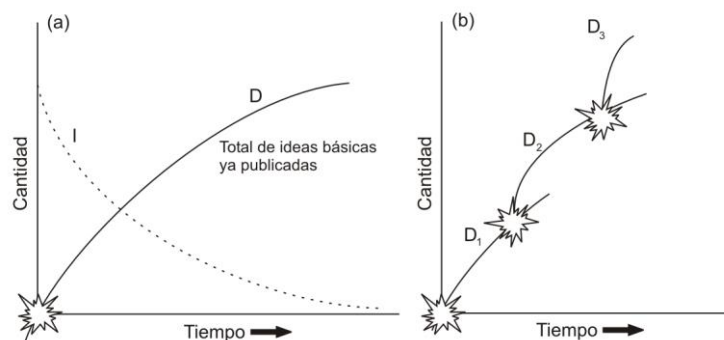


Figura 5: La evolución del descubrimiento (Holton, 1985)

Los mercados privados no suelen ofrecer una cantidad de investigación básica suficiente al ser la ciencia un bien público. La investigación básica debe ser sostenida pródigamente por la sociedad debido a los potenciales beneficios que de ella pudieran fluir en el futuro (Ziman, 1980). De ahí que la financie el Estado, filántropos, fundaciones y organizaciones sin ánimo de lucro. Se lleva a cabo en organismos públicos específicos, como los centros públicos de investigación o las universidades. Esto no evita el creciente temor de que la cantidad gastada en investigación básica no sea la adecuada. Algunas empresas que se ven forzadas a realizar investigación básica como complemento de sus desarrollos tecnológicos tienden a mantener en secreto los resultados, por lo que permanecen inaccesibles a la comunidad científica (Nelson, 1959).

La investigación básica no suele ser importante como fuente directa de la tecnología, pues los adelantos significativos en la ciencia no suelen ser aplicables en forma directa e inmediata a las soluciones de problemas prácticos y, por tanto, no se traducen en nuevos productos y procesos. Sin embargo, juega un papel crítico en la producción de conocimiento y entra en el proceso de innovación tecnológica indirectamente, por medio de la formación (Utterback, 1974). A menudo, los resultados de la investigación básica tienen mayor valor como insumo clave de otros proyectos de investigación que, a su vez, pueden generar resultados de valor práctico y patentable.

3. LA TECNOLOGÍA

El conocimiento tecnológico (tecnología) se conocía en la Grecia clásica como *techné* (arte o conocimiento práctico). La tecnología⁴ es la aplicación sistemática del conocimiento científico u otro conocimiento organizado a tareas prácticas (Galbraith, 1967). La tecnología es más que ciencia aplicada. Incluso puede apoyarse en hipótesis científicas falsas, como ocurrió, por ejemplo, con la invención en Francia del globo aerostático por los hermanos Montgolfier en 1783. Joseph de Montgolfier conocía el hidrógeno, un gas mucho más ligero que el aire, descubierto por Henry Cavendish en 1766. Supuso que el fuego desprendía un gas similar, y que este gas, recogido en un recipiente cerrado, al ser más ligero que el aire, sería capaz de elevar el recipiente. El experimento funcionó, pero el razonamiento era, en parte, falaz. Lo que hizo subir el globo no fue un gas más liviano que el aire, sino el aire mismo, que, una vez calentado, se expandía y reducía su peso específico (Mokyr, 1990). Por tanto, la tecnología puede resolver problemas sin que existan conocimientos científicos que la sustenten. No conviene olvidar que los primeros conocimientos desarrollados por el hombre fueron de tipo tecnológico. En la prehistoria, nuestros antepasados construyeron sus hachas de sílice sin referencias científicas previas.

La tecnología requiere un deseo de actuar en un sentido perfectamente establecido con objeto de solucionar algún problema determinado. Por tanto, los componentes de la tecnología son: 1) un resultado deseado (la solución de un problema) y 2) un conjunto organizado de actividades que permite alcanzar ese resultado. En consecuencia, el núcleo de cualquier tecnología es una relación causa-efecto, que responde a la pregunta ¿cómo hay que actuar para alcanzar el resultado deseado?, y tiene la siguiente estructura gramatical:

Saber + (verbo) + (complemento)

La tecnología es tanto un conocimiento como el resultado de ese conocimiento. Es decir, la tecnología produce artefactos, pero también los conocimientos necesarios para crearlos. Los artefactos son tanto el medio como el fin de la tecnología (Basalla, 1988). En todos los casos se espera que el investigador que crea una tecnología termine cada una de sus investigaciones afirmando no tanto «he descubierto X» como «he descubierto que X puede servir para producir (o impedir) Y» (Bunge, 1982). Como sugiere Gille (1978), la ciencia tiene como objetivo la comprensión de las leyes que rigen la naturaleza, mientras

⁴ Algunos autores distinguen entre tecnología y técnica. Consideran la tecnología como una rama del saber constituida por el conjunto de conocimientos y de habilidades necesarios para la creación, mejora y utilización de las técnicas. La técnica es el conjunto de operaciones que deben ser efectivamente realizadas para la fabricación de un bien, lo que exige a menudo el empleo de herramientas y máquinas. En este libro nos haremos tal distinción y nos referiremos a todo ello como tecnología.

que la tecnología pretende su explotación, es decir, mientras la ciencia se propone conocer el mundo, la tecnología se propone controlarlo. En comparación con el científico, al tecnólogo le interesa primordialmente: a) la eficiencia, no la verdad y b) crear cosas útiles, no una comprensión más profunda de las mismas. Mientras que la ciencia estudia la manera y las razones por las que suceden las cosas, la tecnología se ocupa de hacer que las cosas sucedan. La tecnología está íntimamente relacionada a lo físico y lo material.

La tecnología parece ser un cuerpo de conocimientos mucho menos congruente y de más difícil acceso que la ciencia, dado que no obtiene validez mediante la discusión y la crítica, sino que la obtiene a través de las necesidades que pretende satisfacer. La medida del éxito de la tecnología no es su brillantez conceptual ni su elegancia metodológica ni los conocimientos que la sustentan, sino los beneficios que permite alcanzar. Ahora bien, más allá del resultado inmediatamente pretendido, los efectos de la tecnología influyen siempre sobre el ser humano y su entorno (Rapp, 1981).

El conocimiento obtenido por la ciencia es de un valor ético neutral, pero no así la tecnología, ya que su utilización puede hacerse con finen moralmente aceptables o reprobables. En la creación de la tecnología influye, de forma muy especial, la filosofía y la ideología. La filosofía inspira problemas, métodos y teorías científicas. La ideología proporciona, para bien o para mal, valores y finalidades: es la que determina qué vale la pena hacer y qué es menester evitar. Con ello sugiere un estilo de vida (Bunge, 1982).

La tecnología es un cuerpo de conocimientos acerca de ciertas clases de sucesos y actividades, que ha generado una determinada tasa de progreso económico durante miles de años (Rosenberg, 1982). Muchos de esos conocimientos son de naturaleza explícita, pero otros son de naturaleza tácita. En el campo empresarial, la tecnología es una oportunidad para satisfacer una necesidad sentida por la empresa o por el mercado. La tecnología recibe su impulso de las alteraciones en el modelo de la demanda de los clientes o de las actividades de los competidores en el mercado. También surge de la mente creativa de los inventores como una forma de solucionar viejos o nuevos problemas.

La creación de una tecnología conlleva una elevada incertidumbre, dado que nunca pueden predecirse los resultados a partir de los insumos utilizados. La serendipia y la suerte juegan un papel importante. En el análisis de la tecnología es necesario tener en cuenta una serie de características que afectan a su potencial para crear ventajas competitivas, y que pasamos a comentar.

Complejidad. En general, una tecnología es una unidad de análisis formada por diferentes partes que se interrelacionan entre sí, por lo que debe conceptualizarse como un sistema con límites que impiden su expansión. Así pues, una tecnología es un sistema tecnológico formado por la combinación de un número finito de módulos o componentes (figura 6) que pueden, a su vez, tener la consideración de sistemas tecnológicos (tecnologías) en relación a sus partes. Un automóvil es una tecnología formada por un motor, un sistema eléctrico, un volante, un sistema eléctrico, etc. El motor, a su vez, es un sistema formado por un cartel, el bloque y la culata. De esta forma, una tecnología compleja se genera y desarrolla en campos tecnológicos especializados. Sin embargo, las ventajas que produce sólo pueden conseguirse combinando las diferentes tecnologías especializadas en un sistema bien organizado. Cuantas más tecnologías puedan acoplarse, mayor número de combinaciones se puede obtener, por lo que es posible ampliar el abanico de alternativas en la solución de problemas. El número de componentes (tecnologías especializadas) varía ampliamente de

unas tecnologías a otras, lo que influye en el coste del desarrollo y en la funcionalidad. En suma, el sistema tecnológico se apoya en la propiedad combinatoria, que significa que una tecnología se puede combinar con otras para obtener una tecnología diferente; dándose incluso el caso, por lo demás usual, de que una determinada combinación de tecnologías conocidas da lugar a una nueva tecnología totalmente original (innovación). La investigación sugiere que la mayoría de las innovaciones no se deben al descubrimiento de algo totalmente nuevo, sino que es el resultado de nuevas combinaciones de tecnologías y materiales conocidos (Usher, 1964).

La fusión tecnológica es la aptitud para combinar las bases de conocimiento alojadas en diferentes capacidades tecnológicas que previamente no habían trabajado juntas (Kodama, 1992). Los descubrimientos tecnológicos crean oportunidades para la fusión de las tecnologías. Dichas oportunidades se producen cada vez con mayor frecuencia, creando nuevas clases de productos y procesos a lo largo de la corriente tradicional de actividades que crean valor en una industria. La fusión de tecnologías anteriormente diferenciadas, que dan lugar a otras nuevas, es una fuente tan importante de innovación como la tradicional invención científica.

	Desarrollador Jobmaster de las herramientas Stanley	Patines en línea Rollerblade	Impresora DeskJet de Hewlett-Packard	Automóvil New Beetle de Volkswagen	Avión Boeing 777
Volumen de producción anual.	100.000 unidades / año	100.000 unidades / año	4 millones unidades / año	100.000 unidades / año	50 unidades / año
Ciclo de vida del producto.	40 años	3 años	2 años	6 años	30 años
Precio de venta.	\$ 3	\$ 200	\$ 300	\$ 17.000	\$ 130 millones
Número de partes únicas (número de partes).	3 partes	35 partes	200 partes	10.000 partes	130.000 partes
Tiempo de desarrollo.	1 año	2 años	1.5 años	3.5 años	4.5 años
Equipo de desarrollo interno (magnitud pico).	3 personas	5 personas	100 personas	800 personas	6.800 personas
Equipo de desarrollo externo (magnitud pico).	3 personas	10 personas	75 personas	800 personas	10.000 personas
Coste (\$ USA) de desarrollo.	\$ 150.000	\$ 750.000	\$ 50 millones	\$ 400 millones	\$ 3 mil millones
Inversión (\$USA) para producción.	\$ 150.000	\$ 1 millón	\$ 25 millones	\$ 500 millones	\$ 3 mil millones

Figura 6: Atributos de cinco tecnologías (Ulrich y Eppinger, 2004)

La concepción sistémica de la tecnología obliga a optimizar el sistema como un todo, evitando, de esa forma, incurrir en una productividad negativa, que acontece cuando se maximizan unas tecnologías del sistema en detrimento de otras, dando lugar a un sistema tecnológico (tecnología) inferior. Además, este enfoque permite considerar entre los diferentes módulos o componentes: a) la existencia de relaciones (interdependencias) y b) el surgimiento de desfases y desequilibrios, como resultado inevitable de su dinámica

diferenciada. En este sentido, hay que tener en cuenta que la eficiencia del sistema se verá afectada por la tecnología más débil (lectura 8).

LECTURA 8: LA INSULINA

Los esfuerzos fallidos de Eli Lilly para crear insulina aún más pura demuestran el impacto de la saturación. Muchos diabéticos utilizan insulina todos los días para ayudar a mantener el nivel apropiado de glucosa en la sangre. Históricamente, la insulina era elaborada a partir del páncreas de vacas y cerdos. Durante gran parte del siglo XX, los fabricantes se centraron en el incremento de la pureza de la insulina. En 1925, las impurezas se mantenían en cincuenta mil partes por millón (ppm), reduciéndose a diez mil ppm en 1950. En el año 1980 se había reducido a solo diez ppm, principalmente como resultado de la inversión y el desarrollo realizado por Eli Lilly, el fabricante líder de insulina en el mundo.

A pesar de la pureza que Eli Lilly había sido capaz de lograr, la insulina animal era aún un poco diferente a la humana. Una fracción del 1% de los pacientes diabéticos manifestaron resistencia a sus sistemas inmunológicos cuando se los trató con insulina animal, por lo que Lilly contrató a Genentech para crear bacterias genéticamente alteradas que pudieran producir proteínas de insulina que fueran cien por cien puras y equivalentes a la insulina humana. Después de invertir casi mil millones de dólares en el esfuerzo, Lilly presentó en el mercado “Humulina”, su marca de insulina, con un precio especial un 25% mayor que los otros productos de insulina.

El mercado no estaba entusiasmado con la Humulina. El crecimiento de las ventas fue decepcionantemente lento, y Lilly tuvo problemas para mantener tal precio especial en el producto. “En retrospectiva –indicó un estudioso de Lilly–, el mercado no estaba terriblemente insatisfecho con la insulina de cerdo. De hecho, estaba bastante feliz con ella”. Lilly tuvo que gastar una suma significativa de dinero y recursos organizativos para crear un producto que saturó la demanda. La mayoría de los consumidores de insulina no querían o no necesitaban un producto más fiable y, por lo tanto, no estaban dispuestos a pagar más por una insulina que, si bien era técnicamente superior, no tenía un impacto importante en el tratamiento de sus condiciones.

Mientras que las ventas de Humulina resultaron decepcionantes, un pequeño fabricante de insulina danés, Novo Nordisk, identificó perfectamente una dimensión de rendimiento no saturada en el mismo mercado: la comodidad. La compañía desarrolló una línea de lápices de insulina, que hizo mucho más cómoda su aplicación por parte de los pacientes diabéticos. Tradicionalmente, la gente con diabetes tenía que llevar consigo una jeringa, cargar desde una ampollita la cantidad precisa de insulina, sostener la aguja y apretar la jeringa muchas veces para retirar las burbujas de aire. Por lo general, tenían que repetir el proceso para cargar un segundo tipo de insulina desde una ampolla diferente. Sólo entonces podían inyectarse la insulina.

El lápiz de Novo simplificó el problema. Tenía un cartucho que contenía una mezcla de dos tipos de insulina, y los usuarios simplemente tenían que girar un pequeño dial para determinar la cantidad de insulina que necesitaban inyectarse, poner la aguja del lápiz en la piel y presionar el botón. El lápiz Novo redujo lo que había sido un proceso de uno o dos minutos a diez segundos. Para los pacientes diabéticos que utilizaban insulina todos los días o varias veces al día, esta mejora en la eficiencia representaba un avance importante.

Novo pudo establecer un aumento del 30% en el precio por unidad de insulina. El éxito de los lápices y los cartuchos premezclados ayudó a la compañía a incrementar su cuota en el mercado mundial de manera sustancial y rentable. La conveniencia del lápiz atrajo muchos más pacientes al mercado de consumo de insulina, especialmente en Europa. Ambos, Eli Lilly y Novo, habían satisfecho las necesidades del mercado principal en cuanto a la pureza de insulina. Los reguladores aseguraron la fiabilidad de ambas marcas. Cuando la necesidad del mercado de fiabilidad se hubo saciado, las bases de la competencia se volvieron hacia la comodidad, y la empresa que había brindado un producto más práctico se benefició.

Fuente: Anthony, S. D.; Sinfield, J. V.; Johnson, M. W. y Altman, E. J. (2008): *The Innovator's Guide to Growth. Putting Disruptive Innovation to Work*, Harvard Business School Press, Massachusetts.

Acumulabilidad. Toda tecnología, una vez introducida en el mercado, es susceptible de múltiples mejoras, tendentes a aumentar sus prestaciones y/o reducir su coste y, así, satisfacer mejor las necesidades de los clientes. Una trayectoria tecnológica es el camino de progreso establecido por la elección de un concepto tecnológico central en el comienzo de su desarrollo. Las decisiones sobre el producto, limitadas por las elecciones técnicas anteriores y por la evolución de las preferencias de los clientes, influyen en esa trayectoria (Utterback, 1994). Las decisiones y sucesos tecnológicos previos irrumpen en el presente y dan forma al futuro. Los pasos dados en el pasado no

se pueden desandar; las acciones del presente están influidas por la trayectoria seguida para llegar aquí. La historia importa (North, 1990). La trayectoria tecnológica se representa mediante una curva S: un gráfico sinuoso que relaciona el tiempo (más bien el esfuerzo realizado para mejorar el producto a lo largo del tiempo) con el desempeño tecnológico (figura 7). Al principio, el progreso es muy lento debido a que los fundamentos de la tecnología son poco conocidos. Si la tecnología es muy diferente de las tecnologías previas, no existen rutinas de evaluación que permitan a los investigadores medir su progreso o potencial de mejora. Además, hasta que la tecnología no haya conseguido un cierto grado de legitimación, puede ser difícil atraer nuevos investigadores para participar en su desarrollo. Después, el ritmo se acelera vertiginosamente, conforme los científicos o las empresas consiguen un conocimiento más profundo sobre la tecnología. La tecnología comienza a conseguir legitimidad como un esfuerzo que merece la pena, atrayendo a otros desarrolladores. Igualmente, se crean medidas para valorar la tecnología, permitiendo a los investigadores dirigir su atención hacia aquellas actividades que generen los mayores avances por unidad de esfuerzo, lo que hace que el rendimiento se incremente rápidamente. Por último, cada vez es más difícil y costoso lograr progresos técnicos (Foster, 1986). El desempeño de la tecnología se irá aproximando asintóticamente a un límite físico o natural. Años de experimentación y de mejoras incrementales habrán agotado muchas de las oportunidades para mejorar el coste y la eficacia, y la curva en forma de S se aplana (Schilling, 2008).

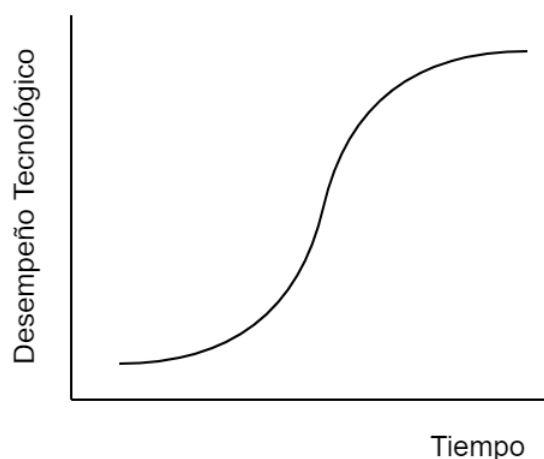


Figura 7: Trayectoria tecnológica (Foster, 1986)

Ahora bien, las mejoras tecnológicas (la acumulabilidad) sólo se producen si tienen valor para el cliente. Una mejora tecnológica que represente un importante avance científico, pero que no contribuya a satisfacer una necesidad del mercado, no se llevará a cabo, ya que representa un coste para la empresa que no añade valor al producto.

La dependencia de la trayectoria se refiere a los recursos tecnológicos que la empresa ha ido acumulando a lo largo del tiempo, a través del conocimiento adquirido y de las inversiones realizadas. La elección de la tecnología pasa a ser el vínculo por el cual condiciones económicas dominantes pueden influir en las dimensiones futuras del conocimiento tecnológico (David, 1975). Esta acumulación histórica integrada configura el potencial de la tecnología. Característica que no está sometida al determinismo, porque todo proceso de desarrollo de la tecnología está sujeto a la incertidumbre de los resultados. La dependencia de la trayectoria es vista como algo que restringe lo que puede ser hecho. Por tanto, el desarrollo de la tecnología está sujeto a un efecto de *lock-in* (amarre), al estar condicionados los conocimientos futuros por los conocimientos

previos, lo que, de alguna manera, puede limitar su apertura a nuevos campos. Es difícil cambiar de una tecnología a otra incompatible debido a los costes de cambio que hay que afrontar.

Especificidad. Lejos de satisfacer necesidades universales, la tecnología obtiene su importancia en el seno de un contexto cultural o sistema de valores específico. La tecnología es específica (idiosincrásica) al entorno, ya que una necesidad para un pueblo, generación o clase social puede carecer de valor utilitario o puede ser un lujo superficial para otro pueblo, generación o clase social. Así, mientras los europeos promovían enérgicamente el transporte rodado, entre los siglos III y XII d.C., las civilizaciones del Oriente Próximo y Norte de África (región de origen de la rueda) abandonaron el transporte con vehículos rodados y adoptaron el camello como animal de carga (Basalla, 1988). En igual sentido, los fabricantes de coches japoneses se especializaron en el desarrollo de coches pequeños, mientras que los norteamericanos se centraron en los coches grandes. Cada tipo de coche siguió su propia trayectoria tecnológica.

Durante la primera mitad del siglo veinte, los observadores ingleses, con frecuencia, señalaban con cierto asombro que los productos norteamericanos se diseñaban para adaptarse no al consumidor, sino a la máquina. Es decir, en Norteamérica el productor de bienes de capital tomaba la iniciativa respecto al diseño de máquinas y suprimió con éxito las variaciones en el diseño del producto que no servían con claridad a un propósito determinado. En otras palabras, logró un alto grado de estandarización en la maquinaria, que simplificó mucho sus problemas de producción y redujo el precio de los bienes de capital. Sin embargo, este tipo de máquinas no eran adecuadas para Inglaterra, que exigía más personalización a los productos. La ausencia de estandarización de los productos exigida por el mercado complicó en gran medida el proceso de adaptación de la tecnología y organización de la producción masiva hasta mediados del siglo XIX (Rosenberg, 1976).

En muchos casos, la especificidad se apoya en el conocimiento tácito. Por ejemplo, cuando el gobierno británico compró una gran cantidad de maquinaria norteamericana para fabricar armas de fuego en la década de 1850 y la introdujo en el arsenal de Enfield, se descubrió que era absolutamente imprescindible emplear a un gran número de mecánicos norteamericanos y personal supervisor del mismo país (Rosenberg, 1976). Por tanto, la asimilación de tecnologías desarrolladas en otros países generalmente obliga a asumir costes, ya que la especificidad de la tecnología significa que las organizaciones que la adoptan tienen que realizar modificaciones tecnológicas y obtener conocimiento adicional (en su mayoría de tipo tácito) si quieren hacerla funcionar.

Transversalidad. Una tecnología es fungible (o transversal), ya que puede aplicarse a diferentes productos y en diversos sectores industriales. Gillette retrasó la introducción de la hoja de afeitar de acero inoxidable porque canibalizaba su hoja «súper azul» de acero de carbono. En ese momento ninguno de sus competidores directos representaba una seria amenaza. Sin embargo, Wilkinson, que se encontraba compitiendo en el sector de menajes de cocina, también desarrolló la tecnología de acero inoxidable, y aprovechó la oportunidad para introducirse en un nuevo sector, el de las maquinillas y hojas de afeitar, y ahí se quedó.

Un buen ejemplo de tecnología transversal (o fungible) sería el láser, acrónimo de *Light Amplification by Stimulated Emission Radiation*. Esta tecnología se puede emplear en diferentes sectores. En el ámbito médico se aplican diferentes tipos de láser en

oftalmología, odontología, dermatología, fisioterapia, cirugía, etc. En la industria se usan para cortar materiales metálicos, en la soldadura, en el taladro de diamantes, en el recorte de componentes microelectrónicos o en la síntesis de nuevos materiales; también se usan en otros sectores, como defensa, arquitectura o investigación, y en productos comerciales como impresoras, lectores de CD, lectores de código de barras, en la limpieza y restauración de obras de arte e, incluso, para animar las fiestas.

Variedad tecnológica. La amplia variedad de estrategias, organizaciones, instituciones y reglamentaciones empresariales, junto a cuanto tiene de tácita y de específica la capacidad tecnológica, hace que, tanto dentro de un mismo país, como en países diferentes, coexisten los más diversos historiales tecnológicos, las más diferenciadas maneras de afrontar un mismo problema y los más variados niveles de calificación, pudiendo todos ellos ser igualmente eficaces en determinadas condiciones. De tales premisas nace el concepto de «variedad tecnológica» (Cimoli y Dosi, 1992).

En el inicio del siglo veinte no era en modo alguno obvio que el moderno motor de automóvil —el motor de combustión interna Otto de cuatro tiempos— ganara a todos sus competidores. En 1900, se fabricaron en los Estados Unidos 4.192 coches. De éstos, 1.681 era de vapor, 1.575 eléctricos y sólo 936 de gasolina. Sin embargo, poco después, el motor de combustión interna empezó a destacar claramente. En la Feria del Automóvil de Nueva York de 1901, se presentaron 58 modelos de vapor, 23 eléctricos y 58 de gasolina. En 1903, el número de modelos de vapor y eléctricos era de 34 y 51, respectivamente, mientras que los coches de gasolina eran 168. En la muestra de 1905, la derrota fue total: los 219 modelos de gasolina exhibidos superaban en una proporción de 7 a 1 al conjunto total de vehículos eléctricos y de vapor (Basalla, 1988).

El coche eléctrico no hacía ruido ni olor, y era muy fácil de arrancar y de conducir. Ningún otro vehículo de motor podía compararse en comodidad y limpieza o simplicidad de construcción y facilidad de mantenimiento. Los primeros vehículos eléctricos fueron producidos en 1894. Entre sus inconvenientes destacan que era lento, incapaz de superar subidas escarpadas y caro de adquirir y mantener. Pero, por encima de todo, tenía un limitado alcance. Sus pesadas baterías de almacenamiento de plomo y ácido habían de recargarse cada 50 kilómetros aproximadamente, por lo que no era apto para adentrarse en el campo o viajar a una ciudad lejana. Ahora bien, como el reducido ámbito de acción del vehículo eléctrico satisfacía las exigencias del servicio de reparto urbano, se construyeron camiones eléctricos para transportes de bienes dentro de las ciudades. Los automóviles de vapor gozaron de una gran popularidad a comienzo de siglo. El aspecto general del automóvil de vapor era similar al de un coche de gasolina de la época. El de vapor no era tan silencioso como el eléctrico, pero su precio de compra y sus costes de mantenimiento eran considerablemente menores, y un potente motor le permitía superar todo tipo de carreteras sin esfuerzo. El motor de vapor tenía muchas menos partes móviles que el de gasolina, lo que significaba menos desgaste del motor y más fácil mantenimiento. Para calentar el agua podían utilizarse destilados de petróleo de baja graduación. El coche de vapor necesitaba repostar cada 50 kilómetros aproximadamente. Otro problema era que el tiempo necesario para generar vapor, a fin de que el coche arrancara cada día por primera vez, era de media hora de espera. Finalmente se redujo a unos cuantos minutos (Basalla, 1988).

Apropiabilidad. Este término se relaciona con la capacidad del propietario de la tecnología para capturar el valor derivado de su aplicación comercial. La tecnología como

conocimiento es un bien público. Por tanto, es necesario privatizarla para excluir a terceros de su uso. La expresión régimen de apropiabilidad se usa para describir aquellas propiedades del conocimiento tecnológico y de los artefactos técnicos, los mercados y el entorno legal que permiten las innovaciones y las protegen, en grados diversos, de las imitaciones de los competidores, manteniéndolas como activos que producen beneficios (Dosi, 1982). El régimen de apropiabilidad es más fuerte si una parte del conocimiento indispensable para configurar y fabricar el diseño es de naturaleza tácita (Teece, 1986) y, por tanto, difícil de decodificar por los competidores. También resulta de difícil acceso el conocimiento que se mantiene en secreto en los procesos productivos. En un régimen de apropiabilidad fuerte, el propietario es capaz de capturar una parte sustancial del valor creado. En un régimen de apropiabilidad débil, otros participantes obtienen la mayor parte del valor. En cualquier caso, resulta muy difícil que el propietario pueda cosechar todos los beneficios que la tecnología genera, ya que existen derrames (*spillover*) tecnológicos, es decir, filtraciones hacia otras empresas de los conocimientos tecnológicos desarrollados internamente, por lo que permiten a los no propietarios obtener una parte de los beneficios que genera la tecnología. La apropiabilidad se halla en relación inversa respecto a la facilidad de imitación o de sustitución de la tecnología.

La tecnología posee la propiedad de no rival. Es decir, puede utilizarse simultáneamente en múltiples aplicaciones y por diferentes usuarios. Esta propiedad permite al propietario de la tecnología conceder permisos de fabricación y uso a otras partes interesadas mediando una contraprestación económica.

Límite. El movimiento a lo largo de una curva *S* es generalmente resultado de innovaciones incrementales que van teniendo lugar dentro de la trayectoria tecnológica existente, mientras que un salto hacia la siguiente curva tecnológica implica adoptar una innovación radical. En consecuencia, toda empresa debe ser especialmente hábil para decidir entre alternativas relacionadas con la eficacia y la eficiencia tecnológica. La eficacia está relacionada con el tipo de curva en *S* que va a seguir la empresa. La eficiencia determina la pendiente. La eficacia atañe a la estrategia a elegir y la eficiencia a la utilización económica de los recursos (Foster, 1986). La tecnología tiene un límite (o frontera tecnológica), no siempre fácil de delimitar. El límite de una tecnología no implica que no pueda existir otra capaz de resolver mejor el problema del cliente. Es, por lo tanto, la mejor señal para plantearse nuevas alternativas empresariales. Este límite representa el logro que la tecnología puede alcanzar, valorado en términos de sus dimensiones técnica, económica y social. Los límites de la tecnología condicionan su desarrollo y puesta en práctica y afectan a su integración en la empresa. Tales límites se pueden clasificar en cuatro grandes grupos (Rapp, 1981):

- La estructura del mundo material, definida por la lógica y las leyes científicas.
- Los recursos intelectuales: la situación del conocimiento científico, la capacidad y el saber tecnológico.
- Los recursos materiales: disponibilidad en cantidad y calidad de materias primas, maquinaria, mano de obra y energía.
- Las condiciones sociales: aceptación del mercado, restricciones políticas, jurídicas y sociales.

Lo que maneja la tecnología dentro de esos límites son los problemas prácticos establecidos por la sociedad. Es decir, la sociedad marca los límites sobre los tipos de soluciones que pueden considerarse seriamente y examina detenidamente aquellos que se prueban. Por tanto, parece haber un elemento social ineludible en la tecnología. La ciencia

sólo define los límites externos de los problemas que pueden intentar abordarse con ella. Por ejemplo, en la batalla librada en el mercado de los detergentes en los años cincuenta, Procter & Gamble (P&G) y sus competidores se afanaban por conseguir un producto que permitiese obtener la ropa más limpia. Lograron el producto, pero, con frecuencia, la ropa recién lavada adquiría un aspecto gris y apagado que los consumidores asociaban con la suciedad. De hecho, este aspecto se debe a fibras rotas y deshilachadas, pero el consumidor no apreciaba este aparente recóndito detalle técnico. Para solucionar el problema, P&G decidió añadir abrillantadores ópticos al detergente, es decir, sustancias químicas que reflejaran la luz. La ropa retenía estas sustancias, que hacían que apareciese más luminosa y, por tanto, más limpia a los ojos de los consumidores, aunque, en el verdadero sentido de la palabra, no estaba en absoluto más limpia. Posteriormente, el departamento técnico se planteó que la ropa quedase más luminosa, pero ahora se había traspasado un nuevo tipo de límite. No el límite del brillo óptico, sino más bien el límite de la percepción del brillo óptico y de la limpieza. En la época actual, los consumidores no quieren sólo brillo óptico, quieren además «aroma fresco». Estos cambios pueden deberse a una modificación del entorno social o económico. Por ejemplo, las nuevas leyes del medio ambiente, una variación en el precio de la energía, o la aparición de un producto completamente nuevo (Foster, 1986). En cualquier caso, el límite del sistema tecnológico se puede expandir a campos diferentes a lo que realmente constituye la verdadera esencia de la tecnología, por lo que resulta difícil de establecer.

3.1. El paradigma tecnológico

Fue Dosi (1982) quien acuñó el término paradigma tecnológico y elaboró una teoría a partir de las ideas de Kuhn (1962). Tanto los paradigmas científicos como los tecnológicos incorporan una perspectiva, un conjunto de procedimientos, una definición de los problemas relevantes, un esquema de indagación y el conocimiento específico relativo a su solución (Dosi, 1982). Un paradigma tecnológico define contextualmente las necesidades que se han de satisfacer, los principios científicos en los que se apoya la tarea y el material tecnológico a utilizar. En otras palabras, un paradigma tecnológico se puede definir como un modelo de solución de determinados problemas tecnoeconómicos basados en principios muy selectos derivados de las ciencias fácticas, junto con reglas específicas orientadas a la adquisición de nuevos conocimientos y a salvaguardarlos, cuando sea posible, de una rápida difusión a los competidores (Dosi, 1988). Los paradigmas tecnológicos son la fuente de los métodos, problemas y normas de resolución de problemas aceptados por los tecnólogos en cualquier momento dado. El paradigma conlleva prescripciones sobre la dirección del cambio tecnológico a seguir y las alternativas a rechazar.

El paradigma tecnológico aporta los conocimientos necesarios y marca la agenda de investigación de los tecnólogos en el desarrollo y mejora de una tecnología que satisface una necesidad del mercado. Un paradigma tecnológico es: a) un ejemplar, es decir, un artefacto que ha de ser desarrollado y mejorado, tal como un coche y b) un conjunto de elementos heurísticos (matriz disciplinar, en terminología de Kuhn), que se apoya en los conocimientos actuales para tratar de dar respuesta a preguntas tales como ¿hacia dónde nos dirigimos?, ¿dónde deberíamos buscar?, ¿de qué tipo de conocimiento nos nutrimos?, ¿cómo hacer las cosas? y ¿cómo mejorarlas? (Dosi, 1988). El término «base del conocimiento», acuñado por Nelson y Winter (1982), denomina el conjunto de *inputs* de información, conocimiento y habilidades del que se nutre los tecnólogos cuando buscan soluciones innovadoras dentro de un paradigma dado. Los elementos heurísticos son

específicos a cada tecnología, aunque en modo variable. En otras palabras, cada paradigma conlleva una tecnología del cambio tecnológico (Dosi, 1988).

La mejora de algunos atributos de comportamiento típicos de los ejemplares subyace en la idea de postes de indicadores tecnológicos de Sahal (1985). Un poste indicador tecnológico sería un artefacto básico cuyas características económicas están mejorando progresivamente. Por otra parte, Rosenberg (1976) resalta la importancia de mecanismos de enfoque, esto es, problemas típicos, oportunidades y objetivos que tienden a enfocar el proceso de investigación en una dirección particular. A menudo ocurre que los modelos de resolución de problemas prototipo, las reglas sobre cómo buscar y en qué objetivos concentrarse y las creencias sobre lo que el mercado quiere se convierten en una visión compartida por la comunidad de tecnólogos (Dosi, 1988).

Los tecnólogos que desarrollan su agenda de investigación dentro de un determinado paradigma son acríticos con la tecnología que desarrollan, al tiempo que excluyen tecnologías alternativas para usos similares. La investigación dentro de un paradigma hace que las actividades de innovación sean fuertemente selectivas, ya que están orientadas en direcciones bastante precisas, al desarrollarse dentro de los límites del paradigma, lo que aboca en un proceso acumulativo en la adquisición de conocimiento para la resolución de problemas. Es precisamente la paradigmática naturaleza acumulativa del conocimiento tecnológico lo que explica el carácter relativamente ordenado de los esquemas observados de cambio técnico. En este sentido, la exploración de tecnologías particulares y el desarrollo de modelos concretos de resolución de problemas aumentan las cualificaciones de las empresas y las industrias en estas direcciones específicas y, por tanto, aumenta el incentivo para hacerlo así en un futuro. Estas formas de rendimientos crecientes dinámicas específicos a la tecnología tienden a amarrar (*lock-in*) los procesos de cambio tecnológico dentro de un paradigma particular, llevando consigo un reforzamiento mutuo (una retroalimentación positiva) entre un cierto esquema de aprendizaje y un esquema de asignación de recursos a actividades innovadoras donde el aprendizaje ya ha tenido lugar en el pasado. En este sentido, los procesos de investigación tecnológica en cada empresa son también procesos acumulativos. Lo que la empresa puede esperar conseguir tecnológicamente en el futuro está estrechamente restringido, por lo que ha sido capaz de hacer en el pasado (Dosi, 1998).

Definamos como una trayectoria tecnológica –avenidas de la innovación en terminología de Sahal (1985)– la actividad del proceso tecnológico junto con las elecciones (*trade-offs*) tecnológicas y económicas definidas por el paradigma (Dosi, 1988). La trayectoria describe los avances tecnológicos acumulativos que se producen en el ejemplar (artefacto) y que contribuyen a satisfacer mejor las necesidades de los clientes. Las decisiones sobre el producto, limitadas por las elecciones técnicas anteriores y por la evolución de las preferencias de los clientes, influyen en las diversas trayectorias (Utterback, 1994). Por ejemplo, cuando surgió el vídeo fueron introducidos varios formatos (ejemplares) en el mercado (*VHS* y *Betamax*, entre otros), cada uno de ellos incrustado en una trayectoria tecnológica incompatible y, hasta cierto punto, inconmensurable. La trayectoria tecnológica son desarrollos tecnológicos y mejoras relativamente menores a lo largo de un ejemplar (producto) creado dentro de un paradigma.

La dependencia de la trayectoria de una empresa significa que sus tecnologías están limitadas no sólo por su historia y posición actual, sino también por sus específicas oportunidades futuras. Las empresas suelen ser lentas en la ruptura de sus rutinas una vez

establecidas. Frecuentemente tampoco es posible para una empresa, de forma individual, cambiar su lógica de desarrollo tecnológico o alterar la demanda del mercado o de la sociedad. Hay una evidencia creciente de que las «avenidas de la innovación» específicas constituyen un aspecto extendido de los esquemas observados de cambio técnico (Sahal, 1985).

Los tecnólogos generalmente tratan de mejorar las características deseables que son específicas a un cierto producto, herramienta o mecanismo considerando las elecciones entre ellas. Relacionado con esto, la evidencia histórica indica que se ha derivado un impulso importante para la innovación a partir de «desequilibrios» entre las dimensiones técnicas que caracterizan una trayectoria (o avenida) tecnológica. Se pueden hallar también argumentos que extienden el alcance de los desequilibrios a la relación entre el cambio tecnológico y el papel social y el comportamiento de distintos grupos de tecnólogos. A su vez, las principales pautas de progreso tecnológicas están relacionadas con (Dosi, 1988): 1) la I+D formal en empresas y laboratorios de investigación; 2) los procesos informales relacionados con la difusión de la información y la innovación; 3) las externalidades de cooperación inter-empresas; 4) las innovaciones adoptadas de otras industrias y 5) los *inputs* de innovación insertados en equipos de capital y bienes intermedios.

El paradigma entra en un período revolucionario cuando surge un paradigma alternativo prometedor. El nuevo paradigma se caracteriza por satisfacer mejor que el actualmente en uso una necesidad del mercado (por ejemplo, cuando el automóvil sustituyó al coche de caballos). Los paradigmas en competencia son inconmensurables e incompatibles. El desenlace de un conflicto entre paradigmas no es fortuito, sino que comporta sus propios criterios de racionalidad. Por encima de todo, el paradigma triunfante debe dar un tratamiento satisfactorio a las necesidades del mercado. En el ámbito industrial son las pequeñas empresas innovadoras (aunque no siempre) las que introducen la mayoría de los nuevos paradigmas, ya que sus investigadores no están condicionados por las tecnologías en uso y pueden criticar las soluciones tecnológicas actuales proponiendo nuevas alternativas. En la etapa preparadigmática varias tecnologías incompatibles compiten entre ellas y con el paradigma establecido. Normalmente una de esas tecnologías se convierte en el nuevo paradigma dominante.

Un cambio en el paradigma implica generalmente rupturas en las trayectorias. Junto con bases de conocimiento diferentes y distintos prototipos de artefactos, las dimensiones tecnoeconómicas de la innovación también varían. Algunas características pueden que se vuelvan fáciles de conseguir, pueden emerger nuevas características deseables y otras pueden perder importancia. De manera relacionada, la visión de los tecnólogos sobre los avances tecnológicos futuros cambia, unida a un énfasis cambiante en las distintas elecciones que caracterizan los nuevos artefactos. De manera más general, también se ha sugerido que los principales paradigmas tecnológicos predominantes conllevan la utilización intensiva de algún *input* crucial disponible en abundancia a bajo coste (Dosi, 1988).

3.2. La investigación aplicada y el desarrollo experimental (I+D)

La tecnología se crea mediante la investigación aplicada siguiendo dos caminos, cuya separación no siempre resulta clara: a) la ciencia y sus orientaciones prácticas y b) las intuiciones y corazonadas sin un apoyo científico claro. Así pues, esta investigación está dirigida fundamentalmente hacia un objetivo práctico específico, y se emprende para

determinar los posibles usos de los resultados de la investigación básica o para determinar nuevos métodos o formas de alcanzar objetivos específicos predeterminados. En este sentido, Francis Bacon (1561–1629) clasificaba las tecnologías en dos categorías: a) las que dependen del estado general de la ciencia y sólo pueden realizarse cuando se dispone de los conocimientos científicos adecuados y b) las puramente empíricas, que, al no sustentarse en la ciencia existente, pueden lograrse en cualquier momento de la historia.

Una parte de la tecnología es el resultado de una investigación aplicada de «prueba y error» más bien tosca, que permite lograr resultados prácticos a pesar de desconocer muchas de las leyes y teorías científicas que los sustentan. De hecho, la concepción y desarrollo de un gran número de tecnologías no ha dependido de los progresos alcanzados previamente en la ciencia. Por ejemplo, la máquina de vapor, la invención por excelencia de la revolución industrial, fue anterior a la ciencia que formalizó los principios y leyes en los que se basaba. A su vez, la fabricación de útiles de piedra, una de las más primitivas tecnologías conocidas, floreció cerca de dos millones de años antes del advenimiento de la mineralogía y la geología (Basalla, 1988). En definitiva, los tecnólogos de muchas industrias resuelven problemas para los cuales no hay explicación científica y la tecnología resultante genera la investigación básica subsiguiente para encontrar los conocimientos científicos que sustentan los resultados alcanzados.

No obstante, mucho conocimiento científico puede ser aplicado en actividades económicas y sociales. Por tanto, la ciencia es una fuente directa de la tecnología. La investigación aplicada implica la consideración de todos los conocimientos existentes y su profundización, en un intento de solucionar problemas tecnoeconómicos específicos. En el sector empresas, la investigación aplicada consiste en preparar un nuevo proyecto para explorar un resultado prometedor obtenido en el marco de un programa de investigación básica. Los resultados de la investigación aplicada dan lugar a un producto único o un número limitado de productos, operaciones, métodos o sistemas. Este tipo de investigación desarrolla las ideas para convertirlas en algo operativo. Los conocimientos o informaciones obtenidas de la investigación aplicada son a menudo patentados, aunque igualmente pueden permanecer secretos (OCDE, 2003).

La investigación aplicada aumenta el stock tecnológico, pero no el científico, salvo por accidente, ya que, si se necesita avanzar en el conocimiento científico antes de resolverse un problema práctico particular y los costes esperados del esfuerzo investigador son muy elevados, no se iniciará la investigación aplicada sobre el problema y no se intentará la invención. En muchos casos la utilidad económica de un invento particular es tan grande que un esfuerzo inventivo resulta económicamente rentable, aunque el conocimiento científico básico sea escaso y, por tanto, el coste esperado de la fabricación del invento sea elevado. El intento de Edison de desarrollar una lámpara incandescente y el de Goodyear por mejorar las características del caucho son dos claros ejemplos. En estos casos, en virtud de que había escasos conocimientos científicos útiles, el procedimiento de la invención fue el de prueba y error, donde un ensayo queda indicado –pero sólo en forma burda– por una teoría muy vaga formulada a medida que prosigue la investigación. De todos modos, aunque los inventores sabían que probablemente resultaría costoso el logro del objetivo, creían que los beneficios, si se alcanzaba el éxito, serían suficientemente grandes para hacer rentable el esfuerzo (Nelson, 1959).

El invento es la tecnología recién creada, a menudo bastante tosca, que requiere ulteriores refinamientos antes de su aplicación o uso comercial. A veces, para estudiar la viabilidad comercial del invento y perfeccionar sus características técnicas se hace necesaria la construcción de un «prototipo», un modelo básico completamente funcional que posee las características esenciales y se comporta de forma equivalente a lo que se intenta producir, pero al que fácilmente puedan incorporársele cambios y modificaciones. Igualmente, para la realización de las actividades de desarrollo puede ser necesario construir una «planta piloto», ya que rasgos esenciales del proceso de producción no pueden ser, posiblemente, resultado exclusivo del conocimiento científico (Rosenberg, 1993). Estas plantas fueron esenciales en el pasado, ya que permitían conocer con precisión las características básicas de la producción, cuando ninguna otra alternativa era factible. Su construcción y funcionamiento tiene como objetivos principales adquirir experiencia y obtener datos técnicos (o de otro tipo) que puedan utilizarse para verificar hipótesis, obtener nuevas fórmulas de productos, establecer nuevas especificaciones para productos terminados, diseñar equipos y estructuras especiales para un nuevo proceso y redactar los modos operatorios o los manuales relativos al proceso (OCDE, 1981).

El desarrollo experimental consiste en trabajos sistemáticos fundamentados en los conocimientos existentes obtenidos por la investigación o la experiencia práctica, que se dirigen a la fabricación de nuevos materiales, productos o dispositivos, a establecer nuevos procedimientos, sistemas o servicios, o a mejorar considerablemente los que ya existen (OCDE, 2003). En esta etapa, el aprendizaje está orientado hacia las dimensiones comerciales del invento con objeto de transformarlo en una innovación: descubrir las necesidades del mercado (y de los submercados apropiados) e incorporar al producto las prestaciones suficientes que permitan satisfacerlas de la forma más eficiente posible. Este es un proceso muy sutil, particularmente cuando estamos tratando con productos que tienen un alto grado de complejidad (Rosenberg, 1982). Esta actividad también se lleva a cabo en la tecnología en uso para adaptarla mejor a las necesidades del mercado, dando lugar a una innovación incremental.

Orientación de la I+D. Una empresa puede seguir dos orientaciones diferentes para generar invenciones: el mercado (tirón de la demanda) y la ciencia (empujón de la ciencia). El empujón de la ciencia apoya que se debe vender todo lo que se puede producir. El punto de partida lo constituye la explotación comercial de los conocimientos generados en el departamento de I+D o afines. El tirón de la demanda considera que cuando el mercado constituye la inspiración básica de la investigación. En este caso se debe producir todo lo que se puede vender. La empresa ha de estudiar las necesidades del mercado con objeto de satisfacer la demanda latente y, así, poder desarrollarse y crecer. En general, una vez que se ha producido una innovación importante, puede aparecer un patrón de innovaciones e inventos secundarios generados por el lado de la demanda, a lo largo de muchas décadas (Freeman *et al.*, 1982).

En un estudio sobre 567 innovaciones de productos y procesos que tuvieron lugar en 121 empresas, juzgadas por los ejecutivos responsables como importantes y, presumiblemente, muy significativas desde el punto de vista comercial, Myers y Marquis (1969) observaron que la gran mayoría (tres cuartas partes) fueron estimuladas por la demanda del mercado o una necesidad de producción (45% por la demanda del mercado y 30% por una necesidad de producción). Sólo el 21 por 100 de las innovaciones fueron la respuesta a una oportunidad tecnológica percibida para un producto o proceso nuevo o mejorado y un 4 por 100 a un cambio administrativo. A su vez, Rothwell (1977) también considera que la

mayoría de las innovaciones con éxito (alrededor del 75% de ellas, como promedio) surgen como respuesta al reconocimiento de una necesidad del mercado (tirón de la demanda) y no del reconocimiento de un nuevo potencial técnico (empujón de la ciencia). En aquellos casos en que el potencial técnico surge primero (lo cual ocurre con más frecuencia en las innovaciones más radicales desde el punto de vista técnico), las empresas innovadoras determinan si existe una necesidad antes de seguir con el proyecto. Al mismo tiempo analizan con detalle las necesidades del usuario, a fin de interpretarlas en el diseño del nuevo producto o equipo. Por otra parte, las innovaciones estimuladas por la demanda se basan en una tecnología más antigua que las estimuladas por la oportunidad tecnológica (Myers y Marquis, 1969). En suma, el reconocimiento de la demanda es un factor de éxito más importante que el reconocimiento de un factor técnico (Marquis, 1982).

Resulta de utilidad la metáfora de las dos hojas de unas tijeras utilizada por Schmookler (1966) para describir el proceso de cambio tecnológico: una de ellas equivale a los descubrimientos científicos y la otra a los cambios en el estado de la demanda del mercado.

Por otra parte, los motivos económicos para reducir costes siempre existen en las operaciones comerciales, pero debido a que son tan difusos y generales no explican mucho sobre la secuencia y el tiempo de la actividad innovadora (Rosenberg, 1976).

Las tecnologías complejas crean presiones internas que enfocan la actividad investigadora en direcciones determinadas. Estas presiones operan tanto a nivel de planta como dentro de los componentes de los productos acabados (Rosenberg, 1976). El impulso lo puede proporcionar una necesidad del proceso productivo. Las nuevas máquinas en los procesos productivos provocan desajustes que deben corregirse mejorando las máquinas antiguas o sustituyéndolas por otras nuevas más eficientes. Por otra parte, cualquier mejora importante en el funcionamiento de un componente del producto, sea éste restrictivo o no, es probable que cree nuevos obstáculos en forma de limitaciones impuestas por otros componentes, para el logro de una mayor funcionalidad y/o eficiencia. Por tanto, se hace necesario eliminar los obstáculos mejorando (o innovando) los componentes que los provocan.

La escasez de materias primas favorece la búsqueda de nuevos materiales o el logro de mejoras tecnológicas que aumenten la productividad de los equipos en uso. Por ejemplo, a medida que la madera aserrada se hizo más cara, hacia mediados del siglo XIX, el esfuerzo inventivo europeo se dirigió cada vez más hacia la reducción del gasto de madera. Las guerras provocaron la búsqueda de productos sustitutivos. Así, la suspensión del suministro de caucho natural proveniente del sudeste asiático debido a la ocupación japonesa a principios de la Segunda Guerra Mundial fue responsable del avance de las investigaciones que culminaron con el nacimiento o creación de la industria norteamericana del caucho sintético (Rosenberg, 1976). También los desastres naturales contribuyen a orientar la actividad innovadora. Los incendios en alta mar fueron casi siempre seguidos por una carrera de aplicaciones de patentes para mecanismos que pretendían prevenir funestas repeticiones.

Si bien una empresa tiene un incentivo para reducir una parte de sus costes o gastos, las huelgas o el temor a ellas han servido, históricamente, como poderoso agente para dirigir las investigaciones de nuevas tecnologías en una determinada dirección (Rosenberg, 1976). Marx en su obra *Miseria de la Filosofía* afirmaba: “En Inglaterra, las huelgas han dado

origen regularmente a la invención y a la aplicación de nuevas máquinas”. La escasez de trabajadores también induce la búsqueda de innovaciones que ahorran mano de obra.

Todos estos mecanismos de inducción surgen como una oportunidad de conseguir beneficios extraordinarios o perder la oportunidad de hacerlo (huelgas). Las empresas atacan lo que, en un momento dado, consideran un «cuello de botella»: la más restrictiva coacción sobre sus operaciones.

4. RELACIÓN ENTRE CIENCIA Y TECNOLOGÍA

La relación entre ciencia y tecnología se contempla como dos corrientes de conocimiento acumulativo con muchas interdependencias y relaciones cruzadas, siendo las conexiones internas más fuertes que sus conexiones cruzadas (Brooks, 1994). Una corriente de conocimientos va desde la ciencia a la tecnología, mientras que otra recorre el camino inverso.

La ciencia es una fuente directa de nuevas ideas tecnológicas. Existe la creencia generalizada de que, antes de la primera mitad del siglo XIX, el progreso tecnológico era más o menos independiente del progreso científico y que, desde entonces, la relación entre ciencia y tecnología ha sido cada vez más estrecha. A partir de 1850, un buen número de tecnologías, desde la energía hidráulica hasta la industria química, dependieron o estuvieron inspiradas en los avances científicos. Esto no significa que se redujera la cantidad de avances tecnológicos puramente empíricos, aunque sí declinó su importancia relativa (Mokyr, 1990).

La ciencia también puede ser fuente de técnicas de diseño de los procesos de transformación. Esta relación ha llegado a ser muy fuerte debido al aumento de los costes de los tests y los procesos de evaluación (Brooks, 1994). Por ejemplo, a finales del siglo XIX, se descubrió que era rentable la aplicación del análisis químico en los procesos industriales a gran escala (tales como acero, cemento o productos alimenticios) como un medio de controlar la calidad y la consistencia dentro de los diferentes estados de la producción, así como para variar las propiedades del producto final (Rosenberg, 1982).

En muchos casos, los instrumentos científicos utilizados por los investigadores en la universidad se transforman en bienes de capital industriales (Rosenberg, 1993). Éstos cambian del estatus de una herramienta de investigación básica al de una herramienta de producción. Esta aplicación constituye, por tanto, una parte de los beneficios de la actividad de la investigación básica, independiente y distinta de los beneficios resultantes del conocimiento científico y de las eventuales aplicaciones de dicho conocimiento.

La preparación científica y el trabajo metódico son cada vez más necesarios en la etapa de desarrollo del proceso tecnológico. Para seguir en activo los inventores deben convertirse en expertos o, en caso contrario, abandonar (Hounshell, 1975). Ello no significa que no hay personas que tienen un papel decisivo en los procesos de invención, cuya inteligencia franca y abierta resulte muchas veces esencial en el origen de la actividad inventiva. Sin embargo, la ciencia es capaz de crear las habilidades y las rutinas necesarias en los procesos de innovación tecnológica. De ello se benefician los estudiantes universitarios que adquieren conocimientos científicos, destrezas, métodos, rigor y una red de contactos profesionales que les ayudarán a solucionar, en el futuro, los problemas de tipo tecnológico que se presentan en las empresas donde llevan a cabo su actividad. De hecho, la

contribución de la ciencia a la tecnología es, sobre todo, indirecta, en la forma de contratación de personas con nuevos conocimientos y habilidades valiosas, más que directa, a través de artículos publicados, aunque éstos también pueden ser muy útiles (Pavitt, 1993).

El conocimiento científico también ayuda a valorar la trascendencia tecnológica y económica de una innovación. Cuanto más compleja es la tecnología, más investigación básica se necesita para disminuir la incertidumbre acerca del funcionamiento correcto del producto, de que no ocasionará riesgos medioambientales ni pérdidas de vidas humanas durante su aplicación industrial (Brooks, 1994). Ello no significa que la ciencia determina la forma final de los artefactos, aunque sí dicta los límites de sus posibilidades físicas. La ley de Ohm no dicta la forma y detalles del sistema de iluminación de Edison ni las ecuaciones de Maxwell determinan la forma precisa que habían de adoptar los circuitos de un receptor de radio moderno (Basalla, 1988).

La ciencia es importante en el desarrollo de todo tipo de innovaciones, incluso en el caso de las innovaciones incrementales, aunque rara vez lo sea la investigación académica. Ésta si será relevante en aquellos sectores en los que la ciencia es prácticamente indistinguible de la tecnología, como ocurre con la biotecnología (Utterback, 1994).

La ciencia también puede ser una fuente de información en el desarrollo de la tecnología, al mostrar los caminos de la investigación aplicada o descubrir mediante investigaciones y análisis los fundamentos científicos y tecnológicos de los productos de la competencia. Las mayores invenciones del siglo XX hubieran sido imposibles sin la previa acumulación de conocimientos científicos. Incluso algunos avances científicos jugaron un papel crítico en la etapa de desarrollo tecnológico (Utterback, 1994). En la invención, como en la mayoría de las actividades dirigidas hacia un objetivo, el actor tiene varios caminos alternativos entre los que elegir. Cuanto mayor sea su conocimiento de todos ellos, más probable le resultará encontrar un camino satisfactorio explorando el menor número de alternativas posibles. Así pues, cuanto mayor sea el conocimiento científico, menor será el coste esperado de la realización de un invento particular (Nelson, 1959).

Por su parte, la tecnología contribuye al desarrollo científico, al menos de dos formas: como fuente de ideas y aportando instrumentos y técnicas de observación y medida. La concepción y desarrollo de un gran número de inventos no ha dependido de los progresos alcanzados previamente en la investigación científica. Por ejemplo, la máquina de vapor, la invención por excelencia de la revolución industrial, fue anterior a la creación de la ciencia que formalizó los principios en los que se basaba. En 1824, el francés Sadi Carnot, después de observar el funcionamiento de una máquina de vapor y preguntarse por qué eran mejores las de alta presión, concibió el núcleo de lo que después se conocería con el nombre de termodinámica (Mokyr, 1990). En relación a este planteamiento, Henssen y sus colegas soviéticos en el Congreso Internacional de Historia de las Ciencias de 1931, expusieron su idea de que la tecnología marcaba el camino y la ciencia lo seguía. Parece existir cierta evidencia empírica para esta hipótesis en el caso de los plásticos, los tintes y los productos químicos intermedios, ya que se produjo una demanda de trabajos científicos adicionales, secundaria y derivada, varios años después de la aparición de los productos en el mercado (Freeman, 1974). Sin embargo, estos ejemplos de ningún modo pueden servir para defender la idea de que la ciencia sigue a la tecnología. Ahora bien, sí es cierto que los tecnólogos de muchas industrias resuelven problemas para los cuales no hay explicación científica y que la solución que aportan genera la investigación básica subsiguiente para

encontrar los conocimientos que sustentan los resultados alcanzados. Es decir, una parte del trabajo de los científicos de hoy en día implica la sistematización y reestructuración del conocimiento y de las soluciones y métodos prácticos acumulados previamente por los tecnólogos.

Las mejoras en la instrumentación, por sus efectos diferenciales sobre las posibilidades de observación y medición en subcampos específicos de la ciencia, han sido desde hace mucho tiempo un determinante principal del progreso científico. Utilizando una larga lista (el microscopio y el telescopio, entre otras), sería fácil demostrar cómo la instrumentación posibilita el progreso científico. Sin embargo, dejar la discusión a este nivel constituiría una forma más bien tosca de determinismo tecnológico. Las relaciones entre tecnología y ciencia son mucho más atractivas (y dialécticas) que lo que este determinismo sugiere, pues la presión para perfeccionar determinados instrumentos reflejará la determinación de hacer avanzar un campo particular de la ciencia, así como una expectativa de que la instrumentación importante está madura para este perfeccionamiento. Finalmente, no puede sobrestimarse que las mejoras de las capacidades de observación tuvieron, por sí mismas, una utilización limitada hasta que se desarrollaron los conceptos y se formularon las hipótesis que dieron un significado potencial a las observaciones (Rosenberg, 1982).

En resumen, la relación entre tecnología y ciencia es interactiva, por lo que la expresión tecnología relacionada con la ciencia resulta, por lo general, preferible al de tecnología basándose en la ciencia o al de ciencia basada en la tecnología, con sus implicaciones de movimientos de ideas unidireccionales y supersimplificadas. Ya los primeros estudios de Carter y Williams (1959) señalan la importancia de los contactos, a veces informales, entre el mundo de la ciencia y el de la industria, así como la creciente interdependencia entre ciencia y tecnología.

5. EL INVENTO Y LA INNOVACIÓN

Un invento es una tecnología recién creada, a menudo bastante tosca, que requiere ulteriores refinamientos antes de realizar alguna función práctica que satisfice alguna necesidad. Enos (1962) identifica la fecha de una invención como “la primera concepción del producto en su forma sustancialmente comercial”. El significado de la palabra innovación la podemos encontrar en su raíz latina, *nova*, o nuevo. Se puede interpretar como la introducción de un objeto o método nuevos en el mercado. En este sentido, la innovación es invención más explotación económica. Un invento puede ser ingenioso o creativo, pero si los clientes potenciales no lo demandan, no se convertirá en una innovación. Por ejemplo, la Oficina de Patentes de Estados Unidos ha concedido hasta la fecha más de 4.000 patentes de trampas para cazar ratones (Hope, 1996) de las cuales solo unas veinte han dado dinero alguna vez⁵. Que se invente algo no significa, naturalmente, que vaya a ser o deba ser utilizado, ni siquiera cuando esté completo todo el desarrollo técnico. Que se use o deba ser usado depende de hasta qué punto es competitivo con los productos y procesos existentes en las condiciones vigentes de oferta y demanda (Nelson, 1968). Por lo tanto, una innovación debe permitir hacer algo que no era posible hacer antes, al menos tan bien o tan económicamente.

A Schumpeter debemos la distinción de extrema importancia entre inventos e innovaciones, que desde entonces ha quedado incorporada con carácter general a la teoría económica. Consideraba que los inventos no contenían ninguna actividad económica y, por

⁵ <http://uh.edu/engines/epi1163.htm>

tanto, los trató como un factor exógeno a la economía, dejándolos al margen de sus estudios, y se centró exclusivamente en la explotación del invento, es decir, consideraba que, a intervalos irregulares, las empresas seleccionaban alguna de las invenciones existentes, las cuales entraban en la escena *schumpeteriana* completamente desarrolladas. Sin embargo, lo hacían no como un objeto o proceso cuyo desarrollo es asunto de explícito interés. Fue Schmookler (1966) el primero en sostener que la actividad inventiva es también un fenómeno esencialmente económico, al considerar que la producción de inventos y de gran parte del conocimiento tecnológico es, en la mayoría de los casos, tan actividad económica como pueda serlo la producción del pan. Aunque los papeles del inventor y del innovador son diferentes, es igualmente cierto que muchos innovadores han sido también grandes inventores. Un ejemplo histórico de un gran inventor que también alcanzó el éxito como emprendedor es Thomas Alva Edison (1843–1931), el mago de Menlo Park.

Así como la creación de un invento supone la solución de un problema técnico y está relacionado con el control de la naturaleza, la innovación es un asunto eminentemente social, ya que, al comercializar una nueva tecnología, el innovador tiene que interactuar con un entorno formado por competidores, clientes, proveedores y la regulación, entre otros (Mokyr, 1990). De esta forma, la invención y la innovación se complementan. A corto plazo la compenetración no es perfecta, incluso es posible que se produzca una innovación sin la presencia de nada que pudiéramos identificar como invención. Así ocurre, por ejemplo, con el descubrimiento y utilización de una nueva fuente de materias primas. En cierto sentido, la innovación es la acción más importante, ya que permite recoger los frutos económicos de la investigación. Podemos considerar que la investigación es la transformación de dinero en conocimiento, y la innovación la transformación de conocimiento en dinero. Por ello, no debe entenderse como un concepto técnico, sino de raíz económica y social. La esencia de la innovación tecnológica es satisfacer una necesidad. Por lo tanto, hay que valorarla siempre desde la perspectiva del cliente. El valor de la innovación debe basarse en su viabilidad económica y en su capacidad para generar riqueza en un mercado competitivo, más que en alguna percepción intelectual de su valor como un nuevo concepto.

La innovación está impulsada por una comprensión sólida, mediante la observación directa, de lo que las personas quieren y necesitan en sus vidas y de lo que les gusta o disgusta respecto de cómo se fabrican, empaican, comercializan, venden y apoyan productos particulares (Brown, 2008). La innovación es, por otra parte, un proceso de solución de problemas que se enfrenta con dos incertidumbres: tecnológica y de mercado. La incertidumbre tecnológica es la información adicional sobre los componentes, las relaciones entre ellos, los métodos y las técnicas que contribuyen a hacer que el nuevo producto o servicio funcione de forma correcta, de acuerdo con unos parámetros previamente especificados. En suma, la incertidumbre tecnológica está relacionada con «cómo crear un nuevo producto y lograr que funcione». La incertidumbre de mercado es la información adicional sobre los clientes potenciales, sus necesidades y expectativas, necesaria para hacer llegar el producto al mercado y conseguir que los clientes lo compren. Está vinculada con «cómo vender el nuevo producto y conseguir que sea un éxito comercial» (Geroski, 1995). Por supuesto, la incertidumbre tecnológica y la de mercado están muy conectadas. Mientras más sepa una empresa sobre lo que desean los clientes, en mejores condiciones estará para tomar decisiones sobre la configuración del producto (Afuah, 1998).

Los inventos, al contrario de lo que normalmente se cree, no se transforman de forma inmediata en innovaciones tecnológicas. De hecho, la mayoría de ellos no lo hacen. Que se invente algo no significa, por lo tanto, que a continuación vaya a ser o deba ser comercializado. Para que un invento llegue a convertirse en una innovación hace falta un trabajo empresarial adicional para desarrollarlo, fabricarlo y comercializarlo. Dejando a un lado las cuestiones de falta de desarrollo y de utilidad quedan muchos obstáculos por superar antes de que un invento pase a formar parte de la vida económica y cultural de un pueblo. En las sociedades modernas, hay que reunir capital, trabajo y recursos naturales, convertir el modelo inventado en un producto aceptable para el consumidor, y fabricarlo de forma que pueda venderse con beneficio. En la figura 8 se recogen los años de retraso entre el invento y la innovación para una serie de productos. La fecha del invento se identifica con la solitud de la patente y la de la innovación con su explotación económica. Este intervalo o retraso varía considerablemente y, a menudo, es de diez o más años para las invenciones importantes. La duración del retraso es difícil de medir con exactitud debido a la naturaleza desordenada e incierta del proceso innovador, aunque puede superar los 50 años. Dicho retraso es producto de la complejidad de los problemas que deben resolverse antes de que un invento sea económicamente viable y/o aceptado por la comunidad, y está condicionado por la industria, el producto, el mercado y los recursos utilizados (Utterback, 1974). Algunos autores han destacado el hecho de que el tiempo que transcurre entre el invento y la innovación ha llegado a ser progresivamente más corto con el paso del tiempo (figura 9). Además, el intervalo se acorta cuando el inventor hace también de innovador en comparación a cuando se conforma con vender a otra empresa el invento para que sea ésta la que lo desarrolle y comercialice (Enos, 1962). Por otra parte, el retraso es menor cuando la innovación se dirige al consumidor en lugar de a la industria o cuando la innovación se desarrolla para el gobierno, como opuesta a la desarrollada para el mercado (Utterback, 1974).

Nuevo concepto	Invento	Innovación
Arrastre automático.....	1904	1939
Bolígrafo.....	1888	1946
Cinerama.....	1937	1953
Fundición de acero continua.....	1927	1952
Locomotora diesel.....	1895	1913
Iluminación fluorescente.....	1901	1938
Helicóptero.....	1904	1936
Insulina.....	1920	1927
Motor a reacción.....	1928	1941
Magnetófono.....	1898	1937
Penicilina.....	1928	1943
Servodirección.....	1925	1930
Radar.....	1925	1934
Radio.....	1900	1918
Detergentes sintéticos.....	1886	1928
Televisión.....	1923	1936
<Terylene>, fibra sintética.....	1941	1955
Titanio.....	1937	1944
Transistor.....	1948	1950
Xerografía.....	1937	1950
Cremallera.....	1891	1923
Acero inoxidable.....	1904	1912

Figura 8: Año estimado del invento y la innovación para determinados productos (Jewkes et al., 1969)

Las razones de los retrasos son múltiples. Puede surgir porque en muchas ocasiones hay que profundizar en el desarrollo de la tecnología y los materiales básicos. Esto resulta esencial antes de poder afirmar que ha llegado el momento oportuno para un invento importante. Existe siempre un vacío entre la capacidad para conceptualizar un artefacto o técnica y la habilidad para materializarlo o ponerlo en práctica (lectura 9). El cuaderno de Leonardo da Vinci está lleno de diseños para nueva maquinaria que no pudo realizarse con las primitivas y rudimentarias técnicas que estaban a su disposición para trabajar el metal (Rosenberg, 1976). Puede ocurrir que la innovación requiera de materiales con un alto grado de pureza o calidad, imposibles de conseguir con las técnicas de fabricación al uso. Por ejemplo, la máquina de vapor *compound* tuvo que esperar la existencia de acero barato de alta calidad (Rosenberg, 1982). A veces, se conocen los nuevos materiales y las técnicas para su fabricación en condiciones de laboratorio, pero se carece de un procedimiento específico que permita producirlos en cantidades suficientes para hacer rentable su comercialización. Así, el método a base de oxígeno para la fabricación del acero es tan antiguo como la patente británica de Bessemer, en 1856. Sin embargo, no fue factible hasta que se logró producir oxígeno puro a gran escala, una posibilidad no realizada hasta tres cuartos de siglo más tarde (Rosenberg, 1976). Los problemas que pueden abordarse con gran precisión en el laboratorio son inmensamente más difíciles de manejar en operaciones a gran escala, especialmente si se requiere un alto grado de precisión. No debe olvidarse que una fábrica de un proceso químico no es, ni mucho menos, una versión a escala del equipo original del laboratorio. Por esta razón, su puesta en marcha para una producción eficiente requiere la solución de múltiples problemas de tipo organizativo y técnico.

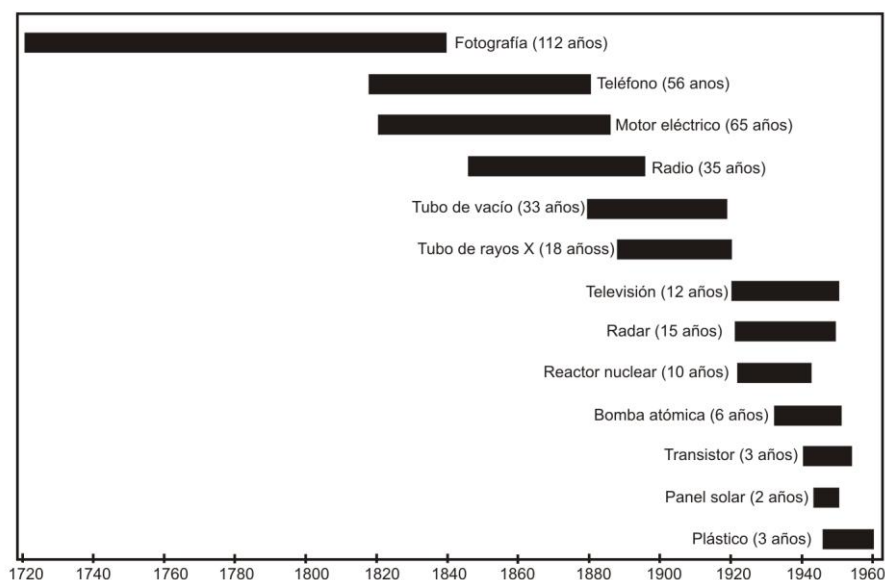


Figura 9: El sendero de la innovación tecnológica (Center for Integrative Studies (1969), citado en Dussauge et al. (1992))

Una gran parte de las nuevas invenciones no tienen importancia económica hasta que las industrias de bienes de capital han resuelto con éxito los problemas técnicos o mecánicos o desarrollado las nuevas máquinas que las invenciones requieren (Rosenberg, 1976).

LECTURA 9: LA HISTORIA DEL SOFÁ-BOCA: DEL SALIVASOFÁ AL DALILIPS

El precedente de todos los muebles dalinianos editados por Bd es su famosa boca convertida en sofá, que aparece por primera vez en el famoso cuadro, pintado en 1936, que transformaba la cara de Mae West en una habitación. Jean-Michel Frank construyó la primera versión del sofá-boca, que fue expuesto en la Gran Exposición Surrealista de 1938. Poco después, Dalí encargó a un tapicero de Figueres una copia para su casa de Port Ligat, y a Green y Abott, en Londres, la fabricación de alguna otra. No hubo más. Su amigo Jean-Michel se suicidó al poco tiempo de trasladarse a Estados Unidos, y los sistemas de fabricación empleados por Green y Abott no permitían conseguir el realismo y las formas sensuales que Dalí soñaba para su obra.

En 1972, Oscar Tusquets colaboró con Dalí en la construcción tridimensional del retrato de Mae West para el Museo de Figueres, con el reto de que el sofá tuviera todo el realismo y la fuerza expresiva de la boca de la célebre estrella de Hollywood. Con la forma, el color y el brillo de los labios humedecidos por la saliva. Incluso imitando las pequeñas grietas de la piel. Lo presentaron un año después en el hotel Meurice de París, bautizado con el nombre de Salivasofá. En aquella época, Tusquets acababa de fundar la editora Bd, desde la que trataron de poner en el mercado 300 ejemplares del nuevo sofá-boca-real, fabricado con espuma de poliuretano moldeado en caliente y acabado con un revestimiento de Polidur en color *rouge* de labios Shoking Pink, de Elsa Schiaparelli, según indicaciones de Dalí. Pero tampoco este modelo, que se mantuvo en la Sala Mae West durante años, llegó a ponerse a la venta. Los materiales de fabricación no garantizaban la calidad y durabilidad de la pieza.

Más de 30 años después, Tusquets y Bd han podido cumplir aquel viejo sueño de Dalí de ser diseñador para convertir industrialmente y con todo realismo una boca en un sofá. Lo han conseguido gracias a las nuevas tecnologías de fabricación con polietileno rotomoldeado y empleando un proceso especial que le da a la pieza una sensación ligeramente mórbida. El flamante diseño, que no es más que una versión tecnológicamente innovadora de aquel Salivasofá, se ha bautizado con el nombre de Dalilips y se ha presentado públicamente en la feria del mueble que se celebró en Milán durante el mes de abril.

Fuente: Úbeda, R. (2004): "Diseñador Dalí", *El País Semanal*, 12 de septiembre, pp. 81-83.

Históricamente, a medida que aumentaban las exigencias de capital fijo en la fabricación, se intentó utilizar la concesión de créditos para crear otra barrera que impidiera la introducción de nuevas ideas y técnicas en el mercado. Por ejemplo, Henry Ford tardó varios años en lograr que un banco financiara el nuevo sistema de producción diseñado para fabricar coches (Mokyr, 1990), ya que los competidores de Ford presionaban a los bancos para que no financiaran su invento de la cadena de montaje.

La falta de un sistema eficaz para proteger los inventos puede, asimismo, retrasar la introducción de innovaciones. Al parecer, Matthew Boulton, el socio de Watt (inventor de la máquina de vapor), invirtió en su proyecto sólo después de saber que el invento estaría protegido por una patente (Mokyr, 1990). También se puede dar el caso contrario. La patente puede restringir el uso de la tecnología. A su vez, las grandes empresas no explotan determinadas patentes, ya que satisfacen una demanda residual del mercado, por lo que no les resulta rentable su comercialización. No obstante, impiden a las empresas de menor tamaño su uso para que, así, no se divulgue la tecnología en el mercado.

El retraso puede acontecer porque la innovación destruye las prácticas en uso, con las graves distorsiones que esto provoca. Deming (1986) recoge el ejemplo de una máquina para mejorar el procedimiento de votación en el Congreso de los Estados Unidos inventada por Thomas A. Edison. Con su dispositivo, cada diputado y cada senador dispondrían de tres botones en el brazo de su butaca: un botón rojo para el no, un botón verde para el sí y un botón blanco para la abstención. Cuando había que votar, cada miembro pulsaría el botón apropiado e instantáneamente aparecería un registro de cómo había votado cada uno

y el total de los votos. Este dispositivo permitía eliminar los errores de la votación por lista y también reducía el tiempo empleado en la misma. Sin embargo, ambos presidentes, el del Congreso y el del Senado, informaron que no deseaban este sistema, al trastocar completamente el funcionamiento normal de las instituciones. La razón era que los retrasos en la votación por lista formaban parte integrante del proceso de deliberaciones con el cual funcionaban.

Otra razón de los retrasos es que librarse de la vieja tecnología puede ser muy caro. Por ejemplo, el cambio de la generación de energía nuclear a la generación producida por gas puede requerir la costosa eliminación de la planta de energía nuclear (Afuah, 1998). En consecuencia, los altos costes de salida pueden impedir el abandono de la tecnología actual para introducir una nueva.

El miedo a la canibalización de los productos, es decir, al proceso por el cual un producto de una empresa quita ventas a otro, también provoca el retraso en la comercialización del invento. La excesiva preocupación por la canibalización puede resultar costosa si implica que se retrasen los nuevos productos o las mejoras en las presentaciones de los productos actuales. Un clásico ejemplo de los peligros que suponen tales retrasos es el de Gillette. Esta empresa fue aplazando la presentación de la nueva hoja de afeitar de acero inoxidable de alta calidad por temor a que dicha hoja, que permitía más afeitados que la hoja «super azul», que proporcionaba cuantiosos beneficios, la canibalizase. Sólo cuando los competidores audaces presentaron sus propias hojas de acero inoxidable se animó Gillette a hacer lo mismo. Pero debido a la falta de decisión a la hora de presentar la innovación y la mejora del producto, la participación de Gillette en el mercado de las maquinillas de doble hoja cayó de un 90 por 100 a un 70 por 100. Nunca llegó a recuperar en su totalidad esta pérdida de la posición competitiva (Hartley, 1990).

Los retrasos pueden estar condicionados por el cumplimiento de la normativa en vigor. En este sentido, muchos nuevos productos y procesos deben cumplir la normativa medioambiental antes de su introducción en el mercado. Por otra parte, los fármacos deben pasar una serie de controles impuestos por los gobiernos antes de su comercialización. Estos controles, en el mejor de los casos, retrasan la innovación y, bastante a menudo, impiden la comercialización de los inventos, al considerar que no son suficientemente benignos.

La innovación puede posponerse de forma deliberada. Ahora bien, retrasar una innovación tiene tanto ventajas como inconvenientes. El momento más oportuno es aquél en que los beneficios y los costes adicionales del retraso resulten equiparables. El coste del aplazamiento equivale al beneficio perdido por no proceder de inmediato al cambio tecnológico. Por otro lado, el beneficio que implica dicho aplazamiento integra dos componentes: a) los intereses generados por el capital durante el tiempo en que se retrasa la inversión y b) los beneficios diferenciales de las futuras mejoras en el rendimiento de la nueva tecnología que pueden relacionarse con el futuro crecimiento del mercado, con cambios en el precio de los factores de producción o con mejoras de la tecnología (Metcalf, 1992).

Los retrasos en las innovaciones se deben igualmente a la lucha por el poder entre grupos con distintos intereses en la empresa innovadora. Las empresas, de hecho, están compuestas de coaliciones políticas, cuyo objeto es proteger y aumentar sus propios intereses (Pfeffer, 1992). Cada coalición influye en las decisiones en función de su

poder, entendido éste como la capacidad para hacer que las preferencias o inclinaciones propias se reflejen en cualesquiera acciones que la empresa lleve a cabo. La coalición dominante logrará imponer sus intereses en las decisiones de la empresa. Por tanto, las innovaciones tienen más probabilidad de desarrollarse si aumentan el poder de la coalición dominante. De la misma forma, la coalición dominante intentará retrasar (o eliminar) las innovaciones que disminuyen su poder (Afuah, 1998). Así pues, los éxitos pasados de la coalición dominante le ayudan a ocupar posiciones influyentes y a controlar recursos clave que emplean para debilitar a quienes defienden el uso de métodos mejores que, si bien beneficiarán a la empresa, amenazarán la hegemonía de la que gozan (Sutton, 2002). Por ejemplo, después de que se inventara el cañón naval de tiro rápido y se demostrara, en numerosos buques británicos y americanos, su superioridad sobre las prácticas vigentes, mediante una evaluación cuidadosamente documentada e informes enviados a Washington para instar a su adopción, esta invención se encontró con la indiferencia, luego con la polémica, después con la hostilidad. Su gran defensor, un teniente de navío de la Armada de EEUU, tuvo que apelar directamente al presidente Theodore Roosevelt para ser oído. Los oficiales superiores de la Armada se oponían a la innovación porque preveían con claridad que a partir de ella se produciría una reorganización turbulenta de la sociedad militar. Entre otros cambios, el oficial de artillería ganó puestos en el escalafón abriéndole caminos hacia los altos rangos (Leonard–Barton, 1995).

Los altos directivos pudieron alcanzar su posición debido a las valiosas contribuciones que hicieron en la invención y comercialización de los productos actuales. Por tanto, tienen apegos emocionales hacia ellos, lo que les ocasionan un coste psicológico muy fuerte si lo abandonaran. De esta forma, tratarán de retrasar la innovación que reemplazará el producto que los hizo ser lo que son (Afuah, 1998). Por ejemplo, Henry Ford se negaba a introducir una nueva línea de automóviles que sustituyera a su Modelo T. De hecho, tardó varios años en aceptar el consejo de sus directivos de introducir el nuevo modelo, incluso a pesar que General Motors llevaba varios años comercializando cinco modelos diferentes.

Los retrasos en las innovaciones también se deben a que la gente se vuelve experta en hacer las cosas a la vieja usanza. Su profunda competencia en la antigua tecnología y su falta de habilidad con la nueva pueden hacer que desarrollen peor su trabajo cuando prueban conceptos, métodos y tecnologías nuevos (además de superiores). Por tanto, rechazan todo aquello que represente un cambio en la forma establecida de hacer las cosas. March (1991) llamó a esto la «trampa del éxito» o la «trampa de la competencia».

Las sociedades conservadoras muestran mayor tendencia a rechazar las innovaciones, como ocurrió en la sociedad china en la época de los mandarines que eran unos celosos guardianes de la tradición. Muchas veces ese rechazo está inducido por grupos de presión bien organizados que se benefician de mantener el *statu quo* frente a una alternativa mejor (Mokyr, 1990). Por ello, intentarán boicotear por todos los medios posibles la innovación. Gremios, sindicatos y asociaciones profesionales han contribuido a reforzar el conformismo dentro de la industria e impedir la entrada del *homines novi* con nuevas ideas. En el siglo XV el gremio de escribientes de París consiguió retrasar la introducción de la imprenta en esta ciudad veinte años⁶. Igualmente han provocado el retraso de la

⁶ Tradicionalmente, hubo fuertes oposiciones a las nuevas tecnologías en el sector textil. Entre 1811 y 1816, los *Midlands* y los condados industriales ingleses fueron el centro de los ataques de los *luddite*, con graves

innovación determinadas tendencias hacia el misticismo y ciertas actitudes dogmáticas y reaccionarias: un ejemplo es el gran retraso social que supuso para Camboya el gobierno de los *jémeres rojos*, que llegaron incluso a prohibir el uso de las medicinas occidentales.

La principal razón por la que actitudes y valores sociales retrasan la innovación se basa en que el progreso tecnológico aumenta los bienes materiales, mientras que los bienes de posición permanecen inalterados (Hirsch, 1976). Por ejemplo, en una sociedad esclavista donde la riqueza y el prestigio social se valoran por el número de esclavos en propiedad, la adopción de maquinaria para ahorrar mano de obra (aun siendo más eficiente) sería rechazada, al convertir los esclavos en innecesarios. Sin embargo, si en esa misma sociedad la riqueza se estimara en función del total de recursos, representando los esclavos sólo una forma de riqueza, el cambio tecnológico tendría las mismas posibilidades de ocurrir que en cualquier otra sociedad (Mokyr, 1990).

El éxito espectacular de Occidente en el cultivo de la ciencia y la tecnología está arraigado en la creencia judeo cristiana de que el dominio de la naturaleza estaba sancionado por la religión. El persistente esfuerzo agresivo de los occidentales por explotar toda posible fuerza y recurso natural se materializó en su liderazgo mundial de la tecnología. Aquellos pueblos cuya religión les enseña a adoptar una actitud más benigna hacia la naturaleza no llegaron a desarrollar la tecnología hasta su máximo potencial (Basalla, 1988).

6. TIPOS DE INNOVACIONES

La innovación es un concepto complejo y, por tanto, susceptible de múltiples interpretaciones. Una forma de allanar su comprensión es clasificándola, si bien son posibles una gran variedad de clasificaciones.

6.1. Innovación radical *versus* incremental

Atendiendo a su originalidad una innovación puede ser radical o incremental (Dewar y Dutton, 1986). La innovación radical hace obsoleta la tecnología en uso y es destructora de competencias (Tushman y Anderson, 1986). En el otro extremo de la dicotomía se halla la innovación incremental, que mejora las prestaciones de la tecnología en uso y contribuye a incrementar las competencias. La mayoría de las innovaciones son incrementales.

Schumpeter (1883–1950) se ha centrado exclusivamente en el análisis de los cambios propiciados por las innovaciones radicales, o sea, en las «nuevas combinaciones» de los factores de producción que dan origen a las incesantes tormentas de «destrucción creativa». En concreto, define la innovación radical como una nueva función de producción, una alteración del conjunto de posibilidades que define lo que es posible producir y cómo, y que no puede reducirse a una serie infinitesimal de modificaciones. La

daños para las máquinas (*luddite* fue el nombre que se dio a los grupos de trabajadores que en esos años provocaron desórdenes y destruyeron máquinas, cuyo nombre se deriva probablemente de Ned Lud(d), un trabajador de fines del siglo XVIII que destruyó máquinas de tricotaje en Inglaterra). En 1769 el Parlamento de Gran Bretaña aprobó una ley que consideraba la destrucción deliberada de maquinaria un crimen que podría ser castigado con la pena de muerte. Contra los motines de los *luddite* de 1811 a 1813, el gobierno británico empleó 12.000 hombres, una fuerza superior en tamaño al primitivo ejército peninsular de Wellington en 1808. Los motines se sofocaban por la fuerza y solían acabar en condenas a la horca y deportaciones (Mokyr, 1990)

innovación radical es algo nuevo para el mundo, constituye una ruptura profunda con las formas establecidas de hacer las cosas, suele abrir nuevos mercados y aplicaciones potenciales, cambia las bases de la competencia, crea grandes dificultades a las empresas establecidas y puede suponer la base para la entrada con éxito de nuevas empresas e incluso la redefinición de la industria. Schumpeter (1934) considera que una innovación radical tiene lugar cuando se produce alguna de las cinco situaciones siguientes:

1. La introducción de un nuevo bien –uno que no es todavía familiar a los consumidores– o de una nueva calidad de un bien; por ejemplo, Manuel Jalón Corominas inventó y comercializó la fregona en el mercado español, una innovación que, más tarde, se extendería por los cinco continentes.
2. La introducción de un nuevo método de producción, de uno no probado por la experiencia en la industria de que se trate, que no precisa apoyarse en un descubrimiento nuevo desde el punto de vista científico, y puede consistir simplemente en una forma nueva de manejar comercialmente una mercancía. Tal es el caso de la cadena de montaje desarrollada por Ford o el sistema *just in time* creado por Toyota.
3. La apertura de un nuevo mercado, esto es, aquel en el cual no haya entrado la industria del país de que se trate, a pesar de que existiera anteriormente dicho mercado; por ejemplo, el uso del gas para calentar espacios, previamente se había utilizado para alumbrado residencial y comercial urbano.
4. La conquista de una nueva fuente de aprovisionamiento de materias primas o de bienes semifabricados, con independencia de si ya existía o de si hay que crearla; por ejemplo, la creación de la industria norteamericana del caucho sintético para sustituir al natural.
5. La creación de una nueva organización de cualquier industria, como la de una posición de monopolio o la anulación de una posición de monopolio existente con anterioridad; por ejemplo, la creación de una franquicia o la estructura multidivisional.

Estas nuevas combinaciones provocan la «destrucción creativa», al destruir el antiguo orden y crear al mismo tiempo nuevos negocios que alimentan el cambio y provocan mayor crecimiento económico (Schumpeter, 1942). Estos nuevos negocios detentan una superioridad decisiva en el coste o en la calidad del producto y atacan, no ya los márgenes de los beneficios y de la producción de las empresas existentes, sino a sus cimientos y a su misma existencia. En suma, Schumpeter reemplazó la noción de la competencia dentro del mercado por la competencia por el mercado defendiendo los beneficios de la destrucción creativa que surge del proceso innovador.

Ahora bien, los proyectos dedicados a la innovación radical son arriesgados, caros y normalmente tardan varios años en producir resultados tangibles –si es que llegan a producirlos–. La investigación de Leifer *et al.* (2000) concluyó que hasta pasados 10 años, por lo menos, no se observaban resultados tangibles. Para tener éxito, las empresas deben ser pacientes y disponer del presupuesto necesario para soportarlo.

La innovación incremental conlleva cambios en la tecnología establecida para adaptarla mejor a las necesidades de los mercados actuales. Se apoya en las destrezas y los recursos que la empresa posee y manifiesta consecuencias significativas sobre las características de los productos, el coste y el precio, por lo que pueden servir para fortalecer y aumentar no sólo la competencia en producción, sino también los vínculos con los clientes y los mercados. La innovación incremental suele ser casi invisible, pero puede tener un efecto

acumulativo muy importante sobre la productividad y la funcionalidad de los productos. De hecho, la mayor parte de la reducción de costes de producción, incluso en las grandes empresas, se debe a un conjunto de pequeños, y a veces olvidados, cambios en la tecnología, introducidos por los ingenieros y otros especialistas técnicos, y no al fruto de grandes innovaciones surgidas de sus departamentos de I+D (Schmookler, 1966). Así pues, las innovaciones incrementales son mejoras que se realizan sobre la tecnología existente, es decir, introducen cambios menores en los productos y procesos actuales, explotan el potencial del diseño establecido y refuerzan el dominio de las empresas que los llevan a cabo (figura 10). Los problemas asociados con el elevado riesgo, los cuantiosos gastos y los largos plazos de desarrollo de la innovación radical animan a muchas empresas establecidas a decantarse por la innovación incremental. Se trata de un proceso más seguro, más barato y con más probabilidad de producir resultados positivos dentro de un período de tiempo razonable. No obstante, al llevar a cabo innovaciones incrementales conviene no padecer el síndrome de la «ornamentación», que consiste en incorporar prestaciones adicionales al producto, a pesar de que los clientes no las necesitan.

<i>Innovación incremental</i>	<i>Innovación radical</i>
La demanda del mercado es conocida y predecible.	Una demanda potencial elevada pero poco predecible. El riesgo de fracasar también es alto.
El reconocimiento y la aceptación del mercado son rápidos.	Puede que la aceptación por parte del mercado sea lenta en un principio, pero se espera una reacción imitativa rápida de los competidores.
Fácilmente adaptable a las ventas existentes en el mercado y a la política de distribución.	Puede exigir unas políticas de distribución, de marketing y ventas exclusivas para educar a los consumidores o a causas de problemas especiales de garantía y reparaciones.
Encaja en la actual segmentación del mercado y en las políticas de producto.	La demanda puede no coincidir con los segmentos del mercado establecidos, distorsionando el control de la empresa, y absorber el mercado de otros productos.

Figura 10: Características de las innovaciones radicales e incrementales (Hayes y Abernathy, 1980)

De acuerdo con Schumpeter, la capacidad o la iniciativa de los empresarios creaban nuevas oportunidades de beneficio, lo que atraía a su vez a un «enjambre» de imitadores, explotando las nuevas vías con una ola de nuevas inversiones que generaban unas condiciones de expansión. El proceso competitivo puesto en marcha por esta «ebullición» reducía entonces los beneficios de la innovación, pero antes de que el sistema pudiera establecerse en una situación de equilibrio, comenzaba de nuevo todo el proceso a través de los efectos desestabilizadores de una nueva ola de innovaciones. Ahora bien, de acuerdo con Rosenberg (1982) el proceso de ebullición no puede considerarse una simple réplica o copia exacta, sino que, a menudo, trae consigo una cadena de innovaciones adicionales de tipo incremental, al verse envuelto un creciente número de empresas que comienzan a aprender la nueva tecnología y se esfuerzan en superar a los competidores. Estas innovaciones incrementales las puede llevar a cabo tanto la empresa que creó el nuevo negocio como cualquier imitador.

Así pues, la innovación no es un artefacto único, sino más bien una secuencia de artefactos o mejoras subsiguientes a la propia innovación, a medida que el diseño original mejora y se extiende a nuevas aplicaciones durante la ebullición. En este sentido, las invenciones subsiguientes a una invención tras su primera introducción podrían ser mucho más importantes económicamente que la disponibilidad de la innovación en su forma original. Por ejemplo, el primer modelo de máquina de escribir comercial en 1874 sólo tenía letras mayúsculas y las teclas golpeaban el papel dentro de la máquina, haciendo imposible que el mecanógrafo pudiera ver su trabajo hasta que las cuatro primeras líneas estaban impresas, momento en que el papel empezaba a salir de la máquina (Utterback, 1994). Sólo las posteriores mejoras contribuyeron a mejorar su eficiencia y funcionalidad. Como apuntan Mansfield *et al.* (1977), las mejoras en el proceso o en el producto pueden ser casi tan importantes como la nueva idea en sí. De igual forma, en su estudio sobre las fuentes del incremento de la eficiencia de las fábricas de rayón de DuPont, Hollander (1965) llega a la conclusión de que los efectos acumulativos de pequeños cambios tecnológicos sobre la reducción de costes fueron más importantes que los efectos de cambios tecnológicos superiores. Enos (1962), que estudió cuatro procesos en la industria del refino del petróleo en Estados Unidos, observó que el ahorro promedio en el coste anual de la adopción inicial de procesos completamente nuevos fue inferior al ahorro que resultó de las siguientes mejoras. Otras investigaciones empíricas confirman que las ventas y los beneficios mejoran, sobre todo con la introducción de innovaciones incrementales (Minasian, 1962; Branch, 1974). En consecuencia, las innovaciones incrementales contribuyen de manera muy significativa al éxito comercial (Marquis, 1982). No obstante, su importancia de ningún modo minimiza el efecto de las innovaciones radicales. En este sentido, resulta muy útil el planteamiento de Hollander (1965), al considerar que hay una interdependencia entre los cambios principales y los menores y que, sin la existencia de algún cambio principal precedente, la corriente potencial de los cambios menores se agotaría. En suma, la empresa, después de introducir una innovación radical en el mercado, debe continuar mejorándola, ya que, si no lo hace ella, lo harán los competidores para mejorar su posición en el mercado.

Mejora continua (*kaizen*). La mejora continua (o *kaizen*) es una innovación incremental menor, cuya responsabilidad compete a todos los trabajadores de la empresa en todos los niveles jerárquicos, incluida, por supuesto, la base de la pirámide organizativa, es decir, el taller y las oficinas. La mejora continua significa reducir el despilfarro y aumentar la calidad en todas las actividades del proceso (Deming, 1986). El objetivo final es la perfección absoluta, que nunca se podrá alcanzar, pero que siempre debe perseguirse.

En general, las empresas apenas valoran la cualificación ni las ideas de los trabajadores (Hayes, 1985). No puede decirse que la innovación suponga un avance en una progresión de escalera, porque, por lo general, no lo hace, ya que un sistema, una vez que ha sido instalado como resultado de una innovación radical, está sujeto a un deterioro uniforme, a menos que se hagan esfuerzos continuos primero para mantenerlo y luego, para mejorarlo (Imai, 1986). Por eso se hace necesario establecer el hábito de la mejora continua como una parte inseparable del trabajo diario de cada uno. La mejora continua puede llevarse a cabo en cualquier actividad empresarial. Estos cambios tienen su origen en los trabajadores de la empresa con independencia de la función que realizan.

Una empresa que desarrolla un método de mejora continua da pocos pasos que sean apreciables o arriesgados, pone en marcha un proceso gradual que apenas precisa recursos, pero sí requiere una gran cantidad de esfuerzo sostenido y dedicación. En lugar de aplicar

recursos masivos a la puesta en práctica de cambios bruscos, esta empresa espera que la mayoría de los cambios aflore a la superficie de una forma intraemprendedora, desde los niveles inferiores de la organización (Hayes, 1985). En pocas palabras, la mejora continua está orientada a las personas, en tanto que la innovación radical está orientada a la tecnología y al dinero (Imai, 1986). La mejora continua es, por definición, lenta, gradual y, a menudo, invisible, con efectos que se sienten a largo plazo.

6.2. Innovación original versus adaptada (imitación)

Muchas innovaciones son nuevas para la empresa, pero no para el mercado. Para Baumol (1993) son también innovaciones las actividades de transferencia de tecnología que aprovechan las oportunidades de introducir una tecnología ya disponible y válida en áreas geográficas, cuya aptitud en el mercado nacional no había sido anteriormente reconocida ni utilizada. Este proceso de transferencia provoca siempre algún tipo de innovación en la empresa receptora, ya que normalmente necesita hacer alguna modificación en la tecnología para adaptarla a las condiciones de su mercado (lectura 10). La adaptación de las innovaciones de otras empresas es también una actividad innovadora que requiere capacidad de absorción, es decir, tener la habilidad para reconocer (identificar) el valor de una tecnología nueva y externa, asimilarla y aplicarla a fines comerciales (Cohen y Levinthal, 1990). Así pues, como es necesario conocimiento previo para absorber conocimiento nuevo, la eficacia de la adaptación está en relación directa con el grado en el que la empresa tiene un conocimiento afín que le permita absorber el conocimiento que sustenta la innovación. Este conocimiento afín incluye habilidades básicas (*know-how*), un lenguaje compartido y conocimiento técnico (Dixon, 2000). En general, Utterback (1974) estima que un número importante de innovaciones (23 al 33 por 100) se han adaptado de otras empresas. Ellas, al igual que las originadas en la empresa, contribuyen significativamente al éxito comercial (Marquis, 1982).

LECTURA 10: ADAPTACIÓN DE LAS GALLETAS OREO AL MERCADO CHINO

Hace mucho tiempo que las galletas Oreo son las de mayor venta en el mercado estadounidense, pero Kraft Foods Inc. tuvo que reinventarlas para que se vendieran bien en el país más poblado del mundo. Las ventas de Oreo en China representan una mínima fracción de los ingresos anuales de la compañía, 37.200 millones de dólares; sin embargo, el viaje de la galleta hasta China es un ejemplo de transformación emprendedora que la CEO Irene Rosenfeld está tratando de difundir a todo este coloso fabricante de alimentos; la CEO ha dado más autonomía a las diversas unidades de negocio de Kraft en todo el mundo e informando a los empleados que las personas de la oficina central, con domicilio en Northfield, Illinois, Estados Unidos, no deberían tomar todas las decisiones sobre los productos de la compañía.

Las galletas Oreo aparecieron en el mercado estadounidense en 1912, pero no fue sino hasta 1996 que Kraft las empezó a vender en China; nueve años después apareció en escena Shawn Warren, un veterano de 37 años en Kraft que llevaba varios de ellos comercializando las galletas dulces y saladas de la compañía por todo el mundo, y quien advirtió que las ventas de China se habían estancado durante cinco años; el problema: Kraft vendía la versión estadounidense de las Oreo en China; según Warren, la definición de demencia de Albert Einstein caracterizaba a Kraft: hacer lo mismo una y otra vez esperando resultados diferentes.

El señor Warren asignó a su equipo un extenso proyecto de investigación, el cual produjo algunos resultados muy interesantes; en primer lugar, Kraft supo que las Oreo tradicionales eran demasiado dulces para el gusto de los chinos y que el paquete de 14, con un precio de venta de 72 centavos, era demasiado caro.

La compañía creó 20 prototipos de Oreo con menos azúcar y realizó pruebas con consumidores chinos antes de llegar a la fórmula con el sabor justo; Kraft también introdujo paquetes con menos galletas a un precio de venta de 29 centavos.

Sin embargo, Kraft sabía que debía hacer mucho más que sólo modificar su receta para obtener una mayor participación del mercado chino; tras más experimentos recreó la Oreo en el 2006, que ahora consistía en cuatro capas de galleta crujiente, rellenas con crema de vainilla y chocolate, y todo recubierto con otra capa de chocolate; la empresa también instituyó un proceso patentado de manejo que garantiza que el producto pueda viajar por todo el país, soportando los climas gélidos del norte y caliente y húmedo del sur, y llegar listo para derretirse en la boca.

Los esfuerzos rindieron frutos: en 2008 y 2009 Kraft duplicó los ingresos generados por Oreo en China; además, con la ayuda de esas ventas, las mundiales excedieron por primera vez la cantidad de los mil millones de dólares.

Fuente: Dess, G. G.; Lumpkin, G. T. y Eisner, A. B. (2010): *Strategic Management. Text and Cases (8th Edition)*, McGraw-Hill, Nueva York.

Una innovación adaptada también puede ser la resultante de adaptar los productos o servicios de la sociedad occidental a las necesidades de los usuarios pobres y deficientemente instruidos de África, Asia y América Latina. Es una forma de explotar las posibilidades comerciales del «terreno de nadie» existentes entre la elaborada tecnología occidental y el rudimentario sistema de producción de los países en vías de desarrollo (Norman y Blair, 1982).

6.3. Innovación tecnológica versus mercado

Las innovaciones pueden clasificarse según el efecto que causan en el conocimiento tecnológico y el mercado objetivo del fabricante (Abernathy y Clark, 1985). Si cruzamos la variable mercado objetivo de la empresa (actual o nuevo) y la variable tecnología (actual o nueva), se obtienen cuatro tipos de innovaciones: creadora de nichos, arquitectónica, regular y revolucionaria (figura 11).

		Tecnología	
		Actual	Nueva
Mercado	Actual	<i>Regular</i>	<i>Revolucionaria</i>
	Nuevo	<i>Creadora de nichos</i>	<i>Arquitectónica</i>

Figura 11: Innovación tecnología–mercado (Abernathy y Clark, 1985)

La innovación regular conlleva cambios incrementales en la tecnología establecida para adaptarla mejor a las necesidades de los mercados actuales. Esta innovación consiste en asentar las destrezas y los recursos que se poseen y manifiesta consecuencias significativas sobre las características de los productos, el coste y el precio, por lo que puede servir para fortalecer y aumentar no sólo la competencia en producción, sino también los vínculos con los clientes y los mercados. Suele ser casi invisible, pero puede tener un efecto acumulativo muy importante sobre la productividad y la funcionalidad de los productos. De hecho, la mayor parte de la reducción de costes de producción, incluso en las grandes empresas, se debe a un conjunto de pequeños, y a veces olvidados, cambios en la tecnología, introducidos por los ingenieros y otros especialistas técnicos, y no al fruto de grandes innovaciones surgidas de sus departamentos de I+D (Schmookler, 1966).

La innovación revolucionaria consiste en una nueva tecnología que se aplica a los mercados actuales. La nueva tecnología, a menudo, invade la industria tradicional introduciéndose en una serie de submercados. Aunque es rudimentaria al principio, suele contar con ciertas ventajas de aplicación para determinados segmentos (Cooper y Schendel, 1976). Esta innovación destruye la tecnología establecida convirtiéndola en obsoleta, desde el punto de vista tecnológico, y se aplica a los mercados y clientes actuales. Esta destrucción debe medirse en función de lo que altera los parámetros de la competencia, así como por los cambios que provoca en la tecnología en uso. La innovación revolucionaria conmociona los datos técnicos, modificando la estructura o los elementos fundamentales del producto y, muy a menudo, la forma de producirlo, sin cambiar la funcionalidad central y, por ende, la clientela o el mercado.

Las reglas de cálculo especialmente diseñadas para los negocios tenían escalas para hacer cálculos de préstamos o determinar cantidades de compra óptimas. Durante los años cincuenta y sesenta, Keuffel & Esser era el fabricante de reglas de cálculo más importante en Estados Unidos, fabricando 5.000 reglas de cálculo al mes. Sin embargo, a principios de los setenta, una innovación relegó la regla de cálculo a coleccionistas y museos en sólo unos pocos años: la barata calculadora de mano. Keuffel & Esser no tenían experiencia en los componentes electrónicos que hacían posibles las calculadoras electrónicas y fue incapaz de hacer la transición hacia la nueva tecnología. En 1976, Keuffel & Esser se retiró del mercado. Aunque la calculadora de mano se agregó a las competencias existentes en compañías como Hewlett-Packard y Texas Instruments (y por tanto para ellas sería potenciadoras de competencias), para Keuffel & Esser la calculadora fue una innovación revolucionaria destructora de competencias (Schilling, 2008).

La innovación creadora de nichos utiliza las tecnologías establecidas para abrir nuevas oportunidades del mercado. Pueden basarse en la propiedad transversal de la tecnología, es decir, los cambios se construyen sobre la competencia técnica establecida y mejoran su aplicación en nuevos segmentos de mercado. Además de mejorar la tecnología continuamente, es necesario adaptarla a la demanda de los nuevos clientes (lectura 11). La gran suerte para la industria del gas, cuando empezó a perder el mercado del alumbrado residencial y comercial urbano frente a la energía eléctrica, fue que encontró y desarrolló nuevos mercados para su producto –el calentamiento de espacios y el calor para los procesos de fabricación– en el mismo instante en que el mercado de gas para iluminación se extinguía debido al alumbrado eléctrico (Utterback, 1994).

LECTURA 11: UN NUEVO NICHOS DE MERCADO

A finales de la década de 1870, un ginecólogo francés llamado Stéphane Tarnier paseando por el zoo de París observó a unos pollitos recién nacidos en una incubadora para pollos, y tuvo una asociación de ideas. Poco después, contrató a Odile Martin, la criadora de pollos del zoo, para que le construyera un aparato que pudiera hacerles un servicio similar a los recién nacidos humanos. En aquella época, en una ciudad tan avanzada como París, uno de cada cinco bebés moría sin haber alcanzado la edad de gatear. Tarnier sabía que la regulación de la temperatura era un aspecto crucial para que los niños no murieran, y sabía que el *establishment* médico francés tenía obsesión por la estadística. Así que, en cuanto tuvo en la Maternité su incubadora para recién nacidos, donde se calentaban a los bebés con botellas de agua tibia instalada bajo la estructura, Tarnier llevó a cabo un estudio rápido sobre unos quinientos bebés. Mientras que el 66% de los nacidos con bajo peso moría a las pocas semanas de nacer, sólo corría esa misma suerte el 38% de los que pasaban un tiempo en la caja de incubar de Tarnier. En conclusión, se podía reducir a la mitad la tasa de mortalidad entre los bebés prematuros si se les trataba como a los pollitos del zoo. Pocos años después, las autoridades parisienses decretaron la obligación de instalar incubadoras en todos los hospitales de maternidad.

Fuente: Johnson, S. (2011): *Las Buenas Ideas. Una Historia Natural de la Innovación*, Turner Publicaciones, Madrid.

La innovación arquitectónica conlleva nuevas tecnologías para nuevos mercados. La nueva tecnología a menudo crea un nuevo mercado no disponible para la antigua tecnología (Cooper y Schendel, 1976). Se trata de un concepto tecnológico original que rompe con los sistemas tecnológicos establecidos, y, a continuación, crea un nuevo mercado. Es la alternativa que conlleva más riesgo, ya que, al ser todo nuevo para la empresa, le resulta difícil de asimilar y gestionar. En contrapartida, es la que puede proporcionar los mayores beneficios a largo plazo; por ejemplo, cuando se introdujo la máquina de escribir manual no existía ninguna tecnología previa. Esta tecnología no sólo crea un nuevo sector, sino también nuevas profesiones, como la de secretaria, que, en su inicio, fue desempeñada fundamentalmente por hombres.

6.4. Innovación modular

De acuerdo con Utterback (1994), los productos pueden ser no ensamblados o ensamblados. Los productos no ensamblados están compuestos de uno o unos pocos materiales, como los cristales de las ventanas y los productos químicos. Estos productos generalmente no admiten mejoras incrementales, sino que, una vez introducidos en el mercado, se comercializan con la misma forma y prestaciones durante su ciclo de vida que, en algunos casos, puede durar siglos.

El producto ensamblado puede considerarse como un todo (sistema) formado por un conjunto de partes (componentes) interrelacionadas (lectura 12). Esta concepción del producto ensamblado como un sistema indica que el desarrollo de un producto con éxito requiere dos tipos de conocimientos (Henderson y Clark, 1990): a) de los componentes y b) de la manera en que los componentes se relacionan para configurar un producto único.

LECTURA 12: ENFOQUE MODULAR VERSUS ENFOQUE PLATAFORMA

Dos políticas para el diseño de productos complejos son el enfoque modular y el enfoque plataforma, si bien las dos forman parte del mismo *continuum*, que es la innovación modular.

Enfoque plataforma. Una plataforma de producto es una serie de subsistemas e interfases que forman una estructura o tecnología central común de la que surgen varios productos derivados que se pueden producir y desarrollar eficientemente (Meyer y Lehnerd, 1997). A la configuración central se le añaden accesorios para que correspondan a las necesidades de diferentes segmentos del mercado. El grupo francés PSA Peugeot Citröen decidió fabricar todos sus modelos de coches en tres plataformas (la estructura de los coches). Esto significa que los vehículos de la misma categoría comparten el bastidor, el motor y todos los elementos que no ven los consumidores; en total, el sesenta por ciento de sus componentes. El cuarenta por ciento restante, todos los componentes que se ven desde el exterior, son distintos y permiten que cada marca mantenga su personalidad. La política de plataforma compartida ha permitido al grupo industrial ser más eficiente, reducir más costes y, al mismo tiempo, aumentar la oferta de productos.

El popular reloj *Swatch* es otro ejemplo de una familia de productos con éxito basada en una plataforma común. La plataforma *Swatch* es una pequeña serie de subsistemas para medir el tiempo conectada a través de unas cuantas interfaces electrónicas. Esta plataforma es, en efecto, la innovación. Casi todos los *Swatch* utilizan la misma plataforma, que es simple, barata de fabricar y capaz de soportar una cantidad ilimitada de variaciones externas.

Las plataformas de producto basadas en la elegancia del diseño y de la producción ofrecen a las empresas oportunidades de bajo coste para personalizar sus productos y dirigirlos a segmentos de mercado diferentes. La plataforma de elementos comunes se puede combinar con ciertos elementos únicos para

producir un producto para un segmento de mercado particular. *Swatch* lo puso en práctica introduciendo su nueva maquinaria innovadora dentro de una amplia variedad de carcasas de diseño exclusivo, produciendo muchos relojes diferentes para diferentes segmentos a partir de una base común. La verdadera innovación estaba en la maquinaria. Esa innovación se aprovechó en el mercado por medio de los derivados de la plataforma.

Enfoque modular. La modularidad se refiere al grado en el que los componentes de un sistema pueden ser separados y recombinados. La esencia de la modularidad consiste en diseñar, desarrollar y producir componentes que puedan utilizarse mundialmente para combinarse en un número máximo de configuraciones de producto. En este caso no existe una plataforma básica más bien rígida, sino componentes susceptibles de modificaciones y sustituciones. Un enfoque modular tiene dos propiedades: a) los componentes constituyen uno o algunos elementos funcionales en su totalidad y b) las interacciones entre los componentes están bien definidas y, por lo general, son fundamentales para las funciones primarias del producto. Por ejemplo, al diseñar todos los componentes de estanterías para que funcionen mediante conectores estándar, IKEA asegura que los componentes pueden ser combinados y ajustados de manera libre. Los componentes individuales pueden combinarse sin necesitar ningún cambio en el diseño de los otros componentes. Además, como la modularidad puede permitir que un componente sea mejorado sin cambiar los demás componentes, puede permitir a la empresa y clientes mejorar sus productos sin reemplazar su sistema completo. El ordenador es un excelente ejemplo (Schilling, 2008). El enfoque modular proporciona diferentes ventajas. La empresa puede minimizar el número de componentes estándares, ensamblarlos en las primeras etapas del proceso y añadir los componentes que diferencian el producto en la etapa final. Los componentes pueden fabricarse en paralelo y de forma independiente, lo que conlleva una reducción del ciclo de producción. La empresa puede diagnosticar con mayor facilidad los problemas de producción y aislar los potenciales problemas de calidad. El número de componentes estándar es menor que el de componentes en un sistema integral. Las reparaciones resultan más sencillas en el caso de sustitución. La canibalización de componentes se simplifica. La adición de productos es sencilla, por lo que resulta fácil llegar a diferentes segmentos del mercado. Se consiguen economías de escala en la fabricación de componentes. Entre las desventajas se encuentran: a) cada componente requiere más piezas de las necesarias y b) la interconexión de componentes puede resultar difícil.

Un componente es una parte del producto que se corresponde con un diseño tecnológico específico (o tecnología) y que realiza una función bien definida. Las relaciones entre los componentes son el elemento fundamental en la estructura del producto, ya que permiten integrar a los diferentes componentes en un todo homogéneo. La esencia de una innovación de componente es que modifica el conocimiento de componente dando lugar a uno nuevo o a una mejora en las características funcionales del componente actual. Por su parte, la esencia de la innovación en las relaciones entre los componentes es la reconfiguración del sistema establecido, dando lugar a una nueva combinación de componentes (Henderson y Clark, 1990). La figura 12 recoge los cuatro tipos de innovaciones que se obtienen al cruzar las variables componentes (actuales y nuevos) y relaciones entre los componentes (actuales y nuevas): radical, incremental, modular y arquitectónica.

		Componentes	
		Actuales	Nuevos
Relaciones entre componentes	Actuales	<i>Incremental</i>	<i>Modular</i>
	Nuevas	<i>Arquitectónica</i>	<i>Radical</i>

Figura 12: Innovación modular (Henderson y Clark, 1990)

La innovación radical consiste en la creación de algo totalmente nuevo, a partir de tecnologías que no existían con anterioridad. Por ejemplo, Francis Bacon observaba en su obra *Novum Organum* que tres grandes inventos mecánicos habían cambiado la faz de la tierra: la imprenta en la literatura, la pólvora en la guerra y el compás en la navegación. La innovación radical permite obtener un producto totalmente nuevo en el mercado, formado por nuevos componentes, que están unidos a través de una configuración original.

La innovación incremental se apoya en la mejora de los componentes actuales del producto, manteniendo la estructura de relaciones actual o mejorada. La innovación incremental no suele ser visible, ya que el aspecto externo del producto y las funciones que realiza apenas cambian. Sin embargo, puede tener una incidencia fundamental en los costes de producción, proporcionando a la empresa que la desarrolla una importante ventaja competitiva en costes, que, por lo general, va a tardar en ser detectada por los competidores.

La innovación modular surge al modificar radicalmente alguno de los componentes del producto, manteniendo la misma estructura de relaciones. En la innovación modular tiene particular importancia la compatibilidad entre los componentes, ya que la modificación de algunos de ellos puede distorsionar el equilibrio previo que existía entre todos los componentes del producto. Esta innovación es muy importante porque facilita la división del producto en partes y, de este modo, permite al usuario adquirir los componentes que considere necesarios o incluso realizar una compra programada para ir sustituyendo todos los componentes del producto viejo. Los automóviles actuales constituyen un ejemplo familiar de diseño modular. Como señala Iacocca (1984): “los coches nuevos son actualmente una mera ilusión. Cada novedad que sale al mercado es una combinación de partes y componentes de modelos anteriores. La originalidad suele ser el laminado del metal, la transmisión o el chasis. Pero nadie, ni siquiera General Motors, puede permitirse el lujo de fabricar un coche a partir de cero”.

La innovación arquitectónica utiliza componentes conocidos en una nueva configuración. Por tanto, no conlleva un avance en los componentes que se están aplicando. Sin embargo, puede revolucionar los mercados, sobre todo si la nueva configuración permite obtener un producto totalmente diferente a los existentes. Por ejemplo, la máquina de escribir *Remington Número 1*, la primera en ser ofertada al público en general en 1874, era una síntesis de muchas tecnologías existentes y de muchos elementos mecánicos de amplio uso en aquel momento. Los aparatos de relojería sugerían la idea de «escape» —es decir, del avance del carro una letra cada vez. Las teclas y sus brazos de conexión eran adaptaciones de la tecla de telégrafo. El pedal de una máquina de coser hacía retroceder el carro, y el piano sugería un modelo para los brazos de libre movimiento y los martillos que golpeaban las letras contra el papel (Utterback, 1994).

Las innovaciones arquitectónicas proceden a menudo de nuevos entrantes —no sólo nuevas empresas, sino también de grandes compañías que entran en nuevas áreas de negocio (IBM creó en el mercado de las máquinas eléctricas sin pertenecer al de las máquinas de escribir) o de filiales de competidores ya establecidos que se han creado para desarrollar una nueva idea (Utterback, 1994). La misión de la innovación arquitectónica consiste en tomar conciencia de las posibilidades de combinación de las tecnologías y adecuarlas al mercado y a los recursos de la empresa.

6.5. Naturaleza funcional de las innovaciones

El concepto de innovación va más allá de los puros cambios técnicos para comprender también la promoción, la distribución y muchas otras formas en las que las empresas pueden conseguir ventajas competitivas sobre sus rivales. Así, la OCDE (2006) define la innovación como la introducción de un nuevo, o significativamente mejorado, producto (bien o servicio), de un proceso, de un nuevo método de comercialización o de un nuevo método organizativo, en las prácticas internas de la empresa, la organización del lugar de trabajo o las relaciones exteriores. Una característica común a todos los tipos de innovación es que deben haber sido introducidos. Un nuevo (o mejorado) producto se ha introducido cuando ha sido lanzado al mercado, mientras que un proceso o un método de organización se ha introducido cuando ha sido utilizado efectivamente en el marco de las operaciones de una empresa.

En este apartado recogemos algunos ejemplos de innovaciones en diferentes áreas funcionales. No obstante, a lo largo del texto nos centraremos fundamentalmente en las innovaciones tecnológicas.

Innovaciones de producto. Las innovaciones de productos suponen la comercialización de un producto (bien o servicio) nuevo, o significativamente mejorado, en cuanto a sus características o en cuanto al uso a que se destina (OCDE, 2006). Tomando como referencia el grado de novedad con el que los productos son percibidos por las empresas y por el mercado, la consultora Booz, Allen y Hamilton (1982), en su estudio sobre aproximadamente 13.000 nuevos productos industriales y de servicios, distinguió seis categorías de nuevos productos (figura 13).

1. *Productos nuevos para el mundo.* Se trata de productos totalmente originales que pueden crear un mercado completamente nuevo. Son productos nuevos para la empresa y para el mercado (10 por ciento). Por ejemplo, el magnetófono de Sony en el momento de su introducción en el mercado.

2. *Nuevas líneas de productos.* La empresa entra por primera vez en un mercado ya establecido. Se trata de un producto nuevo para la empresa, pero no para el mercado (20 por ciento). Por ejemplo, la entrada de Bic en el mercado de las maquinillas de afeitar desechables.

3. *Incorporaciones a líneas de productos ya existentes.* Son productos nuevos que complementan una línea establecida de una empresa. El grado de novedad para la empresa y para el mercado es medio (26 por ciento). Por ejemplo, cuando Burger King introdujo la hamburguesa vegetariana.

4. *Mejoras en los productos ya existentes.* Aportan un mejor desempeño o un mayor valor de percepción y sustituyen a los existentes (26 por ciento). Por ejemplo, los teléfonos móviles cada vez tienen más prestaciones, que los hacen más atractivos para los clientes.

5. *Reposicionamiento.* Productos existentes que son dirigidos hacia nuevos mercados o segmentos (7 por ciento). El ejemplo paradigmático es el *nylon*, que inicialmente tenía una aplicación militar: paracaídas, hilos y cuerdas. A esto, le siguió la entrada del nylon en el mercado del tejido de confección y su dominio posterior del negocio de ropa interior femenina (Levitt, 1965). Actualmente, además de los anteriores tiene múltiples usos: mantas, neumáticos, pines, cepillos, cojinetes, alfombras y muchas más.

6. *Reducciones de costes.* Se eliminan los componentes no relacionados con la función principal del producto y se mejora la eficiencia de los sistemas productivos (11 por ciento). Los fabricantes de automóviles utilizan la misma plataforma (estructura de coche) en varias líneas de productos.

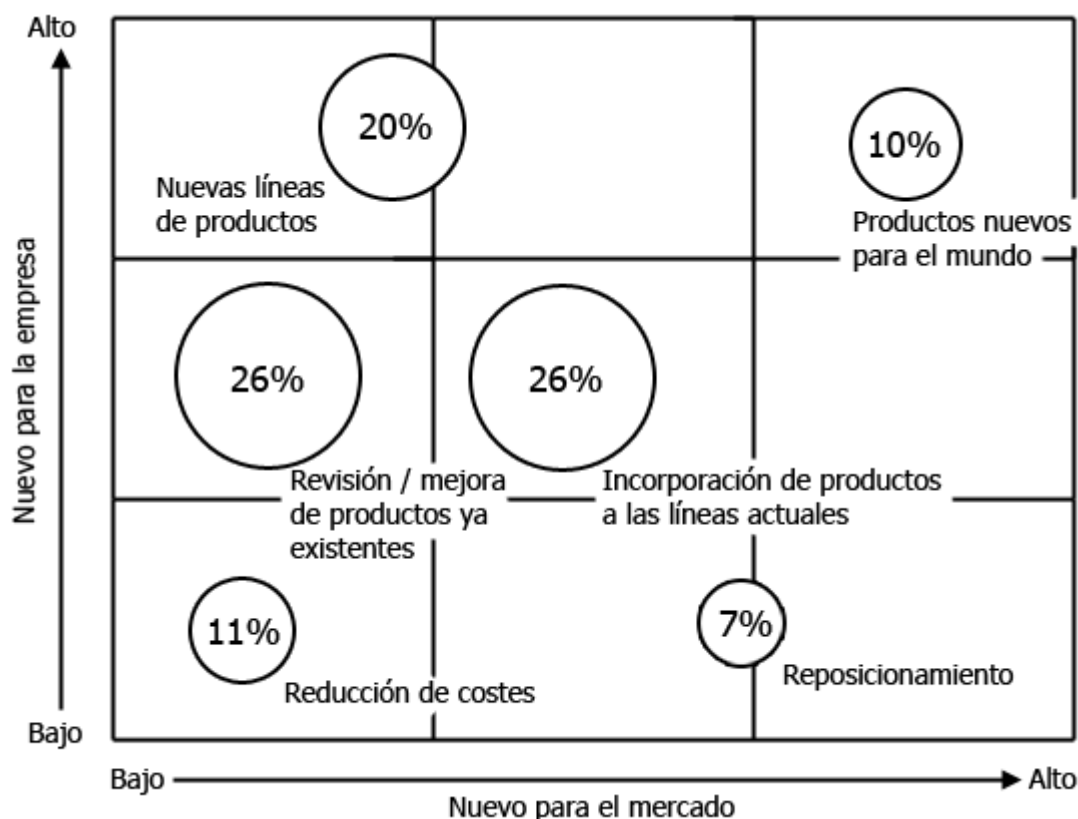


Figura 13: Categorías de nuevos productos (Booz, Allen y Hamilton, 1982)

Se puede observar, por tanto, que un nuevo producto no es necesariamente un producto realmente original, sino que puede ser un producto ya existente en el mercado, al que, sin haber alterado de forma sustancial su base tecnológica, se ha mejorado o ampliado sus prestaciones. Es importante destacar que sólo un 10 por ciento de todos los nuevos productos pertenecen a la categoría de productos absolutamente originales y que la mayoría de ellos (70 por ciento) son, en realidad, extensiones o modificaciones de productos existentes. La categoría de nuevos productos para el mundo supone los mayores costes y riesgos, ya que el producto es nuevo tanto para la empresa como para el mercado.

Innovaciones de procesos. Las innovaciones de procesos son la aplicación industrial de un nuevo, o significativamente mejorado, proceso de producción o de distribución. Ello implica cambios significativos en las técnicas, los materiales y/o los programas informáticos. Algunos avances son de tipo radical, lo que provoca una caída espectacular del coste de producción, mientras que otras innovaciones constituyen mejoras incrementales en los procesos, por lo que las reducciones de coste son más suaves.

Las innovaciones radicales en procesos crean una nueva arquitectura del proceso que provoca una discontinuidad en el mismo. Discontinuidades de procesos son modos esencialmente diferentes de fabricar un producto que se reflejan en mejoras de órdenes de

magnitud en el coste o en la calidad del producto (Anderson y Tushman, 1990). Las discontinuidades se apoyan en la eliminación de etapas o en la introducción de una nueva tecnología de proceso. Así pues, cada nuevo método de producción puede implicar una combinación o eliminación de etapas anteriores, dando como resultado espectaculares ganancias de productividad y costes unitarios inferiores. Lo mismo puede decirse de los casos en que se utiliza una nueva tecnología de procesos, como en el cambio de fabricación de vidrio por el método corona a vidrio en cilindros (Utterback, 1994). Puesto que cada nueva arquitectura del proceso conduce a costes unitarios inferiores, la relación entre tiempo, coste unitario y arquitecturas del proceso se asemeja a una escalera de bajada, representando cada escalón un cambio en la arquitectura del proceso.

Aunque los cambios en la arquitectura del proceso son normalmente pocos y espaciados, el progreso hacia una mayor productividad no se para entre los momentos en que se producen. En realidad, a los cambios discontinuos importantes siguen, normalmente, un cierto número de pequeñas mejoras incrementales. Estas mejoras son el resultado del perfeccionamiento de las diferentes etapas del proceso, de la mejora de las máquinas y del aprender mediante la práctica. De este modo, vemos una caída importante en los costes unitarios producidos por un cambio de arquitectura del proceso, seguida por un cierto número de pequeñas caídas, que resultan de mejoras incrementales (figura 14). La suma de las dos partidas es igual a las ganancias de productividad a largo plazo (Utterback, 1994).

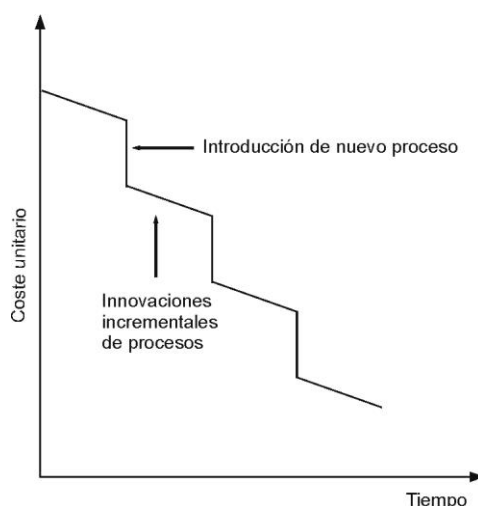


Figura 14: Ganancias de productividad a largo plazo (Utterback, 1994)

Normalmente las innovaciones en proceso comienzan mecanizando algunas actividades manuales. El diseño de máquinas las puede realizar la propia empresa o un proveedor especializado. A continuación, se automatizan grupos de actividades de forma independiente, lo que se conoce como «islas de automatización». Estas islas están conectadas por procesos manuales. Finalmente, se intenta integrar todas las islas en un flujo continuo⁷. En general, la escala aumenta de forma espectacular a medida que entran en juego las innovaciones de proceso. Enos (1962), en el caso del refino del petróleo, atribuye cada una de las oleadas de innovaciones del proceso a una causa diferente: a)

⁷ Una de las innovaciones radicales en proceso más conocida fue la cadena de montaje de Ford. Si bien en las últimas décadas las innovaciones de los fabricantes japoneses les ha permitido desafiar la regla de hierro de la producción en masa: las largas series de fabricación de productos normalizados y los bajos costes unitarios van necesariamente unidos.

extraer más gasolina del petróleo, cada vez más escaso, b) reducir los costes del refino y c) mejorar la calidad.

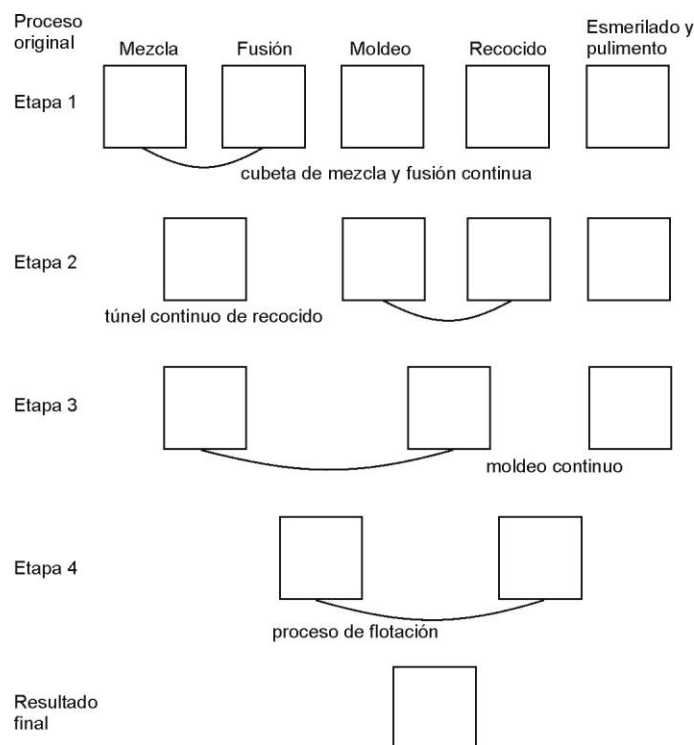


Figura 15: Evolución de los procesos de fabricación del vidrio plano pulimentado (Utterback, 1994)

En el caso de los productos no ensamblados, las innovaciones en procesos a menudo eliminan por completo pasos de producción manuales y automatizados, convergiendo finalmente en un proceso de producción continuo. Por ejemplo, el vidrio ha permanecido más o menos idéntico como tal. Sin embargo, ha cambiado el proceso de fabricación. Es éste último el que cuenta para el vidrio de hoy de alta calidad uniforme, en grandes cantidades y a bajo precio. El rumbo de las innovaciones de proceso ha llevado el negocio de la fabricación de vidrio a un estado altamente específico, en el que todas las etapas importantes de producción se han automatizado y unido en un sistema continuo. La figura 15 resume esta evolución (Utterback, 1994).

Los procesos de distribución están vinculados a la logística de la empresa y engloban los equipos, los programas informáticos y las técnicas para el abastecimiento de insumos, la asignación de suministros en el seno de la empresa o la distribución de productos finales. Las innovaciones en procesos incluyen métodos novedosos de prestación de servicios, así como las técnicas, los equipos y los programas informáticos utilizados en las actividades auxiliares de apoyo como las compras o la contabilidad.

Innovaciones comerciales. Una innovación comercial es la aplicación de un nuevo método de comercialización que implique cambios significativos del diseño o del envasado de un producto, su posicionamiento, su promoción o su tarificación (OCDE, 2006). Estas innovaciones tratan de satisfacer mejor las necesidades de los consumidores, de abrir nuevos mercados o de posicionar en el mercado de una nueva manera un producto de la

empresa con el fin de aumentar las ventas. Drucker (1985) considera que el agente de ventas que consiguiera vender frigoríficos a los esquimales para congelar los alimentos sería, sin duda, tan innovador como el que hubiera desarrollado un nuevo procedimiento o inventado un nuevo producto. Por tanto, deben considerarse también las innovaciones comerciales manifestadas a través de una novedosa presentación del producto, una nueva forma de distribución de un producto dado, una nueva campaña publicitaria o un nuevo envase, entre otros.

Por ejemplo, en 1970 Pepsi Cola concibió un plan para destronar a Coca-Cola y formuló una estrategia básica para hacerlo⁸. Primero, hizo un seguimiento de los hábitos de compra de los consumidores, estudiando en profundidad el comportamiento de 350 familias a las que se les dio la oportunidad de adquirir semanalmente cualquier refresco a precios especiales. Independientemente de la cantidad que compraran, la consumían toda. Ello llevó a Pepsi a crear botellas y cajas fáciles de transportar (permitiendo, así, a los consumidores comprar y llevarse a casa cantidades mayores de Pepsi). Esto se hizo sustituyendo la botella estándar de 6,5 onzas (unos 184 gramos) por otras más grandes de plástico y haciendo que los paquetes de botellas contuvieran ocho en lugar de seis. De esta forma, la empresa incrementó significativamente el volumen de ventas, pero estas actuaciones fueron pronto imitadas por los competidores. Segundo, Pepsi atacó la imagen de Coke a través de la publicidad. Los anuncios de «la generación Pepsi» presentándola como la bebida de la gente joven, relegaban a Coke, por eliminación, para gente mayor. Pepsi trató los *spots* como si fueran películas en miniatura y, para ello, reclutó a alguno de los mejores cineastas de Hollywood para que filmaran con calidad cinematográfica. La mayor parte de las empresas gastaban entre 15.000 y 75.000 dólares para producir un *spot*. Pepsi empezó a gastar entre 200.000 y 300.000 dólares por cada uno de ellos. Era algo nunca visto en aquellos días. Esta campaña pilló a Coca-Cola por sorpresa. Pepsi también atacó a Coke persiguiendo mercados locales concretos y canales de distribución característicos (tiendas de descuento, supermercados, restaurantes, etc.), mientras que Coke tenía que defender todas sus posesiones. Tercero, Pepsi realizó publicidad comparativa: mostrando a gente bebiendo refrescos de cola sin conocer la marca (pruebas a ciegas). En estas pruebas, se observó que, cuando la gente no conocía la marca prefería Pepsi a Coke, mientras que si conocían la marca se decantaban por Coke. Ello se debía a que Pepsi era ligeramente más dulce. Estas tres actuaciones proporcionaron a Pepsi ventajas temporales, pero en su conjunto le dieron una importante ventaja en Estados Unidos, permitiéndole derrotar a Coca Cola por primera vez en 1978 (30,8 por ciento del mercado norteamericano frente a un 29,2 por ciento de Coca-Cola) (Sculley, 1987).

A veces, la introducción de un nuevo producto en el mercado conlleva cambiar las prácticas comerciales actuales. Por ejemplo, una de las razones principales que explica el éxito de IBM en sus comienzos fue que alquilaba los ordenadores en lugar de venderlos. En aquella primera época los ordenadores tenían un coste muy elevado y se volvían obsoletos en pocos años, por lo que no era posible su amortización en un tiempo prudencial, constituyendo un importante *handicap* para los compradores potenciales. Cuando IBM decidió alquilar los ordenadores logró colocarlos con facilidad a sus clientes, aunque éstos sabían que pagaban un alquiler elevado (Drucker, 1985).

⁸ Este ejemplo aparece desarrollado en Sculley, J. (1987): *Odyssey: From Pepsi to Apple*, Harper & Row, Nueva York.

La ventaja de un mejor servicio se convierte en una parte esencial de un escenario muy competitivo. Los clientes demandan mayor velocidad, exactitud y fiabilidad en el servicio. Entre las empresas que han explotado tales demandas de los clientes se encuentran Caterpillar, que ofrece un servicio de repuestos las 24 horas del día en todo el mundo. En teoría, las mejoras en servicio sólo pueden ofrecer ventajas competitivas a corto plazo, ya que otras empresas pueden imitar rápidamente el mismo tipo de ventajas. En la práctica, las ventajas se mantendrán a largo plazo.

Innovaciones organizativas. Una innovación de organización es la introducción de un nuevo método organizativo en las prácticas, la organización del lugar del trabajo o las relaciones exteriores de la empresa (OCDE, 2006). Las innovaciones de organización pueden tener por objeto mejorar los resultados de una empresa reduciendo los costes burocráticos, mejorando el nivel de satisfacción en el trabajo (y, por consiguiente, aumentar la productividad), facilitando el acceso a bienes no comercializados (como el conocimiento externo no catalogado) o reduciendo los costes de los suministros.

Por definición, cualquier innovación tecnológica radical implica algún cambio en la organización de la producción y de los mercados. Las innovaciones organizativas e institucionales están intrínsecamente asociadas a las innovaciones técnicas. Además, en algunas ocasiones las innovaciones de organización y gestión pueden abrir el camino. Algunos ejemplos son la cadena de montaje y los supermercados (Freeman, 1998). Ford creó la cadena de montaje, una estructura organizativa que revolucionó el sistema productivo mundial. Por otra parte, en la década de 1920, algunas de las empresas más grandes de Estados Unidos (General Motors, DuPont Corporation y Sears, entre otras) crearon una revolución organizativa. En lo esencial, abandonaron la estructura funcional que predominaba hasta entonces, y adoptaron una estructura multidivisional (Chandler, 1962). En las últimas décadas han tenido éxito múltiples innovaciones organizativas: la gestión de la calidad total, la reingeniería de procesos y la organización virtual, entre otras.

Las empresas pueden introducir nuevas políticas de recursos humanos, tanto en lo relativo a la selección con objeto de atraer a los mejores trabajadores como en lo relacionado con la motivación y retención del personal para, así, mejorar su productividad y apropiarse parte de los beneficios resultantes. Un ejemplo de prácticas innovadoras lo proporciona la empresa Ford. El 5 de enero de 1914, Henry Ford anunció un nuevo salario mínimo de cinco dólares por ocho horas de trabajo al día, duplicando el salario medio del sector (2,38 dólares por nueve horas de trabajo) uno de los más altos de Estados Unidos (Ford, 1922). Al tomar esta medida fue criticado por la mayoría de empresarios norteamericanos, incluso fue tachado de comunista. No obstante, según el propio Ford esta decisión fue una de las mejores tácticas de reducción de costes que llevó a cabo la empresa. Este empresario había redescubierto una antigua verdad: pagando a los trabajadores más de lo que pueden ganar en cualquier otro lugar, es posible conseguir mano de obra que trabaje con mayor ahínco, sea más leal y tenga unos niveles de absentismo más bajos. Al mismo tiempo, consiguió atraer a los trabajadores más productivos del sector. Todos los días, al abrir sus puertas, la empresa Ford tenía cientos de trabajadores a cola solicitando empleo, pudiendo escoger a los mejores de la industria.

A veces, las políticas públicas y la regulación proporcionan el surgimiento de prácticas empresariales oportunistas. En este sentido, Baumol (1993) incluye las innovaciones en los procedimientos especulativos, tales como el descubrimiento de una táctica legal, no utilizada previamente, eficaz para desviar rentas hacia aquéllos que primero lo explotan.

No obstante, esta última actividad se considera totalmente improductiva, al no crear ningún tipo de riqueza.

6.6. Innovaciones en la base de la pirámide

La distribución de la riqueza y la capacidad de generar ingresos en el mundo pueden representarse en forma de una pirámide económica (figura 16). En la cima de la pirámide están los ricos, con numerosas oportunidades para generar altos niveles de ingreso. Más de cuatro mil millones de personas habitan la base de la pirámide, con menos de dos dólares de ingresos al día.

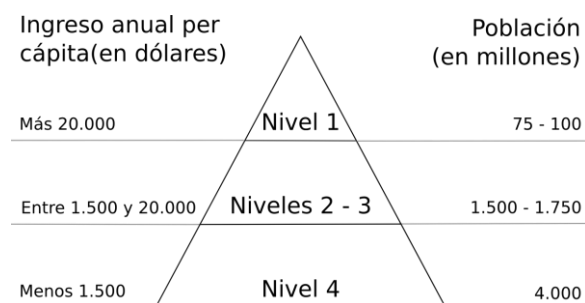


Figura 16: La pirámide económica (Prahalad y Hart, 2002)

La innovación en la base de la pirámide es un enfoque para ayudar a los menos favorecidos; un enfoque que supone asociarse con ellos para innovar y lograr escenarios ganadores en los que los pobres estén activamente comprometidos y donde, al mismo tiempo, las empresas que les suministran productos y servicios sean rentables (lectura 13). Por tanto, la economía básica del mercado de la base de la pirámide se fundamenta en paquetes de unidad pequeños, márgenes bajos, altos volúmenes y bajos rendimientos sobre el capital empleado. Esto es distinto de los paquetes de unidad grandes, altos márgenes, elevados volúmenes y rendimientos razonables sobre el capital invertido. El teléfono móvil 1100 de Nokia, con un precio de 15 dólares, fue diseñado para la innovación frugal, una innovación donde el precio es la prioridad número uno. El teléfono ofrece las posibilidades básicas de un móvil: hacer y recibir llamadas y enviar sms. Sólo incluye una función que no es esencial, que es una pequeña linterna, popular debido a los apagones frecuentes en algunas regiones emergentes, y tiene un revestimiento antideslizante para evitar que resbale de la mano en las regiones indias donde la humedad es muy elevada (Sehgal *et al.*, 2010)

Prahalad (2005) ha identificado doce principios que, tomados en su conjunto, constituyen los componentes esenciales de una filosofía de la innovación para los mercados de la base de la pirámide.

El rendimiento de los precios. El centro de atención debe estar en el rendimiento de los precios de los productos y servicios. Ocuparse de las oportunidades de mercado en la base de la pirámide exige comenzar por una comprensión radicalmente nueva de la relación precio–rendimiento, en comparación con la que se emplea en los mercados desarrollados. No se trata de reducir los precios sino de modificar la relación precio–rendimiento.

La innovación: los híbridos. La innovación exige soluciones híbridas. Las oportunidades del mercado de la base de la pirámide no pueden satisfacerse mediante versiones

«aguadas» de las soluciones tecnológicas tradicionales procedentes de los mercados desarrollados. El mercado de la base de la pirámide puede y debe atenderse con las tecnologías más avanzadas, combinadas de manera creativa con la infraestructura existente (y en evolución).

LECTURA 13: INNOVACIONES EN LA BASE DE LA PIRÁMIDE

Las incubadoras de hoy son aparatos complejos y caros. Una de las que se usan normalmente en Estados Unidos puede costar más de cuarenta mil dólares. Además, el problema no es solo lo que cuestan; los aparatos complicados se estropean, y entonces hace falta un experto para arreglarlos, y hace falta además repuestos. En el curso del año que siguió al tsunami del océano Índico de 2004, la ciudad indonesia de Meulaboh recibió ocho incubadoras, donadas por diversas ONGs internacionales. A finales de 2008, un profesor del MIT llamado Timothy Prestero visitó el hospital de esa ciudad y vio que las ocho estaban fuera de servicio, por culpa de los altibajos de la tensión eléctrica y de la humedad tropical; además, el personal médico no sabía inglés y no podía leer el manual de reparaciones. Y el de las incubadoras de Meulaboh no es más que un ejemplo gráfico; hay estudios que apuntan a que hasta un 95% del equipamiento médico donado a los países en vías de desarrollo queda fuera de servicio durante los primeros cinco años.

Prestero tenía interés personal en esas incubadoras estropeadas, porque era el fundador de una empresa, Design that Matters (Diseño que Importa), que llevaba varios años trabajando en el desarrollo de una incubadora más resistente y menos cara, un aparato que tuviera en cuenta el hecho de que la tecnología médica compleja probablemente requeriría un mantenimiento distinto en un país en vías de desarrollo que en un hospital europeo o norteamericano. Diseñar una incubadora a la medida de un país poco desarrollado no era solo cuestión de crear algo que funcionara; también se trataba de diseñar algo que, al estropearse, no lo hiciera de manera irreparable. Como no se podía garantizar que fuera a haber repuestos, ni técnicos con la formación necesaria para arreglar los aparatos, Prestero y su equipo decidieron construir una incubadora hecha de materiales que ya existieran en abundancia en esos países. La idea original fue de un médico de Boston llamado Jonathan Rosen, que había observado que incluso en los pueblos más pequeños de algunos países del Tercer Mundo parecían arreglárselas para conservar los automóviles en razonable buen estado. En aquellos pueblos quizá no hubiera aire acondicionado, ni ordenadores portátiles ni televisión por cable, pero se las ingeniaban para que un Toyota 4x4 siguiera funcionando. Así que Rosen le fue a Prestero con su idea: ¿por qué no se hacía una incubadora con componentes de automóviles?

Tres años después el equipo de Design that Matters presentó un prototipo denominado *NeoNurture* que, por fuera, parecía una incubadora moderna corriente, de líneas sencillas, pero que por dentro era capaz de autogenerarse. El calor, que es lo principal, lo fabricaban unos faros sellados; la circulación de aire se hacía con un ventilador de motor, y las alarmas sonaban con las campanillas de las puertas. El aparato se podía enchufar con un mechero de coche adaptado, o conectarse a una batería de motocicleta estándar. El que la *NeoNurture* se construyera con este tipo de componentes fue eficiente en dos sentidos: en el lugar donde se instalara habría repuestos para ella, como ya sabían, y habría además los conocimientos necesarios, porque eran los mismos que se aplicaban a la reparación de coches. Ambas cosas, como a Rosen le gustaba decir, abundaban en los países en vías de desarrollo. Y no había que ser técnico cualificado para arreglar una *NeoNurture*; no había ni que leerse el manual: si sabías cambiar un faro a un coche, sabrías hacerlo.

Fuente: Johnson, S. (2011): *Las Buenas Ideas. Una Historia Natural de la Innovación*, Turner Publicaciones, Madrid.

La escala de las operaciones. Teniendo en cuenta la magnitud de los mercados de la base de la pirámide, las soluciones que se desarrollen deben poder incrementar su escala y transportarse entre países, culturas e idiomas. Es fácil tener éxito en un experimento restringido, pero las necesidades de mercado de cuatro mil millones de personas sugieren que los experimentos deben tener la posibilidad de ampliar su escala comercial.

El desarrollo sostenible: benigno para el medio ambiente. La reducción de la intensidad de los recursos debe constituir el principio esencial en el desarrollo de productos. Todas las innovaciones deben concentrarse en la conservación de recursos: reducir, reutilizar y

reciclar. Los pobres, como mercado, suman 4.000 millones de personas. Esto significa que las soluciones que se desarrollen no pueden basarse en las mismas pautas en cuanto a uso de recursos que se espera emplear en los países desarrollados. Las soluciones deben ser sostenibles y benignas para el medio ambiente.

Identificación de la funcionalidad: ¿Es diferente la base de la pirámide de los mercados desarrollados?. El desarrollo de productos debe comenzar por una profunda comprensión de la funcionalidad y no sólo de la forma. No bastará con cambios marginales en los productos desarrollados para clientes ricos. Reconocer el hecho de que la funcionalidad que exigen los productos y servicios del mercado de la base de la pirámide puede ser diferente de la que se encuentre disponible en los mercados desarrollados es un punto de partida crucial (lectura 14). Los promotores del desarrollo deben partir de esta perspectiva y buscar las anomalías en sus expectativas anteriores basadas en sus experiencias con los mercados desarrollados.

La innovación de los procesos. Las innovaciones en los procesos son tan cruciales como las innovaciones de producto en los mercados de la base de la pirámide. Con frecuencia, la innovación debe concentrarse en erigir una infraestructura logística, incluida una manufactura sensible a las condiciones imperantes. Llegar a los clientes potenciales y educarlos puede también ser una tarea sobrecogedora para los no iniciados. Existe una significativa oportunidad para la innovación en los mercados de la base de la pirámide en torno a la redefinición de los procesos para adaptarlos a la infraestructura. La innovación de los procesos es un paso crucial para hacer que los productos y servicios sean accesibles para los pobres. Cómo entregar es tan importante como qué entregar.

La subcalificación del trabajo. El trabajo de capacitación es esencial. La mayoría de los mercados de la base de la pirámide son deficientes en cuanto a habilidades. En la mayoría de los mercados de la base de la pirámide hay escasez de talento. El trabajo, por tanto, debe ser subcalificado.

La educación de los clientes. La educación de los clientes (semianalfabetos) en el uso de los productos es esencial. La innovación en los mercados de la base de la pirámide exige inversiones significativas en la educación de los clientes en cuanto al uso adecuado y a los beneficios de productos y servicios específicos. Dada la deficiente infraestructura para el acceso a los clientes, es vital la innovación en el proceso educativo (lectura 14).

LECTURA 14: FRACASO DE UN BUEN PRODUCTO

A finales de los años ochenta Procter & Gamble (P&G) lanzó al mercado el detergente Ariel Ultra. Lavaba mejor la ropa y se requería sólo la mitad para hacer el mismo trabajo. P&G consideró que este era un beneficio importante porque buena parte de los hogares de bajos ingresos cuentan con un espacio limitado para su almacenamiento. Ultra también tenía enzimas que mejoraban el rendimiento en limpieza. P&G estaba tan convencida de que tenía un producto ganador que la mayoría de la producción se dedicó a Ultra y se inició una campaña enorme. Las mujeres mexicanas le informaron a P&G que las cosas eran diferentes. Por una parte, realmente no creían que la ropa quedara bien lavada con tan poco detergente. Por otra parte, Ariel Ultra no hacía espuma. Muchas personas en los hogares de bajos ingresos hacen sus labores de manera manual y tienen una percepción especialmente afinada de los olores; para ellas la producción de espuma indicaba que de verdad a la ropa se le estaba quitando el sudor. “Nos equivocamos completamente en cuanto a la importancia que tienen las señales estéticas y visuales de rendimiento para el consumidor de bajos ingresos”, afirma Herrera Moro de P&G. En cuestión de meses, Ariel Ultra fue retirado del mercado.

Cuando lanzó el producto Downy Single Rinse (Downy de un solo enjuague), que permitía a las amas de casa mexicanas reducir el número de enjuagues durante el lavado de ropa, lo que ahorra tiempo, esfuerzo

y agua, P&G lanzó el producto con la recomendación de la entidad mexicana para el agua y el medioambiente y se hicieron muchas demostraciones en las tiendas para que las mujeres pudieran ver cómo funcionaba.

Fuente: Lafley, A. G. y Charan, R. 2009): *Cambio de Juego*, Granica, Barcelona.

Diseñar para una infraestructura hostil. Los productos deben trabajar en entornos hostiles. No sólo deben soportar el ruido, el polvo, las condiciones antihigiénicas y el maltrato, sino acomodarse a la baja calidad de las infraestructuras, como en el caso de la electricidad y el agua. Los mercados de la base de la pirámide se dan en una infraestructura hostil, y esto debe ser tenido en cuenta en el diseño de productos y servicios.

Interfaces. La investigación sobre las interfaces es fundamental dada la naturaleza de la población consumidora. La heterogeneidad de la base de consumidores en cuanto a idioma, cultura, niveles de habilidad y familiaridad previa con las funciones o las características es un reto para el equipo de innovación. El diseño de la interfaz debe meditarse cuidadosamente. La mayoría de los clientes de los mercados de la base de la pirámide son usuarios de vez primera de productos y servicios y la curva de aprendizaje no puede ser larga ni ardua.

La distribución: el acceso al cliente. Las innovaciones deben llegar al consumidor. Es crucial el diseño de métodos de distribución para llegar a los pobres con bajos costes. Son cruciales los sistemas de distribución que lleguen a la base de la pirámide para el desarrollo de este mercado. Las innovaciones en la distribución son tan decisivas como las innovaciones de productos y procesos.

Los mercados de la base de la pirámide esencialmente permiten poner en tela de juicio la opinión ortodoxa en cuanto a provisión de productos y servicios. Los mercados de la base de la pirámide permiten, y obligan, a poner en tela de juicio los paradigmas existentes. La evolución de las características y las funciones en los mercados de la base de la pirámide puede ser muy rápida. Quienes desarrollan los productos deben concentrarse en la amplia arquitectura del sistema –la plataforma– de manera que las nuevas características puedan incorporarse con facilidad.

6.7. Innovación inversa

Mediante la glocalización, las multinacionales hacen negocios en los mercados emergentes exportando versiones ligeramente modificadas de sus productos globales, desarrollados para los clientes de los países ricos. Son, principalmente, modelos menos costosos, con un menor número de funciones. La glocalización es una mezcla de escala global y respuesta local. Las corporaciones calculan un balance óptimo entre escala global (crucial para minimizar costes) y la adaptación local que se requiere para maximizar la participación del mercado. La glocalización implica una mentalidad orientada hacia el producto. Sony vendía productos y tecnología viejos en China y como consecuencia de esta estrategia perdió terreno ante Samsung (Govindarajan y Trimble, 2012).

Los países emergentes se enfrentan a muchos desafíos, como un poder adquisitivo limitado, infraestructuras incompletas y falta de recursos. Por lo tanto, los productos que están dirigidos a mercados emergentes deben diseñarse para que hagan frente a estos desafíos.

Los compradores de los países en desarrollo no pueden exigir los altísimos niveles de desempeño a los que estamos acostumbrados en los países ricos. Para que las empresas que invierten en este tipo de innovación creen valor, los productos tienen que tener las siguientes características:

- Centrarse en las necesidades fundamentales para simplificar el diseño del producto. Por tanto, hay que eliminar las funciones no esenciales de los productos.
- Coste total reducido: no solo el precio al momento de la compra, también los costes de uso, mantenimiento y reparación.
- Los productos necesitan adaptarse a entornos hostiles. Los productos deben hacer frente a varias deficiencias de infraestructura, como cortes de corriente o temperaturas extremas.
- Facilidad de uso. Los compradores no tienen experiencia previa en el uso de productos similares. Ahorro en el uso de los recursos.

La innovación inversa tiene lugar cuando los productos diseñados para regiones emergentes se abren camino en economías más desarrolladas. Este enfoque contrasta con el tradicional ciclo internacional del producto. General Electric (GE), por ejemplo, diseñó un electrocardiograma y una máquina de ultrasonidos portátiles, ambos a bajo precio, para usarlos en zonas rurales de Asia. La prioridad de esos productos eran el precio, la facilidad de uso y la portabilidad, seguidas por último del rendimiento. Ambos productos se desarrollaron para mercados emergentes, y ambos han sido herramientas de crecimiento importantes para GE en esos mercados. Sin embargo, ahora estos dos artículos se comercializan en EEUU para segmentos del mercado donde la portabilidad es importante, como los equipos de emergencia que prestan asistencia en un accidente de tráfico (Immelt *et al.*, 2009). Las versiones estadounidenses de estos productos, desarrollados originalmente para médicos indios y chinos, se venden al por menor con un precio un 80% menor respecto a otros productos comparables en el mercado (figura 17).

GLOCALIZACIÓN (INNOVACIÓN ADAPTADA)	INNOVACIÓN INVERSA
Optimizar los productos para los clientes del mundo desarrollado.	Encontrar la mejor solución para el cliente del mercado emergente.
Productos ultramodernos, tecnológicamente sofisticados, de alto desempeño, con muchas funciones y nuevas y atractivas aplicaciones.	Productos sobrios, funcionales y de calidad aceptable.
Tomar el camino más corto para diseñar productos destinados a los mercados emergentes: eliminar funciones para reducir costes.	Reinventar el producto a partir de lo básico; innovación desde cero.
Precio alto, orientación hacia los amplios márgenes de ganancias.	Precio bajo, orientación hacia el volumen
Énfasis en la tecnología; orientación hacia el producto.	Orientación hacia el cliente y la necesidad.
Buscar los clientes a los que se les pueda vender los servicios.	Identificar las dificultades del cliente y desarrollar productos para resolver sus

	problemas.
Vender productos a los clientes actuales del producto.	Crear nuevo consumo entre los nuevos clientes.
Ampliar la participación en el mercado.	Crear el mercado.
Aprovechar las competencias centrales existentes.	Crear nuevas competencias centrales.
Visión de explotación para las economías emergentes.	Visión de exploración para las economías emergentes.
Usar productos del mundo desarrollado para transformar los mercados emergentes.	Crear nuevas plataformas de crecimiento basadas en los mercados emergentes.

Figura 17: Glocalización versus innovación inversa (Govindarajan y Trimble, 2012)

6.8. Innovación en el modelo de negocio

Los modelos de negocio no tienen nada de secreto. En el fondo son historias, historias que explican cómo funcionan las empresas (Magretta, 2002). Un modelo de negocio es una forma de crear valor para los clientes y retener una parte de dicho valor en la empresa. De manera más específica Chesbrough y Rosenbloom (2002) reivindican que el modelo de negocio debe cumplir las siguientes funciones:

Propuesta de valor. Articular una propuesta de valor: determinar cómo crear valor para los clientes a través de una oferta basada en la tecnología. Es necesaria una definición preliminar de cómo será el producto final que se lance al mercado y de qué forma lo utilizará el usuario. Un modo útil de imaginar una propuesta de valor es considerar el punto de vista del cliente: ¿qué problema se está solucionando? ¿es muy importante este problema para el cliente?

Segmento de mercado. Identificar un segmento del mercado: especificar la finalidad para la que el producto será utilizado. El modelo de negocio debe tener en su punto de mira a un grupo de clientes, o segmento de mercado, a los cuales la propuesta de valor resultará atractiva y, por tanto, la demandarán.

Las empresas necesitan delimitar un conjunto de clientes para poder decidir en qué atributos tecnológicos se pondrá el énfasis durante la etapa de desarrollo. Es probable que haya que considerar muchas alternativas técnicas, diversos posibles candidatos a los que dirigirse en el mercado y eventuales competidores en el desarrollo de las tecnologías. Plantearse un segmento de mercado proporciona a los científicos y técnicos pistas sobre dónde deben focalizar sus actividades. Si una empresa no consigue centrar lo suficiente su proyecto, corre el riesgo de sobrecargar el producto resultante con demasiados atributos de dudoso beneficio.

Cadena de valor. Hay que definir la estructura de la cadena de valor necesaria, requerida para crear y distribuir el producto, y determinar los activos complementarios que son necesarios para afianzar la posición de la empresa en el mercado. Es la característica más identificada con el modelo de negocio. Esta cadena de valor debe alcanzar dos objetivos fundamentales: debe crear valor (haciendo llegar ese valor al consumidor al final de la cadena) y debe permitir a la empresa retener una porción suficiente del valor creado (Chesbrough, 2009).

Crear valor es necesario, pero a continuación la empresa debe decidir cómo se apropiará de alguna porción de ese valor. Según ha demostrado Michael Porter (1980), la capacidad de reclamar valor dependerá del poder de las cinco fuerzas competitivas. Otras investigaciones han confirmado que reclamar valor depende igualmente de la disponibilidad de activos complementarios, lo cual aumentará el valor de las ofertas de la empresa. La posesión o el acceso a los activos complementarios, como fabricación, distribución y marcas, ayuda a la empresa a retener una parte del valor que genera.

Estructura de costes y márgenes. La empresa tiene que especificar los mecanismos de generación de ingresos y estimar la estructura de costes y los beneficios potenciales de elaborar el producto, dada la propuesta de valor y la estructura de la cadena de valor elegidas. Para determinar la estructura de ingresos hay que tener en cuenta: cómo pagarán los clientes, a qué precios vender y cómo se repartirá el valor creado entre los clientes, la empresa y sus proveedores. Para ello existen muchas opciones, incluyendo la venta sin más, el alquiler, el pago por la transacción, modelos de publicidad y suscripción, licencias, e incluso regalar el producto y vender luego servicios de mantenimiento y asesoramiento (modelo maquinilla y hojas de afeitar). Una empresa puede también utilizar más de un mecanismo de pago, como hacen los periódicos, que cobran a los lectores por los ejemplares en circulación y a los anunciantes por la publicidad.

Una vez que conocemos las especificaciones generales de la oferta, podemos entonces determinar su estructura de coste. Este cálculo preliminar de precio y coste nos proporcionará los márgenes. Estos márgenes justifican los activos reales y económicos necesarios para efectuar una propuesta de valor (Chesbrough, 2009).

Red de valor. Creada alrededor de un negocio dado, la red de valor (también llamado ecosistema) moldea el papel que proveedores, clientes y terceros desempeñan en la creación de valor. Además de aumentar la provisión de productos complementarios desde la perspectiva del suministro, la red de valor puede incrementar los efectos globales entre los clientes desde la perspectiva de la demanda. Construir sólidas conexiones en una red de valor puede aumentar el valor del producto para el usuario. Fracasar en la construcción de una red de valor semejante puede reducir el valor potencial de ese mismo producto, en especial si compite con otro que goce de una sólida red de valor (Chesbrough, 2009).

Estrategia competitiva. Hay que formular la estrategia competitiva mediante la cual la empresa obtendrá y retendrá su ventaja frente a sus rivales. Las tempranas investigaciones de Porter sobre esta cuestión en la década de 1980 enfatizaban la necesidad de competir en el coste, en la diferenciación o en la base del nicho. Estudios más recientes han analizado los fundamentos para sostener el éxito competitivo, que incluyen la habilidad para obtener acceso diferencial a recursos clave, la creación de procesos internos valiosos para los clientes y difíciles de imitar por la competencia, y la experiencia pasada y situación presente de la empresa en el mercado (Chesbrough, 2009).

El 26 de septiembre de 1959, Haloid sacó al mercado la fotocopidora 914, que empleaba la *xerografía* para hacer fotocopias. En contraste con las tecnologías imperantes en esa época, esta nueva tecnología producía copias secas de alta calidad sin necesitar papel térmico. Por entonces, las fotocopias se hacían para uso comercial, bien

mediante métodos fotográficos húmedos o por procesos térmicos secos. Cada método producía imágenes de baja calidad que no sobrevivían en buen estado al paso del tiempo. Los modelos de negocio imperantes para cada proceso incluían cobrar los equipos a un precio relativamente modesto y luego vender por separado suministros y otros complementos imprescindibles, por lo general, con un coste comparativo mucho más alto, un modelo de negocio de “maquinilla y cuchillas de afeitar”. En lugar de vender las máquinas, Haloid les ofreció a los clientes arrendarlas. Un cliente solo debería pagar 95 dólares mensuales por el alquiler de la máquina y comprometerse a pagar cuatro céntimos por copia cuando se superasen las primeras dos mil copias efectuadas cada mes con dicha máquina. Haloid proveería de todos los servicios y respaldo necesarios y el alquiler podría cancelarse con un aviso previo de solo quince días.

El modelo de negocio de Haloid enfocaba sus productos y esfuerzos de venta hacia el mercado de las grandes empresas y las organizaciones gubernamentales. La empresa organizó su cadena de valor para ofrecer sistemas de fotocopiado completamente configurados a través de su propia organización de ventas directas, con amplios servicios de mantenimiento proporcionados por sus propios técnicos. La empresa puso precio a sus productos y servicios para ganar algo de dinero con las máquinas, pero obtuvo la mayor parte de sus ganancias con la venta de servicios y complementos (por ejemplo, cartuchos de tinta y papel). Los clientes solo se comprometían al pago del alquiler mensual y no se les obligaba a hacer ningún otro pago a menos que la calidad y la conveniencia de la fotocopidora 914 les llevasen a hacer más de dos mil copias mensuales. Este modelo de negocio solo rendiría beneficios a Haloid si la fotocopidora 914 conseguía incrementar en gran medida el volumen de copiado.

La fotocopidora 914 demostró ser una apuesta acertada. Una vez que se instaló en las oficinas de los clientes, las ventajas de la máquina resultaron evidentes. Los usuarios superaron las dos mil copias en un día (no en un mes), dada la alta calidad de imagen de las fotocopias (ni manchas de huellas digitales debidas a los procesos de copiado húmedo, ni papeles térmicos amarillentos y ondulados) y la conveniencia de su uso. Este notable incremento en el uso implicó que la mayoría de las máquinas generasen a Haloid ganancias por copia hacia el segundo día de cada alquiler mensual. Este modelo de negocio superó con mucho las previsiones más optimistas de la empresa, con un incremento medio de beneficios del 41% durante doce años. Como consecuencia, la empresa Haloid se transformó en una empresa global (rebautizada como Xerox) con 2.500 millones de dólares de ingresos en 1972.

Dicho modelo de negocio no requería las asociaciones con terceros. De hecho, Xerox decidió proporcionar por sí sola todos los elementos necesarios. Llevó adelante sus propias investigaciones, ejecutó todas las actividades de desarrollo necesarias para lanzar el producto y respaldar otros nuevos, fabricó todos sus productos internamente y los distribuyó mediante sus propios canales de distribución. La empresa proporcionó a sus clientes sus propios métodos de financiación y servicios de mantenimiento y apoyo. Xerox fabricó incluso su propio papel, para lograr una óptima alimentación de sus máquinas, aunque en este aspecto, tuvo que utilizar también papel fabricado por otras empresas.

Este enorme éxito tuvo efectos duraderos en Xerox. El gran triunfo del modelo de negocio de la fotocopidora 914 (que generaba más ganancias cuanto mayor era el número de copias realizadas) estableció la lógica dominante para las posteriores

actividades en la industria del copiado. Xerox desarrolló máquinas todavía más veloces que podían manipular enormes volúmenes de copias con un tiempo de funcionamiento mínimo y gran disponibilidad.

Entretanto, las nuevas empresas japonesas identificaron un talón de Aquiles en el modelo de Xerox. Este sistema funcionaba bien cuando se aplicaba a las grandes empresas, que necesitaban grandes cantidades de copias de alta calidad. Sin embargo, no se adaptaba tan bien a las necesidades de las pequeñas empresas o de los particulares. Estos grupos no tenían necesidad de grandes volúmenes de fotocopias, eran mucho más sensibles al precio de la máquina y estaban dispuestos a reducir la calidad de la imagen si eso les ahorra dinero.

Las empresas japonesas se centraron en este segmento del mercado de las fotocopadoras con un modelo de negocio distinto. Así, diseñaron un producto que se podía mantener sin necesidad de recurrir a técnicos especialmente entrenados por las empresas. Para ello, colocaron las piezas de la fotocopadora más frágiles y que se gastaban con mayor frecuencia en un cartucho reemplazable. De este modo, se permitía a las empresas volver a aplicar el modelo de la maquinilla de afeitar, dado que las fotocopadoras se venderían con un margen de ganancia más modesto, mientras que los cartuchos de repuesto les supondrían grandes márgenes de ganancia. Los fabricantes japoneses crearon un canal de distribución indirecto de vendedores y distribuidores para vender sus máquinas y proveer de servicios y financiación cuando fuese preciso. Esta modalidad les ahorra el coste de crear una fuerza de ventas directa. También les permitía conseguir con gran rapidez una capacidad de distribución en todo el territorio nacional y proporcionaba a los potenciales clientes la ventaja de ir a probar las nuevas máquinas a una tienda local antes de comprarlas.

	Xerox	Fotocopadoras japonesas
Propuesta de valor	Fotocopias de alta calidad a un bajo coste de alquiler mensual.	Máquina de bajo coste; mayor rentabilidad de fotocopias.
Segmento de mercado identificado	Mercado de grandes empresas e instituciones gubernamentales.	Mercado de particulares o pequeñas empresas.
Elementos de la cadena de valor	Desarrollo del sistema de fotocopiado en su totalidad, incluyendo complementos; venta a través de una organización comercial directa en la empresa.	Cartucho y máquina creadas internamente; distribución, servicios, respaldo y financiación externos a cargo de terceros.
Coste definitivo y margen de ganancias	Escasa ganancia con las máquinas, alta ganancia con los complementos y cantidad de fotocopias.	Escasa ganancia con la máquina, altos márgenes de ganancia con los cartuchos; modelo de negocio similar al de “maquinilla y cuchilla de afeitar”.
Posicionamiento en la red de valor	Pionera en el proceso de copiado seco; no requirió la ayuda de socios.	Emplean a terceros como vendedores de las máquinas para expandir las ventas territoriales; cartuchos más duraderos.
Formulación de la estrategia competitiva	Compite en tecnología, calidad del producto, capacidades del producto.	Compite en menor coste de venta de la máquina, puestos de venta convenientemente situados, calidad de la máquina/servicio.

Figura 18: El modelo de negocios de Xerox comparado con el nuevo modelo de los fabricantes japoneses (Chesbrough, 2009).

Cuando las empresas japonesas como Canon y Ricoh comenzaron a vender sus fotocopiadoras para escritorio en 1976, Xerox no les prestó atención, ya que esas máquinas pequeñas y baratas no podían hacer demasiadas copias por minuto. Es más, ni siquiera podían alimentar múltiples hojas de forma automática, recopilar hojas, ni aumentar o disminuir el tamaño de la imagen de una copia. Lo que Xerox no supo ver, sin embargo, era la propuesta de valor tan diferente que ofrecían: en lugar de tener que ir a un centro de reprografía para hacer las fotocopias, uno podía tener su propia fotocopiadora en la oficina, una auténtica ventaja (figura 18).

Las grandes empresas clientes de Xerox no percibieron demasiado valor en las fotocopiadoras de primera generación de Canon, Ricoh y otras empresas apenas llegadas al mercado, pero los pequeños comercios y los clientes particulares hallaron en ellas un enorme valor.

La llegada de las empresas japonesas demostró ser un desafío de enormes proporciones para Xerox. Aunque podían diseñar fotocopiadoras mucho más elaboradas y perfectas, los técnicos de Xerox, en respuesta a este desafío, tuvieron que abandonar la lógica dominante de la empresa tan exitosa que habían creado. Esto implicó que los técnicos que antes habían destacado por desarrollar máquinas que movían más deprisa el papel mediante complejos sistemas mecánicos, ahora tuviesen que crear adicionalmente productos mucho más sencillos y a precios mucho más bajos. El departamento de ventas debió determinar cómo gestionar una fuerza de ventas indirecta en paralelo con la ya existente fuerza de ventas directa, y pasaron incontables horas debatiendo por qué y cuándo un consumidor debía ser atendido mediante canales directos o indirectos. A su vez, en marketing tuvieron que decidir cómo promocionar la marca Xerox en el otro extremo del mercado (que representaba menores ganancias por máquina), pero manteniendo al mismo tiempo el extremo de tecnología punta y excelencia técnica (que representaba altos márgenes de ganancia por máquina y había catapultado a la empresa al éxito). A Xerox le llevó una década sobrellevar la amenaza de las nuevas empresas japonesas llegadas al mercado de las oficinas personales y las pequeñas empresas. En 2001, sufriendo grandes presiones en el sector, Xerox decidió abandonar esa parte del mercado, considerando que no merecía sus esfuerzos ni sus recursos.

7. LA DIFUSIÓN DE LA INNOVACIÓN

Una vez que se ha conseguido la innovación comienza su difusión o expansión en el mercado, lo cual dota de verdadera relevancia económica a la innovación. El extraordinario historial de inventos chinos es bien conocido, incluyendo la impresión, la brújula, la pólvora, la rueda, los naipes y muchas cosas más. No existía una falta de invención, lo que estaba ausente en su economía era un mecanismo efectivo de estímulo de los otros componentes de una innovación: la difusión generalizada de las nuevas tecnologías y productos.

En el proceso de difusión desde el lado de la oferta participan tanto la empresa que introdujo la innovación como los imitadores; en la medida de lo posible el innovador debe levantar barreras a la imitación, no que no resulta fácil en la época actual (figura 19). Desde el lado de la demanda, a la empresa innovadora le interesa que los usuarios le

acepten ampliamente sus productos, para lograr una rápida penetración en el mercado y apropiarse de los beneficios resultantes.

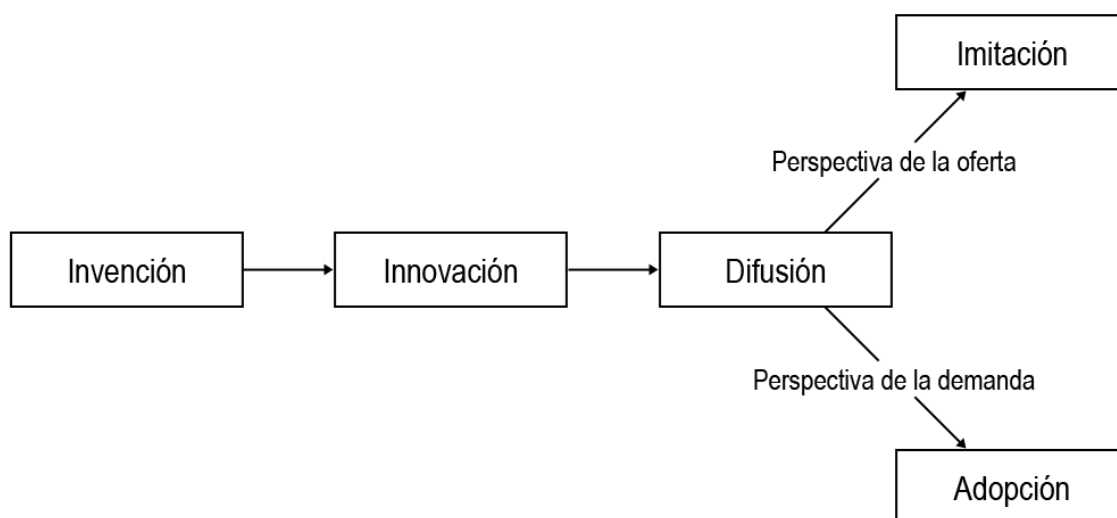


Figura 19: Componentes del proceso innovador

La difusión es un proceso por el que una innovación se comunica a través de ciertos canales, a lo largo del tiempo, entre los miembros de un sistema social (Rogers, 1962). En consecuencia, es tanto un proceso de influencia en los diversos actores del sistema social como un proceso económico, en el que se alteran los costes, los beneficios, la estructura del mercado y las condiciones competitivas.

Schumpeter (1942) considera que la difusión comprende sólo actos de imitación, sin ningún componente innovador, que permiten la divulgación y extensión del nuevo producto o proceso. Sin embargo, Rosenberg (1976) critica esta concepción, por entender que la innovación es sólo el comienzo del proceso de difusión en el mercado, que, a su vez, dependerá de una corriente de mejoras en las características de rendimiento de la innovación: modificaciones progresivas para adaptarla a los distintos segmentos de demanda e introducción de mejoras funcionales complementarias, que los clientes acepten mayoritariamente. El perfeccionamiento de las características de funcionamiento, las posteriores modificaciones técnicas para satisfacer mejor las demandas del usuario y la reducción de los costes unitarios determinan la creciente aceptación de la innovación o, lo que es lo mismo, contribuyen a su difusión. De hecho, para que a partir de un nuevo progreso tecnológico surja un negocio, es necesario que la innovación siga perfeccionándose continuamente y que permanezca en la vanguardia del mercado (Morita, 1986).

La difusión no se produce súbitamente en un instante de tiempo, sino durante un intervalo determinado, cuya longitud, difícil de predecir, varía de un sector a otro y de una innovación a otra. Rogers (1962) considera que el proceso de adopción de innovaciones durante un determinado período de tiempo sigue una curva normal: con forma de campana cuando se utilizan frecuencias (permite clasificar a los innovadores), y con forma de S (figura 20) si se utilizan datos acumulados.

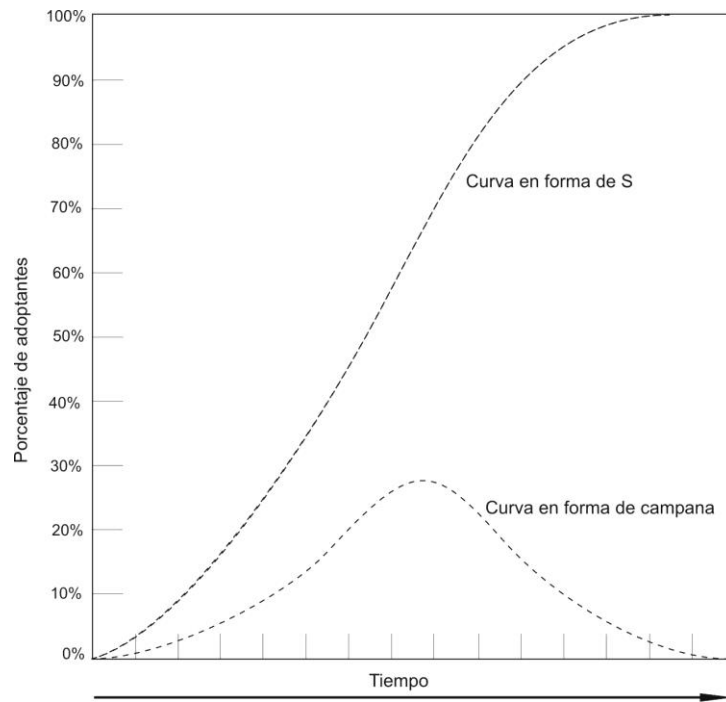


Figura 20: Modelo de difusión (Rogers, 1962)

La curva de difusión tiene forma de S debido a que cuando se introduce en el mercado una tecnología no familiar, la adopción es inicialmente lenta, se acelera conforme la tecnología se convierte en más conocida y es utilizada por el mercado de masas, y finalmente el mercado se satura, por lo que la tasa de nuevas adopciones disminuye, y la curva parece aproximarse asintóticamente a un límite (figura 21). Todas las innovaciones no pasan por este proceso siguiendo el mismo ritmo, por lo que pueden observarse diferentes grados de simetría y de irregularidad. La duración de la fase inicial, en particular, puede variar considerablemente. Durante los últimos años, el tiempo requerido para la completa difusión de una innovación se ha ido acortando. Hoy se considera que puede durar una media de 2 ó 3 décadas, que representa una tercera parte de lo que duraba durante el siglo diecinueve y principios del veinte (Ray, 1990).

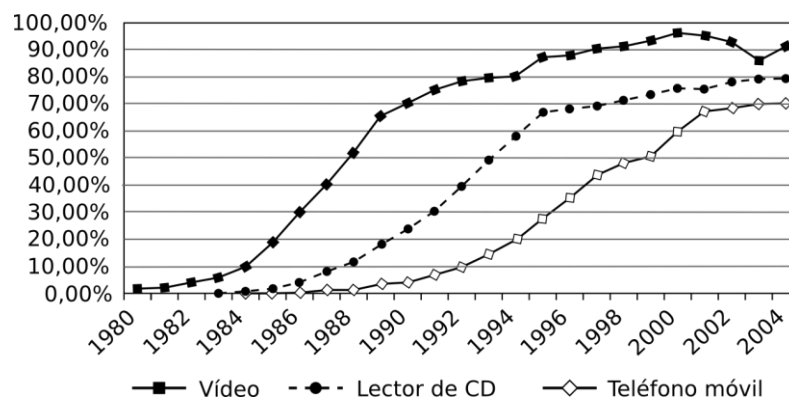


Figura 21: Penetración de la electrónica de consumo (Schilling, 2008)

7.1. Características de la innovación que afectan a la difusión

Hay una serie de características relacionadas con la innovación que contribuyen a explicar su tasa de difusión en el mercado. Entre otras, parecen tener cierta relevancia las cinco siguientes (Rogers, 1962).

Ventaja relativa. Grado en que se percibe la innovación como superior al producto (o proceso) que reemplaza. Pasará un cierto tiempo hasta que los productos en uso se deprecien en tal grado que el cambio al nuevo producto resulte económico. Lo que es más importante: hasta que los clientes potenciales comprendan y aprecien la innovación (Nelson, 1968).

La ventaja será efectiva si se percibe con claridad por los adoptantes potenciales. Para ello, los resultados de la innovación tienen que ser claros. David (1966) afirma que, a comienzos de la década de 1850, a los granjeros en pequeñas plantaciones les resultaba más provechoso utilizar guadañas para cosechar el grano que comprar una segadora mecánica. Según sus cálculos, una granja de 19 hectáreas marcaba el umbral mínimo del tamaño para el uso de la segadora. Sólo los que disponían de una extensión superior a este umbral encontraban productivo invertir en segadoras mecánicas. Los demás siguieron utilizando guadañas. Desde mediados de la década de 1850, aumentaron los costes de la mano de obra y el tamaño de las granjas, y permaneció estable el precio de la segadora, momento en el que empezó a resultar económico utilizarla, lo que contribuyó a su amplia difusión en el mercado.

Compatibilidad. Grado en que la innovación es congruente con los valores al uso y las prácticas existentes de los adoptantes. Cuanto mayor sea la compatibilidad de la innovación con el sistema existente, mayor será la posibilidad de que se adopte. Así, una innovación verá frenado su proceso de difusión si origina cambios significativos en los procesos, de forma que se altere totalmente el sistema de producción o la estrategia competitiva (Sutton, 1980).

Hay que valorar la compatibilidad del sistema tecnológico con el entorno donde la empresa desarrolla su actividad. La tecnología es sólo la «punta del iceberg», lo cual quiere decir que una máquina es inútil sin energía que la ponga en marcha, sin materiales que transformar o sin un contexto en el que funcionar, sin técnicos ni piezas para su mantenimiento y sin un sistema social capaz de proporcionar estos servicios de base y de distribuir y utilizar su producción. El fracaso de la aportación de todo este sistema de sostén equivale, sin más, al fracaso de la difusión de tecnología (Robertson, 1981).

En el caso de una tecnología de proceso, a nivel intra empresarial, es muy importante el equilibrio que debe haber entre el sistema social que la utiliza y la tecnología de proceso. Así, por ejemplo, un estudio realizado por Trist y Bamforth (1951) acerca de las repercusiones que tuvo en la industria del carbón el cambio de una tecnología tradicional por otra nueva, pone de manifiesto la necesidad de valorar siempre la tecnología desde una perspectiva integradora y nunca aisladamente. El gobierno inglés se había propuesto mecanizar de manera masiva las minas de carbón de propiedad pública y para ello realizó un primer ensayo en la mina más productiva. La tecnología tradicional requería grupos pequeños de dos a ocho personas, que formaban un equipo muy interdependiente, que casi siempre actuaba aislado del resto de equipos. Cada equipo era autónomo y tenía la responsabilidad de excavar, cargar y remover el carbón de las capas de la mina. La interdependencia y autonomía del equipo provocaba el desarrollo de unas relaciones emocionales y afectivas muy fuertes entre sus miembros. Los equipos competían para

conseguir ser la mejor sección de excavación, lo que repercutía positivamente en la productividad y el rendimiento. Esta actitud se aceptaba como normal por la comunidad y no creaba ningún tipo de tensión social.

La dirección, para incrementar aún más la productividad de la mina, introdujo una tecnología mecánica para excavar y remover el carbón. En principio todo (dirección, sindicatos y trabajadores) estaban de acuerdo con la nueva tecnología, ya que realizaba las tareas más penosas y era muy productiva, por lo que aumentarían las primas a la producción. La nueva tecnología requería grupos de trabajo grandes, lo que rompió la organización existente. Además, provocó la dispersión de los trabajadores y su división en tres turnos. También creó diferencias en estatus entre los trabajadores en relación con el trabajo que realizaban. Esta ruptura con el sistema social anterior provocó una frustración en los trabajadores y abocó en una disminución de su productividad y en un incremento del absentismo. La aportación más importante de este caso es que un cambio tecnológico dictado por consideraciones racionales desde el punto de vista de la ingeniería puede perturbar la organización social de los trabajadores hasta tal punto que la nueva tecnología no funcione eficientemente (Trist y Bamforth, 1951).

Por otra parte, la difusión se verá afectada por la compatibilidad del nuevo producto con los estilos de vida y la cultura imperante (Rogers y Shoemaker, 1971). Por ejemplo, en el Senegal se introdujo la cocina de gas para reemplazar a la de carbón con objeto de aprovechar el exceso de gas butano que tenía el país y, de paso, evitar la deforestación. Para ello, se hizo una campaña de modernización, que provocó un fuerte incremento de la demanda. No obstante, al año de su introducción, las ventas de cocinas cayeron en picado, debido a que las cocinas de gas no permiten preparar un buen té, aspecto importante de la vida social en el Senegal (Ackermann, 1981).

Complejidad. Grado en que una innovación se comprende y se utiliza con facilidad. Cuanto más difícil resulte su uso menos atractiva será para los clientes potenciales, que no la comprarán. La complejidad de un producto/máquina está relacionada con el número de tecnologías que la sustentan. Las innovaciones complejas son difíciles de comprender por lo que aprender su aplicación resulta un proceso arduo y lento, lo que aboca en una ralentización de la difusión. Igualmente, la complejidad está asociada a altos costes y riesgos elevados, por lo que se reduce el número de clientes potenciales de la tecnología.

Divisibilidad. Grado en que una innovación se puede probar de manera limitada sin realizar mucho gasto. La disponibilidad de diferentes tamaños, permitiendo una primera prueba, con un nivel de riesgos de pérdida más bajo en caso de que la innovación no cumpla con el desempeño prometido, contribuye a incrementar la difusión. En este sentido, la adopción de una innovación no se convierte en una decisión irreversible. La divisibilidad también se puede lograr alquilando el producto en lugar de venderlo.

Observabilidad. Grado en que los resultados sobre el uso de una innovación pueden observarse o describirse por parte de otros. La descripción del producto y sus aplicaciones favorece la transmisión de información entre usuarios potenciales. Davis y North (1971) subrayan que la reducción del coste de obtención de información acerca de las nuevas tecnologías parece haber sido crucial en el ritmo de la difusión de las innovaciones. Saxonhouse (1974) atribuye la difusión «super-rápida» de las técnicas de la industria textil japonesa al bajo coste de la adquisición de información por las empresas del sector. A su vez, considera que estos bajos costes se lograron por la actuación de una asociación de

negociantes, el empleo de un suministrador común de equipo (Platt Brothers), un eficiente departamento de ventas y un alto grado de cooperación tecnológica entre las empresas.

Rogers y Shoemaker (1971) sugieren que la ventaja relativa, la compatibilidad, la divisibilidad y la observabilidad estaban positivamente relacionados con la tasa de adopción de innovaciones, mientras que la complejidad estaba negativamente relacionada con la misma.

7.2. Dimensión temporal

La difusión no es instantánea, sino que se lleva a cabo durante un proceso de tiempo más o menos largo. La percepción del consumidor es crucial para determinar la difusión de una innovación en el mercado. En consecuencia, no es importante el tiempo objetivo que una innovación ha estado en el mercado, sino la reacción del usuario hacia ella (Rogers, 1962).

Proceso de adopción. Es el proceso mental que afronta un individuo desde el momento en que tiene conocimiento por primera vez de una innovación hasta que la adopta o compra el producto. El modelo más socorrido para explicar su estructura identifica cinco etapas (Rogers, 1962):

- **Conciencia.** El usuario potencial se entera de la existencia de la innovación, pero no tiene información suficiente acerca de su naturaleza. Puede ser un acto pasivo y no intencionado o el resultado de investigar conscientemente la solución a una necesidad. Los estudios han mostrado que las fuentes impersonales de información, como la publicidad en medios masivos, son muy relevantes.
- **Interés/información.** El usuario se interesa lo suficiente como para obtener más información sobre la innovación: cómo funciona el producto, cuál es su precio y qué ventajas tiene respecto a los productos actuales, entre otras cuestiones. El usuario investigará y buscará información adicional.
- **Evaluación.** El usuario utiliza la información reunida para tomar en consideración si merece la pena o no probar el nuevo producto. Esta etapa también se denomina «evaluación mental».
- **Prueba.** El usuario ensaya o prueba el producto. La prueba en sí misma es una decisión arriesgada, porque requiere un compromiso de tiempo y dinero, que tiende a alargarse si el resultado no indica de manera inequívoca si se debería adquirir la innovación o no. Obviamente, esta etapa es, en muchos aspectos, crítica, ya que aboca en la aceptación o rechazo de la innovación. Un buen ejemplo, es la prueba de un automóvil antes de comprarlo.
- **Adopción.** Es la decisión final de convertirse en usuario regular del producto o utilizar algún sustituto. Los estudios muestran que, a medida que una persona se desplaza desde la evaluación hasta la prueba y la adopción, las fuentes personales de información son más importantes que las impersonales.

Estas etapas son definiciones conceptuales, no hechos empíricos contrastados, que describen el proceso de toma de decisiones de un individuo en la adquisición de un nuevo producto. El modelo es aplicable tanto a los productos de consumo (el usuario es una persona individual) como a los productos industriales (el usuario es una empresa).

Tasa de adopción. Es la velocidad relativa con que una innovación se adopta por los miembros de un sistema social (Rogers, 1962). La tasa de adopción depende de cuántos

compradores avanzan en el proceso de adopción, de cuántos comienzan y con cuánta rapidez toman la decisión de comprar el producto. Esta tasa varía enormemente dependiendo del producto, el sistema tecnológico y el país donde se comercializa (figura 22). Cuanto mayor sea la tasa de adopción, más atractiva será la innovación para la empresa, ya que toma posiciones en el mercado con mayor rapidez a la par que logra la lealtad de los clientes.

Una rápida velocidad del cambio tecnológico normalmente conduce a una aparente lentitud en la difusión de una innovación. La decisión de comprar ahora puede ser, en efecto, una decisión que hará al usuario cargar con una tecnología que pronto será obsoleta. En este sentido, la empresa debe hacer frente a una verdadera paradoja: intentar persuadir al comprador potencial de la estabilidad tecnológica del producto, al tiempo que destina recursos a la mejora del mismo, a fin de no quedar rezagado ante las mejoras que pone en marcha la competencia. Recíprocamente, cuando la velocidad del cambio tecnológico disminuye y el producto se estabiliza, el ritmo de las adopciones se incrementa, debido a la mayor confianza por parte de los compradores potenciales en que el producto no se verá desplazado por uno mejor en un período de tiempo relativamente corto. Este argumento contribuye a explicar el fracaso de las empresas que funcionan al límite de la «mejor tecnología», debido en gran parte a diferencias en las expectativas empresariales sobre el ritmo futuro de mejoras tecnológicas (Rosenberg, 1982).

Las mejoras continuadas en la tecnología actual retrasarán la difusión de la tecnología nueva. Por ejemplo, la rueda hidráulica continuó produciendo mejoras importantes, al menos hasta un siglo después de la introducción de la máquina de vapor de Watt. Asimismo, el barco de madera de vela estuvo vinculado a mejoras importantes e imaginativas durante mucho tiempo después de la introducción del barco de hierro de vapor (Rosenberg, 1982).

	Año de introducción	Nivel de penetración después de 15 años (%)
<i>Electrodomésticos</i>		
Lavadoras	1908	23
Neveras	1911	3
Congeladores	1929	3
Aire acondicionado	1932	1
Secadoras de ropa	1939	7
Hornos microondas	1957	1
<i>Artículos de uso doméstico</i>		
Aspiradoras	1908	33
Licuadoras	1946	9
Abridor eléctrico	1956	46
Cuchillo eléctrico	1963	42
<i>Electrónica de consumo</i>		
Radio	1920	72
TV en blanco y negro	1939	63
TV en color	1954	38

Figura 22. Tasa de penetración en el mercado de diferentes tecnologías (Bayus, 1992)

La duración del proceso de difusión puede explicarse muy bien en términos económicos. Griliches (1957), al estudiar el comportamiento de los granjeros y de los productores de

semillas híbridas, observó que tanto el retraso en el tiempo de la primera introducción de las semillas híbridas en una región particular como el grado en que las nuevas semillas reemplazaban a las antiguas variedades de polinización abierta giraban alrededor de la rentabilidad. En las áreas donde los beneficios eran importantes y evidentes, la transmisión fue muy rápida. En Iowa, donde el maíz híbrido encajaba particularmente bien –y la rentabilidad de la transición fue, por tanto, alta–, los granjeros emplearon cuatro años en cambiar sus terrenos al maíz híbrido de un 10 a un 90 por 100, es decir, las técnicas nuevas se aportan a las áreas buenas antes que a las malas y las técnicas nuevas se aceptan con mayor rapidez en las áreas «buenas». Por lo tanto, es posible que el cambio tecnológico tienda a acentuar las disparidades regionales en cuanto a los niveles de ingreso y las tasas de crecimiento (Griliches, 1960).

Por otra parte, cuanto menores sean el riesgo o la incertidumbre, mayor será la posibilidad de adoptar una innovación. La tasa de reducción de la incertidumbre inicial respecto a las prestaciones de la innovación, a medida que ésta se difunde, afecta al proceso de difusión (Mansfield, 1968a). En un estudio sobre cambio en la productividad de los embarques oceánicos antes de 1850, North (1961) descubre que un carguero superior, el filibote (o buque volador), que se utilizaba en la ruta del Báltico y en el comercio de carbón inglés, no entró en las rutas del Atlántico durante un siglo debido al predominio de la piratería. Las economías resultantes de la modificación del diseño de buques y la consiguiente reducción del número de tripulantes no fueron factibles, mientras los cargueros tuvieran que embarcar armamento para defenderse de los piratas, ya que necesitaban muchos hombres para el manejo de los cañones.

Probablemente el estímulo más importante para la difusión de una tecnología sea la reducción del precio a medida que lo hace el coste unitario. Con ello se consigue una mayor penetración en el mercado, lo que permite incrementar la capacidad y disminuir los costes (Day, 2001). Probablemente el estímulo más importante para la difusión sea el declive de los precios reales con respecto a los sustitutos. Las principales razones de las bajas de los precios reales son: 1) efectos de experiencia, como un resultado conjunto de aprendizaje acumulativo, economías de escala y avances tecnológicos que dan lugar a incrementos en la productividad y baja de coste y 2) un estrechamiento persistente del tamaño del margen entre los precios predominantes y los costes totales promedios debido a fuerzas competitivas. La tasa de declive en el precio relativo también tiene un impacto directo sobre la expansión del potencial de mercado, a causa del creciente número de nuevos usuarios que entran en el mercado y del estímulo para que se produzca una utilización más intensa entre los usuarios actuales. En cierto sentido la tasa de declive es una profecía que se cumple a sí misma. A medida que los precios más bajos expanden el mercado y estimulan las ventas, el incremento más veloz de la experiencia acumulada hace posible que se abaraten los costes, lo que provoca finalmente precios más bajos, y así sucesivamente (Day, 2000).

Los fabricantes de la tecnología asentada pueden utilizar tácticas políticas y de otro tipo para retrasar la difusión. Por ejemplo, la industria del gas minimizó públicamente la importancia del alumbrado eléctrico, exageró sus peligros y utilizó influencias políticas en relación con las normas de seguridad y las tareas municipales para la distribución de energía con el fin de impedir su progreso (Utterback, 1994).

La aceptación de una innovación se verá dificultada si los clientes a los que se dirige no la conocen, no comprenden completamente sus ventajas, no están convencidos de sus méritos

o no pueden encontrarla. Por ello, cuanto más grande sea el gasto colectivo en publicidad, fuerza de ventas, apoyo promocional y cobertura de distribución, mayor será el impacto del valor percibido del producto, lo que, a su vez, acelera el crecimiento del mercado (Day, 2001).

Un factor no muy estudiado es la influencia de los mercados de segunda mano en el ritmo de difusión tecnológica. Cuanto mayor sea el mercado de segunda mano, mayores serán los costes de capital y el coste de oportunidad de mantener la máquina existente y menor la desventaja de la tecnología alternativa (Metcalf, 1992).

La difusión de una tecnología será más rápida si la empresa que la comercializa tiene una cierta reputación en el mercado y una imagen corporativa consolidada. Si la empresa innovadora no tiene una reputación establecida, puede llevar a cabo una alianza estratégica con una que sí la tenga. El compromiso que la empresa adopte respecto a la innovación también afectará a su difusión.

Los directivos senior de una empresa deliberadamente pueden contribuir al retraso de la difusión. Por ejemplo, si a los mejores vendedores de los productos en uso los asigna a vender los nuevos productos, podrían pasar dos o tres años antes de que empiecen a ser productivos. Sin embargo, si los mantiene vendiendo los productos actuales, seguirán produciendo durante todo el tiempo que el directivo esté en activo y compartirán los beneficios (Leonard-Barton, 1995).

Unos incentivos mal diseñados retrasan la difusión. Así, temerosa de que un excesivo empuje de los ordenadores pudiera producir el colapso del negocio tradicional, NCR no introdujo incentivos para promocionar las ventas de aquellos. De hecho, las comisiones sobre las ventas de ordenadores eran al principio desalentadoramente bajas. Cuando salió el ordenador 315 de NCR, por ejemplo, se ofreció al personal de ventas un 0,006%, una comisión potencial de 1.200 sobre 200.000\$ por ordenador. Se podía ganar mucho más vendiendo una sola máquina contable que por entonces tenían mejor salida en el mercado (Leonard-Barton, 1995).

Hacer más abierta una tecnología (por ejemplo, no protegiéndola intensamente o abriendo parcialmente la tecnología mediante licencias) podría acelerar su adopción al permitir que más productores mejoraran y promocionaran la tecnología y que los desarrolladores de bienes complementarios pudieran apoyar la tecnología con mayor facilidad.

La difusión es más rápida cuando la competencia entre los oferentes es intensa. El progreso en la tecnología está en gran parte impulsado por la necesidad competitiva de igualar lo que los competidores ya han alcanzado, mientras se buscan nuevos límites que los rivales no puedan imitar fácilmente. Cuanto más intensa es la rivalidad, más se gastará en I+D y mayor será la urgencia por llevar los resultados al mercado (Day, 2001).

El retraso en la difusión de una innovación en un determinado país puede estar ocasionado por leyes proteccionistas. Tradicionalmente, e incluso en la época actual, los gobiernos han permitido a las tecnologías nacionales sobrevivir en el mercado nacional mediante la aplicación de tarifas protectoras, a pesar de su menor eficiencia en relación con las nuevas tecnologías. Se utiliza el argumento basado en favorecer la protección de la industria infantil. Algunas empresas no pueden competir con las extranjeras por sus costes actuales, aunque creen que podrán hacerlo más adelante con el aprendizaje adquirido en la

producción. Para ello deben protegerse temporalmente de la competencia extranjera hasta que estén listas para competir.

Los gobiernos de muchas naciones imponen a las multinacionales restricciones a la importación de productos y, por el contrario, incentivan acuerdos de licencia o la creación de *joint ventures*. Como resultado de ello, la tecnología puede emigrar a más lugares y hacerlo con mayor rapidez. Asimismo, una medida habitual de algunos gobiernos de países en vías de desarrollo es prohibir la importación de tecnología intensiva en capital debido al desempleo masivo que ocasionan (Robinson, 1988a), lo que retrasa su difusión en el mercado.

7.3. Sistema social (mercado potencial)

Un sistema social se describe como una población de individuos funcionalmente diferenciados que exhiben una conducta colectiva de resolución de problemas (Rogers, 1962). Por consiguiente, un sistema social es un grupo de personas que tienen algo en común, como problemas, actividades, intereses, lugar de crecimiento o de vida, o ambos. En lo referente al proceso de difusión, los miembros de un sistema social tienen, por lo menos, una característica en común: ser compradores potenciales de la innovación. Constituyen, pues, el mercado potencial hacia el que se dirige la innovación.

En relación con el mercado potencial es necesario realizar una serie de consideraciones. Los usos inmediatos de una innovación no son, en modo alguno, evidentes de por sí. A menudo, resulta difícil determinar con precisión qué hay que hacer con un nuevo artefacto. Éste fue el problema que afrontó Thomas Edison después de haber inventado el gramófono en 1877. Al año siguiente publicó un artículo en el que especificaba diez maneras en las que la invención podía resultar útil para el público. Sugería su empleo para: a) tomar dictado sin uso de un estenógrafo, b) disponer de «libros parlantes» para ciegos, c) enseñar a hablar en público, d) reproducir música, e) conservar dichos y recuerdos importantes de la familia, así como las últimas palabras del moribundo, f) crear sonidos nuevos para cajas y juguetes musicales, g) producir relojes capaces de anunciar la hora y un mensaje, h) disponer de la pronunciación exacta de lenguas extranjeras, i) enseñar a escribir y memorizar y j) registrar llamadas telefónicas. Esta lista es importante porque representa el orden de prioridades del mismo Edison en relación a los usos potenciales de la máquina parlante. La reproducción musical se clasifica en cuarto lugar porque Edison pensaba que era un uso trivial de su invento. Una década después, cuando el inventor entró en el negocio del fonógrafo, aún se resistía a los esfuerzos de comercializarlo como instrumento musical y se dedicó a ofrecerlo a la venta como máquina de dictado (Basalla, 1988).

La innovación puede dirigirse al mercado de masas o a algún segmento específico. Dividir un mercado en segmentos permite que la empresa considere las acciones de los competidores y la fuerza de los productos existentes con respecto a cada grupo bien definido de clientes. La empresa debe decidir la forma de llegar al mercado global. Puede tratar de comercializar el producto en la mayoría de los países al mismo tiempo o entrar país a país de forma secuencial. Es importante señalar que el país de la primera innovación no es necesariamente aquél en el que se produce una difusión mas rápida ni las mayores ganancias de productividad.

La composición del sistema social afecta directamente a la difusión de una innovación. En general, en los países con una población homogénea la velocidad de adopción de nuevos

productos es más rápida que en los países con diferencias culturales significativas. Tres características más de la población afectan a la fase de difusión de una innovación en un país (Gatignon *et al.*, 1989): el cosmopolitismo, la movilidad y el porcentaje de mujeres en la fuerza laboral. Los cosmopolitas son personas que ven más allá de su entorno social inmediato, frente a los más orientados hacia su sistema social inmediato. Cuanto más cosmopolita sea la población del país, mayor será la propensión a innovar. La movilidad es la facilidad con la que los miembros de un sistema social pueden trasladarse e interactuar con otros miembros del mismo. Facilita la comunicación interpersonal y, por tanto, afecta positivamente a la penetración del producto en el mercado. Está condicionada principalmente por la infraestructura del país. Finalmente, una mayor participación de mujeres en la fuerza laboral significa mayores ingresos y, en consecuencia, mayor poder de compra. Los productos que ahorran tiempo (lavadoras o lavaplatos, por ejemplo) son más atractivos y logran una difusión más rápida. Por la misma razón, los productos que consumen más tiempo serán peor valorados y su difusión más lenta.

También se ha identificado una relación positiva entre una cultura individualista y la adopción de innovaciones (Lynn y Gelb, 1996). Las culturas individualistas valoran la autonomía y tienden a ser más hedonistas y materialistas que las culturas orientadas al grupo. Así pues, cuanto más individualista sea la cultura de un país, mayor será su capacidad de innovación nacional. Los datos también muestran que hay una asociación negativa entre la capacidad de innovación de un país y la menor inclinación a asumir riesgos o a experimentar por los ciudadanos de un país.

Categorías de adoptantes. Las categorías de adoptantes son una clasificación de los individuos de un mercado basado en su comportamiento ante una innovación. Hay un grupo de factores que afectan a la adopción y que se relacionan con las características básicas del individuo y son independientes de la propia innovación. Este grupo de factores se compone fundamentalmente de cinco variables: espíritu aventurero, integración social, cosmopolitismo, movilidad social y posición económica (Schewing, 1974). El conjunto de estos factores suele ser, al menos en las sociedades desarrolladas, el reflejo de un talante abierto y están relacionados o incluso integrados bajo el concepto de espíritu innovador, que se define como la disposición de un individuo a adoptar una innovación o el grado con que se anticipa a otros miembros de su sistema social en adoptar nuevas ideas (Rogers, 1962). Una forma de medir el espíritu innovador de un consumidor es a través de la duración de su período personal de adopción, o tiempo que transcurre desde el momento en que se recibe la primera noticia sobre una innovación hasta su adopción. Un consumidor con talante abierto, aventurero, activo y flexible es muy probable que figure entre los primeros compradores de muchas innovaciones.

Los individuos pueden clasificarse en varias categorías de acuerdo con el tiempo que tardan en adquirir la innovación. El modelo de Rogers (1962) asume que la difusión adopta una distribución normal y distingue cinco categorías de adoptantes (figura 23):

- Innovadores: constituyen el primer 2,5 por 100 de los adoptantes.
- Primeros adoptantes: representan el próximo 13,5 por 100.
- Primera mayoría: el próximo 34 por 100.
- Mayoría tardía: el siguiente 34 por 100 de los adoptantes.
- Tradicionalistas o rezagados: forman el 16 por 100 final.

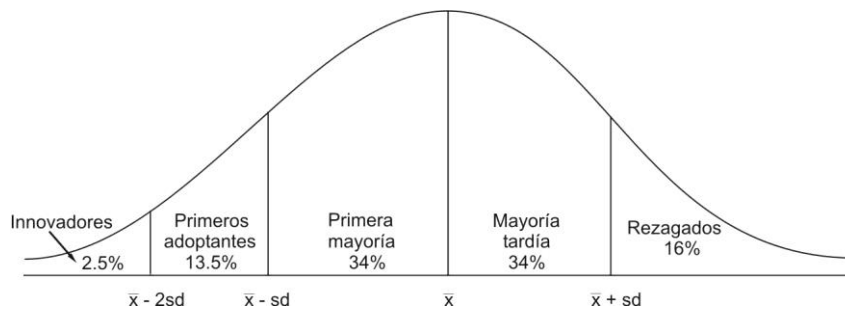


Figura 23: Curva de adopción (Rogers, 1962)

Existen algunas evidencias que sugieren que cada uno de los cinco grupos tiene diferentes escalas de valores. Los innovadores son proclives a la aventura y están dispuestos a asumir riesgos tanto de naturaleza económica como social en la compra de nuevos productos (Rogers, 1962). El riesgo económico está relacionado con la cantidad real de satisfacción derivada del uso o consumo de la innovación (el innovador pierde dinero si el rendimiento del producto no se corresponde con sus expectativas). El riesgo social es mucho más importante para el innovador, ya que perderá prestigio si su grupo de referencia rechaza una innovación que él ya ha adoptado. La influencia del grupo de referencia es especialmente fuerte en relación con la elección de producto y marca en categorías, tales como automóviles, cigarrillos y cerveza (Schewing, 1974). Por otra parte, Donnelly y Etzel (1973) plantean serias dudas sobre la posibilidad real de describir, utilizando parámetros demográficos y de comportamiento, a los innovadores como grupo separado. Señalan que las personas pueden ser innovadoras en función de las semejanzas o desigualdades que presentan los nuevos productos en relación con los utilizados con anterioridad.

Los primeros adoptantes suelen ser personas jóvenes (aunque de más edad que los innovadores), disfrutan de gran prestigio social, tienen una posición financiera favorable y una capacidad mental abierta. Asimismo, están en contacto cercano con el origen de las nuevas ideas, utilizan fuentes de información impersonales, sus relaciones sociales son más cosmopolitas y están considerados como líderes de opinión (Rogers, 1962), por lo que representan un elemento clave en el proceso de difusión. Si la característica distintiva de los innovadores es el espíritu aventurero, la de los primeros adoptantes es el respeto.

La estrategia de marketing para la introducción de una tecnología en el mercado (básicamente, una tecnología de producto) debería centrarse en los primeros adoptantes. Sin embargo, Kotler y Zaltman (1976) señalan que es posible que no constituyan la mejor alternativa para convertirse en el centro de la estrategia, ya que la propensión a la adopción temprana puede no ir acompañada por una fuerte propensión de compra o por una alta disposición a influir sobre otras personas.

La primera mayoría es divulgadora y acoge las nuevas ideas antes que la persona de tipo medio, aunque rara vez son líderes de opinión. Se deciden a adoptar la innovación una vez que sus iguales lo han hecho. Su posición social y sus ingresos también son superiores a la media. El liderazgo de opinión de la primera mayoría se limita a los miembros de su propio grupo y de la mayoría tardía.

La mayoría tardía es escéptica y adopta una innovación después de que gran parte de las personas la hayan probado. Su posición social es baja, por lo que esperan a que el precio de un nuevo producto se reduzca hasta que les resulte asequible su adquisición. La influencia

de opinión se restringe a los miembros de su propio grupo. Poseen poco espíritu aventurero, o mejor expresado, no pueden arriesgarse. Confían poco en los medios de difusión y, en su lugar, recurren a personas de la primera mayoría y mayoría tardía como fuentes de información. Debido a una falta de espíritu abierto y flexible son muy lentas en tomar decisiones.

Finalmente, los rezagados son personas que se aferran a la tradición, es decir, sospechan de los cambios y sólo adoptan la innovación cuando se ha convertido en una tradición en sí misma (Rogers, 1962). Los rezagados suelen encontrarse aislados de su entorno. Principalmente son personas de edad avanzada con un ingreso (a veces, la paga de jubilación) inferior a la media y suelen tener una baja posición social. Su influencia sobre las opiniones de otros es mínima, porque sólo personas en una situación similar les escucharán. Ellos, a su vez, reciben información de vecinos, parientes y amigos. Su resistencia contra todo lo nuevo es difícil de vencer.

Cruzar el abismo. El paso que debe dar el pionero para competir en el mercado de masas, Moore (2002) lo denominó «cruzar el abismo», y conlleva la creación de un nuevo modelo de negocio, por diferentes razones. Primero, mientras que los primeros adoptantes de un nuevo producto son un nicho caracterizado por ser jóvenes visionarios, proclives a la aventura, de mentalidad abierta, dispuestos a asumir riesgos, con una posición financiera sólida y gran prestigio social, que valoran la singularidad del producto (Rogers, 1962), los clientes de un mercado mayoritario son más conservadores, poseen menos recursos y valoran por encima de lo demás la calidad y el precio del producto (Moore, 2002). Segundo, el acceso a los primeros adoptantes se realiza mediante canales especializados de distribución, teniendo singular importancia la promoción personalizada y, en concreto, el boca a oreja (Frenzen & Nakamoto, 1986); sin embargo, llegar a un mercado mayoritario requiere canales de distribución de alto alcance para acaparar el mayor número posible de puntos de venta con la más alta exposición del producto, así como diseñar campañas publicitarias en medios de comunicación masivos (Moore, 2002). Finalmente, los primeros adoptantes son relativamente pocos en número, por lo que el pionero utiliza un sistema de producción en pequeños lotes, mientras que, para atender al mercado mayoritario, se requiere un proceso de producción en masa (Abernathy & Utterback, 1978). La organización tiende a burocratizarse, se especifican las tareas, se establecen procedimientos muy formalizados y se llevan a cabo controles rígidos y un reparto de funciones a través de una línea de autoridad jerárquica; en el mercado predominan las grandes empresas.

De hecho, la investigación demuestra que las capacidades y formas de pensar necesarias para la creación de un producto y su introducción a pequeña escala, como suele ser en el caso de los pioneros, no solo son distintas de las que se necesitan para su comercialización a gran escala, como requeriría el mercado mayoritario (Macher & Richman, 2004), sino que además entran en conflicto con ellas (Robinson *et al.*, 1992; Markides & Geroski, 2005). Las empresas que poseen una buena capacidad inventiva no suelen ser buenas en la comercialización a gran escala, y viceversa (Moore, 2002).

En resumen, el modelo de negocio y las estrategias requeridas para competir en un mercado embrionario repleto de innovadores y primeros adoptantes son muy diferentes de los necesarios para competir en un mercado de masas de alto crecimiento formado por la primera mayoría. Como consecuencia, la transición entre el mercado embrionario y el mercado de masas no es tan suave ni fácil. En cambio, representa un abismo

competitivo que las empresas deben cruzar. De acuerdo con Moore, muchas empresas no desarrollan o no pueden desarrollar el modelo de negocio correcto; caen en el abismo y salen del negocio. Por tanto, aunque los mercados embrionarios suelen estar repletos de un gran número de pequeñas empresas, una vez que el mercado de masas comienza a desarrollarse, el número de empresas disminuye en gran medida.

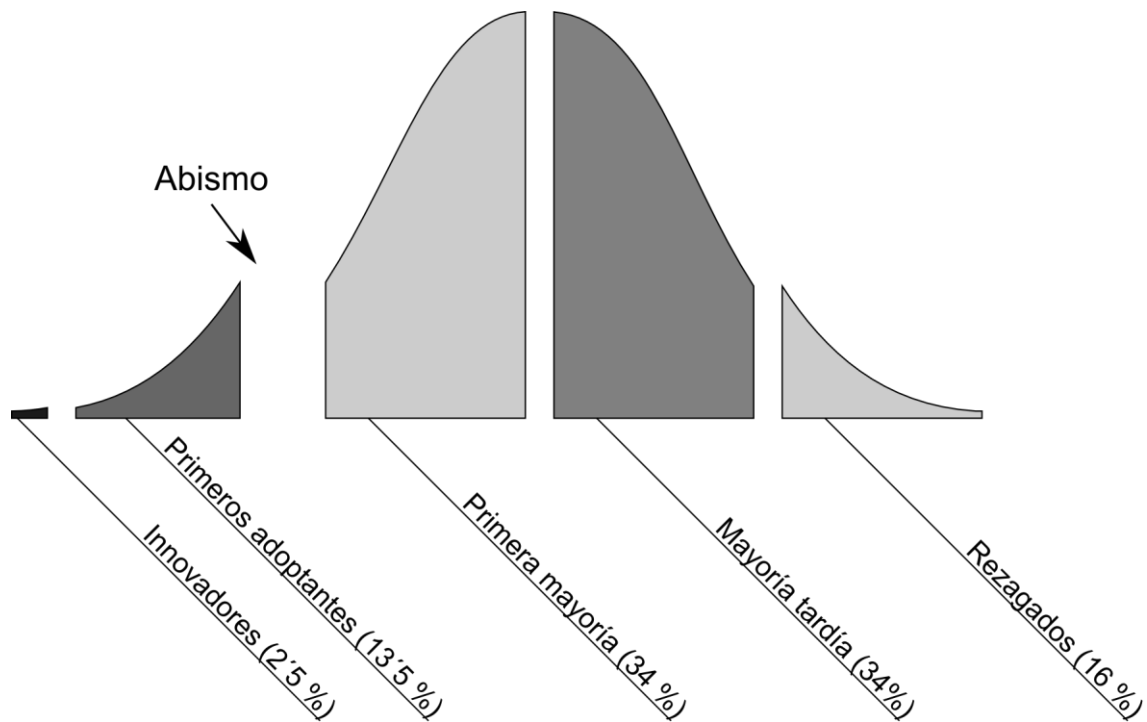


Figura 24: El abismo (Moore, 2002)

La figura 24 muestra que existe un abismo entre los primeros adoptantes y primera mayoría, es decir, entre el mercado embrionario y el mercado de masas en rápido crecimiento. Para cruzar con éxito el abismo, Moore sugirió que los gerentes deben identificar de manera correcta las necesidades de los clientes de la primera mayoría, es decir, la vanguardia del mercado masivo. Después deben modificar sus modelos de negocio para desarrollar nuevas estrategias y rediseñar los productos, y crear canales de distribución y realizar campañas de marketing para satisfacer las necesidades de la primera mayoría. Deben contar con un producto idóneo a un precio razonable que puedan vender a la primera mayoría cuando comiencen a ingresar al mercado en grandes cantidades. Además, los pioneros de la industria deben abandonar su antiguo modelo de negocio enfocado, dirigido tan solo a las necesidades de los innovadores y primeros adoptantes, pues este enfoque ocasionará que ignoren las necesidades de la primera mayoría y el desarrollo de las estrategias necesarias para implementar la diferenciación o modelo de negocio basado en liderazgo en costes y continuar siendo un competidor dominante en la industria.

8. IMITACIÓN DE LA INNOVACIÓN

La imitación es la difusión de una tecnología entre los competidores, lo que contribuye a disipar sus ventajas competitivas. Los competidores están siempre vigilando el mercado y cuando perciben una alta rentabilidad potencial en un nuevo producto (o proceso productivo), comienzan a invertir con fuerza en esta dirección, produciéndose lo que, en terminología *schumpeteriana*, se denomina una «ebullición». Mientras que la pura fogosidad es la que, a menudo, sostiene al innovador de una nueva tecnología, el enjambre de imitadores suele verse empujado por el ejemplo de alguna o algunas operaciones excepcionalmente rentables. Esta ebullición, consecuencia de la imitación creativa del producto innovador, se produce, no al introducir el producto en el mercado, sino cuando su potencial de demanda comienza a ser alto. A veces, incluso hay que esperar algunos años.

De acuerdo con Mansfield *et al.* (1977), cuanto más grande es una empresa, más tiempo tarda en imitar una innovación. De la misma manera, cuantas más personas estén implicadas en el proceso, más tiempo se necesita. También observaron que cuanto más tarde empiece una empresa su propio proceso de imitación, más rápidamente se desplazará a través de las etapas del proceso. Por tanto, las empresas que primero imitan una innovación tienden a ser más reflexivas. Aquéllas que lo hacen después tienden a tomar la decisión con mayor rapidez y tardan menos en sustituir la nueva tecnología por la antigua. De acuerdo con Moch y Morse (1977), las empresas imitan con mayor probabilidad las innovaciones si cuentan con profesionales que las consideren compatibles con sus necesidades e intereses.

Mansfield (1961) encuentra que la velocidad de imitación está correlacionada positivamente con la rentabilidad de la innovación y negativamente con la magnitud de la inversión requerida. Para la innovación en productos esto implica que la velocidad de imitación depende de la política de precios del innovador (Scherer, 1980). Unos precios elevados es una invitación a los competidores para que imiten el producto. Si la imitación es costosa las empresas no la llevan a cabo. Estos costes surgen del carácter irrecuperable de gran parte de las inversiones necesarias para acumular los recursos estratégicos.

La patente puede desvelar demasiada información, ya que pone a disposición de los competidores los conocimientos en los que se apoya la invención, por lo que pueden usarlos para desarrollar variantes de la tecnología básica. El conocimiento de la patente induce a los competidores a buscar procesos y productos sustitutivos, es decir, a tratar de inventar en torno a la patente. De esta forma, los competidores desarrollan invenciones que no violan la patente, pero que realizan las mismas funciones y tienen éxito en apropiarse de algunos de los posibles beneficios de la invención original (Winter, 2001). El objetivo de la investigación es encontrar una forma de evadir la patente y participar en las rentas del dueño de esta (Stiglitz y Greenwald, 2014).

Por otra parte, a veces, una idea desarrollada en una patente puede incitar nuevas ideas, aun si las nuevas ideas no utilizan la antigua idea o no se basan directamente en ella. En ese sentido, actúan como un catalizador: un factor químico que facilita una reacción, pero que no se utiliza en el proceso (Stiglitz y Greenwald, 2014).

Las máquinas y productos son portadores de conocimientos, en cuanto que su configuración determina cómo desempeñar ciertas tareas (Badaracco, 1991). En muchos casos los competidores pueden adquirir este conocimiento mediante la ingeniería inversa: descomponer el producto en partes, con objeto de estudiar sus propiedades

individuales y la estructura de las relaciones que las sustentan. Posteriormente, se mejoran algunas de estas características, e incluso la propia estructura del producto, dando lugar a un producto diferente al original (de ahí lo de imitación o copia creativa): copia porque la imitación realiza idénticas funciones y satisface las mismas necesidades que el original, y creativa en el sentido de que es diferente con objeto de no vulnerar los derechos de propiedad industrial que protegen al innovador. El *software* y los diseños industriales –sobre todo, los del campo de la moda– están continuamente sometidos a la ingeniería inversa.

Los proveedores de equipo ayudan a transferir información clave y conocimientos tecnológicos. Una empresa, al encargar una máquina específica a un proveedor para fabricar un componente para un nuevo producto, le está revelando información. Esta información es fácil que se transfiera a otros clientes del fabricante de máquinas, competidores de la empresa.

Los derrames tecnológicos se producen cuando los beneficios de las actividades de investigación de una empresa se filtran hacia otras empresas. Se trata de una externalidad positiva de los esfuerzos en I+D. Los observadores de la industria, consultores, universidades, asociaciones mercantiles o sociedades profesionales ayudan a transferir la tecnología y otros conocimientos comerciales a las empresas del mercado, favoreciendo, de este modo, la imitación de la tecnología. Divulgaciones no autorizadas de información secretas son comunes y no punibles legalmente. Por otra parte, el espionaje industrial es una práctica bastante recurrente.

Las empresas sólo tardan entre 12 y 18 meses en recabar información de las decisiones que toman sus competidores acerca de los nuevos productos que están desarrollando (Hill, 1997). En la economía global los competidores a veces pueden imitar rápidamente las acciones del innovador. Antes las empresas se protegían con patentes, las cuales eran fuente de su ventaja competitiva, pero ahora la velocidad de difusión de la tecnología ha anulado esa protección. Es frecuente que muchas empresas que compiten en la industria electrónica no soliciten patentes para evitar que sus competidores tengan acceso al conocimiento tecnológico que se debe incluir al solicitarlas. La mayoría de los productos y procesos son copiados antes de los tres años (Levin *et al.*, 1987).

En ciertas industrias los clientes insisten en contar con fuentes secundarias para los componentes que utilizan en sus sistemas. Estimular a otras empresas para que adopten la tecnología propia asegura a tales clientes fuentes secundarias y puede aumentar la cantidad de clientes que están dispuestos a adoptar la innovación. De esta forma, en algunos casos, el innovador es quien revela los conocimientos tecnológicos a terceras partes.

El conocimiento puede extenderse vía contactos. La localización de los contactos y la capacidad de comunicarse son claves. Uno de los beneficios de la globalización es que ha ampliado la exposición de los individuos a nuevas ideas. Internet ha extendido la capacidad de los individuos para acceder a los conocimientos y a nuevos contactos.

Los conocimientos que poseen los individuos se trasladan con ellos cuando cambian de trabajo o de país. A este respecto, en los comienzos de la Revolución Industrial, las mejores copias (imitaciones) de las máquinas textiles británicas fueron obra de mecánicos ingleses inmigrantes en la Europa continental. La difusión de la tecnología se realizó

principalmente persuadiendo a trabajadores expertos a emigrar hacia regiones en las que su pericia no era todavía excesiva. Quizá la contribución más grande de estos inmigrantes no fue lo que hicieron, sino lo que enseñaron (Rosenberg, 1982). Para corregir esta anomalía el Reino Unido promulgó unas leyes que castigaban a los infractores con multas y penas de prisión (Landes, 1969). La fuga de los trabajadores, además de drenar el capital humano de la organización, provee a los competidores de trabajadores ya formados, portadores de conocimientos y/o secretos de relevancia industrial, lo que les permite potenciar sus bases cognitivas.

Ahora bien, para imitar los productos de sus rivales una empresa necesita poseer un cierto nivel de conocimientos afines, también denominados capacidad de absorción, que aglutina las habilidades colectivas para reconocer el valor de nuevo conocimiento, asimilarlo y aplicarlo con fines comerciales (Cohen y Levinthal, 1990). Así pues, la empresa necesita conocimiento previo relacionado para asimilar y utilizar nuevo conocimiento. Este conocimiento previo incluye habilidades básicas y, en ocasiones, un lenguaje compartido, aunque también abarca los más recientes desarrollos científicos y tecnológicos en un campo dado. Se apoya en la capacidad de absorción de los empleados, aunque es superior a su suma. Depende, igualmente, no sólo de la relación de la empresa con el entorno, sino de la transferencia de conocimiento entre unidades organizativas. Woiceshyn y Daellenbach (2005) encontraron que las diferencias en la capacidad de absorción de las empresas se deriva de: a) el compromiso de inversión en proyectos de I+D; b) el aprendizaje de los socios y de las investigaciones propias; c) la selección y aplicación correcta de nuevas tecnologías, el desarrollo y uso de tecnologías complementarias y d) como se comparte la información entre los equipos multifuncionales.

9. APROPIACIÓN DE LOS BENEFICIOS DE LA INNOVACIÓN

La apropiabilidad describe hasta qué punto un innovador puede capturar el valor que se deriva de la innovación en forma de beneficios. La empresa debe evitar que otras empresas imiten la innovación para mantener los beneficios. Como residente o como nuevo participante, hay ciertas cosas que un innovador puede hacer para prolongar el tiempo durante el cual continúe obteniendo beneficios de monopolio. A continuación, comentamos algunos instrumentos y actuaciones que puede utilizar la empresa para proteger sus innovaciones.

Patente. La patente es un título de propiedad otorgado por el Estado, que concede a su poseedor un monopolio temporal, al permitirle excluir a terceros del empleo comercial de una invención tecnológica, sea ésta de producto o de proceso. Como contrapartida, la patente se pone a disposición del público para general conocimiento, lo que en ningún caso conlleva su usurpación. En España, la duración de la patente es de veinte años a contar desde la fecha de la presentación de la solicitud. Para mantenerla en vigor es preciso pagar unas tasas anuales a partir de la concesión.

Transcurrido el tiempo de protección, la invención que protege la patente entra en el dominio público, lo que propicia a su poseedor un declive precipitado de sus ingresos. Por ejemplo, el *Zantac*, un medicamento contra las úlceras del Glaxo Wellcome, se convirtió en el fármaco más recetado del mundo y sus ventas alcanzaron los 3.800 millones de dólares en 1994. Cuando las patentes norteamericanas sobre el fármaco expiraron a finales de julio de 1997, los competidores genéricos se hicieron con más de la

mitad de las prescripciones en un trimestre. De 1996 a 1998 las ventas de Glaxo Wellcome del medicamento descendieron de 3.010 millones a 1.260 millones de dólares (Winter, 2001).

Las patentes tienen una efectividad limitada en algunas industrias, aunque es muy alta en la industria farmacéutica y media en las de alta tecnología, como las de ordenadores y semiconductores (Levin et *al.*, 1987).

Secreto industrial. El secreto industrial es un conocimiento no patentado sobre un producto o proceso de producción que la empresa no revela a otras. Los más difíciles de identificar son los conocimientos utilizados en el proceso productivo, incluyendo las máquinas construidas por la empresa para uso propio. Sin embargo, los productos pueden imitarse utilizando ingeniería inversa. No obstante, hay que tener en cuenta que los tecnólogos forman parte de una comunidad profesional, que mantiene redes de contacto, lo que facilita la comunicación personal y el derrame de conocimientos, lo cual hace vulnerables los secretos industriales.

Activos complementarios. Son activos necesarios para que el bien principal pueda cumplir su función. Controlar los activos complementarios puede ayudar a recoger las ganancias de la innovación. Tales activos incluyen recursos como los canales de distribución, la imagen corporativa o la capacidad de producción, entre otros (Teece, 1986b). El innovador que carece de algunos de los activos complementarios críticos puede tener grandes dificultades para capturar las ganancias de la innovación. Cuando es difícil que las demás empresas adquieran o imiten los activos complementarios, es mucho más probable que la empresa innovadora acapare los beneficios de su innovación.

Durante la década de 1960 y principios de 1970 Harley-Davidson quiso competir con Honda en el mercado de las motos de pequeña cilindrada. Por tanto, fabricó una línea de motos pequeñas de 150 cc. y 300 cc. Harley intentó vender estas motos a través de su red de distribuidores, pero ésta se opuso, por lo que la empresa fracasó en lograr una cuota de mercado aceptable. Los márgenes de rentabilidad eran mucho mayores con las motos de gran cilindrada y temieron que las motos pequeñas afectaran a la imagen de Harley-Davidson con sus clientes principales (Christensen, 1997).

Algunas nuevas tecnologías pueden volver obsoletos los activos complementarios existentes necesarios para el éxito. Por ejemplo, cuando las calculadoras eran electromecánicas, un activo complementario importante era una red de ventas y servicio grande y bien capacitada, que requería más de 1.500 individuos dedicados a esta actividad. Sin embargo, cuando se produjo el cambio hacia las calculadoras electrónicas más fiables, esta red de servicios ya no era necesaria. De hecho, los distribuidores de equipamiento de oficina se convirtieron en un sustituto viable de la organización de servicio y de ventas (Tripsas, 2001).

Los activos complementarios pueden conseguirse mediante alianzas con otras empresas. Las empresas de biotecnología, por ejemplo, se alían con empresas farmacéuticas para la realización de ensayos clínicos, la fabricación y la comercialización. Cuando una innovación y los activos complementarios que la sustentan son suministrados por diferentes empresas, la distribución del valor entre ellos depende del poder relativo de cada una de ellas.

Efecto experiencia. El efecto experiencia afirma que a medida que una empresa duplica el número acumulado de productos fabricados, su coste en términos reales disminuye a una tasa o porcentaje constante. El efecto experiencia se apoya en el principio de que una actividad o tarea siempre puede realizarse mejor y de forma más eficiente a medida que se llevan a cabo repeticiones de la misma (Boston Consulting Group, 1972). Arrow (1962b) observaba que, en el proceso de producir e invertir, se aprende. A medida que producimos e invertimos, mejoramos en lo que hacemos. Aprendemos haciendo. Aprendemos a producir de manera más eficiente, produciendo y, a medida que producimos, observamos cómo podemos hacerlo de forma más eficiente.

El efecto experiencia funciona eficazmente en el inicio del ciclo de vida del producto: la ventaja del pionero es obvia, ya que el liderazgo en experiencia se consigue con mayor facilidad al comienzo, cuando la producción se dobla rápidamente –es decir, la experiencia aumenta diez veces cuando se pasa de la unidad 2 a la 2.000, pero sólo se dobla una vez cuando se pasa de la unidad 2.000 a la 4.000–. La empresa que quiera desarrollarse tendrá que acumular experiencia, de modo que una clave reside en llegar al mercado antes que los competidores, para aventajarlos en costes y obtener, así, una ventaja competitiva en los mercados.

Tiempo de liderazgo (*lead time*). El innovador puede disfrutar de una ventaja natural en la aplicación práctica de una nueva invención, incluso sin patente, pues hace falta tiempo para que los competidores examinen el producto, descubran la tecnología utilizada, construyan las instalaciones y monten el equipo necesario para producirlo ellos mismos y pongan su producto competitivo en el mercado. Mientras tanto, el primer usuario puede haber obtenido suficientes beneficios monopólicos para que le compensen por sus esfuerzos (Machlup, 1968). La simple percepción de que existe un producto nuevo y atractivo y el mero examen de sus especificaciones son a menudo bases inadecuadas para una imitación con éxito. La extensión del tiempo de liderazgo está determinada por la buena suerte, la flexibilidad y las capacidades de la empresa pionera, así como por la mala suerte, la inercia o la incompetencia de las otras. Las investigaciones tienden a calificar el «tiempo de liderazgo» que otorga el mover primero en el mercado como el mecanismo de apropiación más efectivo para la innovación de productos (Levin *et al.*, 1987).

Un riesgo importante está asociado con el tiempo de introducción de la innovación en el mercado. Una innovación sólo tiene valor, en algunos casos, si se introduce en un momento específico o en una secuencia en particular dentro de las operaciones de la organización. Mientras que el innovador *shumpeteriano* experimenta beneficios anormalmente altos hasta que los imitadores lo alcanzan, el innovador impetuoso puede quedarse sin nada, como consecuencia de haber invertido en el modelo prematuro de un invento. Son numerosos los ejemplos de innovaciones que fracasaron por comercializarlas antes que los clientes potenciales estuvieran dispuestos a aceptarlas.

Carrera de patentes. En la estrategia de carrera de patentes, la empresa se mantiene haciendo innovaciones continuamente, a menudo volviendo obsoletas algunas de sus propias tecnologías y canibalizando los productos actuales. La estrategia de carrera de patentes admite que, con frecuencia, las obstrucciones a la entrada son penetrables o caerán con el tiempo, prescindiendo de cuán formidables pudieran parecer.

En la carrera de patentes, la primera empresa en patentar consigue un gran premio y los perdedores no obtienen nada o, al menos, muy poco. En una carrera de patentes, la estrategia óptima de la empresa consiste en adelantarse lo suficiente a los rivales como para que estos no se animen a entrar, y habiendo hecho esto, poder dormirse en sus laureles. Ahora bien, a menudo existen múltiples dimensiones en las características del producto. Un innovador puede desarrollar un producto que es mejor en una dimensión y otro, uno que sea mejor en otra. Aquellos que fracasan en todas las dimensiones son los perdedores, y la competencia del mercado los saca del negocio (Stiglitz y Greenwald, 2014).

Complejidad tecnológica. La proliferación de patentes puede proteger una ventaja basada en la tecnología propia, al limitar las oportunidades técnicas a los competidores. Las grandes empresas pueden invertir para desarrollar extensas carteras de patentes (Mansfield, 1986). Pueden usarlas para crear una maraña de patentes: densas redes de reivindicaciones de patentes que les proporcionan una base plausible para amenazar con demandas en un amplio espectro de propiedad industrial. Pueden hacerlo para impedir que otros introduzcan una innovación superior o para exigir licencias de componentes más débiles en condiciones favorables (Shapiro, 2001).

Los fabricantes pueden hacer que sus productos sean resistentes a la ingeniería inversa. Por ejemplo, las empresas tratan de proteger sus productos de ella mediante lo que se conoce como «empaquetado», que consiste en embalar la innovación, de tal manera que sea muy difícil abrir la carcasa sin destruir su contenido (Rogers y Larsen, 1984). Así, los fabricantes de semiconductores han previsto proteger sus dispositivos de la ingeniería inversa del circuito, encerrándolo en resinas de epoxia, de tal manera que, al eliminar el recubrimiento, se destruye el circuito (Shapley, 1978). Otras empresas retienen el control físico del producto innovador y venden sus servicios, como en el caso de los servicios de reparación de equipamientos realizados por los fabricantes o servicios de equipamientos de apoyo de negocios, como, por ejemplo, el procesamiento de datos (Winter, 2001).

Dependencia de la trayectoria. Dado que las tecnologías se desarrollan lentamente, no se pueden cambiar en poco tiempo: siguen estando ahí. Reflejan, pues, la historia de la empresa, limitando la gama de estrategias que ésta puede plantearse de forma realista (Schoemaker, 1992) y provocan un efecto de *lock-in* (o efecto amarre). La experiencia previa dicta futuras posibilidades de aprendizaje, por lo que este proceso es un fenómeno acumulativo, de manera que las posibilidades del futuro inmediato dependen, en gran medida, de las tecnologías desarrolladas hasta el momento presente.

La dependencia de trayectoria está condicionada por la «interconexión de activos» (Dierickx y Cool, 1989), que significa que el ritmo de acumulación de una tecnología principal puede estar influenciado no sólo por el *stock* existente de esa tecnología, sino también por el nivel de *stocks* de otras tecnologías que se utilizan con la principal. El desarrollo histórico de la tecnología aboca en las denominadas «eficiencias de masa» (Dierickx y Cool, 1989), que suponen que resulta más fácil acumular tecnología cuanto mayor sea el nivel del *stock* tecnológico de partida. Las eficiencias de masa afectan directamente a la «capacidad combinatoria» (Kogut y Zander, 1992), que hace referencia a la posibilidad de generar nuevo conocimiento tecnológico mediante la combinación de conocimientos previos –presentes en las tecnologías actuales–.

Deseconomías de compresión del tiempo. Estas deseconomías se apoyan en la ley de los rendimientos decrecientes, y postulan que, al mantener fijo un *input*, la productividad media decrece a medida que se añaden unidades adicionales del resto de factores productivos (Dierickx y Cool, 1989), es decir, mantener fijo una ratio de inversión en tecnología para un intervalo dado produce un mayor incremento en el *stock* de la tecnología que mantener dos veces la ratio de inversión durante la mitad del período. Por tanto, las empresas que tomaron la delantera son difíciles de seguir a menos que los seguidores realicen inversiones extraordinarias en período de tiempo breves a costa de ser ineficientes.

Ambigüedad causal. La ambigüedad causal se refiere a la incertidumbre asociada a las conexiones causadas entre acciones y resultados (Lippman y Rumelt, 1982). Tal ambigüedad no sólo impide la imitación externa, sino que también explica por qué algunas empresas encuentran grandes dificultades para identificar las causas de sus elevados beneficios o replicar sus plantas de producción más eficientes (Szulanski, 1996). Tres características de la tecnología pueden generar ambigüedad causal (Reed y Defillippi, 1990): el conocimiento tácito, la complejidad social y la especificidad.

El conocimiento tácito se revela mediante la aplicación y se adquiere a través de la práctica. Por lo tanto, los posibles imitadores tienen que hacer frente a tres problemas (Afuah, 1998). En primer lugar, resulta difícil saber qué desea imitar. En segundo lugar, aun cuando una empresa sepa con exactitud qué desea imitar, es posible que no sepa cómo hacerlo, puesto que el conocimiento se aprende acumulativamente a través de los años y está incorporado en los individuos o en las rutinas de la empresa. En tercer lugar, puesto que este proceso de imitación lleva tiempo, los imitadores pueden encontrarse siempre retrasados, ya que mientras imitan a los innovadores, éstos desarrollan niveles superiores de tecnologías.

La complejidad social resulta de tener un amplio número de habilidades y recursos interdependientes trabajando juntos (Reed y Defillippi, 1990). Esta complejidad proporciona una gran ambigüedad causal en el caso de los equipos, las rutinas y la cultura, ya que todos ellos comparten un conocimiento colectivo, que, en gran parte, es de naturaleza tácita. Cuanto más compleja es una innovación, más difícil es de entender y probablemente más difícil de imitar.

Una tecnología es específica o idiosincrásica cuando posee un valor en su uso superior al que poseería en un uso alternativo o bajo cualquier otro usuario. En este sentido, Ghemawat y del Sol (1998) abogan por el uso de tecnologías específicas para la empresa, ya que suelen ser «adhesivas» en el sentido de que existen considerables costes para separarlas de la empresa que las posee. La decisión de invertir en ellas supone, pues, un compromiso irreversible: al ser menos fáciles de revocar, hay que considerarlas más estratégicas, por lo que crean una necesidad de mirar más allá del presente, hacia el futuro.

Disuasión económica. La disuasión supone enviar señales amenazadoras creíbles a los competidores, que les haga pensar que una estrategia de imitación no les será beneficiosa. Debido a que llevar a cabo una amenaza normalmente cuesta dinero, se necesita que esté apoyada en el compromiso de su ejecución, lo que será creíble si la empresa tiene una reputación en el mercado avalada por su historia (Weigelt y Camerer,

1988). Así pues, las represalias deben estar respaldadas tanto por la capacidad como por la disposición de la empresa para tomarlas. Por eso, la disuasión es más eficaz si se utilizan diferentes «colchones», como la liquidez, la capacidad anticipada, pequeñas posiciones en los mercados de los competidores para perturbarlos, lucha de marcas o mejora continua (Ghemawat, 1999). Si la tasa de depreciación de la tecnología es elevada, la credibilidad de la amenaza de una respuesta hostil disminuye, porque la tecnología va a perder valor con rapidez (Eaton y Lipsey, 1980).

Colaboración. Las empresas pueden establecer acuerdos con sus competidores, complementadores, proveedores y clientes para lograr que su tecnología se convierta en el estándar del sector, beneficiarse del efecto red o sencillamente evitar que terceros la copien, sobre todo en aquellos países cuyas leyes de propiedad industrial son débiles. Por ejemplo, es frecuente que las empresas otorguen licencias para seguir manteniendo un cierto control sobre la tecnología. También pueden intercambiar licencias o compartir patentes dentro de un acuerdo de colaboración más integrado.

Esfuerzos en marketing y servicios. El innovador utiliza la publicidad para influir en el consumidor en cuanto a las características del producto y conseguir una lealtad a la marca. Las compañías farmacéuticas gastan más en publicidad y marketing que en investigación. Estos gastos en marketing están diseñados para reducir la elasticidad de la demanda, lo cual permite al titular de la patente elevar los precios y aumentar las ganancias monopólicas (Stiglitz y Greenwald, 2014).

El innovador negocia en exclusiva con los canales de distribución, eligiendo los más eficaces, llegando incluso a imponer sus criterios en lo referente a la colocación de los productos en la estantería. La proliferación de variedades de productos por parte de la empresa líder del mercado puede dejar pocas oportunidades para que los nuevos entrantes y los pequeños competidores establezcan un nicho de mercado (Schmalensee, 1978).

Retención de trabajadores. La movilidad de los trabajadores es uno de los mecanismos por los que los competidores tienen acceso a los conocimientos de la empresa (Arrow, 1996). Para evitar esta difusión de la tecnología la empresa debe desplegar políticas de recursos humanos que faciliten la retención del personal. Igualmente puede parcelar los puestos de trabajo para impedir que un trabajador tenga acceso al conocimiento completo de la tecnología en uso y haciendo que el conocimiento limitado de su parcela no tenga valor para los competidores. La inclusión en el contrato de cláusulas de no competir en el supuesto de que el trabajador abandone la empresa, contribuyen a evitar la migración de tecnología a terceros.

Adecuada redacción de los contratos. Se trata de incluir cláusulas que protejan los conocimientos de la empresa. Los secretos de fabricación se protegen mediante acuerdos de confidencialidad o de no publicación, en los que suele constar el período de vigencia del convenio y las condiciones bajo las cuales debe mantenerse el secreto (Brooking, 1996). El acuerdo de confidencialidad pretende restringir el uso público de determinada información garantizando que sea accesible sólo para aquellos autorizados a tener acceso. Crea una relación comercial entre los participantes para proteger cualquier tipo de información. Mediante la exclusividad el empleado se compromete a no prestar sus servicios para otra empresa competidora. La no competencia representa un compromiso del profesional de no ejercer la misma actividad durante un período de tiempo tras la terminación del contrato.

Competencia desleal. Se pueden ejercer las acciones previstas en la Ley 3/1991, de 10 de enero, de Competencia Desleal (artículos 32 a 36). Esta legislación ofrece protección a posteriori en caso de acciones de imitación, denigración, comparación, explotación o aprovechamiento de los competidores. De este modo, aunque no pueda protegerse anticipadamente la reputación, el modelo de negocio o el *know how*, las acciones contra la competencia desleal, permitirían actuar contra estas acciones desleales para obtener una protección ulterior, solicitando que se declare que el comportamiento es desleal (acción declarativa de deslealtad), que cese tal conducta (acción de cesación de la conducta desleal o de prohibición de su reiteración futura), que se retiren los efectos derivados de la conducta desleal (acción de remoción de los efectos producidos por la conducta desleal), que se rectifiquen las informaciones en caso de que sean engañosas, incorrectas o falsas (acción de rectificación) o que se resarzan los daños y perjuicios ocasionados (acción de resarcimiento de los daños y perjuicios ocasionados por la conducta desleal, si ha intervenido dolo o culpa del agente o incluso acción de enriquecimiento injusto, que solo procederá cuando la conducta desleal lesione un derecho de exclusiva: por ejemplo, una marca o una patente).