

TEMA 2. PALANCAS DEL VALOR

Libro 100. DE.02.FactoresCompetitividad
Esteban Fernández Sánchez
Borrador actualizado: Agosto 2020

1. ECONOMÍAS DE ESCALA
2. EFECTO EXPERIENCIA
3. ECONOMÍAS DE ALCANCE
4. EFECTO DE RED
5. INFORMACIÓN Y CONOCIMIENTO
 - 5.1. La información
 - 5.2. El conocimiento
 - 5.2.1. Conocimiento explícito
 - 5.2.2. Conocimiento tácito
6. CALIDAD
7. INNOVACIÓN

1. ECONOMÍAS DE ESCALA

La escala, el tamaño de las operaciones de la empresa, a menudo es un determinante importante del coste medio de producción. Las economías de escala surgen cuando disminuye el coste medio de fabricación de un producto a medida que aumenta el volumen de producción por período de tiempo. Como muestra la figura 1, la curva de costes medios asociada a una fábrica tiene generalmente forma de 'U'. Existe un volumen de producción por período donde los costes alcanzan el valor mínimo, correspondiendo dicho valor con la capacidad óptima de la instalación (escala mínima eficiente, EME). Esto supone la existencia en la curva de costes medios de una primera zona de rendimientos crecientes a escala: la que abarca los niveles de producción inferiores a la capacidad óptima (desde 0 hasta EME). A la derecha de la capacidad óptima, sin embargo, comienza el tramo de rendimientos decrecientes a escala, que se corresponde con la existencia de deseconomías de escala (lectura 1).

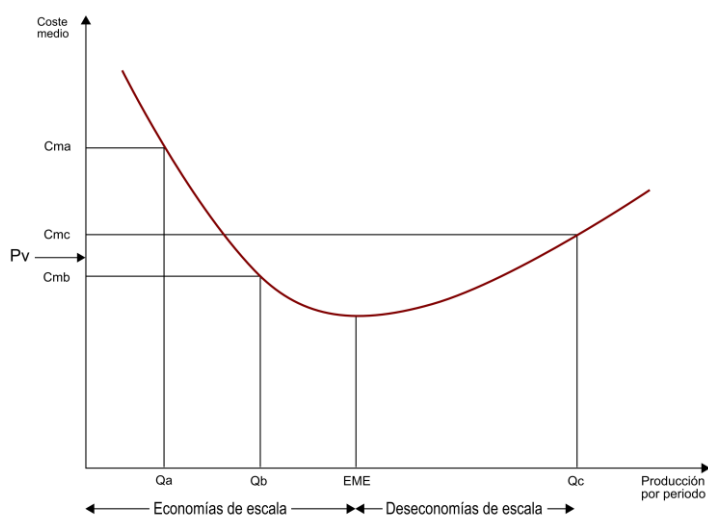


Figura 1: Economías de escala

La fábrica optará siempre por la escala que genere el menor coste medio para la demanda prevista. En la práctica, las fábricas no operan en la capacidad óptima. Normalmente se van adaptando con el tiempo. De ahí que la expansión y contracción de la demanda afecte a los costes. En la figura 2 se observa que, si Q_a es la cantidad de producción prevista por una empresa para satisfacer la demanda esperada, el coste medio en la fábrica FA es C_{ma} . Si en un futuro la demanda prevista es Q_b , a la empresa le interesa más continuar con la fábrica FA que construir la fábrica FB, ya que, a pesar de su menor tamaño, la fábrica FA resulta más eficiente: para una producción Q_b los costes C_{mb} de FA son menores que los costes C_{m1} de FB. En suma, el paso de un nivel de producción Q_a a un nivel Q_b dentro de la misma fábrica supone una reducción de los costes medios de C_{ma} a C_{mb} . Esta reducción de costes, derivada de un mejor aprovechamiento de la capacidad de la fábrica actual, activa las economías de escala a corto plazo.

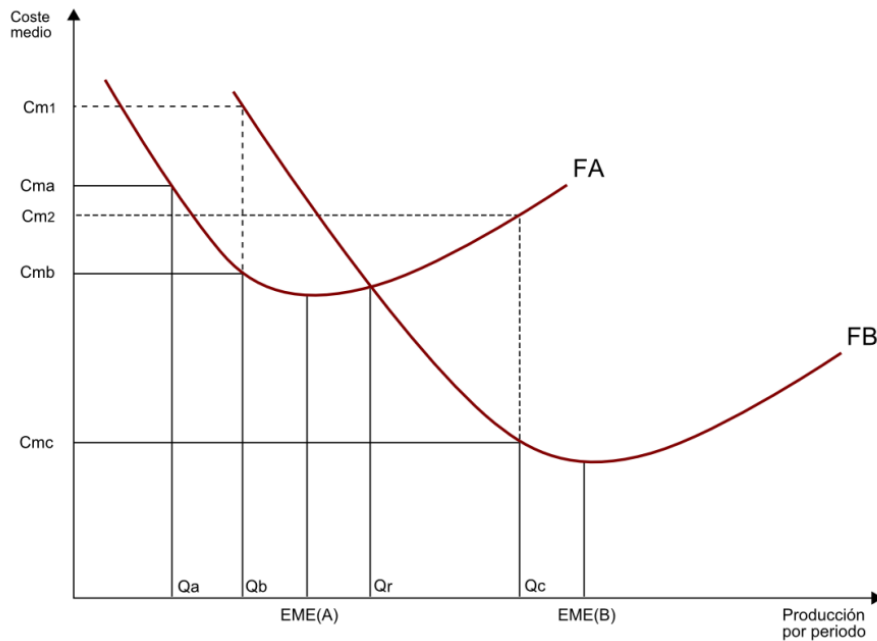


Figura 2: Un ejemplo de economías de escala

Ahora bien, si se espera que la demanda a largo plazo sea Q_c unidades por período, a la empresa le interesa construir la fábrica FB, ya que para ese nivel de producción tendrá unos costes medios C_{mc} inferiores a los costes C_{m2} que obtendría con la fábrica FA. Por tanto, el paso del nivel de producción Q_a en FA a Q_c en FB supone una reducción de los costes medios de C_{ma} a C_{mc} . Esta reducción de costes, derivada del cambio de la fábrica FA a FB debido al aumento de la producción por período, origina economías de escala a largo plazo. En suma, las economías de escala pueden darse al aumentar la producción con una tecnología dada o al cambiar de tecnología al aumentar la producción. Por debajo de un cierto nivel de producción (Q_r en la figura 2) no es eficiente cambiar de tecnología.

LECTURA 1: ECONOMÍAS DE ESCALA EN SONY

El concepto de coste medio lo supo expresar muy bien Morita (presidente-fundador de Sony), cuando intentaba comercializar sus radios en el mercado norteamericano. En sus propias palabras: “mientras visitaba distintos clientes potenciales, me encontré con otro comprador norteamericano, que observó la radio y dijo que le gustaba mucho. Dijo que su cadena tenía alrededor de ciento cincuenta tiendas y que necesitaba grandes cantidades. Esto me complació y, por fortuna, el comprador no me pidió que al producto le pusiera la marca de la cadena: sólo me solicitó que le diera los precios para cantidades de cinco mil, diez mil, treinta mil, cincuenta mil, y cien mil radios. ¡qué invitación!”

Pero cuando regresé a la habitación del hotel, empecé a reflexionar sobre el posible impacto que pedidos tan importantes ejercerían sobre nuestras pequeñas instalaciones de Tokio. Desde que desbordamos la capacidad de la despintada casucha con goteras de Gotenyama, nuestra fábrica se había ampliado considerablemente. Nos habíamos mudado a edificios más grandes y sólidos próximos al emplazamiento originario, y le habíamos echado el ojo a otros inmuebles. Pero, en nuestra pequeña cadena de montaje, no teníamos capacidad para producir cien mil radios de transistores al año sin dejar de hacer las demás cosas: nuestra capacidad era de menos de diez mil radios al mes. Si recibíamos un pedido para elaborar cien mil, tendríamos que contratar y formar empleados nuevos y ampliar aún más nuestras instalaciones. Esto significaría una gran inversión, una importante ampliación y un cierto riesgo.

Yo era inexperto, y todavía un poco ingenuo, pero contaba con capacidad intelectual: sopesé todas las consecuencias que se me ocurrieron y, después, me senté y tracé una curva que tenía el aspecto parecido a una letra u, uno de cuyos lados era mayor que otro. El precio de cinco mil radios sería nuestro precio normal: ése sería el comienzo de la curva; para diez mil equipos habría un descuento, y eso estaba en la parte inferior de la curva; por treinta mil, el precio empezaría a ascender; por cincuenta mil, el precio

unitario sería mayor que por cinco mil; y, para cien mil unidades, el precio unitario tendría que ser mucho mayor que por las primeras cinco mil radios.

Sé que esto suena extraño, pero mi razonamiento era que, si teníamos que duplicar nuestra capacidad de producción para servir el pedido de cien mil radios, y si no podíamos obtener una renovación del pedido para el año siguiente, nos hallaríamos en serios problemas –la bancarrota, quizá–, porque ¿cómo, en ese caso, podríamos mantener todo el personal adicional y pagar las nuevas instalaciones que no se usaban? Era un enfoque conservador y cauteloso, pero estaba convencido de que, si aceptábamos un pedido enorme, debíamos de obtener por él un beneficio lo suficientemente grande como para pagar las nuevas instalaciones durante la duración de ese pedido. La ampliación no es sencilla –obtener más dinero sería difícil– y no pensé que este tipo de ampliación fuese una buena idea, al hacerse sobre la base de un solo pedido. En Japón no podemos contratar gente, y despedirla así como así cada vez que aumentan o disminuyen los pedidos que recibimos. Tenemos un compromiso a largo plazo con nuestros empleados, y ellos lo tienen con nosotros. Naturalmente, también me preocupaba un poco que, si ofrecía un precio muy bajo por cien mil unidades, el comprador pudiera decir que se quedaría con cien mil radios, pero que, al principio, sólo haría un pedido de diez mil, a modo de prueba, al precio de las cien mil unidades y, después, podría ser que ya no mantuviera nuevos pedidos.

Al día siguiente volví con mi oferta. El comprador la miró y parpadeó como si no pudiera creer lo que veía. Bajó el papel y dijo con paciencia: “señor Morita, estuve trabajando como encargado de compras durante casi treinta años y usted es la primera persona que alguna vez se me haya presentado y dicho que cuanto más compro, más elevado será el precio unitario. ¡es ilógico!”.

Le expliqué mi razonamiento y escuchó cuidadosamente lo que yo le tenía que decir. Una vez que superó la sorpresa, se detuvo por un momento, sonrió, y, después, hizo un pedido de diez mil radios al precio correspondiente a las diez mil unidades, que era precisamente lo adecuado para él y para nosotros.

Fuente: Morita, A. (1986): *Made in Japan: Akio Morita and Sony*, E. P. Dutton, Nueva York.

Economías de escala pecuniarias y reales. Las economías de escala pecuniarias se derivan de los descuentos que obtiene la empresa a causa de sus operaciones a gran escala: mayores descuentos en compras, menor coste de financiación, descuentos en publicidad y promoción y reducción de las tarifas de transporte, entre otras. Sin embargo, si estos descuentos únicamente reflejan el poder de mercado del comprador, no son el resultado de verdaderas economías de escala. Sólo cuando es más barato para el proveedor del *input* ofrecer su producto al por mayor se obtiene una verdadera ganancia ocasionada por la escala.

Las economías de escala reales (o técnicas) son las que están ligadas a una reducción de la cantidad física de los *inputs* (materias primas, trabajadores, capital y conocimiento) por unidad de producto. Se logran a nivel de fábrica o en el ámbito de toda la empresa. Aunque el efecto escala aparece, sobre todo, en la fabricación, existe potencialmente en todos los componentes del coste de un producto: compras, investigación y desarrollo, publicidad y distribución, entre otros, si bien la sensibilidad de cada una de esas funciones respecto de la escala varía mucho. Por ejemplo, Chrysler invirtió 150 millones de dólares en publicidad en 1980, lo que supuso 123 dólares por automóvil vendido. Por el contrario, Ford invirtió 280 millones de dólares o 63 dólares por vehículo comercializado y General Motors 316 millones de dólares, lo que supuso tan sólo 44 dólares por vehículo vendido. Los gastos de publicidad deben repartirse entre menos unidades en las empresas más pequeñas, por lo que las economías de escala son menores (Aaker, 1987).

En general, las economías de escala reales se apoyan fundamentalmente en los siguientes mecanismos: indivisibilidades, especialización, relaciones área–volumen, economías estocásticas, control de mercados y densidad.

Indivisibilidades. Con independencia de su tasa de producción, la empresa tiene que incurrir en determinados costes fijos, también denominados indivisibilidades. Las economías de escala surgen cuando los costes fijos totales se reparten entre un mayor volumen de producción. A mayor volumen de producción, menor es el coste fijo unitario. Las indivisibilidades no se presentarían si la escala de las plantas y los procesos fueran susceptibles de aumentar o disminuir en pequeñas cantidades sin sufrir ningún cambio en su naturaleza, esto es, si fueran divisibles de manera perfecta. Las indivisibilidades proporcionan el logro de economías de escala a corto plazo: están originadas por una ocupación más completa de los activos fijos (Abell y Hammond, 1979). Cuando el desarrollo de un producto (una indivisibilidad) es muy costoso, el volumen es esencial para la rentabilidad. Así, el Boeing 747 fue muy rentable, ya que se construyeron 1.415 aviones; el Concorde con sólo veinte aviones, fue un desastre financiero (Grant, 2013).

En el supuesto de una función de costes lineal, donde representan: $CT(Q)$ el coste total, $CM(Q)$ el coste medio, F los costes fijos, v el coste variable unitario y Q la cantidad producida, tenemos que

$$CT(Q) = F + v \cdot Q$$

$$CM(Q) = \frac{CT(Q)}{Q} = \frac{F}{Q} + v$$

Por tanto, observamos que al aumentar la producción por período (Q) disminuye el coste medio de producción $CM(Q)$, pues los costes fijos (F) se reparten entre un mayor número de productos.

Especialización. En realidad, se trata de una indivisibilidad, debido a que es el resultado de la capacidad fija de un trabajador que se emplea de forma óptima cuando se dedica de manera exclusiva a una tarea específica. De acuerdo con Adam Smith, la especialización favorece la habilidad del trabajador, el ahorro de tiempo en el cambio de herramientas cuando se realizan tareas diferentes y la mecanización de los procesos productivos. Las grandes empresas también pueden especializarse en diferentes funciones, tales como producción, finanzas, investigación de mercados y diseño de productos. Una mayor especialización contribuye a reducir el coste medio.

Relaciones área–volumen. Para incrementar el volumen de cualquier recipiente, contenedor o envase en una determinada proporción, hay que aumentar la cantidad de material en una proporción menor. Si la capacidad productiva depende del volumen, en tanto que los costes dependen del área, entonces el coste técnico de la instalación crece menos que proporcionalmente en relación con el volumen productivo instalado. Un típico ejemplo son los contenedores: con 6 metros cuadrados de material se construye un contenedor de 1 metro cúbico de volumen (constituido por seis superficies cuadradas de 1 metro por lado); con 24 metros cuadrados de material un contenedor de 8 metros cúbicos (seis superficies cuadradas de 2 metros por lado). El material empleado aumenta cuatro veces, el volumen conseguido ocho. Estas economías son importantes en las llamadas industrias de procesamiento, como el refino del petróleo, el transporte de gas, la industria química, la del cemento, la del vidrio o la siderúrgica. Así pues, en muchas actividades, los incrementos de producción no requieren incrementos proporcionales de los factores empleados. Un tanque de 10.000 litros de combustible no cuesta cinco veces más que un tanque de 2.000 litros. En estos casos, tiene aplicación la denominada *Regla 0,6–0,8*, según

la cual, al duplicar la capacidad productiva, la inversión requerida no se duplica, sino que aumenta en 2^a , siendo 'a' un número que varía entre 0,6 y 0,8. Esto corresponde a un incremento entre el 52 y el 74 por ciento de la inversión si se duplica la capacidad productiva (Hax y Majluf, 1984). Por ejemplo, instalar una refinería con una capacidad productiva de 90 millones de toneladas por período cuesta sólo $(90/45)^{0,6}$, es decir, 1,52 veces más que instalar una refinería de 45 millones de toneladas de capacidad, lo cual es menos que el doble de la inversión. Por lo tanto, doblando la capacidad de la planta se reduce el coste por unidad de capacidad instalada en casi un 25 por ciento ($1/1,52 = 0,66$ versus $1/1 = 1$). Esto se traduce en un menor coste de amortización por unidad de producto. También se requiere menos del doble de los trabajadores y se puede mejorar el diseño de la instalación, características que no son posibles en una instalación de escala menor.

Economías estocásticas. Surgen porque la posesión de recursos masivos permite dispersar el riesgo. Es más difícil que una flota de veinte camiones opere al 50 por ciento que lo haga una flota de dos. Las empresas mantienen inventarios para reducir la probabilidad de quedarse sin los productos necesarios para satisfacer la demanda o para continuar el proceso de producción. Las economías de inventario tienen por objeto hacer frente a los cambios aleatorios, tanto por el lado de los *inputs* como por el lado de los *outputs*, en el funcionamiento de la empresa. La existencia de inventarios aumenta con la escala, pero en menor proporción.

Control de mercados. La integración vertical 'aguas arriba', que puede ser una estrategia válida para controlar el suministro de ciertas materias primas estratégicas, también suprime posibles comportamientos oportunistas de los proveedores y, en cualquier caso, elimina los beneficios que estos reciben, por lo que pueden disminuir los costes de los *inputs*. Igualmente, la integración 'aguas abajo' permite controlar el sistema de distribución de los productos y, así, atender mejor las necesidades de los clientes. Esta integración también constituye una buena fuente de información acerca del comportamiento de los mercados (Pratten, 1971). Asimismo elimina los beneficios de los distribuidores.

Densidad. Los ahorros surgen cuando una empresa tiene que abastecer un grupo numeroso de clientes agrupados en un área geográfica reducida en relación con ese mismo número de clientes repartidos en una zona geográfica extensa. Igual ocurre cuando los factores de producción están concentrados en un área geográfica determinada. También se las denomina economías de escala externas, y surgen cuando los costes promedio de una empresa disminuyen conforme aumenta la producción de la industria a la que pertenece. Las apalancan los siguientes hechos:

- ❑ La agrupación de las empresas de una misma industria en un área geográfica atrae, por ejemplo, mayores concentraciones de trabajadores especializados que requiere la industria y reduce, por lo tanto, los costes de contratación de la empresa. La necesidad de trabajadores cualificados puede favorecer la creación de un centro de formación.
- ❑ Los conocimientos tecnológicos se transfieren entre las distintas empresas a través de contactos directos entre ellas o cuando los empleados rotan de una empresa a otra.
- ❑ El fácil acceso a los insumos especializados aumenta cuando los proveedores de tales componentes se agrupan cerca del centro de producción. La proximidad de los proveedores disminuye el coste de transporte y favorece las economías de escala.

- ❑ Cuando un país tiene una industria en expansión, se origina crecimiento económico para ese país y gracias a esto, el gobierno puede recaudar mayores impuestos. En reconocimiento, el gobierno puede invertir en mejores instalaciones para la I+D en las universidades locales con el propósito de que algunas empresas de esa área se beneficien.

Deseconomías de escala. Cabe señalar que las economías de escala no se consiguen de forma ilimitada a medida que aumenta el volumen de producción. Una vez que la producción por período supera la capacidad óptima, los costes medios se incrementan (aparecen rendimientos decrecientes a escala), surgiendo las denominadas deseconomías de escala. En la figura 1 se observa que se producirán deseconomías de escala si la producción supera la escala mínima eficiente (EME). En la práctica, los aumentos de costes por las deseconomías de escala pueden producirse por múltiples circunstancias, entre las que cabe destacar las siguientes: el incremento en el coste de los factores de oferta constante (escasez de inputs), el aumento de los costes de transporte, sindicación de los trabajadores y la complejidad en la gestión (costes burocráticos).

La escasez que se origina en un factor al aumentar su demanda por la mayor escala productiva ocasiona un incremento desproporcionado del coste. Así pues, las deseconomías de escala en la función de compras pueden sobrevenir si la demanda se enfrenta a una oferta inelástica, lo cual eleva los precios de los factores productivos. Tal es, por ejemplo, el caso de la oferta de trabajadores en una zona geográfica determinada o con una cualificación concreta o un espacio físico disponible fijo, entre otros. Por otra parte, algunos recursos críticos son no replicables, como el talento de los directivos o las competencias extraordinarias de algunos ingenieros, por lo que su oferta siempre será limitada.

La búsqueda de economías de escala conduce a instalar un número reducido de grandes fábricas. Ahora bien, una instalación grande induce mayores costes de transporte que dos instalaciones pequeñas situadas más cerca de sus mercados. Puede ocurrir que el aumento de los costes de transporte supere el ahorro conseguido por la vía de la escala y estemos incurriendo, entonces, en deseconomías de escala (Pratten, 1971).

Los costes burocráticos aumentan como resultado de lo compleja que resulta la coordinación de actividades a mayor escala y la dificultad que entraña diseñar incentivos eficaces. En las fábricas muy grandes, donde trabajan miles de operarios, las estructuras son muy burocráticas y las relaciones puramente formales, ya que se tienen que utilizar normas y procedimientos para lograr la estabilidad y el funcionamiento eficaz de la organización, y existe una pirámide organizativa muy pronunciada, lo que provoca un distanciamiento entre la dirección y los trabajadores, causando apatía y desmotivación del personal, además de un enrarecimiento de las relaciones laborales y un retraso en la toma de decisiones. Todo ello genera enormes ineficiencias.

Por otra parte, el crecimiento de la fábrica acaparando la producción de un volumen excesivo de productos y componentes la empuja hacia la mediocridad. La fabricación se congestiona, crecen las existencias en las estanterías y almacenes, la planificación de la producción y los sistemas de control se hacen necesariamente más complicados, lentos y expuestos a errores. De igual manera, el incremento de la producción por período se apoya, a menudo, en políticas de fabricación que fomentan la contratación de

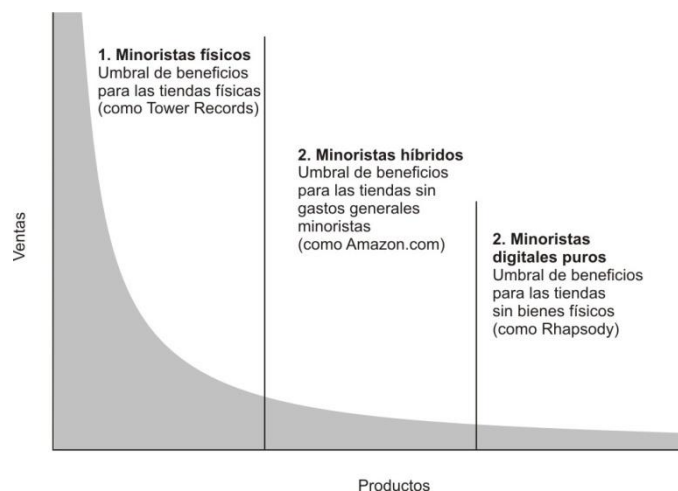
trabajadores a tiempo parcial, los horarios sostenidos de tiempo extra, el inicio de turnos adicionales (que conducen a tasas mayores de rotación de personal), la supervisión estricta de los trabajadores y un liderazgo controlador. Estas políticas crean tensión, estrés y absentismo entre los trabajadores, causando un impacto negativo en la productividad, que puede llegar a neutralizar los esfuerzos por aumentar la escala de la fábrica, favoreciendo la aparición de las deseconomías de escala.

En igual sentido, los sueldos de los trabajadores son más altos en las grandes empresas que en las pequeñas. Una explicación es que resulta más agradable trabajar en una pequeña empresa, por lo que la gran empresa debe pagar una compensación diferencial para atraer trabajadores. Otra explicación es el mayor poder que tienen los sindicatos en las grandes empresas en la negociación colectiva (Buzzell y Gale, 1987).

La agrupación de empresas puede generar congestiones de tráfico, costes de la vivienda excesivamente elevados, salarios altos y mayores concentraciones de polución.

LECTURA 2: LA ECONOMÍA DE LARGA COLA

En la actualidad, el mercado masivo se está convirtiendo en una masa de nichos. El nuevo mercado de nichos no está reemplazando el mercado tradicional de productos de éxito, sólo está compartiendo la escena con él, por primera vez. Durante un siglo hemos seleccionado sólo los *best sellers* para hacer un uso más eficiente del costoso espacio de exhibición, las pantallas, los canales y la atención del público. Ahora, en una nueva era digital de consumidores conectados, el aspecto económico de esta distribución está cambiando radicalmente, mientras Internet disminuye el coste tradicional de cada industria que afecta: tiendas, cines y canales de distribución.



En una era sin limitaciones de espacio físico y otros obstáculos que dificulten la distribución, los bienes y servicios especializados pueden ser económicamente tan atractivos como los artículos generalistas. Todo esto se traduce en seis elementos que caracterizan la era de la larga cola.

1. En casi todos los mercados hay más bienes de nicho que productos de éxito. Esta proporción crece exponencialmente a medida que las herramientas de producción llegan a ser más baratas y más ubicuas.
2. Los costes de acceso a esos nichos están bajando notablemente. Gracias a una combinación de fuerzas que incluyen la distribución digital, las poderosas tecnologías de búsqueda y la gran penetración de la banda ancha, los mercados *online* están transformando la economía del comercio minorista. En consecuencia, ahora muchos mercados pueden ofrecer una enorme variedad de productos.
3. Sin embargo, ofrecer simplemente más variedad no cambia la demanda. Hay que ayudar a los consumidores a encontrar los nichos que se adaptan a sus necesidades e intereses particulares. Para ello, se puede utilizar una gama de herramientas y técnicas que van desde las recomendaciones hasta las clasificaciones. Estos 'filtros' pueden orientar la demanda hacia la larga cola.

4. Una vez que se ha expandido la variedad y se aplican los filtros para navegar por ella, la curva de demanda se aplana. Todavía hay productos de gran popularidad y nichos, pero los primeros son relativamente menos populares, y los nichos relativamente más conocidos.

5. Todos estos nichos se suman. Si bien ninguno vende grandes cantidades, hay tantos productos de nicho que colectivamente pueden crear un mercado que rivaliza con los éxitos.

6. Una vez que se establecen estas condiciones, se revela la forma natural de la demanda, liberada de las dificultades de distribución, la escasez de información y la limitación del espacio de venta. Al contrario de lo que nos han hecho creer, de esta forma está mucho menos orientada a los productos de éxito. Es tan diversa como la población misma.

Una característica verdaderamente asombrosa de la larga cola es su gran tamaño. Si se reúnen suficientes temas no populares, es posible establecer un mercado de nichos que rivaliza con el mercado tradicional de productos de éxito. Consideremos los libros: en promedio, las librerías Borders poseen aproximadamente 100.000 títulos diferentes. Sin embargo, casi un cuarto de la venta de libros de Amazon no proviene de sus 100.000 libros más populares. Si tomamos las estadísticas de Amazon como guía, el mercado de libros que ni siquiera se vende en las librerías corrientes representa un tercio del volumen del mercado existente, y además va en aumento. Si esta tendencia continúa, el mercado potencial de libros podría ser un 50 por ciento más grande de lo que es hoy, si conseguimos superar la economía de la escasez.

Las empresas de mayor éxito en Internet sacan provecho de la larga cola de un modo u otro. Por ejemplo, Google obtiene la mayoría de sus ganancias no de las grandes empresas anunciantes, sino de las pequeñas. eBay también es esencialmente una empresa Long Tail. Al superar las limitaciones de la escala y la geografía, estas compañías no sólo han expandido los mercados existentes, sino que también han descubierto nuevos mercados. Además, en cada caso, estos nuevos mercados que están fuera del alcance del minorista real han resultado ser mucho más grandes de lo que todos esperaban, y su crecimiento no ha hecho más que empezar.

Fuente: Anderson, C. (2006): *The Long Tail. Why the Future of Business Is Selling Less of More*, Hyperion, Nueva York.

Las economías de escala determinan en parte el máximo número de empresas que pueden competir en una industria de manera rentable. Cuando la escala mínima eficiente es grande en relación con la demanda, la industria puede soportar relativamente pocos actores.

2. EFECTO EXPERIENCIA

El efecto experiencia es una generalización del efecto aprendizaje. Ambos se basan en el mismo principio: una actividad o tarea siempre puede realizarse mejor y de forma más eficiente a medida que se repite (aprendizaje mediante la práctica). Su importancia estratégica se deriva del hecho de que ambos efectos (experiencia y aprendizaje) son cuantificables y predecibles.

Los resultados del efecto aprendizaje se observaron por primera vez en 1925, en la industria aeronáutica, cuando se detectó que el número de horas de mano de obra directa utilizadas en el montaje de un avión disminuía un 20 por ciento –lo que es lo mismo, tenía una curva del 80 por 100– a medida que la producción acumulada de aviones se duplicaba (Wright, 1936). Posteriormente, a mediados de los años 60, el Boston Consulting Group (BCG) comprobó empíricamente que este fenómeno era más amplio y que la reducción constante no solamente afectaba al número de horas trabajadas, sino a cada una de las partidas que componen el coste unitario del producto. El rango de reducción típico varía entre un 10 y un 30 por ciento. Así, pues, mientras el efecto aprendizaje sólo incluye el número de horas trabajadas (o el coste de la mano de obra), el efecto experiencia abarca el coste unitario del producto.

El efecto experiencia indica que, a medida que una empresa duplica su experiencia en la fabricación de un producto, sus costes en términos reales disminuyen en una tasa o

porcentaje constante. En este contexto, la experiencia tiene un significado muy preciso: el número acumulado de unidades producidas hasta la fecha. La generación de experiencia puede ser lo suficientemente regular como para que resulte predecible el coste futuro del producto que estemos considerando. Ahora bien, las reducciones de coste no son automáticas, sino que es necesaria una gestión eficaz que persiga tales logros.

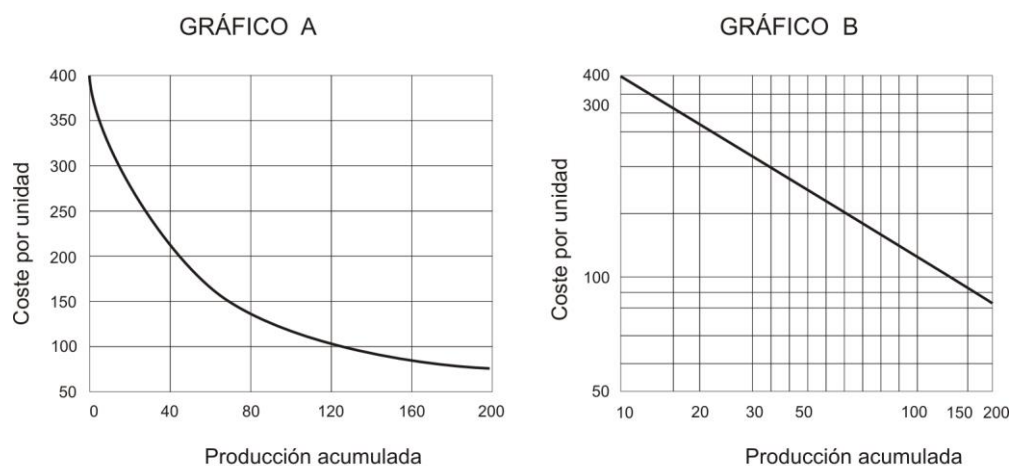


Gráfico A: representación de la relación coste por unidad-producción acumulada en una escala lineal.

Gráfico B: representación de la relación coste por unidad-producción acumulada en una doble escala logarítmica o escala bilogarítmica.

Figura 3: Representaciones gráficas de la curva de experiencia

La forma estándar del efecto experiencia se formula como $y = ax^{-b}$, donde 'y' es el coste unitario de la x-ésima unidad, 'a' es el coste unitario de la primera unidad, 'x' es el número de unidades producidas acumuladas y 'b' es la tasa de experiencia. La representación gráfica del efecto experiencia, en un sistema de coordenadas lineales, con la producción acumulada en el eje horizontal y el coste unitario en el eje vertical, es una curva (curva de experiencia) que muestra en el inicio una caída pronunciada de los costes para, a continuación, aplanarse. En un gráfico logarítmico, en cambio, el coste es una recta decreciente, lo que obedece a una velocidad de reducción constante. Esta relación lineal es más fácil de representar y más manejable para realizar previsiones de costes (figura 3). Ahora bien, la mejora continuada 'sin fin' que aparece en la figura debe matizarse. La base de la curva no es el tiempo, sino la producción acumulada, y el número de productos necesarios para duplicar la producción, por lo que, conseguir un porcentaje de mejora dado, resulta enormemente difícil con el transcurso de la experiencia. De esta forma, en la mayoría de los casos, el aprendizaje llega a ser tan lento que parece estático.

Una vez que se ha calculado el efecto experiencia, se pueden proyectar los costes para cantidades futuras. De ahí su importancia estratégica. Tener en cuenta el efecto experiencia permite determinar con mayor precisión el coste futuro del producto, lo que resultará útil a la dirección de la empresa en la política de fijación de precios. También resulta eficaz para determinar los descuentos por volumen cuando los clientes demandan pedidos significativos. Al mismo tiempo, se puede utilizar para negociar los precios con los proveedores, en caso de conocer su estructura de costes.

Fuentes del efecto experiencia. Para obtener provecho del efecto experiencia, es preciso conocer las fuentes de experiencia. Abell y Hammond (1979) enumeran las siguientes:

1. *Eficiencia de los trabajadores:* Al repetir los trabajadores una tarea específica de producción, se vuelven más habilidosos y aprenden mejoras en los procedimientos que aumentan la eficiencia colectiva, máxime si están motivados para hacerlo.

El aprendizaje lo realizan todo tipo de trabajadores, cualquiera que sea la labor que desempeñan en la empresa, y no sólo la mano de obra directa de fabricación. Para ello, se requiere personal estable y motivado, lo que se consigue con políticas de personal eficaces. Las nuevas contrataciones reducen el efecto experiencia, ya que los trabajadores experimentados deben dedicar tiempo a enseñar las habilidades del puesto a los nuevos empleados, quienes, a su vez, deben también dedicar tiempo a aprender la cultura y rutinas organizativas de la empresa. El abandono de trabajadores experimentados reduce la experiencia de la empresa, ya que son portadores de los conocimientos adquiridos durante su permanencia en la empresa.

2. *Especialización del trabajo y mejora de métodos:* La división del trabajo en tareas elementales y la correspondiente especialización de los operarios contribuyen positivamente al efecto experiencia, ya que éstos adquieren habilidades con mayor rapidez. Los esfuerzos tendentes a mejorar los métodos y procedimientos para la realización de determinadas tareas pueden ser extremadamente eficaces en la reducción de costes, sobre todo cuando se apoyan en programas formales y estructurados. La alienación que esto representa para el trabajador es un efecto negativo que también habría que cuantificar.

3. *Nuevos procesos de producción:* Las innovaciones y mejoras en los procesos productivos pueden ser una fuente fundamental de reducción de costes, sobre todo en las industrias intensivas en capital. Las mejoras tecnológicas que afectan positivamente al efecto experiencia pueden ser exógenas o endógenas. Las exógenas son fácilmente accesibles para los competidores. En cambio, las mejoras endógenas dependen más del *know-how* de la empresa en cuestión y, por lo tanto, deben protegerse contra las posibles filtraciones al exterior, con objeto de seguir acumulando experiencia a un ritmo superior al de los competidores.

4. *Mayor eficiencia del equipo productivo:* El nuevo equipo tiende a utilizarse por debajo de su potencial en el inicio de su funcionamiento. A medida que la experiencia se acumula, los responsables aprenden a utilizarlo en su capacidad óptima y, más adelante, pueden incluso modificarlo para mejorar su funcionamiento. Las mejoras incrementales para aprovechar mejor la capacidad de los procesos productivos contribuyen especialmente a reducir los costes.

5. *Cambios en la mezcla de recursos:* A medida que se acumula experiencia, la empresa puede reemplazar recursos más caros por otros más eficientes; por ejemplo, las máquinas pueden reemplazar a la mano de obra directa.

6. *Estandarización del producto:* La especialización de los trabajadores es eficaz solamente si el producto no sufre modificaciones sustanciales. La estandarización del producto fomenta, asimismo, la mecanización de los sistemas productivos y, por lo tanto, el incremento de su eficiencia.

7. *Rediseño del producto*: A medida que se obtiene experiencia con un producto, tanto el fabricante como el consumidor tienen una mejor comprensión de sus requisitos de desempeño. Esta comprensión permite que el producto sea rediseñado para reducir material, una mayor eficiencia en la fabricación, y una sustitución con recursos y materiales menos costosos, mientras que, al mismo tiempo, mejora el desempeño en las dimensiones importantes. El diseño industrial es también un factor determinante en los costes de producción y contribuye positivamente al efecto experiencia. Un buen diseño permite una mejor configuración de los componentes de un producto, así como la eliminación de los elementos superfluos.

El listado anterior ilustra la observación de que las reducciones de costes debidas a la experiencia no ocurren por inclinación natural; son resultado de un mayor esfuerzo y una mejor coordinación de los factores productivos para lograr menores costes.

Precios y experiencia. La habilidad para prever la evolución de los costes proporciona el potencial para establecer una política de precios adecuada. En mercados competitivos estables, se espera que los precios disminuyan de forma similar a como se reducen los costes (Abell y Hammond, 1979). De esta forma, tal como refleja la figura 4, el margen de beneficios es un porcentaje constante de los precios. En este caso, el efecto experiencia puede obtenerse, a menudo, aunque no siempre, a partir de la información que proporcionan los precios.

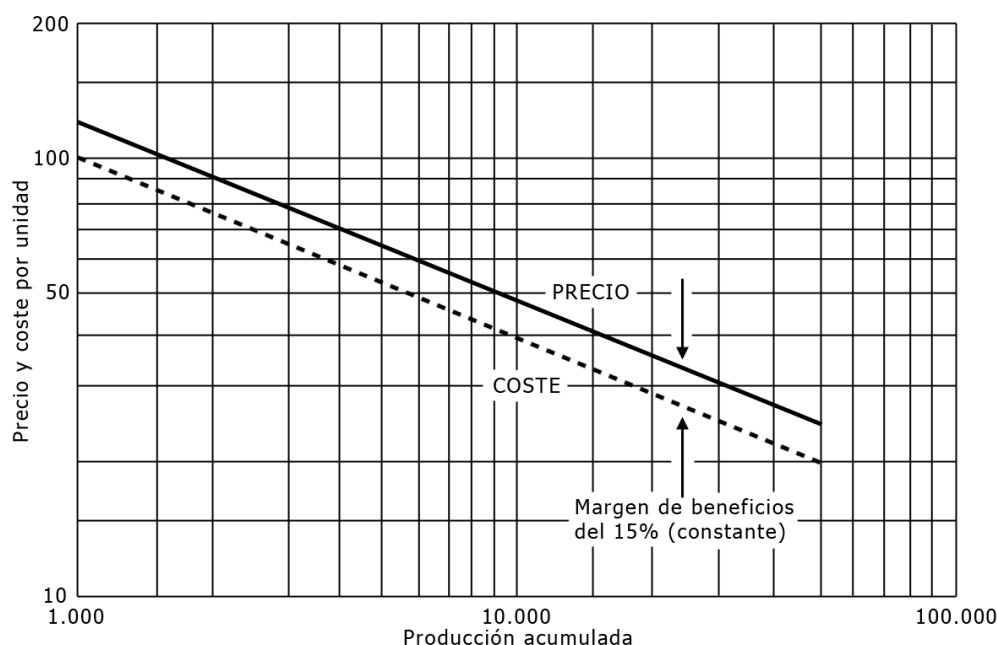


Figura 4: Curva de precios de la industria

La empresa que sigue esta estrategia toma la iniciativa y determina los precios del mercado que se impondrán al conjunto de competidores. Al mantener sus márgenes a un nivel constante, se dificulta mucho la llegada de nuevos entrantes al mercado y se elimina a los competidores más débiles.

Cabe destacar, por otra parte, que el comportamiento de los costes y los precios varía durante las etapas del ciclo de vida del producto, tal y como se muestra en la figura 5 (Abell y Hammond, 1979). Durante la fase de introducción es bastante habitual fijar los precios por debajo de los costes de producción, con objeto, por un lado, de favorecer la penetración del producto en el mercado y, por otro lado, de disuadir la entrada de competidores potenciales en el sector. La etapa de crecimiento se caracteriza porque el precio permanece firme bajo un ‘precio paraguas’ (de protección), apoyado por el líder del mercado, con los costes disminuyendo paulatinamente al ir acumulándose la experiencia. Posteriormente, surge un período turbulento, cuando ante la estandarización del producto, los competidores, con una gran capacidad instalada, comienzan a disminuir los precios de forma agresiva para posicionarse rápidamente en el mercado. En la etapa de madurez los márgenes de ganancia retornan a niveles normales, y los precios comienzan a seguir los costes de la industria en su caída de acuerdo con la curva de experiencia.

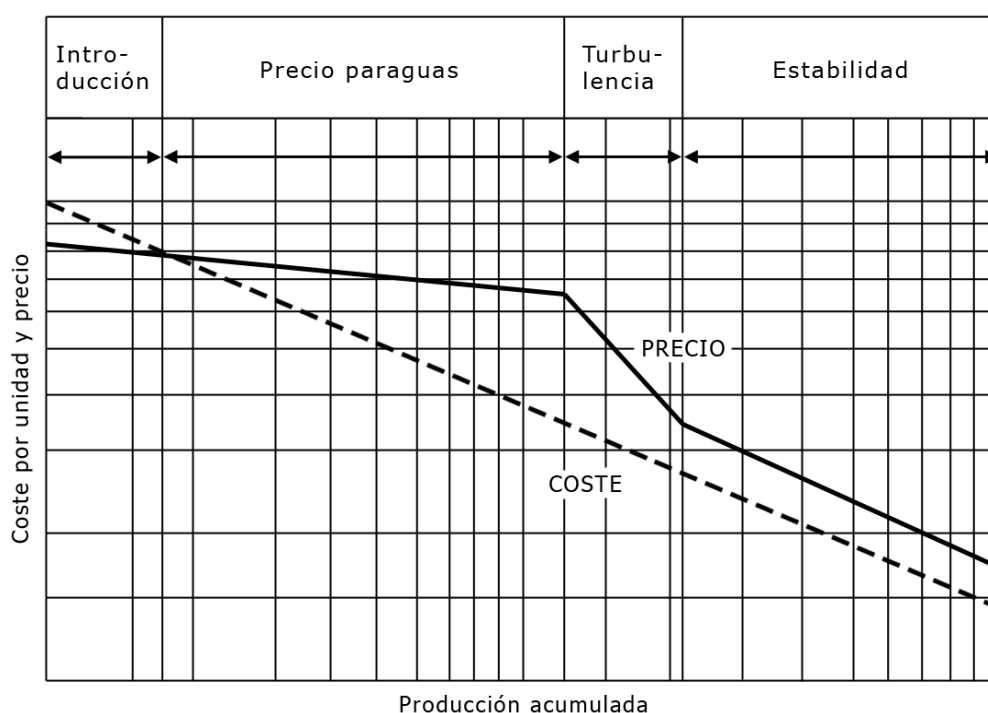


Figura 5: Precio paraguas

Relación entre escala y experiencia. Tanto las economías de escala como el efecto experiencia suponen reducciones en los costes al aumentar el volumen de producción. Sin embargo, ambos fenómenos presentan diferencias notables. El efecto experiencia surge primordialmente debido al ingenio, inteligencia, capacidad y destreza que obtienen los trabajadores con la experiencia; como señala el adagio: ‘la práctica hace al maestro’. Las economías de escala dependen del volumen de producción por período y no del volumen acumulado, como es el caso de la experiencia. Empresas con poca experiencia usan la escala para obtener costes bajos con relación a empresas con mayor experiencia. Son dos efectos independientes. La diferencia clave entre las economías de escala y el efecto experiencia es que la escala, en principio, puede replicarse rápidamente construyendo una gran planta, mientras que la experiencia requiere tiempo para acumular la producción. La figura 6 muestra la diferencia entre las economías de escala y el efecto experiencia. Las economías de escala implican una reducción del coste medio a medida que aumenta la cantidad producida en un período de tiempo. El

efecto experiencia conlleva un desplazamiento de la curva de coste medio: el coste medio de producir una cantidad dada en un período de tiempo disminuye conforme aumenta el volumen acumulado. Si se mantiene constante la producción por período, las economías de escala no variarán. No obstante, la empresa seguirá reduciendo costes a medida que acumule experiencia (que duplique la producción acumulada).

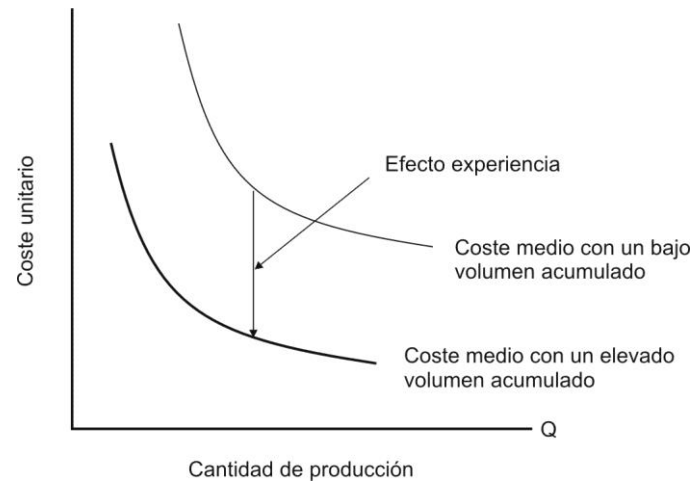


Figura 6: Diferencia entre experiencia y escala

Aunque son fenómenos diferentes, pueden operar de forma simultánea. Realmente, el efecto experiencia actúa siempre, desde que se fabrica la primera unidad, porque cada unidad adicional producida aumenta la producción acumulada y, en consecuencia, se presentan las condiciones necesarias para la reducción del coste. Por el contrario, las economías de escala operan únicamente al variar el volumen de producción por período. Con frecuencia el solapamiento entre los dos efectos es tan grande que resulta difícil (y no muy relevante) el separarlos (Abell y Hammond, 1979).

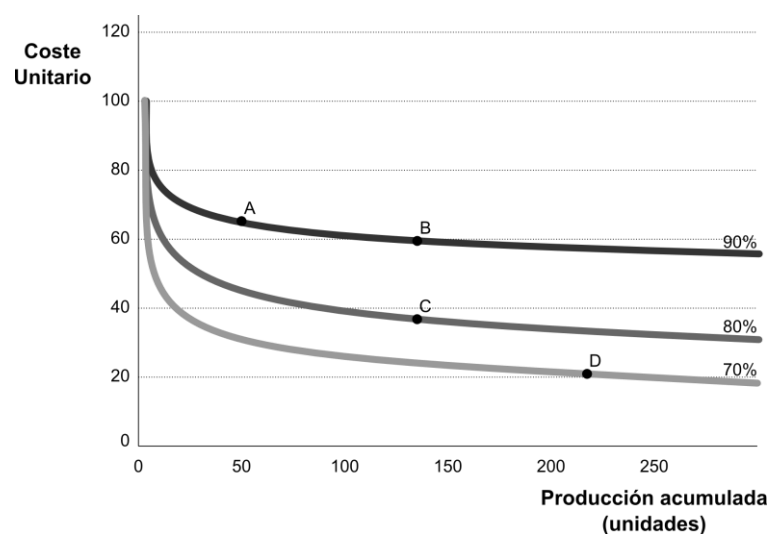


Figura 7: Curvas de experiencia con diferentes pendientes (Rothaermel, 2013)

La producción acumulada provoca un desplazamiento hacia abajo en la curva de experiencia. A medida que se fabrican más productos por período se acumula más experiencia. La empresa B al tener una escala mayor se mueve más rápido hacia abajo que A, aunque ambas tienen el mismo efecto experiencia. La empresa C introduce una nueva tecnología con mayor efecto experiencia, por lo que aunque tiene idénticas escala y experiencia acumuladas que B, alcanza una posición de costes más favorable. La empresa D captura escala y experiencia (figura 7).

3. ECONOMÍAS DE ALCANCE

Las economías de alcance (también denominadas economías de variedad, de ámbito, de enfoque o de gama) surgen cuando es más barato fabricar dos productos en una misma instalación que hacerlo en plantas separadas (Panzar y Willig, 1981). La expresión economías de alcance es, simplemente, una nueva denominación de la sinergia (o efecto $2+2=5$, que significa que el valor conjunto de dos recursos es superior a la suma de sus valores individuales). Las economías de alcance justificarían por sí mismas la existencia de fábricas multiproducto. Casi todos los ejemplos de economías de alcance se basan en el hecho de que varios productos comparten un activo, tangible o no (lectura 3). Así pues, las economías de alcance se logran cuando una fábrica consigue ahorros en coste al incrementar la variedad de bienes que produce. Ello requiere desplegar los activos sobre una mayor variedad de productos.

LECTURA 3: ECONOMÍAS DE ALCANCE EN LA FÁBRICA

En la década de 1990, un nuevo *software* hizo posible la transmisión de instrucciones digitales (codificación), y que de este modo las empresas pudieron proporcionar a sus proveedores una información exacta acerca de cómo ejecutar fases de la producción.

Las fundiciones fabricantes de microplaquetas (sus instalaciones suelen llamarse *fabs*) utilizan los mismos equipamientos para procesar los distintos ficheros que transmiten los clientes, y que contienen las especificaciones. Una *fab* que produce circuitos integrados para ordenadores, aparatos de telecomunicación, cámaras digitales, teléfonos celulares y sistemas destinados a la automoción, resiste los altibajos cíclicos de cada una de estas actividades en particular mejor que las compañías que fabrican los chips para uso propio, como se hacía en la década de 1980. La clientela amplia y diversa también permite soportar la fuerte inversión de capital que se necesita para construir una *fab* nueva. A menudo, las operaciones de este género reciben la contribución de varios grandes clientes, por lo general a cambio de comprometer la reserva de una parte de su capacidad. Las *fabs* deben ofrecer a los clientes la garantía de que están en condiciones de proteger la propiedad industrial de estos. También hay que gestionar la actividad de una manera que impida filtraciones de la PI de unos clientes hacia otros.

Fuente: Berger, S. (2005): *How We Compete –What Companies Around the World Are Doing to Make It in Today's Global Economy*, Currency/Doubleday/Random House, Nueva York.

Los economistas tradicionalmente han supuesto que, en la medida que existan economías de alcance, estas se encontrarían principalmente en los procesos de producción de la empresa. Por ejemplo, una fábrica que produce televisores puede ser más eficiente si fabrica también pantallas de ordenador. Sin embargo, las empresas se dirigen hacia una cartera más diversa de plantas cada vez más homogéneas, lo que significa que las economías de alcance técnicas parecen desempeñar un papel mínimo en la explicación de la diversificación empresarial (Gollop y Monahan, 1991).

Aunque las de tipo técnico no sean relevantes, las economías de alcance también pueden proceder de actividades que no se encuentran directamente relacionadas con el proceso

de producción, incluyendo áreas como investigación y desarrollo, ventas y marketing, distribución, transporte y gastos generales. Así, por ejemplo, desde el 2 de agosto de 2001, PepsiCo, al culminar la adquisición de The Quaker Oats Company, creando una compañía de aperitivos y bebidas por un valor de 25.000 millones de dólares, se está beneficiando de las economías de alcance, debido a que todos los productos van a aprovechar, entre otros recursos, las redes de distribución de la empresa.

El valor de los activos intangibles en las economías de alcance ha recibido recientemente gran atención. Los activos intangibles, a diferencia de lo que ocurre con los activos tangibles, se pueden utilizar simultáneamente en más de una actividad sin que disminuya su valor. Ello es posible porque la base del recurso intangible es conocimiento, que, a su vez, cuenta con dos propiedades que posibilitan las economías de alcance: puede ser utilizado simultáneamente y no se deteriora con el uso (más bien aumenta su valor). Una tecnología aplicable a diferentes líneas de productos se encuentra sujeta a economías de alcance. Igual ocurre si la reputación de la empresa la extendemos a un amplio rango de mercados. No obstante, la empresa que utiliza la misma marca en múltiples productos –marca paraguas– arriesga las ventas futuras de todos sus productos en la medida en que el fallo por mala calidad de uno afectará negativamente a todos los demás, al estar cubiertos por la misma marca.

En definitiva, mientras que las economías de escala se relacionan con el volumen, las economías de alcance lo hacen con la variedad. Ahora bien, ambas se apoyan en el mismo fundamento: el reparto de los costes fijos entre mayores cantidades de productos. Sin embargo, en las economías de alcance, el ahorro se consigue fabricando lotes de una variedad de productos, mientras que en las economías de escala se logra fabricando un gran volumen del mismo producto.

La diferencia entre las economías de escala y de alcance se basa esencialmente en que, para las economías de escala, la ganancia derivada del uso de los recursos en los altos niveles de producción se ve reflejada en un movimiento a lo largo de la curva de costes medios, mientras que en el caso de las economías de alcance, la ganancia derivada del uso compartido de los recursos se ve reflejado en un desplazamiento hacia abajo de la curva de costes medios para cada uno de los productos que comparte recursos.

Al igual que sucedía con las economías de escala, puede hablarse de deseconomías de alcance, destacando que, en ocasiones, la fabricación conjunta de diversos productos puede conllevar un aumento en tres tipos de costes (Porter, 1985): a) coste de la coordinación, b) coste del compromiso y c) coste de la inflexibilidad. La coordinación entraña costes de tiempo, personal y, quizás, de dinero. Estos costes diferirán ampliamente con cada actividad a compartir. Dos negocios que comparten una actividad pueden realizarla de una manera uniforme que quizá no sea la óptima para ninguno. Por tanto, las unidades de negocio casi siempre deben aceptar un compromiso de mínimos en la satisfacción de sus necesidades al compartir una actividad. Por ejemplo, un componente complejo común a dos líneas de producto, del que un producto requiere las funciones A y B, y el otro producto las funciones A y C, forzosamente tendrá que proporcionar las tres funciones a pesar de que no sea lo óptimo para cada producto. La inflexibilidad adopta dos formas: 1) dificultad potencial de responder a las tácticas competitivas y 2) las barreras de salida. Los lazos entre los negocios de una empresa crean vínculos entre sus resultados. Por ejemplo, si disminuye la demanda del producto de un negocio, abandonar el negocio conlleva cargar todo el coste fijo compartido al

otro negocio, lo que lo hace inviable. Por tanto, igual hay que continuar con el negocio aunque no sea rentable. El coste de la inflexibilidad es potencial, ya que dependerá de la probabilidad de que haya que responder o salir. Por otra parte, el simple incremento del volumen de un negocio en absoluto garantiza una mejor estructura de costes de la empresa; combinar dos líneas de negocio similares no significa mejorar la calidad del producto o reducir los costes.

4. EFECTO DE RED

Hay una diferencia fundamental entre la nueva y la vieja economía. La vieja economía industrial estaba impulsada por las economías de escala por el lado de la oferta (con rendimientos decrecientes), la nueva economía de la información está impulsada por los efectos de red (con rendimientos crecientes).

Los efectos de red (externalidades de red, economías de escala del lado de la demanda o rendimientos crecientes del lado de la demanda) surgen del aumento del valor del producto para el usuario a medida que aumenta el número de usuarios (Katz y Shapiro, 1985). Dos fuentes de valor sustentan los efectos de red: el efecto directo y el efecto indirecto.

Efecto directo. El efecto de red directo requiere compatibilidad horizontal. Típicamente ocurre en las redes reales de comunicación bidireccionales, caracterizadas porque los eslabones entre los nodulos (de la red) son conexiones físicas. El efecto directo surge del beneficio que obtiene el usuario del producto al poder ‘conectarse’ con otros usuarios para, por ejemplo, intercambiar información como ocurre con las redes telefónicas. ~~Ser el único propietario de un teléfono no tiene mucho valor. El valor de un fax depende de su capacidad para permitir la comunicación con otros usuarios, que también tengan uno en propiedad, y su valor aumenta conforme crece el número de propietarios.~~ En esta red y otras construidas de forma similar, un usuario está conectado directamente a cada uno de los demás usuarios. Por tanto, el valor de la conexión para un usuario depende de las conexiones con otros usuarios. Además, un nuevo usuario del producto *per se* confiere beneficios a los otros usuarios. Los beneficios de la red son proporcionales al número de personas con las que un usuario puede comunicarse. En general, si hay N personas en una red, y el valor de la red para cada una de ellas es proporcional al número del resto de usuarios, entonces el valor de la red (para todos los usuarios) es proporcional a $N \cdot (N-1) = N^2 - N$. Se denomina ‘regla de cuadrados’ o ley de Metcalfe (Shapiro y Varian, 1999).

Efecto indirecto. Nos encontramos en la misma red informática si podemos usar el mismo software y compartir los mismos ficheros (Shapiro y Varian, 1999). El efecto de red indirecto emerge de la interdependencia en el consumo de productos complementarios. El valor del producto de red emerge indirectamente del volumen y diversidad de los productos complementarios disponibles (Venkatraman & Lee, 2004). Muchos productos tienen poco o ningún valor aisladamente, pero generan valor combinados con otros (David & Greenstein, 1990). Los complementarios pueden ser físicos (las cintas de vídeo para las videograbadoras) o intangibles (habilidades de escribir del mecanógrafo en un teclado QWERTY) (Gallagher & Park, 2002). Al producto de red y sus productos complementarios se les denomina sistema: una colección de dos o más componentes conectados por un interfaz que les permite trabajar juntos (Katz y Shapiro, 1994). Un sistema es la combinación de un bien durable (producto de red, central o base) con bienes o servicios que realizan una función determinada (productos complementarios). En este caso es importante la compatibilidad vertical.

A un sistema también se le denomina paradigma hardware/software: los dos componentes actúan o son combinados para proporcionar valor. El caso más interesante es cuando una unidad de hardware (producto de red) es compatible con una gran variedad de software (productos complementarios). Los compradores del *hardware* aportan beneficios a otros usuarios del mismo *hardware*, porque expanden la base instalada del producto de red, estimulando la demanda para el *software* compatible. Los fabricantes de productos complementarios pueden aprovechar la ampliación de la base instalada para proporcionar *software* más barato y aumentar la variedad, lo que, a su vez, refuerza el valor del *hardware* para los propietarios. El producto de red que cuenta con una gran base instalada es probable que atraiga más desarrolladores de productos complementarios contribuyendo a aumentar la rivalidad entre ellos.

En general, un producto de red (hardware) se configura alrededor de una tecnología nuclear. Empresas competidoras pueden fabricar productos que utilizan la misma tecnología nuclear, cada una con su propia marca. Por ejemplo, una videograbadora JVC y otra Grundig comparten la tecnología nuclear *VHS*, mientras que la videograbadora Sony utiliza una tecnología nuclear incompatible: la tecnología *Betamax*. Las cintas de video que utilizan tecnología *VHS* pueden reproducirse tanto en las videograbadoras JVC como en las Grundig, pero no en las Sony. Ello se debe a que las tecnologías *VHS* y *Betamax* son incompatibles. Así pues, un sistema está formado por a) un producto de red (hardware) que se apoya en una tecnología nuclear, con independencia de que lo comercialicen diferentes empresas cada una con su propia marca y b) diversos productos complementarios (software) compatibles con el producto de red. En suma, un mismo producto complementario no puede utilizarse en dos tecnologías nucleares incompatibles. En el efecto de red indirecto el aumento del valor del producto de red lo proporciona la disponibilidad de productos complementarios compatibles. El valor de una videograbadora depende de la disponibilidad de cintas (películas, documentales, series, etc.) en el mercado. Los productos complementarios son compatibles con el producto de red e incompatibles con el resto de tecnologías nucleares que soportan otros productos de red. La escasez de productos complementarios genera un mayor riesgo de *lock-out* (Schilling, 2002).

El efecto indirecto puede surgir también en otros tipos de redes (Saloner *et al.*, 2001). En las redes estrella, los usuarios interactúan mediante un nexo central. Una oficina central para el intercambio de dinero y documentos bancarios es una red estrella. La red tiene valor debido a que la gente cambia de ubicación. Por ejemplo, cuanto mayor sea el número de lugares donde una persona pueda encontrar un cajero automático compatible, mayor será el valor de la red de cajeros para esa persona. A veces, ninguna red física conecta a los miembros. Un club de personas que intercambian sellos disfruta de los beneficios de la red. De modo similar, el valor que los vendedores potenciales confieren a una subasta aumenta con el número de compradores que participan en ella. El efecto red también surge en los distritos industriales conforme aumenta el número de empresas que contribuyen al fondo común de conocimientos tecnológicos, y les permite aprender unas de otras, así como beneficiarse de las economías de escala de los factores productivos comunes.

En las redes reales, los eslabones entre los nodulos son conexiones físicas, mientras que en las redes virtuales (relacionadas con el efecto de red indirecto) los eslabones entre nodulos son invisibles, pero no por ello son menos importantes. Tanto si son reales como virtuales, las redes presentan una característica económica fundamental: el valor de conectarse a una red depende de cuantas otras personas estén conectadas (Shapiro y Varian, 1999).

Base instalada. A los usuarios de una tecnología nuclear o del sistema se les denomina colectivamente base instalada. Al aumentar la base instalada se producirá una bajada del precio del producto de red (debido a las economías de escala). El producto de red que cuenta con una gran base instalada es más probable que atraiga más desarrolladores de productos complementarios. A su vez, la disponibilidad de productos complementarios influirá sobre la elección de los usuarios entre productos de red incompatibles (Schilling, 1999).

La base instalada en relación con los productos complementarios es el problema del huevo y la gallina (Gupta *et al.*, 1999): una gran base instalada atrae a los fabricantes de productos complementarios y una diversidad de productos complementarios contribuye a aumentar la base instalada. En este sentido, Stremersch *et al.* (2007) concluyen que la base instalada induce la oferta de productos complementarios.

Por razones de mercado, a los proveedores de complementarios les interesa alinearse con el producto de red que tenga una base instalada mayor (figura 8). A medida que aumenta el número de proveedores de productos complementarios para la misma base instalada, lo hace simultáneamente la competencia entre ellos. No obstante, puede que su cuota de mercado no se reduzca si la base instalada aumenta de tamaño. Dado que los complementarios son muy importantes para el éxito de un producto de red, una empresa que intente proporcionar un producto de red deberá potenciar igualmente la creación de complementarios.

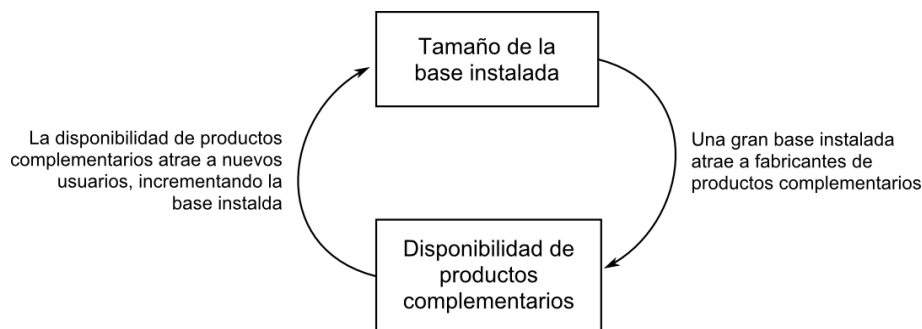


Figura 8: Ciclo de retroalimentación de la base instalada y la disponibilidad de bienes complementarios (Schilling, 1999)

Límites de la red. El efecto de red probablemente tiene rendimientos decrecientes una vez que el número de usuarios que adopta la tecnología se hace muy grande. Conforme las redes crecen, cualquier restricción de los recursos necesarios para mantenerlas puede causar congestión (Westland, 1992). La falta de ancho de banda causa congestión en la Web, disminuyendo el tiempo de acceso y retrasando las transacciones. Este caso muestra que la tasa a la cual el valor de la red aumenta puede declinar con el paso del tiempo debido a los costes de congestión (incluso hacerse negativos si los costes de congestión superan a los beneficios de red). Así pues, la congestión puede hacer que el valor de un nuevo usuario sea negativo, ya que dificulta al resto de usuarios sus actividades con los recursos compatibles (Saloner *et al.*, 2001).

Los límites de la red se entienden mejor al comparar los retornos de las externalidades de red derivados de la cuota de mercado con los costes correspondientes. Estos retornos se

refieren al valor que los consumidores obtienen conforme una mayor porción del mercado adopta el mismo sistema (por ejemplo, es probable que exista una mayor disponibilidad de productos complementarios, mayor compatibilidad entre usuarios y mayores ingresos canalizados hacia un mayor desarrollo de la tecnología). Al aumentar el tamaño de la red se incrementa el ritmo de aprendizaje, haciendo más valioso el sistema. Los costes de monopolio se refieren a los costes que soportan los usuarios conforme una mayor parte del mercado adopta el mismo sistema (por ejemplo, la empresa puede cargar mayores precios al producto central, existe una menor variedad de productos complementarios y la innovación en sistemas alternativos puede ser sofocada). La relación entre retornos de las externalidades de red y cuota de mercado normalmente adquiere la forma de ‘S’. Sin embargo, los costes normalmente aumentan exponencialmente respecto a la cuota de mercado. Al representar ambos en el mismo gráfico (figura 9), se revela como los beneficios de las externalidades de red y los costes de monopolio se compensan entre sí (Schilling, 2008).

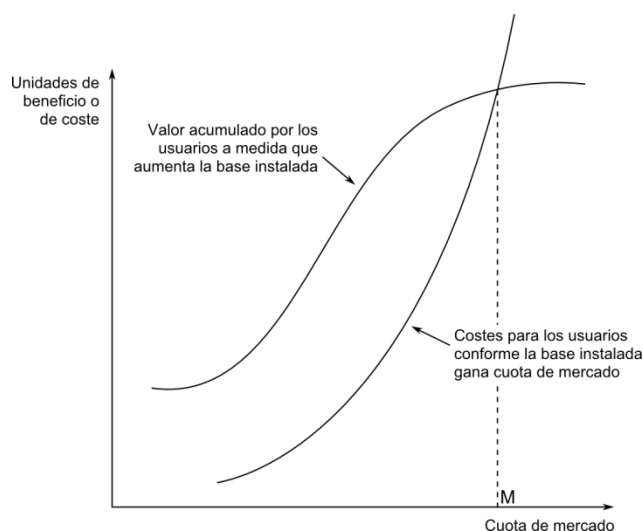


Figura 8: Beneficios y costes del efecto red (Schilling, 2008)

En la figura 9 mientras la base instalada sea menor que M, los beneficios del efecto de red superan a los costes, incluso si M representa una cuota de mercado muy grande. Sin embargo, conforme la base instalada se incrementa más allá de M, los costes del efecto de red superan el valor de los beneficios. Una serie de factores pueden alterar el punto en el que ambas curvas se cortan. Si la utilidad del sistema fuera superior, las curvas se cortarían en un punto mayor que M. Si la curva de retornos de la externalidad de red comienza a reducir su pendiente con una menor cuota de mercado, las curvas se cortarían para una cuota de mercado menor que M (Schilling, 2008).

Lock-in y costes de cambio. El lock-in tiene lugar siempre que los usuarios invierten en activos múltiples, complementarios y duraderos propios de un determinado sistema. Por lo general, al reemplazar un sistema antiguo por uno nuevo incompatible, puede resultar necesario duplicar todos los componentes del sistema.

Cuando llegue el momento de reemplazar el Seat que hemos estado conduciendo durante varios años, no hay ninguna razón imperativa para elegir otro Seat en lugar de un Toyota u otra marca. En nuestro garaje entra igual de bien el Seat como otra marca;

no nos va a llevar tiempo aprender a conducir otra marca; y podemos engancharle el mismo remolque a cualquiera de los vehículos. Para abreviar, se puede transferir fácilmente la inversión en ‘infraestructura automovilística’ a cualquier otro tipo de coche (Shapiro y Varian, 1999).

Cuando los costes de cambiar de un tipo de sistema a otro son sustanciales, los usuarios se enfrentan al efecto lock-in, es decir, quedarán amarrados al sistema. Al estudiar la evolución del teclado QWERTY, David (1985) observó tres factores que ataron (lock-in) a los usuarios: la interrelación técnica, las economías de escala y la cuasi-irreversibilidad de la inversión.

La interrelación técnica conlleva la necesidad de compatibilidad entre el hardware del teclado y el software representado por la memoria de los mecanógrafos respecto a una determinada ordenación de las teclas. La adquisición por parte de un empleador de una máquina con teclado QWERTY transfería una externalidad positiva a los mecanógrafos con formación compatible. Por tanto, más mecanógrafos aprendían mecanografía memorizando la ordenación de las teclas QWERTY. Igualmente, el coste del uso de un sistema de mecanografía QWERTY tendería a disminuir conforme ganase en aceptación respecto a otros sistemas de ordenación de las teclas. Los fabricantes de máquinas de escribir que utilizan un solo teclado, en lugar de dos o más, obtienen economías de escala en la producción de máquinas, lo que contribuye a reducir su precio. Finalmente, el lock-in parece que se debe también en parte al elevado coste de la ‘conversión’ del software (mecanógrafo) y la consiguiente cuasi-irreversibilidad de la inversión en destrezas específicas. La persona que aprende a escribir en un teclado QWERTY encontrará lento e irritante tener que aprender otra colocación de letras diferente.

Ahora bien, no nacemos atados a ninguna tecnología; solamente nos atamos en virtud de las decisiones que tomamos a lo largo del tiempo (figura 10). El punto de partida es la selección de la tecnología. Esta elección puede referirse, por ejemplo, a la compra de un tocadiscos. La primera vez que un determinado usuario elige una tecnología, sus preferencias no van a estar condicionadas por ningún lock-in, ya que no tiene ninguna restricción para elegir entre las diferentes tecnologías que se ofertan en el mercado (puede adquirir un magnetófono o un reproductor de CD, por ejemplo).

A la selección de la tecnología le sigue la fase de muestreo, durante la cual el usuario saca partido de cualquier incentivo que se le hubiera ofrecido para que la probara (permitirle tenerlo a prueba durante una semana). Los usuarios que no se limitan a probar la tecnología avanzan a la fase de afianzamiento. Aquí es cuando el usuario se familiariza realmente con la nueva tecnología, desarrolla su preferencia por esa tecnología por encima de las demás, y quizás se amarra a ella haciendo inversiones en productos complementarios (compra masiva de discos de vinilo). La fase de afianzamiento culmina en lock-in cuando los costes de cambiar a otra tecnología incompatible (reproductor de CD) se hacen prohibitivamente altos. Dado que el usuario amarrado al producto central tiene una demanda muy inelástica, el vendedor (o vendedores) puede elevar los precios de los productos complementarios para extraer una mayor parte del excedente del consumidor.

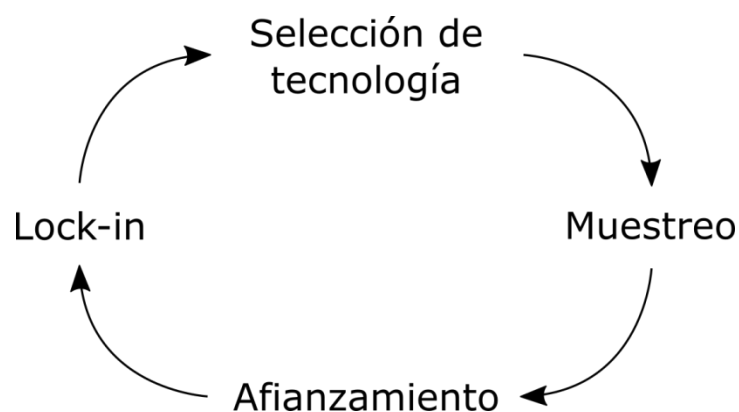


Figura 10: Ciclo de lock-in (Shapiro y Varian, 1999)

Normalmente, el cambio de un sistema a otro incompatible acarrea unos costes, que pueden desglosarse en varios componentes (Porter, 1980; Lieberman y Montgomery, 1988; Klemperer, 1995): a) inversiones en formación (el usuario debe aprender a manejar la nueva tecnología y desarrollar rutinas de funcionamiento), b) inversiones en activos físicos (producto de red y productos complementarios) en que incurren los usuarios al cambiar de tecnología; c) costes contractuales que puede intencionadamente crear el vendedor (por ejemplo, penalizaciones por ruptura del contrato); estos costes crean adherencia, ya que los clientes no están dispuestos a romper el contrato, sobre todo si las penalizaciones son muy altas; d) programas de fidelización de los clientes, como los realizados por las compañías aéreas, que proporcionan un viaje gratis a los clientes que acumulen determinadas horas de vuelo viajando con la misma aerolínea; también podrían incluirse los costes psicológicos por el cambio de marca, que justifican en gran medida la lealtad del cliente; e) la incertidumbre sobre la calidad de las aportaciones de la nueva tecnología provoca inercia en las empresas, ya que no están dispuestos a incurrir en costes de búsqueda de nuevos proveedores; f) coste de conversión de los ficheros a un nuevo formato y g) dificultad para encontrar un proveedor especializado que sustituya al actual (figura 11).

Tipo de lock-in	Costes de cambiar
Aprendizaje de una marca concreta	Aprender un nuevo sistema, costes directos y pérdida de productividad; tienden a crecer con el tiempo
Compras de equipos	Sustitución de equipo; tiende a declinar a medida que envejece
Compromisos contractuales	Indemnizaciones
Programas que favorecen los clientes leales	Pérdidas de ventajas que ofrece el proveedor más posible necesidad de reconstruir la confianza
Costes de búsqueda	Costes de búsqueda combinados de comprador y vendedor; incluyen aprendizaje de la calidad de las alternativas
Información y bases de datos	Convertir los datos en nuevo formato; tienden a aumentar a la larga cuando el volumen de datos crece
Proveedores especializados	Encontrar un nuevo proveedor; pueden aumentar con el tiempo si los proveedores son difíciles de

	encontrar/mantener
--	--------------------

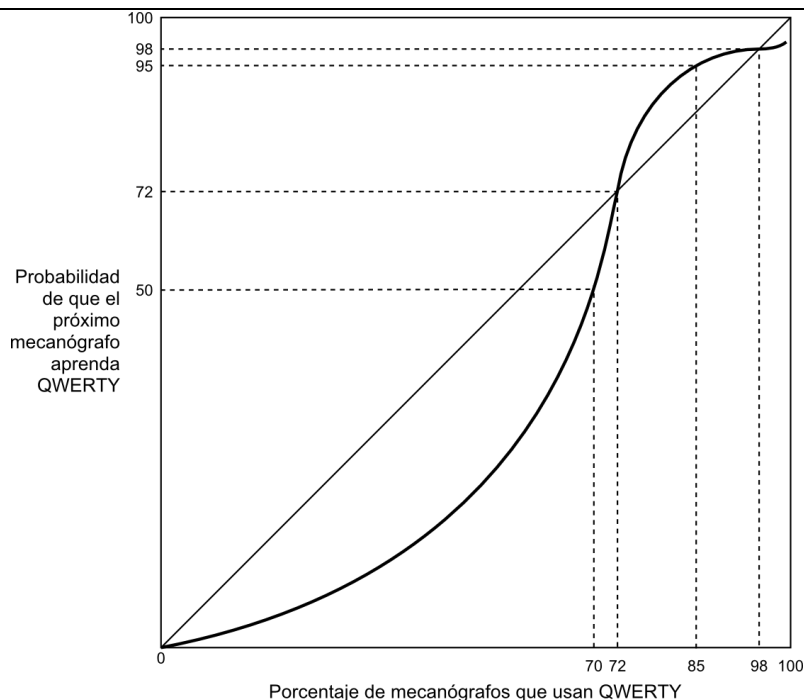
Figura 11: Tipos de lock-in y sus costes de cambiar (Shapiro y Varian, 1999)

Existe un exceso de inercia, en terminología Farrell y Saloner (1985), cuando resulta complicado, lento y costoso sustituir una tecnología por otra, incluso aunque sea mejor. Este exceso de inercia puede producirse en aquellos casos en que la información no es perfecta, y obedece a: 1) los elevados costes de cambio y 2) la existencia de una base instalada amplia. Como resultado, el sistema con un coste de cambio elevado y una gran base instalada retendrá a los usuarios y atraerá a la mayoría de los clientes potenciales. Aparece el efecto *bandwagon* (o arrastre) que afirma que las preferencias de la gente por un producto aumentan a media que crece el número de personas que lo adoptan (Leibenstein, 1950). Esto potencia una ventaja competitiva al sistema con la base instalada más grande, ya que aumenta el poder de permanencia o viabilidad a largo plazo del sistema, lo que resulta particularmente importante para los clientes (lectura 4). Un nuevo comprador preferirá adquirir el sistema con la mayor base instalada. Si los usuarios pueden cambiar de sistema fácilmente, la base instalada será poco importante. Sin embargo, si el sistema tiene costes de cambio altos induce a los compradores a conservar su elección inicial.

LECTURA 4: EL EFECTO BANDWAGON (EFECTO IMITACIÓN O ARRASTRE).

Los consumidores están observándose unos a otros. Si la adopción de un nuevo sistema gana impulso, ganarán confianza en que muchos otros lo adoptarán. Conforme se generaliza esta noción, la masa de consumidores potenciales abandona sus temores y se deja arrastrar. Esto, por supuesto, hace que otros obtengan más confianza lo cual fortalece el efecto arrastre.

A mediados del siglo XIX, no existía una distribución de letras que prevaleciera en los teclados de las máquinas de escribir. En 1873, Christopher Scholes ayudó a diseñar una distribución nueva y mejorada que se vino a conocer con el nombre de QWERTY, tomado del orden de las letras de la línea superior izquierda. En 1904, la empresa de Nueva York, Remington Sewing Machine Company, producía máquinas de escribir con este teclado, que se convirtió en el estándar del sector (norma del sector o diseño dominante). Se han desarrollado nuevas distribuciones de letras en los teclados como el DSK (Dvorak's Simplified Keyboard), que reduce en más del 50% la distancia que tienen que recorrer los dedos al escribir. Un mismo texto se puede escribir con un ahorro de tiempo del 5 a 10% si se utiliza DSK. Pero QWERTY es el sistema establecido. Casi toda la gente que escribe a máquina utiliza este sistema, por lo que es el que todos aprendemos y somos reacios a aprender otros teclados. Por lo mismo, los fabricantes de máquinas de escribir y de teclados siguen con QWERTY. El círculo vicioso se ha cerrado. Detrás del QWERTY hay muchas máquinas de escribir, muchos teclados de ordenador, y muchos mecanógrafos.



Ningún usuario individual estaría dispuesto a pagar el coste de cambiar de convención social, de manera que la falta de coordinación de las decisiones individuales nos mantiene ligados al sistema QWERTY. El problema se conoce con el nombre de efecto *bandwagon*, que podríamos traducir como el efecto imitación (o arrastre) y se puede ilustrar con el gráfico adjunto. En el eje horizontal se muestra la fracción de mecanógrafos que utilizan QWERTY. En el eje vertical se ve la probabilidad de que un mecanógrafo nuevo aprenda QWERTY en lugar de DSK. Como se observa en el dibujo, si el 85% de los mecanógrafos utilizan el sistema QWERTY, hay una probabilidad del 95% de que un mecanógrafo nuevo elija aprender QWERTY y sólo un 5% de que elija DSK. La mayoría de los mecanógrafos nuevos aprenderán DSK siempre y cuando QWERTY tenga menos del 70% del mercado.

Cuando las proporciones de usuarios que utilizan cada uno de estos sistemas son constantes, nos hallamos ante un equilibrio del juego. Demostrar que este juego converge hacia un equilibrio no es tarea fácil. La elección aleatoria de cada mecanógrafo nuevo perturba el sistema constantemente. Recientemente, potentes herramientas matemáticas, concretamente la teoría de la aproximación estocástica, han permitido a los economistas y estadísticos demostrar que este juego sí que converge hacia un equilibrio dinámico (Dixit y Nalebuff, 1991).

Pasamos a describir los posibles resultados. Si la proporción de mecanógrafos que utilizan QWERTY excede el 72%, se da la expectativa de que una fracción todavía mayor aprenda QWERTY. El espacio de QWERTY se expande hasta alcanzar el 98%. En este punto, la proporción de nuevos mecanógrafos que aprenderán QWERTY iguala a su predominancia en la población, el 98% y, por tanto, no se produce ningún crecimiento adicional. Si el número de mecanógrafos que utilizan QWERTY es superior al 98%, la expectativa es que el número decaiga otra vez de vuelta a 98%. Siempre habrá un número reducido, no superior al 2%, de mecanógrafos nuevos que elegirán aprender DSK porque están interesados en la tecnología superior y no les preocupa la cuestión de la compatibilidad.

A la inversa, si la proporción de mecanógrafos que usan QWERTY cae por debajo del 72%, surge la expectativa de que el DSK acabe dominando. Menos del 72% de los mecanógrafos nuevos aprenden QWERTY, y la consiguiente caída en el uso de este sistema da, a los mecanógrafos nuevos, un incremento todavía mayor para aprender el mejor teclado el DSK, no habrá razón alguna para aprender QWERTY y este sistema desaparece del todo.

La aritmética sólo nos dice que terminaremos en uno de estos dos resultados: todo el mundo utilizando DSK o el 98% utilizando QWERTY. Pero no nos dice cuál de los dos ocurrirá. Si partiésemos de cero, todo jugaría a favor de que el DSK fuese el teclado dominante. Pero no partimos de cero, la historia cuenta. Existe la posibilidad de conseguir que todo el mundo mejore, pero eso requiere de una acción coordinada.

Fuente: Dixit, A. y Nalebuff, B. (1991): *Thinking Strategically: The Competitive Edge in Business, Politics and Everyday Life*, WW Norton & Co., Nueva York.

Feedback (retroalimentación) positivo. Las tecnologías sometidas a efectos de red tienen un feedback positivo: a medida que la base de usuarios aumenta, mayor es el número de usuarios a quienes les merece la pena adoptar la tecnología. El feedback positivo hace a los fuertes más fuertes y a los débiles más débiles. Con el tiempo, el producto adquiere masa crítica y se hace con todo el mercado.

El feedback positivo no es completamente nuevo; prácticamente todos los sectores industriales atraviesan una fase de feedback positivo al comienzo de su evolución. Ford fue más eficiente que las pequeñas empresas automovilísticas debido en gran parte a su escala. Esta eficiencia propició su crecimiento. Esta fuente de feedback positivo se conoce como economías de escala en la fabricación: las grandes empresas tienden a tener menores costes unitarios (al menos hasta cierto punto). Desde la perspectiva actual podemos referirnos a estas economías de escala tradicionales como economías de escala por el lado de la oferta.

A pesar de sus economías de escala por el lado de la oferta, Ford nunca creció tanto como para apoderarse de todo el mercado automovilístico, ya que las tradicionales economías de escala basadas en la fabricación se habrían generalmente agotado a niveles muy por debajo del dominio total del mercado. En otras palabras, el feedback positivo basado en economías de escala por el lado de la oferta se topó con sus límites naturales, punto en el cual el feedback negativo asumió el control. Con frecuencia estos límites surgían por las dificultades de controlar empresas muy grandes.

En la economía de la información, el feedback positivo ha aparecido de una manera nueva, más virulenta, basada en el lado de la demanda del mercado. Los clientes de Microsoft valoran su sistema operativo porque es utilizado ampliamente, es de hecho el estándar del sector de los PC. Los sistemas operativos rivales sencillamente no tienen la masa crítica suficiente para suponer ni siquiera una amenaza.

Los productos con feedback positivo aportan rendimientos crecientes, que se apoyan en tres factores (Arthur, 1996a): los costes sumergidos, el efecto de red y el aprendizaje. En las industrias tradicionales, los productos acarrean una gran parte de materias primas y una parte relativamente pequeña de conocimiento o de información. En los productos con feedback positivo, ocurre lo contrario.

Los productos con feedback positivo contienen muchos conocimientos. El principal componente de los costes fijos en la producción de productos con feedback positivo (por ejemplo, los de alta tecnología) son costes sumergidos (*sunk costs*) en investigación y desarrollo, es decir, costes no recuperables. Si invertimos en un nuevo edificio de oficinas, y decidimos después que no lo necesitamos, podemos recobrar parte de los costes vendiendo el edificio. Pero si el fármaco en el que invertimos decenas de millones de euros en su desarrollo fracasa, ¿dónde está el mercado de segunda mano para los fármacos fracasados? Los costes sumergidos generalmente hay que pagarlos de entrada, antes de que comience la producción. Los costes variables de producción de información también tienen una estructura poco corriente: el coste de producir copias adicionales generalmente no aumenta, incluso en el caso de que se hagan muchísimas copias. En condiciones normales, no hay límites naturales a la producción de copias

adicionales de información. Además, el vendedor de información no paga prácticamente nada por distribuir una copia adicional gratuita.

Las redes son otra importante consideración de los productos con feedback positivo o basados en el conocimiento. El producto se beneficia al aprovechar los efectos de red. Una tercera razón que explica la existencia de rendimientos crecientes es que el uso de muchos de estos productos es complicado y difícil de aprender. El aumento del número de usuarios con conocimientos específicos sobre el producto, favorece el feedback positivo, ya que un nuevo usuario logrará un mejor servicio postventa además de los consejos de otros usuarios experimentados. Los usuarios aprenden unos de otros y se benefician de activos o instituciones comunes.

Tanto las economías de escala por el lado de la demanda como de la oferta han estado presentes durante mucho tiempo. Pero la combinación de ambas que ha surgido en el sector de la información es algo nuevo. El resultado es una ‘doble fuerza’ en la cual el crecimiento por el lado de la demanda reduce el coste por el lado de la oferta y hace que el producto sea más atractivo para los demás usuarios –acelerando aún más el crecimiento de la demanda–. El resultado es un feedback positivo especialmente fuerte, lo que propicia que se creen o destruyan industrias enteras con mucha mayor rapidez que durante la era industrial.

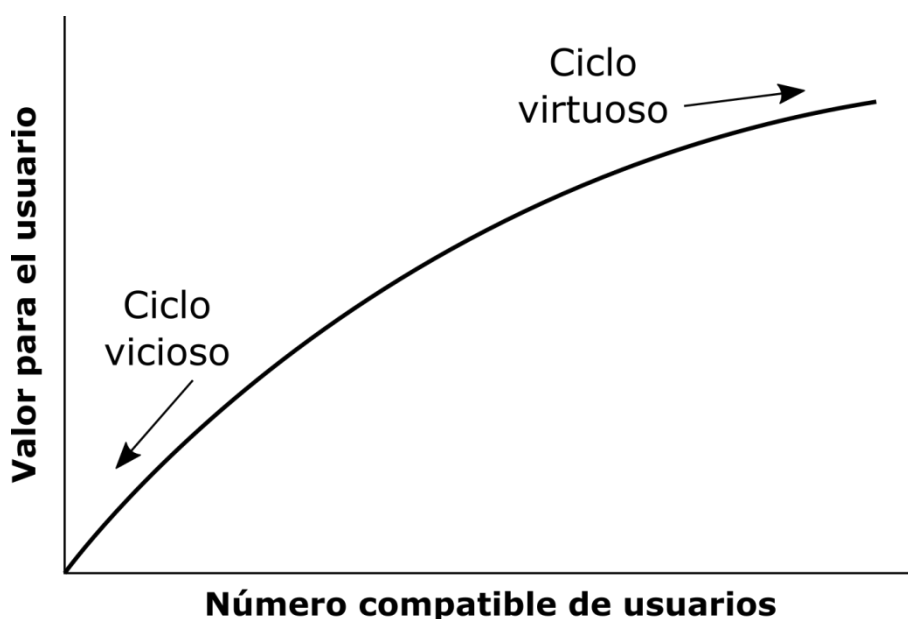


Figura 12: La popularidad añade valor en la economía de redes (Shapiro y Varian, 1999)

La relación positiva entre la popularidad y el valor queda ilustrada en la figura 12. La flecha en la parte superior derecha de la curva representa un ciclo virtuoso: el producto popular con muchos usuarios se hace cada vez más valioso para cada usuario, atrayendo a más usuarios. Por otra parte, al aumentar el número de usuarios, más conocido resulta el sistema y mejor se comprende su base de conocimientos, lo que contribuye a eliminar la incertidumbre atrayendo a aquellos clientes aversos al riesgo (Arthur, 1996a). La flecha en la parte inferior izquierda de la curva representa un ciclo vicioso: una espiral de la muerte en la cual el producto pierde valor cuando es abandonado por los usuarios y que,

eventualmente, va a dejar abandonados a esos intransigentes que se resisten hasta el final, debido a su preferencia especial por el producto o a sus altos costes de cambiar.

La masa crítica de usuarios es el tamaño mínimo que debe tener la base instalada para que a los potenciales usuarios les compense incorporarse a la misma (Rohlf, 1974), es decir, el tamaño mínimo requerido para iniciar el feedback positivo y que el proceso, por tanto, se autosostenga. La propiedad de feedback positivo también indica que las acciones estratégicas o sucesos casuales que confieren ventaja a una tecnología nuclear al comienzo de su vida se amplían para otorgarle una posición dominante. Cuando dos o más empresas compiten por un mercado donde hay un fuerte feedback positivo, solamente una de ellas puede erigirse en vencedora. Los economistas dicen que un mercado como ese es tippy, queriendo indicar que puede inclinarse en favor de un participante o de otro. Es muy improbable que todos sobrevivan. En su forma más extrema, el feedback positivo puede derivar en un mercado donde el vencedor-se-lleva-todo (Frank y Cook, 1995), en el cual una sola empresa o tecnología derrotan a todas las demás (Schilling, 2002).

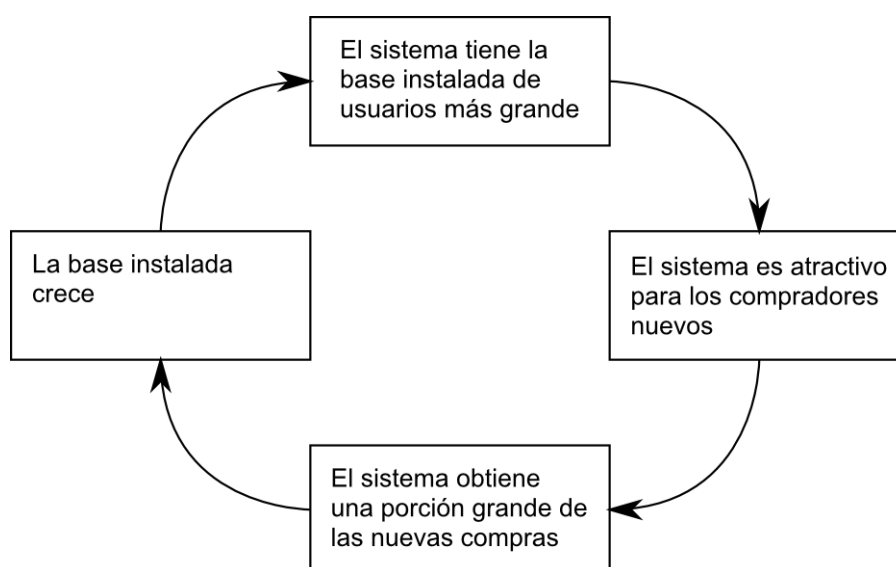


Figura 13: Retroalimentación positiva (Saloner et al., 2001)

El efecto de red puede provocar que la tecnología nuclear B derrote a la tecnología A, con la que compite, aun cuando A sea de una calidad superior a B. Esto ocurrió con el vídeo VHS, que acabó derrotando a *Betamax*, a pesar de la superioridad tecnológica de este último (Robertson y Gatignon, 1986). Así pues, en el momento en que un producto de red se adelante lo suficiente, tiende a seguir haciéndolo todavía más (figura 13). No obstante, al tratarse de un problema de ‘masa crítica’, el fenómeno de las externalidades de red es evidentemente reversible: la disminución de la base de usuarios de una base instalada dada supone un decrecimiento de su utilidad.

No obstante, la heterogeneidad en las preferencias de los consumidores, esto es, la valoración de los distintos sistemas incompatibles puede hacer que varias redes coexistan de forma simultánea. Ello se debe a que determinados usuarios puedan preferir las ventajas intrínsecas de un producto, aunque esto implique pertenecer a una red de menor tamaño. En el caso de los ordenadores personales, Apple es un buen ejemplo (Van Hove, 1999).

VALOR	Costes de cambio	Valor intrínseco del nuevo producto central
	Productos complementarios	
	Base instalada (conectividad)	
	Valor intrínseco del producto central	
SISTEMA EN USO		NUEVO SISTEMA, INCOMPATIBLE

Figura 14: Componentes del valor (adaptado de Schilling, 2008)

A pesar de las enormes ventajas que proporciona la tecnología nuclear cuando se convierte en el estándar del sector, no siempre se puede evitar la entrada de nuevos participantes; por ejemplo, los discos compactos sustituyeron a los de vinilo. Si la nueva tecnología tiene un valor intrínseco que supera ampliamente el valor total de la antigua, algunos usuarios innovadores la adoptarán y sacrificarán los beneficios de la compatibilidad con la actual base instalada (figura 14). Si muchos nuevos usuarios toman esta decisión, construirán su propia base instalada para la nueva tecnología. La adopción masiva de la nueva tecnología nuclear deja atrás a los compradores que poseen la tecnología antigua. Si la base instalada se estanca y a la larga declina, los obliga a asumir los costes de cambio en los que incurren al pasar a la nueva tecnología o vivir con beneficios reducidos por las externalidades de red.

Establecer estándares. Los estándares son descripciones técnicas detalladas, elaboradas con el fin de garantizar la interoperabilidad entre elementos construidos de forma independiente, así como la capacidad de replicar un mismo elemento de manera sistemática y lograr economías de escala. Por tanto, un estándar incluye un conjunto de especificaciones técnicas a las que los productores se adhieren cuando fabrican un producto o componente. El estándar denota las especificaciones técnicas para calidad, referencia, compatibilidad, adaptabilidad y conectividad que son requeridas para el ensamblaje apropiado de productos complejos (Grindley, 1995).

Los estándares técnicos son especificaciones que determinan la compatibilidad de diferentes componentes en un sistema. Un estándar puede ocurrir a nivel de producto (ordenador personal) o a nivel de componente del producto (ranura USB), así como a nivel abstracto (motor de gasolina) o a nivel más concreto (motor de combustión interna V-8) (Murmann & Frenken, 2006). En general, un estándar (a nivel de producto complejo) puede incluir varios estándares (subsistemas y componentes); por ejemplo, el PC agrupa, entre otros, los siguientes: Wintel, teclado QWERTY, ranuras USB y módem TCP/IP. En su forma más simple, un estándar define el encaje físico de dos componentes compatibles. No obstante, los estándares son necesarios en los productos complejos (o diseños) que tienen componentes interdependientes y también entre los usuarios de productos con externalidades de red que requieren conectarse entre sí directa (teléfono) o indirectamente (ordenadores). De esta forma, en un nivel elevado de agregación, podemos identificar estándar técnico con un sistema, como en el caso del PC, que es un estándar del sector informático, al margen de estar constituido por decenas de componentes que también son considerados estándares, aunque a un nivel menos agregado o más elemental. También puede ocurrir que un producto que se mantenga en el mercado con una cuota reducida esté formado por varios estándares técnicos individuales sin que se pueda decir que el producto sea un estándar del sector (por ejemplo, el Mac de Apple).

El surgimiento de un estándar es importante por varias razones (Lee *et al.*, 1995; Hamel & Prahalad, 1994; Hill & Jones, 2009). En primer lugar, la falta de un estándar es un obstáculo para su difusión, ya que los fabricantes de componentes y productos complementarios deben desarrollar distintos modelos para las diferentes variedades de estándares, lo que se traduce en precios altos y un mercado lento en despegar (Gates, 1995). Un estándar garantiza la compatibilidad o interoperabilidad de los componentes, generando un mayor valor en beneficio de los usuarios. En segundo lugar, la existencia de una variedad de sistemas incompatibles confunde a los clientes y los hace mostrarse menos proclives a comprar, pues prefieren ‘esperar y observar’ (Farrell y Saloner, 1986) hasta que surja un claro vencedor. En tercer lugar, ayuda a reducir los costes de producción. Una vez que surge un estándar, el producto y sus complementarios pueden fabricarse en masa, lo que permite aprovechar las economías de escala y el efecto aprendizaje. En cuarto lugar, la existencia de un estándar disminuye el riesgo de productores y clientes de comprometer recursos en un sistema que, al final, abandona el mercado por falta de demanda. En quinto lugar, cambia las prioridades competitivas pasando de las características y funciones del sistema a refinamientos en las prestaciones del producto, la fiabilidad y el coste.

La guerra de estándares (o de formato) se levanta de dos formas. Primero, una nueva tecnología se introduce en el mercado y es incompatible con la vieja. Segundo, los nuevos productos podrían incluir tecnologías incompatibles (VHS y Beta). El ganador se convierte en el estándar *de facto*. La batalla por lograr que un producto sea el estándar de la industria puede enfrentar al pionero con los productores de la tecnología en uso o a los pioneros entre sí. En cualquier caso, los contendientes tienen varias opciones para elegir. Ninguna garantiza el éxito. Pero el conocimiento de los diferentes factores que componen cada opción resulta esencial para lograr la mejor coordinación y la máxima coherencia al integrarlos en la estrategia del negocio. Esperar por un estándar resulta en un exceso de inercia, lo que reduce su probabilidad de emerger (Farrell y Saloner, 1986). La falta de productos complementarios también retrasa la emergencia de un estándar (Gupta *et al.*, 1999).

En la guerra de estándares, las dos primeras decisiones que deben tomar los contendientes desde la perspectiva tecnológica es si abren o cierran la tecnología y si van solos o con aliados (**Tabla 1**).

Tabla 1. Decisiones tecnológicas sobre el nuevo diseño

	SOLO	ACOMPAÑADO
DE JURE (público)	Gobierno	Comisión técnica
DE FACTO (privado)	Competir por el mercado	Competir dentro del mercado

En un proceso abierto regulado (o de jure), el estándar lo fija el gobierno o una comisión técnica. En el primer caso, el pionero puede presionar al gobierno para que regule su tecnología y sea el estándar del sector (Islas, 1999). En este supuesto, el estándar cae dentro del ‘dominio público’, lo que significa que cualquier empresa puede incorporar libremente a sus productos el conocimiento y la tecnología que lo sustentan. Si el estándar regulado no cae en el dominio público, el gobierno exige a la empresa conceder

licencias de cualquier patente que sea esencial para la fabricación del producto en términos justos, razonables y no discriminatorios (Saloner *et al.*, 2001).

Las ventajas de que la tecnología propia tenga que ser aceptada obligatoriamente por toda la industria son relevantes. El pionero que lo ha desarrollado posiblemente tenga sus activos productivos coespecializados con el estándar (Teece, 1986), lo que le proporcionará una importante ventaja competitiva. Los activos de los competidores serán, sin embargo, específicos al producto que estaban comercializando, lo que les ocasiona costes de cambio. No solamente pierden las inversiones específicas, sino que necesitarán tiempo y recursos para fabricar y comercializar el estándar. Además el pionero va por delante en el efecto experiencia.

Cuando el estándar lo fija un comité técnico, donde participan competidores, proveedores y clientes, es difícil alcanzar un consenso, por lo que se retrasa su aparición (David & Greenstein, 1990). No obstante, estos comités tienen reglas de juego limpio y franqueza, así como incentivos para escoger el mejor producto y traerlo al mercado (Farrell y Saloner, 1988). A cambio de aprobar un estándar, los comités exigen que ninguna empresa o grupo de empresas mantenga su propiedad.

El estándar de facto emerge de un proceso de mercado, es decir, de la competencia entre varios sistemas. Si bien el dominio se logra más rápido (Farrell & Saloner, 1988), esta guerra de formatos se caracteriza por la incertidumbre del resultado y los costes de cambio (Burnham *et al.*, 2003).

Todo proceso de mercado requiere tener protegido el producto de red. Para ello, su tecnología nuclear debe estar cubierta por un régimen de apropiabilidad, entendiendo por tal aquellas propiedades del conocimiento tecnológico y de los artefactos técnicos, de los mercados y del entorno legal que permiten las innovaciones y las protegen, en grados diversos, de las imitaciones de los competidores, manteniéndolas como activos que producen rentas (Dosi, 1982). No es necesario que todas, o la mayor parte, de las tecnologías (componentes) del producto de red estén protegidas, pero resulta fundamental que lo esté la tecnología nuclear. El régimen de apropiabilidad es más fuerte si una parte del conocimiento indispensable para configurar y fabricar el estándar es de naturaleza tácita (Teece, 1986) y, por lo tanto, difícil de decodificar por los competidores. También resulta de muy difícil acceso el conocimiento que se mantiene en secreto en los procesos productivos (Levin *et al.*, 1987). En lo que sigue, a la protección del estándar la denominaremos de forma genérica patente, con independencia de que la apropiabilidad sea un constructo multifuncional que, además de la patente, incluye varios mecanismos de protección, tales como el secreto industrial o el tiempo de liderazgo, entre otros (Levin *et al.*, 1987).

Una patente robusta impide inventar alrededor de la misma y protege al producto de forma eficaz, permitiendo a los pioneros construir los activos complementarios necesarios (Pisano, 2006). La historia confirma el poder de la protección de muchas patentes, particularmente para algunas invenciones que contribuyen a crear nuevas industrias basadas en la tecnología. Por ejemplo, las patentes básicas sobre telefonía de Alexander Graham Bell permitieron a Bell Telephone Company alcanzar el dominio rentable de la industria telefónica de Estados Unidos hasta que expiraron en 1894, después de lo cual, el número de competidores independientes creció más de cien veces en menos de diez años (Barnett, 1995). Por otra parte, el proceso de acumulación de patentes forma, en

cierta medida, un efecto de bola de nieve (Scherer, 1980). Por ejemplo, cuando los científicos de DuPont inventaron el nylon no solo se limitaron a patentar los superpolímeros, sino que investigaron todas las variaciones moleculares con propiedades potencialmente similares, cubriendo sus descubrimientos con cientos de patentes y, así, tratar de evitar que otras empresas desarrollasen productos alternativos. Además, cuando una empresa domina un campo tecnológico acumulando una cartera masiva de patentes, no solamente impide a las rivales entrar en el mercado, a no ser con su aquiescencia, sino que también se convierte en el comprador lógico de los nuevos conceptos relacionados y patentados por investigadores independientes. Asimismo, el monopolio que otorga la patente permite al propietario mantener el liderazgo del mercado mayoritario incluso cuando la patente haya expirado, debido a la fuerte penetración de la marca durante el período de monopolio de la misma (Gorecki, 1986).

En un proceso de mercado, además de proteger el producto, la empresa debe decidir si optar por una estrategia de apertura (competir dentro de un producto abierto a otros rivales) o por una de control propietario (competir para convertir el producto en un estándar privado). La de apertura pasaría por abrir la tecnología a otros competidores para, así, aumentar el tamaño de la base instalada. Esto es lo que ocurrió con JVC al licenciar su tecnología *VHS* a otros competidores (Hitachi, Sharp y Philips, entre otros) y ser fabricante de equipo original (para, por ejemplo, RCA, GE o Zenith). La estrategia de control propietario se basaría en explotar en exclusiva el producto, como lo hizo Sony al no ceder su tecnología *Betamax* a empresas rivales y, así, intentar acaparar todos los beneficios (Cusumano *et al.*, 1992). La estrategia de apertura incide en el valor añadido total del sector (competencia en el mercado) y la de control propietario en la cuota del valor del sector (competencia por el mercado).

No existe evidencia empírica suficiente que demuestre que una estrategia es superior a la otra (Shapiro & Varian, 1999). La estrategia de control, en caso de éxito, implicaría unos mayores beneficios para la empresa propietaria. Sin embargo, la negativa a abrir la tecnología puede significar su fracaso al limitar la difusión del producto y la extensión del mercado (Casadesus-Masanell & Ghemawat, 2006) si los usuarios temen un *lock-in* (Arthur, 1989) o si la empresa se enfrenta a poderosos rivales cuyos productos no están monopolizados y ofrecen unas prestaciones comparables. La tecnología *Betamax* de Sony, aunque de mejor calidad, perdió frente a la tecnología *VHS* y con el tiempo fue expulsada del mercado (Cusumano *et al.*, 1992; Liebowitz & Margolis, 1995). Así pues, la tecnología perdedora corre el riesgo de quedar excluido del mercado (*lock-out*), al no ser compatible con el estándar imperante (Schilling, 1998, 2002).

La apertura es una opción más prudente que el control (Shapiro y Varian, 1999). El propósito de esta alternativa consiste en renunciar al control sobre la tecnología para conseguir una mayor movilización del sector. Esta estrategia es de suma importancia cuando ninguna empresa tiene el poder suficiente como para conseguir que su tecnología domine el mercado. Igualmente, una empresa puede escoger compartir su tecnología si va retrasada en la carrera tecnológica con su rival y cree que su tecnología dominará el mercado únicamente si la comparte. Una tecnología abierta tiene mayores probabilidades de éxito, ya que, al estar patrocinada por un grupo de empresas, atraerá a más usuarios, pues saben que tienen al menos una segunda fuente de suministro.

Dentro de la categoría de apertura, conviene distinguir entre una estrategia de apertura total y una estrategia de alianzas para establecer el estándar (Shapiro & Varian, 1999).

En la situación de apertura total, cualquiera tiene derecho a fabricar productos que cumplan con la tecnología, tanto si contribuyeron a su desarrollo como si no; pero, en este caso, la tecnología no está regulada por el gobierno y, por tanto, no goza del paraguas legal para su difusión. La emergencia del PC de IBM como estándar se debió a utilizar una estrategia de apertura total (Khazam & Mowery, 1994). IBM, para sacar al mercado su ordenador personal en menos de un año, decidió construir su producto a partir de componentes disponibles en el mercado, como los microprocesadores de Intel y el sistema operativo de Microsoft. Al utilizar un sistema abierto que podían copiar otras compañías, consiguió que se convirtiese en el estándar del mercado (Gates, 1995). Ahora bien, los estándares verdaderamente abiertos se enfrentan a dos amenazas fundamentales (Shapiro & Varian, 1999). En primer lugar, si no hay un claro patrocinador, ¿quién se ocupará de fijar el rumbo de evolución del estándar? En segundo lugar, sin un patrocinador, ¿quién va a invertir los recursos necesarios para conseguir mejoras e impedir que el estándar se estanque?

Las alianzas con el propósito de promocionar un producto abarcan toda la distancia entre la apertura total y el control. En un extremo del espectro están las alianzas construidas alrededor de un patrocinador, una empresa central que cobra royalties a las demás licenciándoles su tecnología, si bien conserva sus derechos de propiedad y mantiene el control sobre la evolución de la misma. Así lo hizo Microsoft con el MS-DOS, permitió a IBM usar el software libremente, pero sin darle una licencia exclusiva ni dejarle el control sobre las mejoras futuras (Gates, 1995). Otorgar licencias a los competidores asegura al propietario un cierto control de la tecnología, al desincentivar la inversión en tecnologías alternativas. Además, la licencia es la forma que tiene de asegurarse la colaboración de otras empresas (Somaya *et al.*, 2010; Yoffie, 1996), incluso competidores, en la mejora de la tecnología, con objeto de conseguir que se convierta en estándar (Carpenter & Nakamoto, 1989). Del mismo modo, la licencia contribuye a la expansión del mercado. Al utilizar los canales de distribución adicionales de sus socios (Bekkers *et al.*, 2002), proporciona ingresos extra y permite amortizar más rápidamente las inversiones en investigación y desarrollo.

Algunas alianzas operan como un consorcio, formado por empresas independientes que cooperan para lograr que la tecnología se convierta en estándar, creando un fondo de patentes (*patent pools*) común que utilizan para coordinar interfaces, protocolos y detalles técnicos (Baumol, 2004). También intercambian información confidencial sobre la tecnología bajo el compromiso de no revelarla. El fondo de patentes permite solucionar el problema de la aglomeración (maraña) de patentes (*patent thickets*) que presentan los productos complejos. Un aparato complejo, como un ordenador, está compuesto generalmente de muchos componentes cubiertos por las respectivas patentes. Dichas patentes pertenecen normalmente a empresas diferentes, muchas de las cuales son, además, competidores directos en el mercado del producto final. Este hecho sitúa a cualquiera de esas empresas en una posición legal que le permite parar la producción de todas las demás (Baumol, 2004). La manera más efectiva de prevenir las catastróficas consecuencias de que esto suceda es crear un fondo de patentes en el que las empresas aliadas tengan garantizado (generalmente de forma gratuita) el acceso al mismo (Ouchi, 1984), y, sin embargo, a las empresas que no forman parte del acuerdo se les pueda impedir que fabriquen el producto o se les pueda cobrar por ello. Así pues, cada empresa coligada hace alguna contribución a la tecnología y, a cambio, se le permite vender con su marca la tecnología que contribuyó a crear.

El fondo de patentes también ayuda a proteger la entrada de nuevos competidores. Por ejemplo, situémonos en una industria de cinco empresas de idéntico tamaño, cada una con un departamento de I+D con fondos y personal similares a los de las otras. Cada empresa del consorcio tendrá acceso no solo a los descubrimientos de su propio departamento de I+D, sino también a los de las otras cuatro empresas. Ahora supongamos que una sexta empresa quiere entrar en el mercado, pero no es invitada a formar parte del consorcio para compartir tecnología. Teniendo a su disposición solo los productos de su propio departamento de I+D, frente a las otras empresas que obtienen los resultados de cinco, la empresa entrante se encuentra en una situación de clara desventaja competitiva (Baumol, 2004). Así pues, dichos fondos de patentes pueden servir como barreras de entrada efectivas, haciendo más difícil que las empresas que no forman parte del acuerdo entren en el mercado (Stiglitz y Greenwald, 2014).

Todas estas alianzas, por lo general, vienen precedidas de largas negociaciones sobre tres activos clave (Shapiro & Varian, 1999): el control de la base instalada, la superioridad técnica y los derechos de propiedad industrial. Cuantas más empresas participan en el consorcio o colaboran en el desarrollo de la tecnología, más incentivos tienen para impulsar la tecnología compartida y convertirla en el estándar del sector (Gomes-Casseres, 1994; Wade, 1995), si bien los costes de agencia tienden a ser altos (Bergen *et al.*, 1992).

Sin embargo, la estrategia de alianzas también conlleva varios riesgos (Shapiro & Varian, 1999; Schilling, 2008). En primer lugar, si una empresa abre su tecnología, otros productores pueden reducir el precio del producto resultante hasta el límite en el que la empresa propietaria sea incapaz de recuperar su inversión en investigación y desarrollo. Si como resultado de la competencia se reduce el precio, ningún productor obtiene márgenes suficientes sobre el producto, por lo que no tendrán suficientes incentivos para continuar desarrollando la tecnología. Por otra parte, abrir la tecnología puede producir su fragmentación conforme diferentes productores la alteren de acuerdo a sus necesidades, resultando en pérdidas de compatibilidad entre productores y la posible erosión de la calidad de la tecnología. Con el tiempo, hay fuertes incentivos entre los productores para diferenciarse entre sí con el desarrollo de extensiones propias y mantener al mismo tiempo algún grado de compatibilidad con las versiones anteriores. Finalmente, los estándares abiertos también pueden ser secuestrados por empresas que buscan extenderlos con diseños de los que poseen la patente exclusiva y, de ese modo, conseguir un control sobre la base instalada.

Estrategias para un producto de red (cómo ganar la carrera del estándar). Ya sea que una empresa enfrente o no tecnologías competidoras, su deseo es construir una participación de mercado notable. Si no logra suficiente penetración de mercado, incluso la tecnología que promueve un monopolista se estancará y probablemente desaparecerá. Si una tecnología tiene rivales para ser el estándar del mercado, tal vez deberá lograr una gran participación para evitar una espiral de muerte. Esta batalla puede adoptar varias formas:

1. *Construir una ventaja temprana (mover primero).* Una buena manera de lograr una base instalada amplia es teniendo una presencia temprana en el mercado (Schilling, 1998), ya que el pionero obtiene las ventajas de mover primero (Lieberman y Montgomery, 1988). Los clientes son leales a la primera marca (Schmalensee, 1982), máxime si tienen información imperfecta (Kerin *et al.*, 1992) y el pionero logra una buena reputación (Carpenter & Nakamoto, 1989; Gallagher & Park, 2002). Un pionero que ofrece un

diseño de calidad elevada apoyado por una campaña publicitaria agresiva puede ganar una alta cuota de mercado (Schmalensee, 1982). Los gastos en publicidad de los pioneros son más efectivos, contribuyen a la difusión del diseño y crean barreras a la entrada (Comanor & Wilson, 1979). Ahora bien, salir primero generalmente significa sacrificar algunos avances técnicos, lo que confiere a los rivales mucho más espacio para llevar a cabo un producto revolucionario incompatible (Shapiro & Varian, 1999).

2. *Antes de ir a la guerra, buscar aliados.* Si una empresa trata de apropiarse de la mayor parte del valor creado podría fracasar en tener suficiente efecto de arrastre (no podría beneficiarse del efecto *bandwagon*) como para obtener el liderazgo en el mercado. En general, se necesita el apoyo de los consumidores, de los proveedores de componentes e incluso de los competidores (Shapiro y Varian, 1999), ya que ayudan a reducir la incertidumbre sobre la aceptación del producto por el mercado (Leiponen, 2008; Gulati *et al.*, 2000). Un buen socio puede ser el compromiso de un cliente fuerte (*big fish*) de adquirir el producto (Suarez y Utterback, 1995; Suarez, 2004). Por ejemplo, IBM encargó a Microsoft que diseñara su sistema operativo, que luego utilizaría bajo licencia no exclusiva (Gates, 1995).

El deseo de incrementar el tamaño de la base instalada puede animar a las empresas a cooperar (David & Greenstein, 1990). Cuantas más empresas estén implicadas en la creación de valor, más rápido emerge el estándar (Wade, 1995). Un estándar respaldado por muchas empresas goza de mayor credibilidad. Esta credibilidad atrae más adeptos que, al provocar un incremento de la base instalada, logra verse aún más reforzada. Las probabilidades aumentan si además las empresas tienen poder de mercado (Axelrod *et al.*, 1995). De igual forma, el temor al *lock-in* por parte de los potenciales usuarios será menor, ya que existirá competencia futura dentro del mercado, lo que elimina incertidumbre.

3. *Atraer proveedores de productos complementarios.* El atractivo de muchos productos se ve afectado por el número de productos complementarios que también están disponibles (Bessen y Farrell, 1994). Una máquina de juegos Nintendo es mucho más atractiva cuando Nintendo u otras empresas amplían el número de juegos que se pueden utilizar con ella. Una tecnología nuclear útil para los consumidores depende de su disponibilidad y de si encaja con los productos complementarios. En el inicio, IBM logró el control de los ordenadores personales al atacar el eslabón más débil de Apple: el *software* de aplicaciones. IBM logró vender un número suficientemente alto de ordenadores como para justificar el *software* escrito específicamente para su sistema operativo, incompatible con el de Apple.

4. *Gestionar las expectativas.* Cuando se compete para conseguir una gran base instalada, las expectativas de los consumidores son cruciales (Shapiro & Varian, 1999). La decisión de compra viene fuertemente condicionada por las expectativas sobre la evolución del producto y por la consideración de que el mejor producto disponible en el momento actual seguirá siendo aceptado en el futuro (Choi, 1996). En cierto sentido, el producto que se espera que se convierta en estándar lo acabará logrando. De acuerdo con van Hove (1999), la estrategia del pionero únicamente será eficaz si se generan suficientes expectativas de éxito, ya que, en caso contrario, los potenciales consumidores, temerosos de los costes de cambio asociados a una elección incorrecta, no adquirirán el producto y no se logrará el necesario volumen de ventas para alcanzar un lanzamiento exitoso.

Las expectativas son importantes porque son una de las manifestaciones de los fenómenos de *feedback* positivo: a medida que la base de usuarios aumenta, mayor es el número de usuarios a quienes les merece la pena adoptar el nuevo diseño. El éxito engendra más éxito (Arthur, 1996). Con el tiempo, el producto adquiere masa crítica y domina el mercado. La clave para gestionar un *feedback* positivo está en convencer a clientes, proveedores y fabricantes de productos complementarios de que se va ganando la partida (Shapiro & Varian, 1999). En su forma más extrema, el *feedback* positivo puede lograr que una sola empresa o diseño derroten a todos los demás (Schilling, 2002), derivando en un mercado donde “el ganador se lo lleva todo” (Hill, 1997). Si las expectativas son bajas, el mercado no despegará, lo que se agrava si el diseño tiene poco valor en sí mismo o es de baja calidad.

Las expectativas se pueden crear con declaraciones grandilocuentes (Shapiro & Varian, 1999), preanuncios del producto (Besen & Farrell, 1994) y gastos de promoción y publicidad (Eisenmann, 2007). La manera más directa de gestionar las expectativas es mediante declaraciones grandilocuentes sobre la popularidad presente o futura del producto. Es por ello, que las empresas también revelarán y posiblemente exagerarán sus cifras de ventas, con el fin de convencer a posibles compradores de que ya tiene una gran base instalada de usuarios (Besen y Farrell, 1994). Las declaraciones para influir en las percepciones de clientes y consumidores necesitan apoyarse en la reputación del productor para ser efectivas (Porter, 1980).

El propietario de una tecnología particular puede tratar de hacer más lento el crecimiento de un producto rival, al ‘preanunciar’ con mucha antelación la introducción del suyo (Besen y Farrell, 1994). Esta estrategia se suele conocer como *vaporware*, porque el nuevo producto puede no existir o estar lejos de su introducción en la práctica. Por ejemplo, Microsoft prometió Windows 2000 a finales de 1997, pero el sistema operativo no llegó hasta tres años después. Estos preanuncios del producto pueden haberse hecho de buena fe, y los retrasos posteriores ser la consecuencia de complicaciones no previstas. Sin embargo, también puede ser que se hayan anunciado conforme a una estrategia tendente a atraer clientes que en otras condiciones podrían haber comprado un producto competidor, o simplemente para impedir que el competidor desarrollara ese nuevo producto. Preanunciar también conlleva riesgos, porque puede hacer que los clientes actuales de la empresa retrasen sus compras. De todos modos, como Haan (2003) constata, la estrategia de *vaporware* de prometer nuevos diseños que no existían en realidad puede ser un método efectivo y es ampliamente usado para impedir nuevos ingresos.

En una encuesta realizada a directores de marketing en distintos sectores de Estados Unidos, el 51 por 100 afirmaba que su empresa había anunciado con antelación el lanzamiento del último producto o servicio de la empresa. La antelación varía bastante: desde un mes antes de que el producto estuviese disponible hasta veinticuatro meses antes. La media se sitúa entre tres y cuatro meses. La regla general es que, si se decide anunciar con antelación el lanzamiento de un nuevo producto, debe hacerse con una antelación equivalente a la duración del proceso de adopción de un producto (Robertson, 1993).

5. *Compromisos con el precio*. En un escenario de captura de una amplia participación en el mercado, seguir una estrategia de precios agresiva también puede ser determinante. Así, el pionero puede plantearse una estrategia de precio de penetración, la cual suele ser muy eficaz (Besen & Farrell, 1994; Farrell & Saloner, 1986) para atraer a un número de usuarios suficientemente elevado y crear una base instalada amplia (Oren & Dhebar,

1985). Ese precio puede ser incluso cero (gratis total), buscando entonces los ingresos por otras vías como la publicidad, la adición de opciones o una versión de más calidad (Raju y Zhang, 2010). Asimismo, para los productos que requieren complementarios, puede resultar eficaz la estrategia de precio cautivo. Es el caso de las maquinillas y cuchillas de afeitar. En 1902, Gillette decidió vender la maquinilla de afeitar a un precio muy bajo para, así, incentivar las ventas de sus cuchillas desechables, protegidas por una patente, que vendía a un precio ligeramente más alto que el de sus competidores (Drucker, 1985). Por otro lado, ofrecer descuentos para atraer clientes importantes, notorios o influyentes es prácticamente inevitable en una guerra de diseños. Por ejemplo, Microsoft licenció a IBM el MS-DOS a un precio muy bajo. No quería hacer dinero directamente de IBM, sino sacar provecho de la concesión de licencias del MS-DOS a compañías informáticas que deseaban ofrecer máquinas compatibles con el PC de IBM (Gates, 1995). También pueden utilizarse promociones, en cualquiera de sus múltiples formas (cupones, regalos, muestras gratuitas,...). En este contexto, los pioneros pueden usar, además, contratos a largo plazo para asegurar a los clientes que no serán víctimas de aumentos de precios una vez que estén atados al producto, incluyendo también cláusulas de reducirles el precio si así se hiciese a los nuevos clientes (Katz & Shapiro, 1986). Esto será especialmente valioso si la empresa que ofrece el compromiso sabe que hay significativas economías de escala o de aprendizaje en la fabricación del producto. En estas circunstancias, la construcción temprana de una gran base de usuarios genera reducciones de costes que permiten que la empresa cumpla con su promesa de mantener o bajar el precio al mismo tiempo que mejora su rentabilidad (Besen & Farrell, 1994). En algunos casos, una buena estrategia es vender el nuevo diseño en un paquete con un producto que tiene una amplia base instalada (Eppen *et al.*, 1991).

Las condiciones que favorecen este subsidio cruzado de precios son las siguientes (Porter, 1985): a) suficiente sensibilidad al precio en el producto central; b) suficiente insensibilidad al precio en el bien rentable (producto complementario); c) estrecha conexión entre el bien rentable y el bien central y d) barreras a la entrada en el bien rentable. Entre los riesgos del subsidio cruzado se encuentran los siguientes (Porter, 1985): a) clientes selectivos, que compran el producto central por su calidad o prestaciones; b) sustitutivos del bien rentable; c) integración vertical de los clientes y d) competidores especializados que venden el bien rentable a un precio más bajo.

5. INFORMACIÓN Y CONOCIMIENTO

El conocimiento se deriva de la información y la información se extrae de los datos (figura 15). La cantidad de datos que hay en el mundo está creciendo desahoradamente, desbordando no solo la capacidad actual de las máquinas, sino también la propia imaginación de las personas.

Los datos son fragmentos de información sin procesar, la información son datos organizados en un determinado contexto y el conocimiento es la asimilación de la información y la comprensión de cómo alcanzarlo (Machlup, 1983). 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 aislados que las personas realizan por sí mismas a través de los sentidos o utilizando los instrumentos adecuados. Un hecho objetivo sería aquel que fuese verdadero para cualquiera en cualquier momento, con independencia de su perspectiva particular sobre el universo. Ahora bien, la realidad concreta del observador

–su evidencia personal– no cambia la verdad de las cosas, pero tiñe con tonalidades propias la forma en la que es percibida.

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. Están altamente descontextualizados. 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.

5.1. La información

La información son datos procesados (organizados) dentro de un 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).

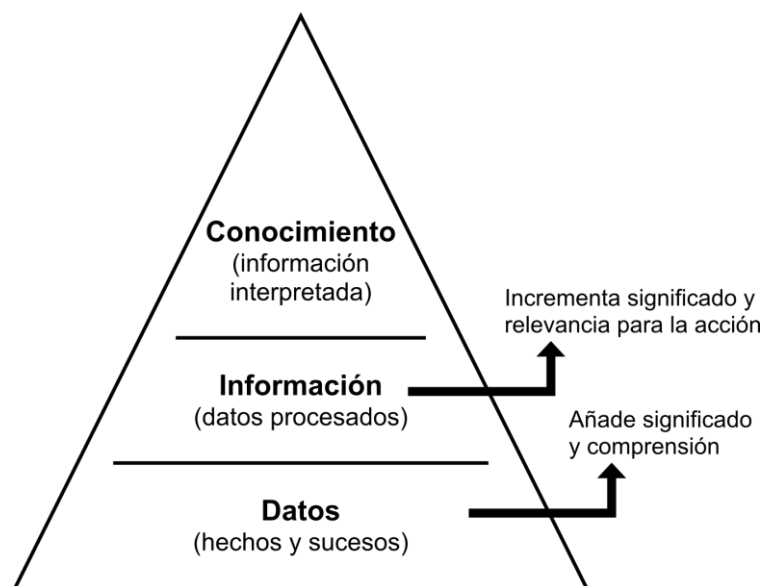


Figura 15: Datos, información y conocimiento (Jasimuddin, 2012)

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 su experiencia, las creencias que tiene, el contexto donde se realiza y la propia 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 habían percibido nada en absoluto (Kuhn, 1962). Por otra parte, la información adquiere significado dentro de las prácticas sociales de la comunidad que la utiliza (lectura 5). En suma, una característica de la información es su subjetividad, lo cual genera discrepancias acerca de su validez, utilidad e importancia.

LECTURA 5: 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 relevantes. 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 podrá obtener 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 (figura 16) que la convierten en imprescindible en el proceso de toma de decisiones. De hecho, es su principal *input*.

Características	Definiciones
Exacta	La información exacta carece de errores. En algunos casos se genera información inexacta porque se introducen datos erróneos en el proceso de transformación. Si entra basura, sale basura.
Completa	La información completa contiene todos los datos importantes.
Económica	La producción de la información debe ser relativamente económica. Los responsables de la toma de decisiones siempre deben evaluar el valor de la información con el coste de producirla.
Flexible	La información flexible es útil para muchos propósitos.

Fiable	La información fiable dependerá de algunos factores. En muchos casos, la fiabilidad de la información depende del método de recolección de datos; en otros, de la fuente de información.
Pertinente	La información pertinente es la realmente importante para el responsable de la toma de decisiones.
Simple	La información debe ser simple, no excesivamente compleja. Por lo general, no se precisa de información sofisticada y detallada. Un exceso de información puede provocar sobrecarga de información, caso este en el que el responsable de la toma de decisiones tiene tanta información que le es imposible identificar la verdaderamente importante.
Oportuna	La información oportuna es la que se recibe justo cuando se le necesita.
Verificable	Esto significa la posibilidad de comprobar que es correcta, quizá mediante la consulta de muchas fuentes al respecto.
Accesible	La información debe ser de fácil acceso para los usuarios autorizados, quienes deben obtenerla en el formato adecuado y el momento correcto para satisfacer sus necesidades.
Segura	La información debe estar protegida contra el acceso a ella de usuarios no autorizados.

Figura 16: Características de la información valiosa (Stair y Reynolds, 1999)

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 explicita 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). El verdadero valor de un proveedor de información se apoya en su labor para localizar, filtrar y comunicar información útil al cliente. 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!

5.2. El conocimiento

Una persona puede aseverar que conoce una cosa, 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 6).

LECTURA 6: 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 y bien sustentada». 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 7). 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 7: 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 8: 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.

Bien sustentada (justificada). Nuestra creencia es 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. 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 instintivamente 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 a 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 9).

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 e 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 9: 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 resta 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. No progresan de ‘certeza en certeza’, como pretendía Descartes, sino por ‘conjeturas y refutaciones’, como propuso Popper, y así progresan mejor. Por eso, no existe aquí, como observa Karl Popper, ninguna razón para desesperar de la razón. Poner el acento en el carácter falible de la humanidad no equivale a renunciar al progreso de los conocimientos, 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 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.

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. 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 (saber cómo) de la explícita (saber qué) 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.

5.2.1. Conocimiento explícito

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 se concibe como información, básicamente cuantitativa, articulada que permite

a los individuos expresar sus opiniones y creencias y confrontarlas. Ese conocimiento explícito también se le conoce como «saber qué» y se crea mediante un proceso que Dutton y Thomas (1985) lo etiquetan como *learning-by-studying*. 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. El conocimiento explícito se almacena en diferentes soportes, tales como libros, revistas, *software*, manuales de procedimiento o en los derechos de propiedad industrial.

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. El potencial de copia 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).

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. Sin embargo, al hacerlo, también conseguimos que sea más fácil para los rivales tener acceso al mismo. Promover la reproducción interna al tiempo que se limita la externa supone un desafío clave para las empresas (Kogut y Zander, 1992).

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).

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

abandonan la empresa. Los individuos pierden el conocimiento no por comunicarlo a los demás, sino por el olvido que provoca su falta de uso.

Por otra parte, el conocimiento también forma parte de los objetos (conocimiento insertado). El conocimiento insertado está incrustado en los medios 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. 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. 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.

La producción de conocimiento difiere mucho de la producción de 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. El proceso de exclusión encuentra obstáculos, debido a la deficiente definición de 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.

	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

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. Otras dificultades surgen de la posibilidad de que se utilice el conocimiento vendido de manera que escapa al marco de derechos de propiedad, como la imitación, la ampliación o la adaptación. De la misma manera, 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). En este sentido, el conocimiento es un «bien de experiencia», ya que los consumidores tienen que pasar por la experiencia de probarlo para poder evaluarlo (Nelson, 1970). 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.

5.2.2. Conocimiento tácito

Tácito significa no articulado y no codificado. 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. Se aprende por el mecanismo ‘aprender haciendo’ y no requiere el uso del lenguaje. 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 explicárselo.

Conocimiento tácito individual (*know-how*). 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 *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. 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). 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. 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 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. 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. 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 colectivo 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 colectivo 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, siguiendo unas pautas marcadas por sus propias rutinas dinámicas de aprendizaje y cambio.

6. CALIDAD

La calidad es una fuente de ventajas competitivas a la que los clientes otorgan cada vez mayor valor. Por ello se utiliza en la publicidad de los bienes y servicios, e incluso en la promoción de la imagen corporativa de la empresa. En muchos casos, la calidad se ha convertido en una restricción para competir, dado que los productos defectuosos son expulsados del mercado. En general, la calidad consiste en satisfacer (a ser posible superar) las expectativas de los clientes. La calidad de un producto no depende tan sólo de las características técnicas que posea, sino del grado en que responda a las necesidades y expectativas del cliente. Por tanto, la calidad es el resultado de comparar una realización y una expectativa. De esta forma, el éxito de una empresa depende de la precisión con la cual establezca las necesidades o expectativas del consumidor en requisitos del producto y de su habilidad para salvar la brecha entre los requisitos y sus propias capacidades para fabricar el producto esperado. Si el rendimiento del producto iguala o supera las exigencias del usuario, éste se sentirá satisfecho; en caso contrario, no.

La obtención de productos de elevada calidad es un proceso lento y no exento de dificultades. Sin embargo, un fallo importante deteriora en poco tiempo la imagen que tanto costó conseguir, y repercute de forma negativa en los beneficios de la empresa más allá del tiempo en que aconteció.

Cabe destacar que la calidad influye positivamente en la rentabilidad y en la cuota de mercado de la empresa. La relación entre calidad, cuota de mercado y rentabilidad se analizó a partir de los datos del PIMS¹. Como puede apreciarse en la figura 17, la rentabilidad de la inversión antes de impuestos (*Return on Investment*, ROI) más elevada corresponde a productos de calidad superior y cuota de mercado por encima del 28 por ciento. También se observa que la media del ROI para las empresas con productos de mayor calidad y elevada cuota de mercado es tres veces superior al ROI del competidor con pequeña cuota de mercado y producto de baja calidad. En general, para cada intervalo analizado, los productos de alta calidad proporcionan las mayores rentabilidades. La calidad también parece tener una importante influencia sobre los beneficios en casi toda clase de mercados y situaciones competitivas (Buzzell, 1983a).

Rentabilidad de la inversión ↙		Cuota de mercado		
		Menos del 13%	Del 13 al 28%	Más del 28%
Calidad del producto	Baja	11	17	23
	Media	11	17	26
	Alta	20	26	35

Figura 17: Calidad, cuota de mercado y rentabilidad de la inversión (Buzzell, 1983a)

Dimensiones de la calidad. La calidad conviene articularla en conceptos más concretos, con objeto de lograr una mejor comprensión y, así, poder gestionarla de forma más eficaz. En este sentido, Garvin (1984) sugiere ocho dimensiones de la calidad: prestaciones, peculiaridades, conformidad, fiabilidad, durabilidad, servicio, estética y percepción.

Prestaciones: Son las características funcionales primarias del producto que contribuyen a satisfacer una necesidad básica del mercado. Dos productos que realizan la misma prestación y satisfacen idéntica necesidad pueden diferir en sus características funcionales y, por lo tanto, en su calidad. En este sentido, aunque un Twingo y un Mercedes cumplen la misma función tienen distinta calidad, ya que difieren en cuanto a seguridad, confort, aceleración, consumo de combustible, resistencia a la corrosión y dureza de los materiales, entre otras características. En suma, poseen diferentes prestaciones.

Peculiaridades: Son características secundarias del producto que, aunque no contribuyen a satisfacer las necesidades básicas del cliente, sirven de complemento a las prestaciones. Ofrecen el conjunto de características adicionales que contribuyen a completar el producto que compra el cliente. En un mundo sin defectos son los toques finales los que impresionarán al consumidor. Por ejemplo, en la industria del automóvil, los clientes valoran muy positivamente los detalles de elegancia y acabado.

Conformidad: Es el grado en que las características operativas de un producto satisfacen las normas establecidas en el diseño. Un Twingo y un Mercedes fabricados de acuerdo con las especificaciones de diseño y que funcionan correctamente son coches de elevada calidad.

¹ La evaluación de la calidad en la metodología del PIMS (*Profit Impact of Market Strategies*) se hace subjetivamente: los directivos deben estimar el porcentaje de ventas de los productos de calidad superior a los de los tres mayores competidores, de los productos de calidad similar y de los de calidad inferior.

Fiabilidad: El producto ha de ser fiable, es decir, debe asegurar al consumidor confianza en su utilización durante toda su vida útil. La fiabilidad adquiere más importancia a medida que aumenta el coste de los tiempos muertos por avería o exigencias de mantenimiento, y se consigue diseñando componentes más seguros y reforzando mediante elementos redundantes los componentes más débiles del producto. Ordenadores, tractores y fotocopiadoras son algunos productos que compiten sobre esta base.

Durabilidad: Está relacionada con la vida útil del producto. Todo producto tiene una duración técnica y otra económica. La vida técnica finaliza cuando un producto averiado no se puede reparar. La vida económica termina cuando resulta más barato sustituir el producto por uno nuevo que repararlo. Esto acontece si surge un producto en el mercado que hace obsoleto y antieconómico el que ahora se utiliza o, sencillamente, cuando los gastos de mantenimiento y de reparación desaconsejan seguir utilizando el producto actual.

Para muchos productos el transcurso del tiempo no es la unidad de medida más adecuada para determinar la vida media útil. Por ejemplo, las expectativas de vida útil (medida en años) de un automóvil crecieron durante la década de los ochenta. La principal causa no fue la mejora de los componentes y del diseño de vehículo (que también contribuyeron), sino los cambios económicos. Las subidas en el precio de la gasolina y una economía más débil redujeron el recorrido medio anual de los vehículos (Garvin, 1987), lo que, a su vez, provocaba un menor desgaste y, en consecuencia, los coches se mantenían en circulación durante más años.

Servicio: Las ventas de los productos, en especial cuanto más refinados son, se sustentan de manera importante sobre la extensión y calidad de los servicios prestados a los clientes: rapidez, cortesía, asesoramiento, formación del consumidor (para evitar un uso inadecuado del producto), sistemas de garantías, financiación, asistencia en la instalación y, cómo no, reparaciones y mantenimiento. Desde esta perspectiva, conviene recordar que, cuanta más atención se dispense en la producción a la calidad del producto, menos serán las demandas del servicio para corregir los problemas subsiguientes (Takeuchi y Quelch, 1983).

Estética: La estética es una actividad de embellecimiento que complementa la estructura funcional y que está claramente vinculada a los aspectos externos del producto, con objeto de adecuarlo a los gustos imperantes o promover nuevas preferencias. No debe utilizarse indiscriminadamente para cambiar la forma de los productos, haciendo creer a los clientes potenciales que son productos diferentes, pues a medio plazo deteriorará la competitividad de la empresa en el mercado.

La estética parece, en ocasiones, evolucionar negativamente. Así, de constituir un ropaje superficial del producto, ha pasado, en muchos casos, a afectar sensiblemente a la funcionalidad del mismo. Algunos productos no sólo no aportan nada nuevo, ya que son el mero resultado de una simple operación cosmética, sino que sacrifican aspectos funcionales para no mermar su capacidad de sorprender y de estar de moda, que es, en suma, lo que se busca en estos tiempos. Esta cirugía estética se lleva a cabo algunas veces amputando algún órgano útil del objeto, con lo cual estos 'objetoides' suelen carecer de algunas importantes cualidades prácticas habituales (Ricard, 1987). En cualquier caso, la estética (aspecto, tacto, sonido, sabor u olor de un producto) es un tema de valoración personal que refleja las preferencias de un determinado individuo (Garvin, 1987).

Percepción: Los clientes no siempre tienen información completa sobre los atributos de un determinado producto. En estas circunstancias, los productos se evaluarán en mayor medida de acuerdo con su imagen, su publicidad o su nombre de marca (suposiciones sobre la calidad más que la realidad de la misma). La percepción se apoya en la experiencia personal del cliente y en la información que le llega a través de múltiples fuentes: empresa, amigos y personas de referencia, entre otros. Una consecuencia de esta diversidad de fuentes es que los clientes tienen diferente información, por lo que su percepción de la calidad del producto es distinta. De ahí que las percepciones de la calidad puedan ser tan subjetivas como las evaluaciones de la estética.

Determinar cómo los clientes definen la calidad no es tarea fácil. Las necesidades y expectativas de los clientes cambian con el tiempo y la mayor parte de ellos tienen dificultades para articular una definición (Ross y Shetty, 1985). Por ejemplo, examinando el caso de los automóviles, los datos del mercado recopilados por SRI International sugieren que las preferencias de los clientes en el mercado norteamericano evolucionaron del siguiente modo: la línea y potencia en 1970, el ahorro de combustible en 1975, y la calidad de diseño y el rendimiento en 1980 (Takeuchi y Quelch, 1983), y un mayor grado de seguridad en los años 90.

La empresa no necesita acometer las ocho dimensiones de la calidad al mismo tiempo. Es más, resulta poco menos que imposible, a no ser que pretenda cobrar precios desorbitados por sus productos. Pueden existir también limitaciones de tipo tecnológico. Por otra parte, los clientes tienden a valorar más unas dimensiones que otras. En consecuencia, hay que asignar la importancia relativa a cada una de ellas, en función de las valoraciones que éstos les otorguen; sino, el resultado será un producto mediocre (Ishikawa, 1985). Un producto puede tener un valor elevado en una dimensión de la calidad y bajo en otra. En cualquier caso, la empresa ha de preocuparse por lograr que el producto destaque en las dimensiones de la calidad que realmente valoran los clientes. Obviamente, sin descuidar el resto, en las que deben alcanzar el valor estándar del mercado.

7. INNOVACIÓN

Un invento es una tecnología recién creada capaz de realizar alguna función práctica. 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. Por tanto, la innovación es invención más explotación económica. Que se invente algo no significa, naturalmente, que vaya a ser o deba ser usado, ni aun 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.

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 poner en práctica una nueva tecnología, el innovador tiene que interactuar con un entorno formado por competidores, clientes, proveedores y el propio gobierno, entre otros (Mokyr, 1990). A corto plazo 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. Por ello, no debe entenderse como un concepto técnico, sino de raíz económica y social. La esencia de la innovación 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 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, esto es, ‘cómo hacer un nuevo producto y lograr que funcione’. La incertidumbre de mercado es la información adicional sobre los clientes potenciales, sus necesidades y expectativas, cómo hacer llegar el producto al mercado y conseguir que los clientes lo compren. Es ‘cómo vender el nuevo producto y hacer que sea un éxito comercial’ (Geroski, 1995). Por supuesto, la incertidumbre tecnológica y la de mercado están muy relacionadas. Mientras más sepa una empresa sobre lo que desean los clientes, en mejores condiciones estará para tomar decisiones sobre qué debe conformar el producto (Afuah, 1998).

Los inventos no conducen necesariamente a innovaciones tecnológicas. De hecho, la mayoría de ellos no lo hacen. Que se invente algo no significa, por lo tanto, que vaya a ser o deba ser comercializado de modo inmediato. Para que un invento llegue a convertirse en una innovación hace falta un trabajo empresarial adicional para desarrollarlo, fabricarlo y comercializarlo. En la figura 18 se recogen los años de retraso entre el invento y la innovación para una serie de productos. Este retraso varía considerablemente. La duración del retraso es difícil de medir con exactitud debido a la naturaleza desordenada e incierta del proceso innovador, aunque puede llegar a más de 50 años. Dicho retraso depende de la complejidad de los problemas que deben resolverse antes de que un invento sea económicamente viable y/o aceptado por el mercado. No obstante, varían con la industria, el producto, el mercado y los recursos utilizados (Utterback, 1974). Algunos investigadores 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. De todas formas, el intervalo se acorta cuando el inventor intenta ser también el innovador en comparación a cuando se contenta meramente con revelar a otra empresa el concepto generado 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 18: Año estimado del invento y la innovación para determinados productos (Jewkes et al., 1969)

Innovación radical versus incremental. Atendiendo a su originalidad una innovación puede ser radical o incremental. Una innovación es radical si el conocimiento tecnológico necesario para explotarla es muy diferente del conocimiento existente y lo hace obsoleto. Se dice que tales innovaciones son ‘destructoras de competencias’ (Tushman y Anderson, 1986). En el otro extremo de la dicotomía se halla la innovación incremental. En esta, el conocimiento necesario para ofrecer un producto se basa en el conocimiento en uso. Esta innovación ‘incrementa 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 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, la comercialización del magnetófono o del *walkman* por parte de Sony.
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, la industria del gas pudo sobrevivir después de perder los mercados de alumbrado residencial, comercial y urbano frente a la industria eléctrica, al encontrar un nuevo mercado en el calentamiento de espacios y en el suministro de calor para los procesos de producción.
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.

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. Las nuevas combinaciones dan lugar, pues, a nuevos negocios que detentan una superioridad decisiva en el coste o en la calidad del producto y ataca no ya los márgenes de los beneficios y de la producción de las empresas existentes, sino a sus cimientos y a su misma existencia (Schumpeter, 1942).

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. Provoca, por tanto, una acumulación creativa. El efecto de los cambios 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 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 lo comercializan (figura 19). Los problemas asociados con el riesgo, los gastos y los plazos de tiempo largos 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 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 19: 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. Rosenberg (1982) afirma que 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 los competidores que participan en el proceso de ebullición.

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 operador pudiera ver su trabajo hasta que las cuatro primeras líneas estaban mecanografiadas, en cuyo momento 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.