

TEMA 2: DESARROLLO INTERNO DE LA INNOVACIÓN

Libro 100. EI.02.DesarrolloInterno
Esteban Fernández Sánchez
Borrador actualizado: Noviembre 2020

1. PROCESO DE INNOVACIÓN TECNOLÓGICA

- 1.1. Modelo cadena–eslabón
- 1.2. Secuenciación *versus* concurrencia
- 1.3. Innovación abierta
- 1.4. Usuarios de avanzada
- 1.5. Sistema de innovación

2. MODELOS DE SEGUIMIENTO DE LA INNOVACIÓN

- 2.1. Los prototipos físicos como objetos para cruzar fronteras
- 2.2. Diseño para producción
- 2.3. CAD/CAM
- 2.4. Coordinación técnica de las funciones (casa de la calidad)
- 2.5. Proceso de evaluación
- 2.6. Cuadro de mando integral

3. DINÁMICA DE LA INNOVACIÓN

- 3.1. Interacción producto/proceso
- 3.2. Destrucción creativa
 - 3.2.1. Tecnología discontinua
 - 3.2.2. Tecnología disruptiva
 - 3.2.3. Inercia estructural

4. FACTORES ECONÓMICOS QUE INDUCEN LA INNOVACIÓN

5. CLAVES DEL ÉXITO DE UNA INNOVACIÓN

1. PROCESO DE INNOVACIÓN TECNOLÓGICA

No existe un único modelo explicativo del proceso innovador. El punto de vista tradicional considera que la innovación se lleva a cabo a través de tareas independientes, aisladas unas de otras, que se inicia con la investigación básica y finaliza con la innovación. El propio orden de los términos Investigación y Desarrollo (I+D) refuerza este planteamiento conceptual. Desde esta perspectiva la innovación como un proceso lineal, es susceptible de planificación, programación y control, que puede desagregarse en actividades independientes para simplificar su gestión (figura 1). La idea de linealidad del proceso está igualmente implícita en la distinción entre innovaciones basadas en el tirón de la demanda (*demand-pull*) o en el empuje de la tecnología (*technology-push*), ya que estar en una situación de tirón o de empuje implica la existencia de un proceso con un comienzo y un final, así como alguna conexión directa entre ambos (Kline, 1985). En este modelo la ciencia adquiere prioridad, no sólo como una condición necesaria, sino también suficiente, para el éxito de la innovación tecnológica, lo que dio lugar a que se igualase la I+D organizada (es decir, departamento de investigación y desarrollo) con el proceso de innovación mismo (Brooks, 1994). De aquí se derivan dos importantes consecuencias: 1) una empresa con un gran centro de investigación tendrá una ventaja respecto a sus rivales de menor tamaño y 2) el cambio tecnológico dependerá de los progresos realizados en la ciencia, por lo que en las áreas con una base científica estable el cambio tecnológico será más lento que en aquellas en las que está creciendo (Kamien y Schwartz, 1982). Por otra parte, el modelo niega la influencia de instituciones, comportamientos de otras empresas o los factores relacionados con la demanda y la educación.

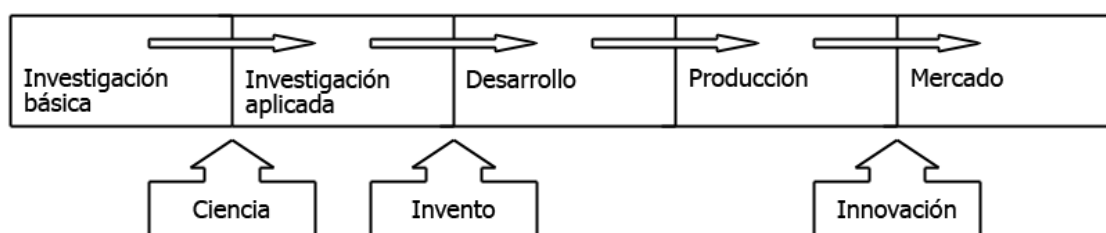


Figura 1: Modelo lineal

Ahora bien, la historia durante los dos últimos siglos de países como Inglaterra, Francia, Estados Unidos, Japón y Rusia pone de manifiesto que unas instituciones científicas de elevada calidad y un alto grado de creatividad no han sido ni una condición necesaria ni suficiente para el dinamismo tecnológico (Rosenberg, 1982). Con demasiada frecuencia la tecnología va antes que la ciencia. Igualmente, han sido las PYMEs, y no las grandes empresas, las que han introducido en el mercado algunas de las innovaciones más importantes, que resultaron trascendentes para la humanidad (Jewkes *et al.*, 1969).

La separación de las diferentes actividades que participan en el proceso innovador retrasa la introducción de innovaciones en el mercado, a la par que merma su eficiencia (Clark *et al.*, 1987). Se hacen necesarios la proximidad y la interacción entre las diferentes etapas para transferir conocimiento tácito y afrontar los problemas con una visión global del proceso. La esencia del proceso innovador es, por un lado, el solapamiento de las distintas actividades (lo que hace difícil identificar cada una de ellas con precisión y, más aún, desagregarlas en etapas independientes) y, por el otro, las frecuentes retroalimentaciones

entre las diferentes etapas. La retroalimentación resulta necesaria para evaluar los resultados y diseñar los próximos pasos.

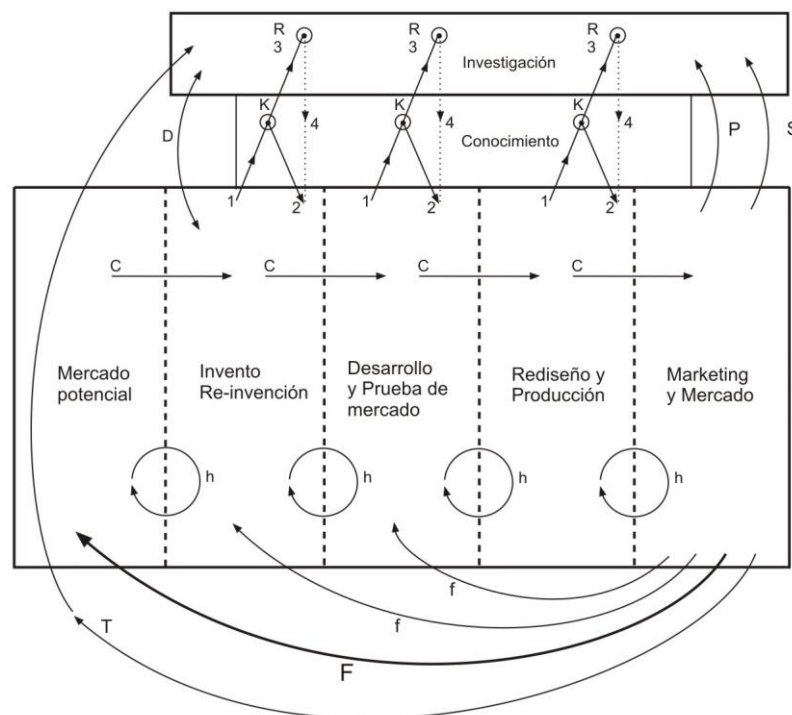
Esta concepción lineal del proceso innovador ha llevado, por otra parte, a políticos, investigadores, directivos y economistas a concentrarse exclusivamente en las innovaciones radicales. Mientras, se ignoraba las innovaciones incrementales, que, aunque no constituyen avances tecnológicos significativos, tienen una importancia económica excepcional. De hecho, la ventaja competitiva en muchas empresas proviene principalmente de una sucesión de mejoras menores en la tecnología existente realizada fuera del departamento de I+D (Mansfield, 1968b). Algunas de estas innovaciones las aportan los técnicos de producción, el personal de mantenimiento y los operarios de las respectivas áreas funcionales. Además, no debe olvidarse la relevancia de las innovaciones no tecnológicas, cuyo origen puede estar en cualquier área o departamento de la empresa.

Si asumimos una relación secuencial entre la ciencia y la tecnología, la ciencia se puede convertir en predictor de la aparición de nuevos productos y procesos (Utterback, 1974). No obstante, si la ciencia pudiera proporcionar una base predictiva más barata para lograr configuraciones óptimas de diseño, los costes de desarrollo, que constituyen unas dos terceras partes de los gastos totales en I+D, no serían ni mucho menos tan elevados. Estos costes son tan altos porque los técnicos y los diseñadores de productos siguen topando con la necesidad de efectuar actividades de comprobación muy caras antes de tener la confianza suficiente en el desempeño de un producto (Rosenberg, 1993). Muchas de estas comprobaciones se llevan a cabo mediante ensayos programados, con objeto de verificar el comportamiento de determinado componente. En estos casos, posiblemente transcurra bastante tiempo, incluso años, hasta que se logre la respuesta científica a determinados comportamientos. Por ejemplo, mediante sucesivos ensayos se puede llegar a determinar si un producto resiste ciertas condiciones del entorno, pero se ignorará el porqué durante un cierto tiempo. En suma, en el proceso de innovación tecnológica, además de la ciencia y la investigación científica, se utilizan algunos determinados conocimientos y una forma de hacer la investigación que, bajo ningún concepto, se pueden considerar científicos.

Finalmente, al considerar el impacto de la ciencia sobre la tecnología, resulta también útil plantear dos observaciones. En primer lugar, la ciencia esencial para determinados adelantos tecnológicos es simplemente la antigua ciencia, incluso algunos investigadores ya no la consideran como tal (Mokyr, 1990). En segundo lugar, cuando la investigación científica abre un campo de posibilidades tecnológicas totalmente nuevas, generalmente se trata de un proceso de etapas múltiples, ya que no es posible pasar directamente del nuevo conocimiento científico a la producción: a menudo se necesita el descubrimiento de alguna nueva información adicional sobre el estado de la naturaleza, que contribuya al posterior desarrollo de un nuevo producto (Rosenberg, 1993). Sin esta información adicional, no sería posible finalizar con éxito el invento.

1.1. Modelo cadena-eslabón

Entre otros investigadores, Schmookler (1966) sugiere que el proceso de innovación tecnológica puede involucrar la síntesis simultánea de varios aspectos de ciencia y tecnología. En esta línea, Kline (1985) ha propuesto un modelo denominado «modelo de cadena-eslabón», con cinco cursos de actividad innovadora. Dichos caminos son vías que conectan las tres áreas relevantes del proceso de innovación tecnológica: la investigación, el conocimiento y la cadena central de innovación tecnológica (figura 2).



C = cadena central de la innovación, h = círculos cortos de retroalimentación, f = retroalimentación del mercado a las fases precedentes, F = bucle de retroalimentación relevante, K-R = conexiones desde el conocimiento y la investigación a la innovación y cursos de retorno (si el problema es resuelto en el módulo K, no se activa el eslabón 3 a R, el retorno de la investigación -eslabón 4- es problemático, de ahí la línea a trazos), D = eslabón directo entre la investigación y los problemas de invención, P = apoyo a la investigación científica desde instrumentos, máquinas y procedimientos tecnológicos, T = conexión indirecta entre mercado e investigación, S = conexión directa entre mercado e investigación.

Figura 2: Proceso de innovación tecnológica (Kline, 1985)

1. Cadena central de la innovación [C]. Existe abundante evidencia empírica que concluye que el mercado constituye la inspiración básica de la innovación. Una vez detectada la necesidad hay que buscar, entre los inventos existentes, alguno que permita satisfacerla. Por ejemplo, Henry Ford observó que el automóvil era una necesidad latente en el mercado y que tendría una demanda muy alta si lograba venderse a un precio asequible. Por ello, en 1906 proclamó la máxima «fabricaré un coche para todo el mundo» con objeto de fabricar un producto estandarizado, en grandes volúmenes, para el mercado de masas. Tal concepto era revolucionario, ya que hasta ese momento el automóvil se consideraba un símbolo de posición social cuidadosamente ensamblado por artesanos. Éstos lanzaban cada año nuevos modelos a la par que modificaban los modelos fabricados hasta la fecha haciéndolos diferentes, para que los propietarios de modelos antiguos desearan deshacerse de los que tenía y, así, adquirir los nuevos.

LECTURA 1: UN INVENTO ESPAÑOL: LA FREGONA.

El mundo se fregaba de rodillas hasta que en 1956 un español colocó un palo, una bayeta, un cubo y un escurridor de rodillos con pedal en un escaparate de Zaragoza. Valorada en 395 pesetas, era un artículo de lujo que nadie sabía para qué servía. Manuel Jalón Corominas, su creador, le puso por nombre lavasuelos, palabra que considera más elegante, pero su primer vendedor la llamó fregona y así ha pasado al diccionario. Antes, esa palabra se aplicaba a la mujer que fregaba y ahora se refiere también al utensilio.

Aunque parezca que el palo y el cubo de fregar son cosas de toda la vida, el primer problema de sus promotores fue enseñar para que servía. Para darla a conocer, organizaban demostraciones en los escaparates de las tiendas. La gente se aglomeraba, se paraba el tráfico y a veces hasta intervenían los guardias para disolver al grupo de curiosos. “Aún así, había mujeres que la devolvían porque no sabían como funcionaba”, rememora su creador.

Su primera crítica escrita data de 1958, cuando se presentó en la Feria Internacional de Muestras de Barcelona. La prensa la llamó ‘escoba-bayeta ultramoderna para ser usada por cualquier miembro de la familia’. El cronista añadía: “incluso nosotros no nos avergonzaríamos de usarla”.

A pesar de su aparente sencillez, la fregona es hija de una época en la que todo resultaba difícil. Como costaba hallar los materiales, los primeros modelos se fabricaron con unos cubos sin galvanizar comprados en Barcelona, Jalón les soldaba unas piezas en Zaragoza, los enviaba a galvanizar en Manresa y los vendía cuando regresaban a Zaragoza. “Al principio sólo hacía uno o dos cada día, hasta que llegó un pedido de cien y compré una prensa con 15.000 pesetas prestadas”. Una trastienda realquilada era su laboratorio de investigación y desarrollo. Otro problema fue encontrar un algodón torcido que absorbiese el agua y que no se deshilara. Le costó meses lograr los cabos adecuados.

En los años sesenta, los turistas y los emigrantes se iban de España cargados de fregonas. En 1980 ya vendía fregonas a 30 países, tenía 150 empleados y facturaba 1.300 millones de pesetas anuales. Fue cuando vendió su empresa a una multinacional que exporta fregonas a 60 países y las fabricará en China. Pero antes de convertirse en uno de los objetos españoles más exportados de la historia, hubo que perfeccionar el invento. En el año 1959, Jalón sustituyó los rodillos de escurrir por un embudo lleno de agujeros. “Se han fabricado más de 50 millones de unidades y no se ha logrado mejorar. Lleva 35 años en el mercado y sigue funcionando bien”, explica. En 1965 cambió la forma del cubo. Por entonces, Manuel Jalón ya había abandonado su profesión de ingeniero aeronáutico militar. “Mi cubo tiene mucho de aeronáutico. Es el primero del mundo con generatriz curva, que es la que da resistencia al fuselaje de los aviones. Eso supone un 25% menos de peso, un 25% menos de material y simplifica el transporte, porque se colocan unos dentro de otros, mientras que antes había que apilarlos”. Su cubo se ha copiado en todo el mundo. La llegada del plástico dejó su invento como hoy se conoce y fue su segunda era.

En los libros y museos se ve que los más representativos del diseño industrial español es la moto Montesa, el Talgo y la fregona. De esos tres frutos del ingenio y la ingeniería de España, sólo la fregona circula por los cinco continentes.

Fuente: Roglan, J. (2000): “Historia de una escoba mágica”, *Magazine*, 6 de febrero, pp. 56–60.

La invención puede existir en cualquier lugar, siendo la misión de la empresa transformarla en una innovación. La disponibilidad de inventos no explotados forma parte del conjunto de posibilidades de innovaciones abiertas a empresarios potenciales, individuos que se encuentran en situación de intentar introducir algo nuevo en el sistema económico (Nelson, 1968).

A veces, la innovación parte de una «reinención», es decir, mejoras y ajustes que se realizan sobre la tecnología actual con objeto de mejorar sus prestaciones y satisfacer mejor las necesidades del mercado. Por ejemplo, los ingenieros de la empresa Sunbeam Appliance descubrieron que había amplias variaciones en el número de componentes de las planchas de vapor que competían entre sí y que se vendían en los mayores mercados de todo el mundo. El número de componentes variaba entre 147 y 74, el de dispositivos de fijación entre 30 y 16 y el número de tipos de dispositivos entre 15 y 9. Los propios productos de Sunbeam caían entre la zona media y el extremo inferior de estos intervalos. Al reinventar su línea de planchas de vapor, Sunbeam prestó una cuidadosa atención tanto al número de componentes como al coste de la mano de obra empleada en el montaje. La nueva línea de productos de la empresa, lanzada en 1986, se basaba en un diseño con sólo 51 componentes y 3 elementos de fijación en sólo dos tipos. Los costes de fabricación de Sunbeam cayeron significativamente, aportándole una ventaja competitiva (Lehnerd, 1987).

El invento se concreta en uno o más prototipos, que permitan probar y analizar las características del futuro producto. Este prototipo incluye la especificación completa del producto, de sus componentes y, en el caso de un artículo industrial, de los esquemas de montaje que constituyen la base para la fabricación en serie. El objetivo del desarrollo es subsanar todos los inconvenientes posibles que tiene el prototipo con objeto de conseguir un producto fiable y rentable.

El test de mercado consiste en comercializar el producto en un mercado de prueba en condiciones similares a las definitivas para estudiar la reacción de los clientes. La gran desventaja del test de mercado se puso de manifiesto con la entrada de Procter & Gamble en la comercialización de una de sus innovaciones: las galletas blandas *Duncan Hines*. El nuevo producto se había probado, con gran éxito, en el mercado de Kansas City. Pero los competidores, antes de que Procter & Gamble alcanzara la capacidad de producción y distribución adecuadas, imitaron el producto. De esta forma, las empresas Nabisco y Keebler lograron ser los pioneros en el mercado (introdujeron sus marcas antes que *Duncan Hines*), ya que disponían de suficiente capacidad de producción. Procter & Gamble llevó a sus rivales a los tribunales por infringir la ley de patentes y ganó el juicio. Sin embargo, el producto *Duncan Hines*, por haber llegado al mercado después de Nabisco y Keebler, no consiguió nunca los resultados que eran previsibles tras el éxito de la prueba de mercado en Kansas (Robertson, 1993).

A veces, los resultados de la prueba de mercado obligan a un rediseño del producto para acceder al mercado de masas. Por ejemplo, los fabricantes de hornos microondas habían concebido el producto como sustituto de la cocina eléctrica o de gas, pero esta idea no tuvo éxito en el mercado. Posteriormente, rediseñaron el producto para reducir el tamaño de forma que pudiera colocarse sobre las encimeras de la cocina y, así, ofrecerlo como un sistema de cocina complementario, con el correspondiente éxito alcanzado (Robertson, 1993). En general, una tecnología nueva es costosa y relativamente tosca en el inicio de su comercialización, por lo que se hace necesario su rediseño para que pueda satisfacer una necesidad real. Rediseñar no significa maquillar, sino actualizar y mejorar, aprovechando al máximo lo ya existente y sustituyendo o añadiendo aquellas partes, materiales y acabados que hagan una aportación rentable. El rediseño contribuye a mejorar las características funcionales de los productos y procesos y, en muchos casos, suele ir acompañado de un *styling* (acabado formal o actividad de embellecimiento) original que satisfaga los gustos imperantes en el mercado.

En la etapa de producción, el invento aún es susceptible de más mejoras, a menudo individualmente pequeñas, pero acumulativamente muy grandes, que pueden identificarse como resultado de su implicación directa en el proceso productivo. Son una consecuencia del denominado *aprender mediante la práctica* o *aprender haciendo* (Arrow, 1962b): la mejora viene determinada por la experiencia acumulada por los fabricantes con el tiempo. Las mejoras en esta etapa son el efecto de desarrollar una habilidad creciente en la producción, que contribuye a aumentar el contenido tecnológico de los productos y procesos y, al mismo tiempo, a reducir los costes medios. A veces, un nuevo producto requiere nuevas máquinas e incluso un nuevo proceso productivo (por ejemplo, Henry Ford diseñó la cadena de montaje para fabricar el Modelo T).

La comercialización del producto también conlleva actividades innovadoras. Por ejemplo, la maquinilla de afeitar de Gillette no era mejor que muchas otras y su producción costaba bastante más. Pero Gillette no vendía la máquina de afeitar. Prácticamente la regalaba,

cobrando sólo 55 centavos y hasta 20 centavos, lo que representaba la quinta parte del coste de producción. No obstante, la diseñó para que pudieran usarla solamente con sus hojas patentadas. Hacer las hojas le costaba menos de un centavo cada una y las vendía por cinco centavos. Como las hojas podían usarse seis o siete veces, cada afeitado costaba aproximadamente un centavo: la décima parte del coste de ir a la barbería (Drucker, 1985).

Las etapas mostradas en la figura 2 representan la cadena más larga. Pertenecen a la industria pesada típica. Sin embargo, para determinados productos algunas etapas no son necesarias, con lo que la cadena se acortará. En cualquier caso, la diferencia en el número de etapas depende mucho del tiempo necesario para el desarrollo del producto y su coste.

2. Eslabones de retroalimentación [h, f, F]. Hay tres tipos de mecanismos de retroalimentación: a) el círculo pequeño de retroalimentación –representado por h–, conecta cada fase de la cadena central con su fase previa (por ejemplo, marketing y mercado con producción); b) el círculo grande de retroalimentación –representado por f– proporciona a cada fase información sobre las necesidades del mercado y c) la retroalimentación proveniente del mercado –representado por F– que proporciona información sobre la posibilidad de desarrollar nuevas aplicaciones industriales. El elemento común a todas estas retroalimentaciones es la comunicación directa entre diferentes áreas o etapas del proceso de desarrollo tecnológico (lo que mejora la creatividad), así como su interacción con el mercado (lo que permite comprender mejor las necesidades de los clientes). Todo ello da lugar al «aprendizaje mediante la interacción».

Las diferentes etapas necesitan retroalimentación inmediata sobre sus resultados, con objeto de solucionar lo más rápidamente posible aquellos fallos detectados por la fase siguiente del proceso (flecha h). Al mismo tiempo, cada etapa debe valorar su actividad desde la perspectiva del mercado (flecha f). Este comportamiento es básico para evitar desarrollar productos que representen un avance espectacular desde el punto de vista tecnológico, pero que no satisfagan ninguna necesidad del mercado. Por otra parte, los usuarios son los que primero detectan fallos en el producto o por qué no alcanza a satisfacer sus necesidades o, incluso, nuevas aplicaciones para el mismo, y su contribución al desarrollo tecnológico se conoce como *aprender usando* (Rosenberg, 1982). Esta contribución se encuentra vinculada a la imposibilidad de predecir, antes de su uso, qué niveles de prestaciones podrá ofrecer la innovación, así como qué resultados podrá alcanzar. Por tanto, el uso del producto actúa como indicador permanente de los resultados posibles, siendo también un inductor de la creatividad. Las contribuciones del usuario acontecen especialmente en las primeras etapas del desarrollo del mercado. En una economía con tecnologías nuevas y complejas, algunos aspectos del aprendizaje no son función de la experiencia adquirida al desarrollar y fabricar el producto, sino de su utilización por el usuario final. En consecuencia, las empresas que dispongan de unos vínculos estrechos y eficaces con los clientes han obtenido ventaja, aunque no hayan sido introductores del producto en el mercado (Rosenberg, 1993).

La aplicación creativa del aprendizaje por el uso puede ser importante en las industrias de alta tecnología. Esto ocurre, por ejemplo, en la industria de los ordenadores, donde el desarrollo de un *software* eficaz depende, en buena medida, de la experiencia del usuario (Rosenberg, 1982). De igual modo, la aplicación creativa del aprendizaje por el uso es particularmente importante en los bienes de equipo, pues permite determinar las características óptimas de funcionamiento de una máquina en lo que afecta a la extensión de su vida útil. Lógicamente, las innovaciones propiciadas por los clientes sólo presentan

ventajas para los fabricantes si los costes de búsqueda y adquisición son menores que los gastos en investigación dentro de la empresa más el coste correspondiente al test del mercado para un producto con una posibilidad de éxito comercial equivalente (Von Hippel, 1978).

Los resultados obtenidos en el aprendizaje por el uso son de dos tipos (Rosenberg, 1982): incorporados (si se traducen en cambios en la concepción del producto o en su propia materialidad técnica) y no incorporados (si únicamente influyen sobre los procedimientos y las reglas de utilización de los productos). En el primer caso, la experiencia inicial de una nueva tecnología conduce a una mejor comprensión de la relación entre las características específicas de funcionamiento del diseño que permiten las mejoras subsiguientes. El resultado es la modificación apropiada del diseño. En el segundo caso –conocimiento no incorporado– el conocimiento generado conduce a ciertas alteraciones del uso que no requieren modificaciones en el diseño (o sólo modificaciones triviales). En este supuesto, la experiencia prolongada con el producto proporciona información sobre su funcionamiento y características operativas, que, a su vez, aboca en nuevas prácticas que aumentan la productividad del producto, bien alargando su vida útil o bien reduciendo sus costes operativos.

De todas formas, en el caso de una innovación radical es normal que el innovador sólo perciba una aplicación limitada (Cooper y Schendel, 1976). Posteriormente se descubren mercados potenciales muy importantes (flecha F). De hecho, las innovaciones tecnológicas frecuentemente consiguen el primer espaldarazo para fines en los que no se pensó al principio o que se consideraron secundarios. Por ejemplo, Univac, la empresa que diseñó el primer ordenador, consideraba que esa máquina sólo se podía destinar a los trabajos científicos. IBM había diseñado su ordenador para los cálculos astronómicos, pero aceptó pedidos de empresas comerciales. Diez años después, alrededor de 1960, Univac tenía aún la máquina mejor y más avanzada, pero IBM poseía el mercado de los ordenadores (Drucker, 1985). Por lo tanto, los promotores de una nueva tecnología deben adoptar una estrategia de exploración de nuevas aplicaciones, con una mente abierta hacia nuevos usos y dispuesto a satisfacer necesidades del mercado inesperadas. En resumen, la comercialización de una innovación permite descubrir nuevos mercados potenciales de mayor dimensión.

3. Conexión entre conocimiento e innovación [K–R]. El eslabón entre el conocimiento y la innovación no es sólo unidireccional, como prevé el modelo lineal. Siempre que acontece un problema en una actividad de la cadena central de la innovación tecnológica, se acude al conocimiento existente. Este vínculo es la base para nombrar al modelo de Kline como modelo de cadena–eslabón. El punto de partida lo constituye la explotación de los conocimientos existentes.

Las innovaciones, incluso cuando hay un acuerdo general sobre como han de utilizarse, no funcionan como prometen. Al principio, suelen ser modelos muy toscos, que materializan ideas nuevas necesitadas de ulterior refinamiento. Por ejemplo, los aparatos de televisión de los años veinte mostraban imágenes pequeñas (de 3'7 por 5 cm), borrosas y parpadeantes. El primer ordenador electrónico ocupaba casi 167 m² de espacio y pesaba 30 toneladas (Basalla, 1988).

La invención tratará, en primer lugar, de conectarse con el conocimiento, que, de no ser suficiente o satisfactorio, obligará a llevar a cabo una investigación con objeto de crear el

conocimiento tecnológico necesario para desarrollar la innovación. Algunas tecnologías apropiadas no podían, de ningún modo, derivarse ni deducirse del *stock* de conocimiento existente. Por el contrario, hubo que llevar a cabo nuevas investigaciones, que incluso permitieron desarrollar y generar ramas de conocimiento claramente diferenciadas, antes de comenzar la producción (Rosenberg, 1993). Por ejemplo, en su trabajo sobre el sistema eléctrico, Edison se vio forzado a contratar a un matemático para que trabajara en el diseño teórico de circuitos paralelos.

La llamada al conocimiento se indica mediante la flecha 1, que liga la invención al pequeño círculo etiquetado con la letra K en la capa conocimiento. Si éste proporciona la información necesaria, se traslada al invento, lo que se refleja mediante la flecha 2. Si no existe tal información, entonces hay que realizar una investigación, lo que se indica mediante la flecha 3, que va desde el círculo K al círculo R. El retorno se refleja en la flecha 4. Los resultados de la investigación se aplican a la tecnología y se añaden, asimismo, al *stock* de conocimientos. Un conjunto similar de conexiones acontece a lo largo de toda la cadena central de la innovación.

El conocimiento se coloca entre la investigación y los otros elementos de la cadena de la innovación, ya que hace de intermediario. La investigación es la herramienta utilizada en la creación del conocimiento, pero es el conocimiento almacenado y el sistema construido por ese conocimiento el que proporciona el *input* primario para la innovación corriente. Cuando el conocimiento necesario no está disponible, se abandonará la idea o se iniciará la correspondiente investigación (Kline, 1985). El conocimiento está almacenado en libros, artículos, patentes y en el cerebro de las personas, entre otros lugares menos comunes.

4. Conexión directa entre investigación e invención [D]. Cuando la innovación tiene su origen en la investigación que se realiza en el departamento de investigación y desarrollo, la mayor dificultad para la empresa consiste en encontrar una aplicación comercial al invento, ya que al ser original no tiene un mercado definido. Para este tipo de innovaciones la investigación de mercados no resulta fiable, al no poder expresar los clientes una opinión clara sobre algo que desconocen (Morita, 1986). Es el empresario quien realmente cree en las posibilidades comerciales del futuro producto, a pesar de no tener ningún elemento de referencia.

Después de crear el magnetófono, Akio Morita (fundador de Sony) hizo demostraciones a todo el mundo en Japón. A la gente le gustaba el artefacto pero nadie quería comprarlo. No sabían que uso darle a un aparato que pesaba treinta y cinco kilos y tenía un precio de ciento setenta mil yen (un graduado universitario ganaba menos de diez mil yen al mes en una empresa). Contar con una tecnología singular no es suficiente, hay que encontrarle un mercado. Cuenta Morita (1986): “Habíamos observado –o mejor dicho, Tamón Maeda había observado– que durante el período inmediatamente posterior al final de la guerra era muy difícil encontrar algún taquígrafo, ya que mucha gente había abandonado los estudios para dedicarse al esfuerzo bélico. Hasta que se pudiera corregir esa escasez, los tribunales de Japón estaban tratando de apañárselas con un pequeño cuerpo de taquígrafos, sobrecargado de trabajo. Con ayuda de Maeda pudimos hacer una demostración de nuestra máquina en el Tribunal Supremo de Japón.... ¡Y vendimos veinte máquinas de modo casi instantáneo! Esa gente no tuvo dificultades para darse cuenta de cómo dar uso práctico a nuestro dispositivo. De inmediato vieron el valor del magnetófono y se dieron cuenta de que no era un juguete... Me pareció un paso lógico el pasar de los tribunales a las escuelas de Japón.”. El resto de la historia es bien conocida.

5. Conexión directa entre producto e investigación [S, P, T]. Algunos resultados de la innovación, tales como instrumentos, máquinas herramientas y procedimientos tecnológicos, se utilizan para apoyar la investigación científica (un ejemplo actual es el uso del ordenador en el laboratorio). Esta conexión se representa mediante la flecha P. Hay otra conexión cualitativamente diferente desde los mercados a la investigación, que puede ser directa (flecha S) o indirecta (flecha T), a través de la conexión con el mercado potencial.

Las características de los mercados en que actúa la empresa, así como el tipo de producto que demanda, inciden sobre la importancia y orientación del esfuerzo innovador (flecha S). En este sentido, las agencias del gobierno y las grandes empresas de los países desarrollados buscan áreas de innovación potencial y apoyan la investigación a largo plazo con el fin de resolver graves problemas sociales y militares (la investigación sobre el SIDA y el programa de la Guerra de las Galaxias, son dos ejemplos) o para sugerir mejoras en la calidad y el desempeño de los productos en uso (la investigación en superconductores y nuevos materiales).

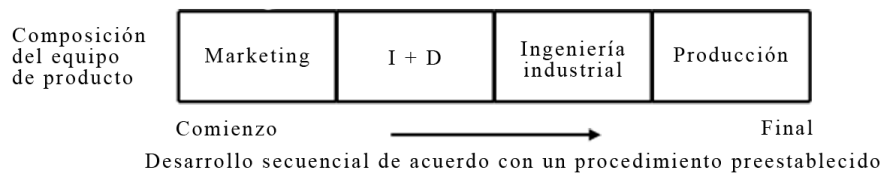
La conexión indirecta (flecha T) considera que la innovación comienza en el departamento de marketing, que detecta una necesidad del mercado (mercado potencial), mientras que la respuesta proviene del personal investigador. Dicho de otra manera, las personas de la empresa que tratan directamente con los clientes plantean un problema y el personal de investigación proporciona la solución (Kamien y Schwartz, 1982). Por ejemplo, las iniciativas comerciales llevadas a cabo por Marconi, cuyos primeros clientes fueron el ejército y la armada inglesa, descubrieron nuevos mercados para la radio y plantearon desafíos tecnológicos que tuvo que investigar con objeto de solucionar los problemas planteados, y para los que aún no existían conocimientos científicos (Basalla, 1988).

Los cinco caminos son importantes: ninguno describe por sí solo todas las fuentes de la innovación o todas las funciones necesarias para el éxito de la I+D industrial (Kline, 1985).

1.2. Secuenciación versus concurrencia

Existen dos enfoques alternativos para la puesta en práctica del proceso de desarrollo tecnológico. Por un lado, el enfoque tradicional se basa en el desarrollo secuencial de las innovaciones. Por otro lado, un enfoque más innovador se apoya en el desarrollo paralelo y en una fuerte cooperación interfuncional, denominada, en el límite, ingeniería concurrente. Takeuchi y Nonaka (1986) describen las empresas japonesas como un equipo de rugby, cuyos jugadores atacan y defienden conjuntamente. En contraste, sus competidores estadounidenses hacen carreras de relevos, donde cada uno de los corredores finaliza su tramo de carrera cediendo su testigo al siguiente corredor, único momento en que entran en interacción los corredores. En el modelo existe una clara división del trabajo a través de toda la secuencia, en la que participan diferentes agentes que actúan atendiendo lógicas también distintas (figura 3).

A. Enfoque secuencial



B. Ingeniería concurrente

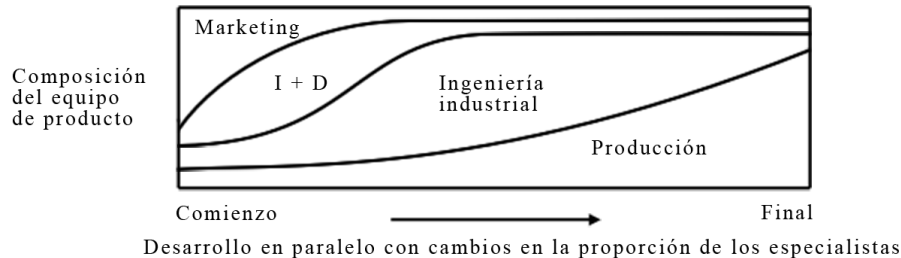


Figura 3: Enfoque secuencial versus concurrente

Enfoque secuencial. El método tradicional de puesta en práctica de una innovación se basa en la utilización de una perspectiva funcional, donde las actividades se realizan secuencialmente, de forma que las responsabilidades de creación de innovaciones se van repartiendo entre los diferentes departamentos de la empresa que participan en el proceso. El proyecto de innovación se traslada de departamento en departamento a lo largo de una especie de cadena de desarrollo, que va de un extremo a otro de la empresa. Esto ocurre, por ejemplo, de la siguiente forma:

- Marketing determina lo que quiere el cliente y lo comunica a Investigación y Desarrollo (I+D).
- I+D interpreta las necesidades del cliente y las traduce en especificaciones realistas, desde el punto de vista técnico, así como en planos, que le sirven para comunicarse con las áreas que le siguen en la secuencia.
- Ingeniería se preocupa de hacer técnicamente factible la fabricación del producto diseñado por I+D y, si no es posible, le devuelve el proyecto.
- Compras recibe los planos y la lista de materiales y su labor consiste en intentar encontrar a los proveedores adecuados para que suministren los materiales y componentes de acuerdo con las especificaciones marcadas por Ingeniería. En el caso de no encontrarlos, la lista de especificaciones volverá de nuevo hacia atrás para que Ingeniería tome la decisión de modificar el producto.
- Producción recibe los planos y especificaciones y debe diseñar el nuevo sistema productivo con la tecnología que posee o puede adquirir. En ocasiones, esto es difícil o incluso imposible. En estos casos, es frecuente devolver los planos al departamento de Ingeniería, para que modifique sus propuestas iniciales y el proceso vuelva a empezar con unas perspectivas más acordes con los recursos materiales que posee la empresa.
- Ventas intenta comunicar los resultados al cliente y vender el producto en las condiciones que lo entrega Producción. Si los resultados no son los deseados, habrá que volver atrás nuevamente para modificar alguna de las decisiones tomadas previamente.

Este proceso secuencial es el que tradicionalmente se conoce como «tirar los planos por encima de la pared», debido a que cada departamento, al finalizar su trabajo, lo pasa al

siguiente para que tome el relevo y continúe el desarrollo de la innovación. De esta forma, el proyecto evoluciona paso a paso y no entra en la etapa siguiente sin antes haber cumplido todos los requisitos de la anterior. Aunque estos trámites de verificación permiten controlar los riesgos, el método secuencial dificulta la integración y la transferencia de conocimiento tácito. Cualquier problema que surja en una etapa puede llegar a frenar, e incluso a deteriorar, el conjunto del proceso de desarrollo (Takeuchi y Nonaka, 1986). La interacción entre los departamentos o funciones es mínima: todo se realiza siempre dentro del contexto de la jerarquía funcional. Esta tendencia implica el peligro de que cada función involucrada en el proceso se limite a trabajar aisladamente en su tarea, sin tener en cuenta la necesidad de interrelación y coordinación entre los diferentes departamentos. Por ello, este enfoque se traduce en continuos retrocesos en cada una de las diferentes etapas para corregir los errores cometidos. Ello genera desarrollos muy largos y muchos problemas de calidad, debidos a la falta de comunicación y comprensión entre el diseño del producto, las capacidades de producción y las necesidades de los clientes.

Lo cierto es que esta forma de desarrollo de una innovación se ha empleado con éxito durante años en entornos tecnológicamente estáticos, caracterizados por largos ciclos de vida, tecnologías homogéneas y exigencias de productos estandarizados y poco complejos (Shenas y Derakhshan, 1994). Sin embargo, en un entorno donde la rapidez y la flexibilidad se convierten en factores clave de supervivencia, es necesario poner en práctica nuevos planteamientos, que aceleren y mejoren la calidad y eficiencia del proceso de innovación.

En un nivel más avanzado, las empresas establecen un enlace entre las diferentes áreas funcionales que facilite el flujo de información entre los expertos. Para llevar a cabo su tarea de forma eficaz, el enlace debe estar familiarizado con los diferentes aspectos del proceso y con el modo en el que éstos deben integrarse para lograr la ingeniería concurrente. Sin embargo, el conocimiento de este enlace y su capacidad para procesar información suelen ser limitados, con lo que, cuando se hace frente a actividades no rutinarias y sofisticadas, la información puede perderse entre las diferentes funciones durante la transmisión. Este camino es el menos deseable para la ingeniería concurrente y suele utilizarse con el fin de evitar cambios organizativos drásticos.

LECTURA 2: UNA VERTIGINOSA LUCHA CONTRA EL TIEMPO

En los últimos años, los teóricos y la experiencia de las empresas han puesto de relieve la necesidad de adoptar cuatro medidas fundamentales para el desarrollo de un nuevo producto:

- Dirigir de forma paralela las diversas fases de los procesos de diseño y de desarrollo –no sólo de los nuevos productos, sino también de las máquinas necesarias para fabricarlos–. Tradicionalmente, el proceso de desarrollo en cualquier empresa que no fuera japonesa –y siempre que tuviera ciertas dimensiones– ha sido secuencial, de forma que un departamento entregaba un proyecto al siguiente, como si se tratara de una carrera de relevos. Con demasiada frecuencia, el relevo se caía al suelo o era necesario volverlo atrás para ser reexaminado –ya fuera por razones de coste, rendimiento o de fabricación–.
- Conseguir que los especialistas de marketing, diseño, ingeniería, producción y el resto de los departamentos de la empresa colaboren de forma más eficaz, reuniéndolos a todos en equipos de trabajo dedicados a la tarea específica que en ese momento interesa a la empresa.
- Limitar en lo posible el grado de variación de los productos de una generación a la siguiente. Esta manera gradual de concebir la innovación es típica de los japoneses y, hasta cierto punto, también de las industrias alemanas. Este sistema facilita y acelera enormemente el desarrollo del producto y la puesta en marcha del proceso de producción y resulta un progreso frente al sistema anglosajón de dar un gran salto hacia delante en cada generación de productos.

- No ajustarse a una sola fórmula de desarrollo del producto, sino adaptar la estrategia a cada situación competitiva, económica y tecnológica concreta y conforme al tipo de industria, de empresa y de línea de productos de que se trate.

Fuente: Lorenz, C. (1991): “Una vertiginosa lucha contra el tiempo”, *Expansión*, 6 de julio, p. XVI.

En este sentido, cabe señalar que se observa la necesidad de una nueva metodología que acorte el proceso técnico del desarrollo y establezca las conexiones adecuadas entre las actividades de los distintos departamentos evitando, así, los continuos retrocesos del enfoque tradicional. Precisamente con estos objetivos surge el concepto de ingeniería concurrente o simultánea, que busca un planteamiento integrado del desarrollo del producto, con una planificación interfuncional que permita y facilite la realización en paralelo de muchas de las tareas que antes se realizaban secuencialmente.

Ingeniería concurrente. Para mejorar cada paso del proceso de desarrollo en un esfuerzo por mejorar el desempeño, una empresa puede yuxtaponer las fases. De esta forma, la ingeniería concurrente surge en oposición al enfoque tradicional. Tiene como objetivo, a través de la consideración simultánea de todas las actividades necesarias para la introducción de un nuevo producto, mejorar el rendimiento del proceso de innovación en sus tres dimensiones: 1) menor duración del proyecto de desarrollo, 2) desarrollo más eficiente y 3) calidad y prestaciones más altas (figura 4).

Para alcanzar tales objetivos, la ingeniería concurrente se apoya en dos pilares básicos: 1) el desarrollo paralelo de diferentes actividades que hasta ahora se desarrollaban secuencialmente y 2) la implicación desde el inicio de todas las funciones que contribuyen al desarrollo de la innovación formando parte de un equipo autónomo para el desarrollo de un nuevo producto.

El desarrollo paralelo o concurrente de las actividades no reduce la duración de cada actividad, pero sí disminuye el tiempo de desarrollo global. Ello es así porque propicia la existencia de frecuentes cambios de información bidireccionales entre las mismas, lo que mejora la eficiencia del desarrollo.

La implicación temprana de todas las funciones supone que los diferentes especialistas funcionales expresen sus opiniones y proporcionen sus *inputs* de información desde el principio. De esta manera, se eliminan lagunas informativas y se evitan pérdidas de tiempo inherentes a la necesidad de tener que comunicar funciones físicamente separadas.

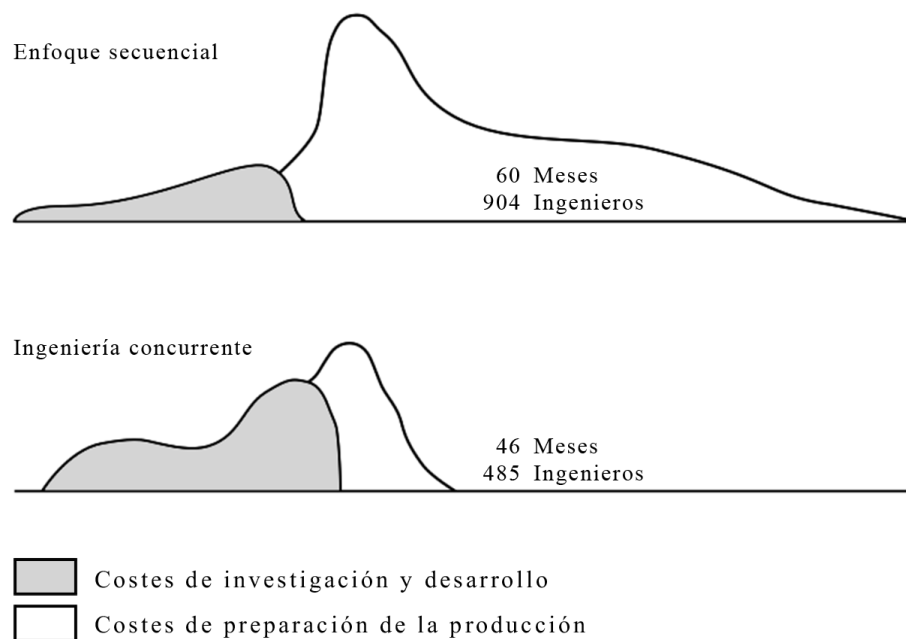


Figura 4: Duración y costes en los enfoques secuencial y concurrente

La utilización de ingeniería concurrente requiere grandes cambios en la organización interna, que no todas las empresas están dispuestas a asumir y, como consecuencia, puede observarse un continuo de caminos para el despliegue de esta metodología (Shenas y Derakhshan, 1994).

El camino más integrador consiste en la creación de un equipo autónomo responsable del producto y del proceso. El trabajo en equipo se erige como un pilar fundamental de la ingeniería concurrente. Así, los participantes en el proceso de desarrollo no sólo deben estar implicados desde el comienzo del proyecto, interactuando e intercambiando información abiertamente, sino que, además, deben colaborar estrechamente, reforzándose unos a otros en el logro de unos objetivos compartidos. Es decir, una de las características más relevantes de la ingeniería concurrente es el uso de equipos multifuncionales en los que las distintas disciplinas deben coordinar sus esfuerzos en la solución de los problemas relacionados con la innovación del producto. Este trabajo en equipo debe entenderse como un proceso caracterizado por unos intereses comunes, un alto grado de transparencia, unos riesgos compartidos y múltiples sinergias.

El equipo autónomo facilita la ingeniería concurrente. También se produce un incremento de las destrezas técnicas e interpersonales de los miembros del equipo al intercambiar conocimiento codificado y tácito. Cada individuo conoce los requisitos, no solamente de su propio campo funcional, sino también de otros campos técnicos. La oportunidad de experimentar y las sinergias tecnológicas se incrementan de forma notable gracias a la fusión de información de las competencias heterogéneas de los diferentes miembros. Por otro lado, los miembros del equipo están sujetos a una jerarquía organizativa uniforme, disminuyendo, así, la posibilidad de conflictos entre departamentos. Esta solución es la más avanzada. Por tanto, puede concluirse que la máxima expresión de la ingeniería concurrente se alcanza con el uso de equipos autónomos.

Aunque parece que la ingeniería concurrente conduce siempre a mejores resultados, se ha podido observar que, en determinadas circunstancias, su uso no es muy adecuado. Así, determinadas investigaciones se han encaminado a analizar qué situaciones o circunstancias hacen más eficaz la ingeniería concurrente, aunque existen importantes contradicciones al respecto. Algunos autores, tales como Shenay y Derakhshan (1994), consideran que, para empresas que desarrollan productos complejos, definidos por una alta incertidumbre, es eficaz la ingeniería concurrente y el uso de equipos autónomos totalmente dedicados y con su propia ubicación. Sin embargo, para empresas que desarrollan productos relativamente simples, es improbable que este tipo de equipos proporcione una solución eficaz.

No obstante, en contra de lo expuesto, existe evidencia empírica que sugiere que, mientras que la ingeniería concurrente puede ser apropiada para productos incrementalmente nuevos, esta práctica puede provocar una serie de costes ocultos que hacen que su uso no sea apropiado para productos radicalmente nuevos. Asimismo, determinados estudios recomiendan restringir el uso de la concurrencia a entornos de baja incertidumbre, ya que se ha observado que una comprensión del proceso de desarrollo, a través de la concurrencia de actividades, requiere una situación con incertidumbre limitada, donde los cambios tecnológicos sean previsibles y puedan mantenerse bajo control. De otra manera, el solapamiento puede causar un reprocesamiento sustancial del producto, que pesaría más que el tiempo ganado por la concurrencia.

1.3. Innovación abierta

De acuerdo con el viejo paradigma de la innovación cerrada, las empresas deben generar por sí solas sus propias ideas y luego desarrollarlas, construirlas, sacarlas al mercado, distribuirlas, mantener los servicios que generan, financiarlas y respaldarlas. Este paradigma aconseja a las empresas mantenerse tan independientes como les sea posible, pues nadie puede estar seguro de la calidad, la disponibilidad y la capacidad de las ideas de otros. Si quieres que se haga bien, debes hacerlo tu mismo. Algunas de las reglas implícitas de la innovación cerrada son las siguientes (Chesbrough, 2009):

- Debemos contratar a los mejores y más brillantes investigadores si queremos que el equipo más capacitado de la industria trabaje para nosotros.
- A fin de sacar al mercado nuevos productos y servicios, nosotros mismos debemos descubrirlos y desarrollarlos.
- Si los descubrimos nosotros, debemos ser los primeros en sacarlos al mercado.
- La empresa que llega primero al mercado con una innovación es en general la vencedora.
- Si superamos a los demás con nuestras inversiones en I+D, descubriremos más ideas, y las mejores de todas ellas, y así podremos liderar el mercado.
- Debemos controlar nuestra propiedad industrial, a fin de que nuestros competidores no saquen provecho de nuestras ideas.

La lógica de la innovación cerrada creaba un círculo virtuoso. Las compañías invertían en I+D, lo que conducía a numerosos descubrimientos revolucionarios; descubrimientos que permitían a dichas compañías sacar al mercado nuevos productos y servicios, incrementar sus ventas y obtener mayores beneficios, que luego reinvertían en más I+D interna, lo que a su vez conducía a nuevos descubrimientos. Y dado que la propiedad

industrial que se deriva de esta I+D interna es celosamente protegida, nadie más puede explotar tales ideas en su propio provecho.

La figura 5 despliega este paradigma de innovación cerrada. Las líneas más gruesas muestran los límites de la empresa. Las ideas fluyen hacia el interior de la empresa a través de los proyectos de investigación. Son examinadas y filtradas, y las ideas supervivientes pasan a su desarrollo y posteriormente salen al mercado. El proceso está diseñado para eliminar falsos positivos, proyectos que en un principio parecen atractivos pero luego resultan ser decepcionantes. Se espera que los proyectos supervivientes, habiendo superado una serie de inspecciones internas, tengan mayores posibilidades de éxito en el mercado.

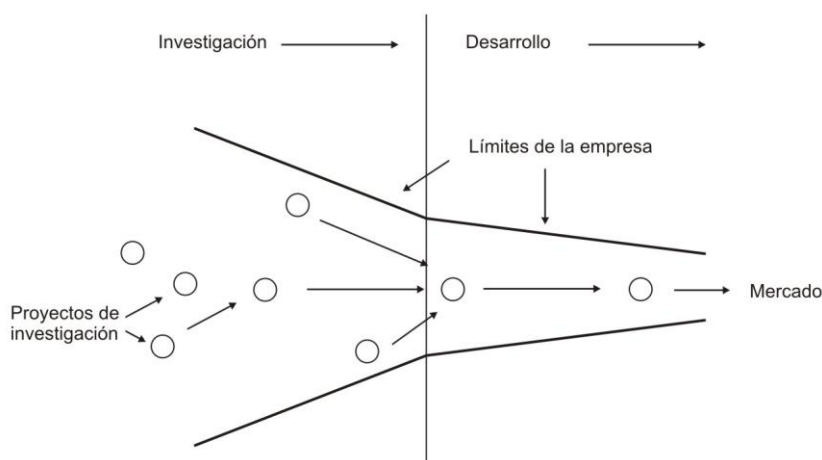


Figura 5: Paradigma de la innovación cerrada (Chesbrough, 2009)

En los últimos años del siglo XX, numerosos factores se combinaron para socavar los supuestos de la innovación cerrada (Chesbrough, 2009). Uno de esos factores fue la creciente movilidad de personas altamente cualificadas y experimentadas. Cuando esta gente abandonaba una empresa tras trabajar en ella durante muchos años, se llevaba consigo buena parte de los saberes tan duramente obtenidos en su anterior empleo, que iban a parar a su nuevo empleador que, no obstante, se negaba a pagar compensación alguna al ex empleador por semejante capacitación. Otro factor corrosivo fue el creciente número de personas formadas gracias a cursos universitarios, másteres o posgrados. El hecho de que hubiese cada vez más personas con estas características permitió que el conocimiento escapase de los silos del saber de los laboratorios de investigación central de las grandes empresas, para llegar a empresas de todos los tamaños en diferentes industrias. Un factor adicional fue la presencia de capital riesgo, que se especializó en la financiación de nuevas empresas de base tecnológica para convertirlas en empresas prósperas y bien financiadas. A menudo estas nuevas empresas se convertían en formidables competidoras para las grandes empresas ya establecidas que previamente habían financiado la mayoría de la I+D en la industria (cuyas ideas alimentaban estas nuevas empresas en su competencia por el liderazgo industrial). La lógica de la innovación cerrada fue asimismo puesta en jaque debido a la necesidad acuciante por sacar al mercado muchos productos y servicios, lo que reducía

significativamente el periodo de incubación de una tecnología específica. Además, clientes y proveedores cada vez mejor informados han cuestionado la capacidad de las empresas de aprovechar su silo de saber. Por otra parte, las empresas no estadounidenses han aumentado a su vez las fortalezas competitivas. Cuando estos factores de corrosión han impactado en una industria, los argumentos y la lógica que alguna vez hicieron de la innovación cerrada una estrategia eficaz ya no pueden ser aplicados.

La innovación abierta es el uso por parte de las organizaciones de las entradas y salidas de conocimiento para acelerar la innovación interna y expandir el mercado para el uso externo de la misma (Chesbrough, 2009). La innovación abierta combina ideas internas y externas para crear estructuras y sistemas cuyos requerimientos son definidos por un modelo de negocio. Ese nuevo modelo de negocio genera valor al tiempo que define mecanismos internos para reclamar alguna porción de ese valor. La innovación abierta presupone que las ideas internas también pueden ser reconducidas al mercado a través de canales externos, por fuera de los negocios actuales de la empresa, a fin de crear un valor adicional. Así pues, la innovación abierta implica que las ideas valiosas pueden provenir tanto de dentro como de fuera de la empresa, y pueden salir al mercado también desde dentro o fuera de la empresa. En suma, la innovación requiere, cada vez con más frecuencia, respuestas coordinadas de varias empresas. Las empresas innovadoras necesitan identificar y hacer un mapa de su ecosistema de innovación para gestionar las diferentes interdependencias (Adner, 2006). La figura 6 ilustra el proceso de innovación abierta.

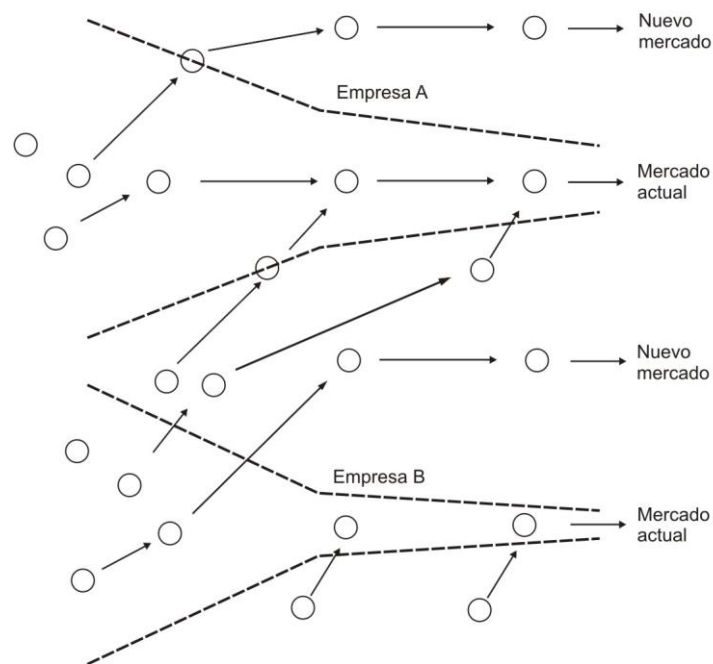


Figura 6: Paradigma de la innovación abierta (Chesbrough, 2009)

En la figura 6 algunas ideas se originan dentro del proceso interno de investigación de la empresa, pero otras ideas pueden generarse también fuera de los propios laboratorios de la empresa y luego ingresar en ellos, sea en la fase de investigación o, más tarde, en la de desarrollo. Existen fuera de la empresa muchas valiosas ideas potenciales. En la figura 6, las líneas que representan los límites de la empresa aparecen punteadas,

reflejando unos límites cada vez más permeables, la interfaz entre lo que se hace dentro de sus muros y lo que se extrae de fuera. Las ideas también pueden fugarse fuera de la empresa, siendo una de las vías principales la creación de una nueva pequeña empresa a menudo fundada con ex miembros del personal de la antigua empresa con apoyo de la misma. Otro mecanismo de fuga es la licencia. Algunas empresas incentivan a sus empleados para encontrar licenciatarios (Dávila y Esptein, 2014). No conviene olvidar que los empleados abandonan la empresa para trabajar con la competencia.

Aunque el proceso de innovación abierta elimina falsos positivos (ahora provenientes tanto de fuentes internas como externas), a la vez permite la recuperación de los falsos negativos, es decir, proyectos que en un principio parecían casi carentes de valor, pero que han resultado después sorprendentemente valiosos.

Principios de la innovación cerrada	Principios de la innovación abierta
Los mejores trabajan para nosotros Para beneficiarnos de la I+D debemos realizar nosotros mismos los descubrimientos, desarrollarlos y distribuirlos Si somos los descubridores, entonces debemos ser los primeros en sacarlos al mercado La empresa ganadora es la que saca primero una innovación al mercado Si generamos la mayor cantidad y calidad de ideas de la industria, venceremos Debemos controlar nuestra propiedad intelectual, a fin de que nuestros competidores no saquen provecho de nuestras ideas	No todos los mejores trabajan para nosotros. Necesitamos trabajar con gente brillante tanto de dentro como de fuera de nuestra compañía La I+D externa puede crear un valor sustancial; la I+D interna es necesaria para reclamar una parte de ese valor No es necesario que generemos las investigaciones para sacar provecho de ellas Edificar un mejor modelo de negocio es preferible a ser los primeros en el mercado Si hacemos el mejor uso de las ideas internas y externas venceremos Debemos sacar provecho del uso que otros hacen de nuestra propiedad intelectual, y debemos comprar la propiedad intelectual de otros cada vez que haga progresar nuestro propio modelo de negocio
Características de La innovación cerrada	Características de la innovación abierta
Ejemplos: reactores nucleares, servidores de informática Por lo general ideas internas Baja movilidad laboral Poco capital riesgo Poco y muy débil desprendimiento de nuevas empresas Poca relación con las universidades	Ejemplos: ordenadores personales, cinematografía Muchas ideas externas Alta movilidad laboral Capital riesgo activo Frecuente desprendimiento de nuevas empresas Fecunda relación con las universidades

Figura 7: Principios contrastados de innovación abierta y cerrada (Chesbrough, 2009)

La industria cinematográfica de Hollywood, por ejemplo, hace varias décadas que innova a través de una red de asociaciones y alianzas entre estudios de producción, directores, agencias de talentos, actores y actrices, guionistas, subcontratistas especializados (por ejemplo, proveedores de efectos especiales) y productores independientes. Las diferentes industrias pueden colocarse en un continuo, uno de cuyos extremos incluye industrias en las cuales prevalecen por completo condiciones de innovación cerrada, mientras que el otro extremo contiene industrias que trabajan íntegramente en condiciones de innovación abierta (figura 7).

1.4. Usuarios de avanzada

Los usuarios, según empleamos el término en este apartado, son empresas o clientes particulares que esperan beneficiarse del uso de un producto (bien o servicio). Los estudios empíricos evidencian que los usuarios están asumiendo una parte importante del desarrollo y la modificación de productos en numerosos campos (von Hippel, 2011). Adam Smith fue un observador temprano del fenómeno y señaló la importancia de “la invención de un gran número de máquinas, que facilitan y abrevian el trabajo, capacitando a un hombre para hacer la labor de muchos”. Smith añadió que “una gran parte de las máquinas empleadas en esas manufacturas, en las cuales se haya muy subdividido el trabajo, fueron al principio invento de artesanos comunes, pues hallándose ocupado cada uno de ellos en una operación sencilla, toda su imaginación se concentraba en la búsqueda de métodos rápidos y fáciles para ejecutarla”. Rosenberg (1976) exploró la cuestión centrándose en las empresas usuarias y no en los trabajadores individuales. Estudió la historia de la industria de maquinaria de Estados Unidos y descubrió que algunos tipos de máquinas básicos e importantes, como los tornos y las fresadoras, fueron desarrollados y contruidos inicialmente por empresas usuarias que tenían una gran necesidad de ellos. Las empresas de fabricación textil, los fabricantes de armas y los de máquinas de coser fueron importantes usuarios–desarrolladores de maquinaria en aquella primera época.

Enos (1962) señaló que casi todas las innovaciones de peso introducidas en el refinado del petróleo fueron desarrolladas por empresas usuarias. Von Hippel (1988) constató que los usuarios eran los desarrolladores de aproximadamente el 80% de las innovaciones más importantes realizadas en el instrumental científico y también de la mayoría de las grandes innovaciones introducidas en la fabricación de semiconductores.

Los usuarios y los productores tienden a desarrollar distintos tipos de innovaciones. Los desarrolladores de productos deben contar con dos tipos de información para tener éxito en su trabajo: información sobre la necesidad y el contexto de uso (generada por los usuarios) e información genérica sobre la solución (en la mayoría de los casos proporcionada inicialmente por los productores especializados en un tipo concreto de solución). Aunar estos dos tipos de información no resulta fácil. Tanto la información sobre la necesidad como la información sobre la solución suelen ser muy adherentes, difíciles de transferir desde el lugar en el que se han generado a otros emplazamientos (Von Hippel, 1994).

Cuando la información es adherente (pegajosa), los innovadores suelen recurrir fundamentalmente a la información de la que disponen. Una consecuencia de la asimetría típica resultante entre los usuarios y los productores es que los usuarios suelen desarrollar innovaciones que son novedosas desde el punto de vista funcional, lo que requiere una gran cantidad de información sobre la necesidad del usuario y el contexto de uso. En cambio, los productores desarrollan sobre todo innovaciones que son mejoras de las soluciones existentes para necesidades conocidas y que para su desarrollo exigen un profundo conocimiento de la información sobre la solución. Del mismo modo, los usuarios suelen disponer de una información superior a la que tienen los productores sobre los modos de mejorar las actividades relacionadas con el uso, que se apoya en el «aprender usando». En este sentido, Riggs y von Hippel (1994)

descubrieron que la probabilidad de desarrollar innovaciones que permitan que los instrumentos realicen tipos de cosas cualitativamente nuevas por primera vez es mucho más alta para los usuarios que para los productores. Por el contrario, los productores suelen desarrollar innovaciones que permiten a los usuarios hacer lo mismo que ya hacían pero de una forma más cómoda o fiable.

De los usuarios innovadores sabemos lo siguiente. En primer lugar, los estudios muestran que pocos usuarios intentan proteger sus innovaciones de los imitadores; sus innovaciones son libres para que cualquiera las utilice. En segundo lugar, la mayoría de las innovaciones de los usuarios no son adoptadas por otros usuarios ni por los fabricantes de productos de consumo. En tercer lugar, una cifra significativa sí es adoptada por terceros. En conjunto, estos hallazgos significan que las empresas que fabrican productos de consumo tienen un inesperado «frente» de diseños de la innovación libre que pueden utilizar como una fuente importante de materia prima para los procesos de innovación comercial en una gran variedad de áreas (von Hippel *et al.*, 2011). Este nuevo paradigma de innovación, en el que los usuarios de los productos desempeñan un papel central, tiene tres fases (figura 8).

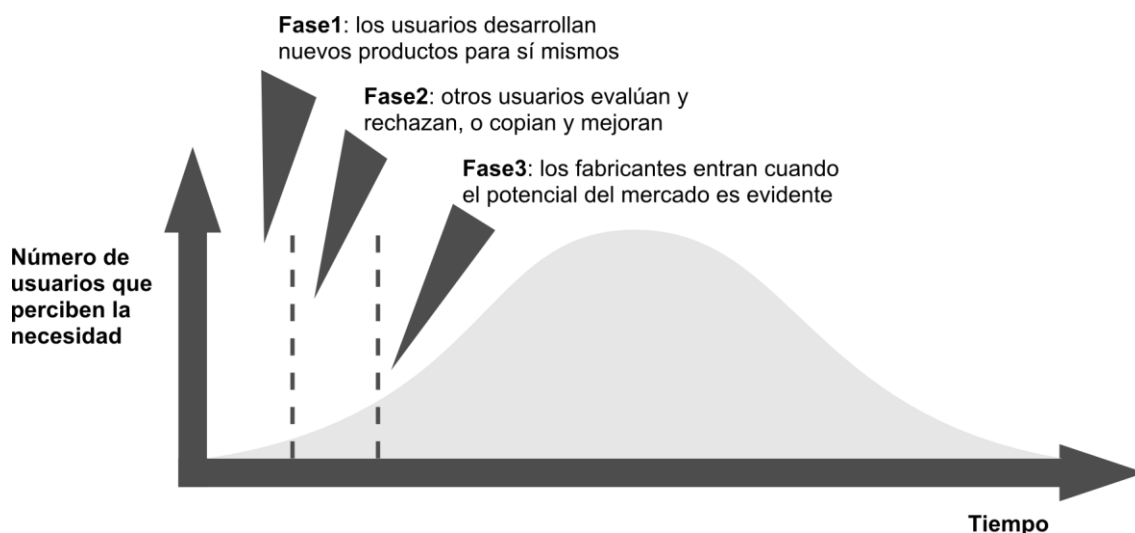


Figura 8: Un nuevo paradigma de innovación (von Hippel *et al.*, 2011)

Fase 1. Inicialmente, los mercados de productos y servicios con nuevas funcionalidades son pequeños e inciertos. Por ejemplo, al principio, nadie sabía si habría un mercado rentable para el primer *skateboard*, ya que fue desarrollado y construido por jóvenes para su propio uso. Lo hicieron desmontando un tipo de patines sobre ruedas que sujetaban a los zapatos y clavando las ruedas de los patines a una madera (de ahí el término *skateboard*).

Fase 2. La mayoría de las innovaciones desarrolladas por los usuarios son de interés sólo para quien las origina; sin embargo, algunas de ellas tienen un potencial mayor. Dado que muchos de los diseños están por lo general disponibles libremente, otros usuarios pueden realizar copias con total libertad, probándolas y quizá también mejorando los diseños. En qué medida esta difusión viral se transmite a otros usuarios, tanto a través de comunidades en la Web como a través de otras comunidades, emite una señal a los fabricantes con respecto a cuáles de los nuevos diseños y funciones constituirán la base para un nuevo producto comercial o línea de producto rentable. En otras palabras, los usuarios no solo están desarrollando nuevos productos, sino que

también están proporcionando datos de estudios de mercado a cualquier fabricante que esté lo suficientemente alerta como para reunirlos y evaluarlos.

Fase 3. Los productores empiezan a decidir que la información sobre el diseño y la función del nuevo producto y sobre cuántos individuos estarán dispuestos a comprarlo ha alcanzado niveles aceptables para su perfil de riesgo. Por ejemplo, solo cuando la popularidad del *skateboard* empezó a extenderse entre los jóvenes, las empresas empezaron a interesarse por fabricar tablas de *skateboard* comercialmente. Por lo general, los pequeños productores son los primeros en entrar porque atienden nichos de mercado. Algunos de ellos son empresas de nueva creación fundadas por los propios usuarios–innovadores. A continuación entran las grandes empresas, por lo general mediante adquisiciones, si el mercado sigue creciendo. Los productores aunque no hayan desarrollado las ideas y los prototipos iniciales de innovaciones novedosas en términos funcionales, también aportan su granito de arena. Pueden mejorar los diseños desarrollados por los usuarios con el fin de hacerlos más fiables y fáciles de utilizar –y, por lo general, rediseñarán los productos con el fin de que se adapten mejor a la producción en masa–.

Las innovaciones de interés general impulsadas por los usuarios pueden transferirse a los productores para su difusión a gran escala, utilizando tres métodos. En primer lugar, los productores pueden buscar activamente las innovaciones desarrolladas por los usuarios líderes que constituyan una buena base para un producto comercial rentable. En segundo lugar, los productores pueden animar a los usuarios innovadores a participar en interacciones de diseño conjuntas proporcionándoles kits de herramientas para la innovación impulsadas por los usuarios. En tercer lugar, los usuarios pueden convertirse en productores para lograr una difusión amplia de sus innovaciones (von Hippel, 2011).

1.5. Sistema de innovación

El concepto de sistema de innovación se basa en la idea de que las innovaciones son el resultado de la interacción social entre actores económicos. Un sistema de innovación es una red de instituciones y organizaciones, del sector privado y público, cuyas actividades e interacciones inician, importan, modifican o divulgan nuevas tecnologías (Freeman, 1987). Una serie de características permite hablar de un marco conceptual de sistemas de innovación (Edquist, 1997).

La innovación y el aprendizaje en el centro. Las innovaciones son nuevas creaciones de significado económico realizadas normalmente por las empresas. Dichas creaciones pueden ser totalmente originales, pero frecuentemente consisten más bien en nuevas combinaciones de elementos existentes. Hay diversas clasificaciones de innovaciones. Una de gran relevancia es la que distingue entre innovaciones de producto y de proceso. Otra clasificación es la que hace referencia al grado de novedad de la innovación: innovaciones radicales e innovaciones incrementales.

Las innovaciones son el resultado de procesos de aprendizaje. Siguiendo a Lundvall y Johnson (1994) podríamos distinguir los siguientes:

- Exploración (*exploring*), es el derivado de actividades de investigación básica que llevan a cabo universidades y organizaciones sin ánimo de lucro para descubrir leyes y teorías científicas. En ocasiones producen resultados imprevistos o no perseguidos, que dan lugar a un nuevo paradigma tecnológico.

- Investigación (*searching*), es una actividad de investigación aplicada que persigue aumentar el conocimiento con objeto de estimular la innovación.
- Aprendizaje (*learning*), que integra los las tres actividades siguientes: a) aprender haciendo (*learning by doing*) al que se refería Arrow (1962b) al indicar que los operarios aumentaban su conocimiento al realizar la actividad productiva, b) aprender usando (*learning by using*) al que se refería Rosenberg (1982) al considerar que los usuarios encuentran formas más eficientes de usar el producto, y c) aprender interactuando (*learning by interacting*) al que se refería Lundvall (1988) al indicar que los usuarios y los productores interactúan y de ello resultan innovaciones de producto.

Estos procesos de aprendizaje abocan en cuatro categorías diferentes de conocimiento (Lundvall y Johnson, 1994):

- Saber qué (*know-what*): conocimiento acerca de hechos; se identifica con la información, ya que puede descomponerse en bits y comunicarse con facilidad.
- Saber por qué (*know-why*): conocimiento científico acerca de los principios y leyes de la naturaleza, de la mente humana y de la sociedad. El acceso a este conocimiento mejora la tecnología y reduce los errores. Es muy importante en el desarrollo tecnológico de ciertas áreas basadas en la ciencia, como la industria química.
- Saber cómo (*know-how*): habilidades para hacer algo. Es algo práctico más que teórico. Puede relacionarse con los artesanos y los trabajadores de producción, aunque desempeña un papel clave en todas las actividades económicas.
- Saber quién (*know-who*): involucra información acerca de quién sabe qué y quién sabe hacer qué. También involucra la capacidad social para cooperar y comunicarse con diferentes clases de personas y expertos. El acceso a diferentes y variadas fuentes de conocimiento resulta esencial, sobre todo si tenemos en cuenta que los productos son cada vez más complejos y se combinan con muchas tecnologías, cada una enlazada con distintas disciplinas científicas.

El conocimiento es el recurso fundamental del sistema de innovación y el aprendizaje interactivo el proceso más importante, ya que coordina todos los elementos del sistema con objeto de dinamizarlo (Lundvall y Johnson, 1995). El énfasis en el aprendizaje interactivo indica que la cooperación entre los diferentes agentes sociales es una estrategia importante para promover las innovaciones (Asheim, 1996). Este aprendizaje tiene lugar: a) entre las diferentes etapas del proceso de innovación, implicando la movilización de diferentes formas de conocimiento e información (por ejemplo, conocimiento científico, información de mercado y habilidades técnicas, entre otros); b) con diferentes empresas, lo que conlleva colaboraciones con proveedores, complementores, competidores y clientes; c) con diferentes instituciones y centros de investigación, tanto públicos como privados y d) entre diferentes departamentos empresariales, que implica una interacción permanente entre los diferentes especialistas con diferentes formas de conocimiento (Asheim y Isaksen, 2001).

Considerando la empresa como un sistema abierto, el aprendizaje interactivo aboca en una organización dinámica y flexible, denominada ‘organización que aprende’ (*learning organisation*), que puede definirse como una organización que promueve el aprendizaje de todos sus miembros y tiene la capacidad de transformarse a sí misma continuamente, al adaptarse rápidamente a los entornos cambiantes, por medio de la adopción y desarrollo de innovaciones (Asheim, 1996). Tales organizaciones aceptan que la

innovación puede ocurrir en cualquier lugar de la empresa y que no es responsabilidad exclusiva de los investigadores, sino que todo trabajador puede ser una fuente de innovación. Por ello, estas empresas descentralizan la toma de decisiones, fomentan las comunicaciones horizontales, definen ampliamente el contenido de las tareas y fomentan el trabajo en equipo, entre otras decisiones estructurales. También están abiertas al entorno, con el que intercambian continuamente información y conocimientos.

Un enfoque holístico e interdisciplinar. El enfoque holístico trata de englobar todos los determinantes de la innovación que son importantes. Además tiene en cuenta factores no sólo económicos, sino también institucionales, organizativos, sociales y políticos. En este sentido es una aproximación interdisciplinar.

La capacidad innovadora de un sistema no solamente depende de su esfuerzo cuantitativo en I+D y de su infraestructura tecnológica, sino que también depende de la generación de externalidades mediante la interacción entre los distintos agentes del sistema como las empresas o las administraciones públicas.

Los distintos agentes de un sistema de innovación se agrupan en cuatro subsistemas (Buesa et al., 2002): 1) las empresas, relaciones interempresariales y las estructuras de mediación; 2) la infraestructura pública y privada de apoyo a la innovación (centros tecnológicos, universidades,); 3) las actuaciones públicas en relación con la innovación y el desarrollo tecnológico (marco legal y políticas tecnológicas), y 4) el entorno global (capital riesgo, capital humano, demanda).

La perspectiva histórica es natural. Las innovaciones se desarrollan a lo largo del tiempo (es necesario que transcurra un cierto período de tiempo desde la invención técnica hasta su transformación en una innovación económicamente importante y a su amplia difusión en el mercado) y las innovaciones, organizaciones e instituciones, tecnologías, regiones e incluso países son dependientes de la trayectoria (*path dependent*). La historia cuenta.

Diferencias entre sistemas y no optimalidad. Hay que reconocer las diferencias existentes de unos sistemas de innovación a otros y de la inexistencia de un sistema óptimo.

Los sistemas de innovación de varios países pueden ser diferentes. Pueden diferir en la estructura de producción, ya que en algunos son importantes las materias primas y en otros los conocimientos. Las instituciones también son diferentes en los países. Por otra parte, los sistemas al basarse en procesos de aprendizaje están sujetos a continuos cambios. El sistema nunca alcanzará un equilibrio óptimo, ya que es un sistema abierto y sujeto a la dependencia de la trayectoria.

Énfasis en la interdependencia y una visión no lineal del proceso innovador. El nuevo conocimiento tiene su origen en diferentes actores y agentes como empresas y universidades. Las empresas nunca innovan aisladamente, sino que interactúan con otras organizaciones. Esas relaciones son muy complejas y a menudo se caracterizan por la reciprocidad, la interactividad y los mecanismos de retroalimentación. Claramente no está caracterizada por una relación lineal causal y unidireccional.

La unidad primaria de la innovación no es el individuo; una persona no es un componente básico. Más bien es la red que se expande hacia el interior (departamentos de investigación, marketing y producción) y hacia el exterior (incluyendo clientes, proveedores, socios y otros). La innovación requiere el desarrollo y el mantenimiento de esta red como si fuera una fuerza abierta y colaboradora; una tarea nada fácil si tenemos en cuenta las complejidades de las relaciones, motivaciones diferentes y objetivos distintos (Davila *et al.*, 2006).

Engloba tanto tecnologías de producto como innovaciones organizativas. La innovación puede ser de producto, proceso y en otras áreas funcionales. La innovación en proceso conlleva cambios organizativos y viceversa.

Las instituciones son centrales. Las instituciones son elementos centrales en el proceso de innovación. Edquist y Johnson (1997) propugnan que se distinga claramente entre instituciones y organizaciones. Las organizaciones serían estructuras formales con un objetivo explícito, que han sido creadas conscientemente. Serían, pues, agentes o actores. Ejemplos de organizaciones relevantes en los sistemas de innovación serían las empresas (que pueden ser proveedores, clientes o competidores), las universidades, las sociedades de capital riesgo y las agencias públicas de política de innovación. Las instituciones, por su parte, serían conjuntos de hábitos comunes, rutinas, prácticas establecidas, reglas o leyes que regulan las relaciones e interacciones entre individuos, grupos y organizaciones. Serían, pues, las reglas del juego. Ejemplos importantes de tales instituciones serían, en los sistemas de innovación, las leyes de patentes y las normas que rigen las relaciones entre la universidad y la empresa.

Es un concepto difuso. El sistema de innovación nacional es un concepto difuso, ya que existen diferencias en la definición de los elementos que lo constituyen. La aproximación de sistema de innovación está asociada con una serie de ambigüedades. Por ejemplo, la definición de instituciones. Tampoco queda claro que elementos deberían ser incluidos en un sistema de innovación y las relaciones lógicas entre ellos.

Es una estructura conceptual más que una teoría formal. No es una teoría formal, sino una aproximación teórica para realizar trabajos empíricos. Proporciona interpretaciones y guías para futuras exploraciones. Proporciona una base para formular conjeturas.

La cooperación cercana con los proveedores, los complementores, los competidores, los clientes y las instituciones de apoyo en la región, basado en las relaciones humanas, puede aumentar el proceso de aprendizaje interactivo y crear un entorno o *milieu* favorable a la innovación y a la mejora continua (Maillat, 1995). Esto influye en la actuación y resultados de las empresas y refuerza la competitividad de los sistemas regionales, y se percibe, cada vez en mayor medida, como un factor importante de la ventaja competitiva regional.

El aprendizaje interactivo, que tiene lugar entre las empresas e instituciones que configuran un sistema de innovación, permite crear un conocimiento único y pegajoso (*sticky*) de naturaleza tácita. Parte de este conocimiento no es ‘propiedad’ de ninguna empresa en particular, sino que pertenece al propio sistema, en su conjunto, como uno de sus recursos intangibles (Gottardi, 1996).

El conocimiento tácito requiere utilizar medios informales en su transferencia. Ejemplos de estos medios de transferencia son la comunicación cara a cara, la formación de los trabajadores y la movilidad del personal. En todos estos casos, la proximidad geográfica facilita profundamente su transferencia. Esto es debido a que estas formas de transferencia de conocimiento son muy sensibles a la distancia entre los agentes implicados. Es por ello que, para algunos autores, la innovación es fundamentalmente un proceso geográfico, facilitado, aunque necesariamente limitado, por el agrupamiento espacial de las partes implicadas dentro de la misma región (Porter, 1998). Las razones son varias. La proximidad espacial facilita una interacción frecuente y cercana, la mayoría de las veces cara a cara. Tal interacción, tanto formal como informal, fomenta y permite el aprendizaje. En segundo lugar, las empresas agrupadas en la misma región frecuentemente comparten una «cultura común» que tiende a facilitar el proceso de aprendizaje. Estas empresas tienen un lenguaje o código de comunicación común producto de repetidas interacciones a lo largo del tiempo. Por otra parte, como el conocimiento más importante transmitido entre los agentes del proceso de innovación es tácito más que codificado, las empresas de la región que participan en las redes de intercambio tienen una ventaja crucial sobre el resto (Patel y Pavitt, 1994). Finalmente, este lenguaje o código de comunicación se apoya, de forma complementaria, en la creación de instituciones regionales, que ayudan a producir y reforzar un conjunto de normas y convenciones que gobiernan el comportamiento de las empresas locales y la interacción entre ellas.

Así pues, determinados recursos sólo se pueden encontrar en algunos lugares y no pueden transferirse rápida y fácilmente ni «copiarse» en otras localizaciones. En estos lugares, se produce un aprendizaje interactivo y «derrame» (*spillover*) de conocimiento desde diferentes instituciones y empresas, estableciendo como resultado una información asimétrica y un conocimiento único creado y absorbido de forma que promueven la competitividad de las empresas locales.

En definitiva, la proximidad de las empresas que forman parte de un sistema de innovación territorial les permite: a) compartir los costes fijos de recursos comunes, tales como infraestructuras y servicios; b) disponer de una base de proveedores especializados; c) tener acceso a un *stock* de trabajadores cualificados y d) aprovecharse de una rápida difusión y filtración de los conocimientos técnicos entre empresas. Estos factores, que facilitan las innovaciones y provocan un aumento de la productividad y que se encuentran fuera de las empresas individuales, reciben el nombre de «economías de aglomeración» (Saxenian, 1995). Las economías de aglomeración pueden representar importantes condiciones básicas y un estímulo a las innovaciones incrementales a través de las prácticas informales del «aprendizaje mediante la práctica» y el «aprendizaje mediante el uso», basadas básicamente en el conocimiento tácito (Asheim, 1996).

2. MODELOS DE SEGUIMIENTO DE LA INNOVACIÓN

En este apartado desarrollamos varios modelos que permiten hacer un seguimiento de las innovaciones, con objeto de contribuir a mejorarlas y verificar su eficacia en el mercado.

2.1. Los prototipos físicos como objetos para cruzar fronteras

En su forma más simple, los prototipos son hojas de cálculo, mapas de procesos o simulaciones; cualquier cosa simple que nos permita visualizar y entender mejor dónde

reside nuestra ignorancia. En una forma más elaborada, un prototipo es una versión preliminar del producto fabricado de forma artesanal para observar el aspecto externo y las prestaciones que realiza. Los prototipos deben ser una competencia central del equipo de innovación. Para el equipo, tiene que ser un modo primario de pensar y de actuar. La competencia de crear prototipos requiere gente que sea capaz de trabajar con y mediante lo incompleto. Construir, probar, perfeccionar, rebatir y corroborar prototipos son actividades esenciales porque (Davila *et al.*, 2006):

- Desafían los modelos mentales que tiene el equipo.
- Descubre patrones de resultados gracias a los cuales puede aprender el equipo.
- Unen al equipo y crean un lenguaje y una visión común y compartida.
- Generan un nivel de entusiasmo que los medios tradicionales no pueden igualar.
- Generan una manera de pensar que, al final, se convierte en innovación radical.

Los prototipos solo deben demandar el tiempo, esfuerzo, e inversión necesarios para generar un *feedback* útil y desarrollar una idea. Cuanto más terminado parezca el prototipo, es menos probable que sus creadores presten atención al *feedback* y se beneficien de él. La meta de crear prototipos no es concluir el proyecto, es aprender sobre la fortalezas y debilidades de la idea e identificar nuevas direcciones que otros prototipos podrán tomar (Brown, 2008).

Los prototipos son una forma útil de experimentar. En las primeras fases, no hace falta que sean elaborados ni caros, solo lo suficientemente desarrollado como para responder a la pregunta entre manos. A medida que progresa el aprendizaje, se vuelven más detallados para abordar preguntas más específicas. Los prototipos no están restringidos a los productos físicos; también se pueden usar para poner a prueba hipótesis sobre el modelo de negocio: qué segmentos del mercado están potencialmente interesados, cómo se puede introducir el producto o el servicio o qué servicios complementarios se podrían ofrecer. Las empresas de *software* e Internet utilizan el concepto de producto mínimo viable para describir la versión más básica de un producto que puede sacarse al mercado. Es frecuente que el producto mínimo viable carezca de muchas de las prestaciones que al final formarán parte del mismo, pero ofrece la suficiente funcionalidad como para descubrir las reacciones de los clientes, y proporciona una plataforma para añadir nuevos diseños y prestaciones. Su objetivo es entrar rápido en el mercado para acelerar el mensaje gracias a la interacción con los clientes (Dávila y Esptein, 2014).

El equipo de innovación debe ser capaz de cambiar, desarrollar y jugar con el prototipo para cosechar información valiosa (Davila *et al.*, 2006). Para comercializar las innovaciones radicales varias empresas han usado prototipos y el enfoque investigar y aprender con muy buenos resultados. Acometer cualquier tarea innovadora requiere fracasar mucho. Lo que hay que hacer es fracasar rápidamente y de una forma barata. Esta es la función de los prototipos. Un prototipo exige prestar atención a tres importantes reglas, que se describen en la figura 9.

Regla 1 Piense de forma modular	No intente resolver todas las piezas. Construya prototipos que aclaren una o dos de las incertidumbres más importantes. Como saben los buenos experimentadores esto da una información muy valiosa acerca de la naturaleza del problema, así como de la posible solución.
------------------------------------	---

Regla 2 Fracase de un modo rápido y barato	Defina pruebas prácticas que pueden hacerse de un modo barato. Construya un prototipo y pruébelo rápidamente. A menudo, es mejor trabajar con un socio, como un cliente importante o un proveedor, para compartir los costes, los riesgos y la información. Obtenga los resultados y determine qué ha aprendido y qué preguntas nuevas ha identificado. Modifique el prototipo.
Regla 3 Fracase a menudo para tener éxito más rápido	Use el enfoque de “Preparados, apunten, ¡fuego”! ... y vuelta a empezar. Es vital superar el viejo síndrome de “Preparados, apunten, apunten, apunten ...” Recuerde que el plural de anécdota es datos.

Figura 9: Las tres reglas del prototipo (Davila et al., 2006)

Wheelwright y Clark (1995) consideran que los prototipos ofrecen una oportunidad maravillosa de reunir las diferentes funciones, determinar el grado de progreso alcanzado hasta la fecha y considerar, en una fase intermedia, cómo podrían funcionar juntas las soluciones alternativas. En esencia, la realización de prototipos puede ser un vehículo importante para la discusión transversal entre las funciones de la organización, para la resolución de problemas y la integración. En este sentido, Schrage (1993) observa que dentro de algunas culturas innovadoras, los prototipos se convierten en el medio esencial para la información, la interacción, la integración y la colaboración abierta en la organización. Cuando los prototipos se usan sólo para probar los conceptos técnicos y no como vehículos de comunicación para la resolución de problemas a través de las fronteras, los desarrolladores pasan por alto enormes posibilidades para la abrasión creativa y la integración. Además, como los prototipos son objetos físicos y visuales, sirven para comunicarse con personas que no tienen una formación especializada, pero cuyo ojo profano quizá pueda predecir la respuesta general del público mucho mejor que el juicio de los expertos.

Al recordar sus días como diseñador de General Motors, Hirshberg comenta que todo en GM giraba en torno al coche en la plataforma de exhibición como primera prioridad. Sin embargo, jamás se permitía a nadie externo al equipo de desarrollo que viera o hiciera comentarios sobre el prototipo de arcilla del coche. Por el contrario, cuando en Nissan Design International un nuevo producto alcanza la fase de prototipo, se invita a todo aquel que esté libre y tenga interés a acercarse y hacer comentarios. Las secretarías, el personal de mantenimiento, el personal de la fábrica, absolutamente toda persona del edificio que está interesada se une a la crítica del diseño. Esta visión igualitaria de la validez del gusto de todo el mundo, incluso de los que no son expertos, puede tener como resultado observaciones dolorosas y útiles. Hirshberg cuenta aquella vez que una secretaria, Cathy Woo, llegó, taza de café en mano, a un coche nuevo que estaba en exhibición y observó con franqueza: “Creo que tiene un aspecto estúpido”. Los miembros del equipo de diseño, que habían estado votando a favor del diseño, se miraron entre sí y supieron que se habían estado engañando a sí mismos; todavía no lo teníamos. Se retiraron a la fábrica, volvieron a las mesas de dibujo y tres meses más tarde llamaron a la Sra. Woo para darle las gracias e invitarla a celebrar con ellos el nuevo diseño (Leonard– Barton, 1995).

Los mejores prototipos son aquellos con los que clientes y proveedores pueden jugar para que podamos reducir nuestra ignorancia acerca de lo que ellos ven en el concepto propuesto. En los primeros prototipos, el objetivo debería ser la simplicidad extrema y el aprendizaje de lo básico. Después, en el segundo prototipo y los siguientes se puede pasar a representaciones de la innovación más robustas y menos modulares. Los prototipos más avanzados deberían centrarse en la esencia del diseño. Los mejores prototipos son los que ofrecen a los diseñadores, clientes y proveedores la oportunidad de hacer mejoras secuenciales. Un prototipo centra al equipo alrededor de un concepto común y en su desarrollo. Conecta al equipo con un modelo fuerte que aclara tanto el problema como las soluciones posibles. Además, un prototipo da pistas a la dirección sobre los pasos siguientes que han de dar en innovación (Davila *et al.*, 2006).

Sony y Panasonic instalaron «tiendas antena» en algunos centros comerciales y aeropuertos seleccionados, donde compran los clientes más exigentes. En estas tiendas se presentan prototipos de los nuevos productos y los ingenieros de desarrollo o los directores de los mismos conversan con los clientes para obtener información de primera mano sobre su reacción frente al nuevo producto y lo que ellos realmente quieren (Mendelson y Ziegler, 1999).

2.2. Diseño para producción

El diseño para producción (*Design For Manufacturing*, DFM) trata de facilitar la integración entre diseño y producción. Esta técnica pretende analizar en la etapa de diseño de un producto los requisitos para su fabricación, y de esta manera desarrollar productos que sean fáciles de fabricar. Cuanto más fáciles de fabricar sean los productos, mayor es la calidad, más alta es la productividad y, por tanto, más bajos los costes unitarios (Schilling y Hill, 1998). Asimismo, también se logra evitar modificaciones posteriores en el producto como resultado de problemas en el área de producción. Esta técnica permite minimizar retrasos costosos que pueden ocurrir debido, por ejemplo, a desajustes entre las especificaciones formuladas en la fase de diseño y las características que puede acometer la fase de producción. Por tanto, se acorta el tiempo de desarrollo del producto.

El diseño para producción examina si los diseños del producto serán fáciles de ensamblar y estimula la simplificación del producto. Esto genera una reducción del número de componentes, que conduce a una reducción tanto de los costes de materiales como de los tiempos de ensamblaje. La simplicidad fluye hasta cubrir todas las actividades, incluido el servicio *in situ*. Esta reducción de componentes facilita la fiabilidad del producto, disminuye los costes del ciclo de vida del producto, reduce el número de horas de ingeniería de diseño necesarias, reduce las compras, los inventarios y el espacio para almacenar los componentes. Por ejemplo, al reducir a la mitad el número de partes para armar el parachoques trasero del Seville, Cadillac disminuyó un 57 por ciento el tiempo de ensamblaje y ahorró más de 450.000 dólares en costes laborales anuales (Schilling, 2008).

Por otra parte, mantener en un mínimo la cantidad de modificaciones hechas a las partes (por ejemplo, clase y tamaño de tornillos) ahorra tiempo y dinero. Emplear componentes estandarizados mantiene bajos los niveles de inventario, los costes de compra y el tiempo de entrega de los pedidos. También evita preparar las especificaciones para un solo componente, encontrar un proveedor y probar el producto

ya terminado. Asimismo, la estandarización facilita la reparación de los productos y posibilita que los repuestos se encuentren con mayor rapidez. Por ejemplo, el 60 por ciento de las piezas empleadas en fabricar el Domani, un modelo de Honda, también se utilizan en otros modelos de la compañía. Con anterioridad, sólo entre el 10 y el 15 por ciento de las piezas eran comunes a los diversos modelos (Noori y Radford, 1995).

Una forma de simplificar el ensamblaje consiste en diseñar el producto de tal manera que sus piezas se ajusten unas a otras sin necesidad de tornillos o tuercas, tal y como lo hizo, por ejemplo, NCR con su registradora electrónica 2760. En la factura de un proveedor, los tornillos y las tuercas pueden representar tan sólo centavos, y colectivamente, sólo constituyen el 5 por ciento del listado de materiales de un producto típico. Sin embargo, si se agregan todos los costes asociados, como el tiempo para alinear los componentes mientras se insertan y se aseguran los tornillos, el coste de utilizar estas piezas tan comunes puede ascender hasta el 75 por ciento del coste total del ensamblaje. Los elementos sujetadores son, pues, lo primero que se debe eliminar al diseñar un producto para producción (Port, 1989).

Al desarrollar componentes que pueden ensamblarse de diversas maneras (diseño modular), la empresa puede ofrecer una amplia variedad de productos mientras se mantiene en un mínimo el número de piezas que deben comprarse o producirse. En Hewlett-Packard (HP) el diseño modular posibilita que la empresa pueda comercializar un flujo continuo de impresoras mejoradas que cuestan menos, tienen más funciones y funcionan mejor que cualquier otra en el mercado. El diseño modular también permite que HP ofrezca paquetes de calidad mejorada para sus ordenadores personales, los cuales pueden emplearse para agregar muchas de las nuevas funciones a los modelos viejos (Noori y Radford, 1995).

La implementación con éxito del diseño para producción requiere cambios culturales que permitan mejorar las comunicaciones entre todas las funciones de la empresa, promuevan el trabajo en equipo e integren los esfuerzos del personal involucrado en las decisiones del producto y del proceso. La figura 10 resume un conjunto de reglas de diseño para producción comúnmente utilizadas, junto con su impacto esperado en el rendimiento.

Regla de diseño	Impacto sobre el rendimiento
Minimizar el número de componentes	Simplificar el montaje; reduce la mano de obra directa; reduce los costes de manipulación de materiales y de inventario; estimula la calidad del producto
Minimizar el número de referencias de componentes (utilizar componentes comunes en una misma familia de productos)	Reduce los costes de manipulación de materiales y de inventario; mejora las economías de escala (incrementa el volumen mediante elementos comunes)
Eliminar los ajustes	Reduce los errores de ensamblaje (incrementa la calidad); permite la automatización; incrementa la capacidad y el rendimiento
Eliminar los cierres	Simplifica el ensamblaje (incrementa la calidad); reduce los costes de mano de obra directa; reduce chirridos y ruidos; mejora la durabilidad; permite la automatización

Eliminar plantillas ¹ y accesorios	Reduce los costes de cambio de línea; reduce la inversión requerida
---	---

Figura 10: Reglas de diseño para el montaje de productos fabricados (Schilling, 2008).

El propósito de estas reglas de diseño normalmente es reducir los costes y potenciar la calidad del producto al asegurar que los diseños de producto son fáciles de fabricar. Cuanto más fáciles de fabricar sean los productos y menores los pasos de ensamblaje requeridos, la productividad del trabajo será mayor, resultando en menores costes unitarios. Además diseñar productos que sean fáciles de fabricar reduce la probabilidad de incurrir en errores en el proceso de montaje, lo que resulta en una mayor calidad del producto.

2.3. CAD/CAM

El diseño asistido por ordenador (*Computer Aided Design*, CAD) es un sistema electrónico que permite el diseño de nuevas piezas o componentes o la modificación de los actuales, sustituyendo el diseño tradicional a mano. Se trata de un proceso de diseño que emplea sofisticadas técnicas gráficas de ordenador, apoyadas en paquetes de *software*, con el fin de ayudar en los problemas analíticos, de desarrollo, de coste y ergonómicos asociados a las actividades de diseño (Hawkes, 1989).

Una de las ventajas más importantes del CAD es su capacidad de archivar resultados en soportes que permiten su fácil reproducción y manipulación. Asimismo, cabe destacar los siguientes beneficios derivados de su utilización (Hawkes, 1989): 1) agilizar y potenciar el diseño de piezas y otros elementos de ingeniería, 2) visualizar y generar las mejores perspectivas de un diseño, 3) obtener distintas secciones del objeto diseñado, 4) simular el comportamiento real del elemento a crear frente a situaciones o problemas específicos, 5) introducir modificaciones en el diseño, 6) controlar y revisar las especificaciones que deben cumplirse, 7) generar la documentación técnica necesaria para el proceso de fabricación, 8) mayor precisión en los dibujos, 9) dibujos más limpios, 10) técnicas especiales de dibujo (como el *zooming*), 11) análisis y cálculo de diseño más rápidos, 12) menores requisitos de desarrollo y 13) integración del diseño con otras disciplinas. Por todo ello, el CAD permite reducir el tiempo y el coste de desarrollo de los productos y, en consecuencia, acorta el tiempo necesario para comercializar nuevos productos.

La fabricación asistida por ordenador (*Computer Aided Manufacturing*, CAM) se ha utilizado desde hace tiempo en los procesos de producción continuos, si bien su aplicación a la fabricación por lotes es posterior. Bajo la denominación CAM se suele incluir: a) los robots industriales, b) las máquinas-herramientas de control numérico, el control numérico por ordenador y el control numérico directo, c) los sistemas de fabricación flexible, d) los sistemas automatizados de manejo de materiales, incluyendo el control de existencias, y e) los sistemas de verificación automatizados, cuyas aplicaciones cubren un amplio espectro.

La integración de las operaciones de diseño y producción se convierten en un elemento de crucial importancia, que se realiza a través de los llamados sistemas CAD/CAM. La

¹ Una plantilla es una «plancha» cortada con los mismos tamaños, ángulos y figuras que debe tener la superficie de una pieza y que, puesta sobre ella, permite cortarla y labrarla.

integración del CAD y el CAM hace posible diseñar el producto y especificar el proceso de modo simultáneo, lo que permite reducir de modo drástico los errores y pérdidas originados al pasar de la fase de diseño a la de fabricación (Thompson y Paris, 1982). Esto significa, por ejemplo, que puede dibujarse cualquier componente sobre una pantalla de representación visual (*Visual Display Unit*, VDU) y transferir los gráficos por medio de señales eléctricas a través de un cable que lo enlace a un sistema de fabricación, donde los componentes se puedan producir automáticamente sobre una máquina de control numérico por ordenador. Los sistemas CAD/CAM permiten observar las interacciones existentes entre los diferentes elementos que integran una pieza o un producto, sin necesidad de construir un prototipo.

LECTURA 3: COTAGUA CORTA DE TODO MEDIANTE CHORROS DE AGUA

El sistema, cuya tecnología es de origen alemán, está en el mercado desde hace siete años, pero todavía no es demasiado conocido. Su funcionamiento impresiona. Cotagua estudia primero las piezas que el cliente desea y las diseña en el ordenador para obtener la máxima cantidad de piezas posible en la plancha del material que le entrega el cliente.

Una vez realizado el diseño en el ordenador, se envía la orden a la mesa de corte –de dos metros de ancho por cuatro de largo–, donde un brazo metálico controlado por ordenador dirige el chorro de agua según lo diseñado. El chorro de agua sale a una presión de 3.800 atmósferas y una velocidad que duplica la del sonido –680 kilómetros por segundo– a través de un tubo que se coloca a un milímetro de distancia del material que se desea cortar. Después, el agua corta el material con un margen de error de una décima de milímetro.

Para evitar que el chorro de agua rompa también la mesa de corte, la superficie es de rejilla y está llena de agua, lo que amortigua la potencia del chorro. Y es que el chorro de agua es capaz de cortar todo tipo de materiales de grosores de hasta treinta centímetros: cristal, corcho, madera, fibra de vidrio, granito, titanio, mármol o acero al carbono, entre otros. Según la dureza del material, Cotagua utiliza un polvo abrasivo que ayuda al corte.

El sistema tiene ventajas: se aprovecha mejor el material, corta no sólo en recto, sino que reproduce diseños con exactitud y no deforma ni temple las piezas, ni cambia su estructura molecular, al no estar sometidas al calor.

Por poner algún ejemplo, el sistema de corte con agua se utiliza en medicina para seccionar órganos humanos –el láser dañaría los tejidos– o en empresas como Bimbo, que realiza el corte de sus rebanadas de pan de molde con este sistema, aunque, en vez de con agua, con aceite de soja que luego recicla.

Fuente: Errea, G. (1997): “Cotagua corta de todo mediante chorros de agua”, *Expansión*, 16 de octubre, p. 12.

Básicamente, las condiciones que deben reunir los sistemas CAD/CAM podrían resumirse en (Hawkes, 1989): 1) el sistema debe ayudar al diseñador a realizar su trabajo mediante realizaciones mutuamente efectivas, es decir, el ordenador debe realizar aquellas tareas en las que es más eficiente que el operador humano, 2) el sistema debe ayudar en todos los procesos, desde el diseño conceptual al control numérico y 3) en la etapa de diseño conceptual, el sistema deberá facilitar una presentación efectiva del objeto diseñado.

2.4. Coordinación técnica de las funciones (casa de la calidad)

Convertir los requisitos del cliente en características técnicas de diseño bien detalladas puede ser una tarea difícil. Con frecuencia, los requisitos de los clientes son vagos y, en algunos casos, contradictorios. Como las características técnicas del producto se expresan en un lenguaje bastante diferente del utilizado por el cliente, a menudo la voz del cliente no se escucha y el resultado final es un producto que no satisface por completo sus necesidades.

El despliegue de la función de calidad (*Quality Function Deployment*, QFD) es un proceso que incorpora la voz del cliente en el desarrollo del producto (Griffin y Hauser, 1993). Es

decir, se trata de un proceso para identificar las necesidades del cliente y convertirlas en características de ingeniería. Este proceso se lleva a cabo empleando un diagrama de dos dimensiones: los requisitos de calidad necesarios para satisfacer al cliente en el eje horizontal y las características de ingeniería en el eje vertical. Es una especie de mapa conceptual que facilita la comunicación y la planificación interfuncional. En la figura 11 se representa un gráfico de despliegue de la calidad, también denominado, por su forma, «casa de la calidad».

El despliegue de la calidad comienza con el descubrimiento de las necesidades de los clientes, que se colocan en la columna izquierda, de forma que cada fila horizontal refleja un único requisito del producto o atributo valorado por el cliente (AC), y se reproducen con las mismas palabras que este utiliza. Esta etapa es la más difícil, porque requiere obtener y expresar lo que el cliente realmente quiere, y no lo que la empresa piensa que el cliente espera (Sullivan, 1986). La voz del mercado incluye, no solamente a los usuarios del producto, sino también a otras partes interesadas que, de alguna manera, pudiesen verse afectadas, como la fuerza de ventas o las instituciones reguladoras, entre otras. Los requisitos del cliente con relación al producto indican «*qué*» hay que hacer. Posteriormente, se priorizan los requisitos en virtud de la experiencia directa de los miembros del equipo que construye la casa de la calidad con los clientes. Para ello se pueden utilizar diferentes métodos: desde otorgar a cada requisito una puntuación de 0 a 10 hasta asignar un peso expresado como porcentaje.

Para tener una referencia respecto a los competidores, es necesario saber en qué posición se encuentra la empresa. La parte derecha de la casa compara cómo satisface la competencia los deseos del cliente y cómo lo hacemos nosotros, lo que permite identificar oportunidades de mejora. Esta sección de la casa de la calidad proporciona una conexión natural de la concepción del producto con la visión estratégica de la empresa (Hauser y Clausing, 1988). Además, es un punto de partida para buscar la forma de conseguir una ventaja competitiva.

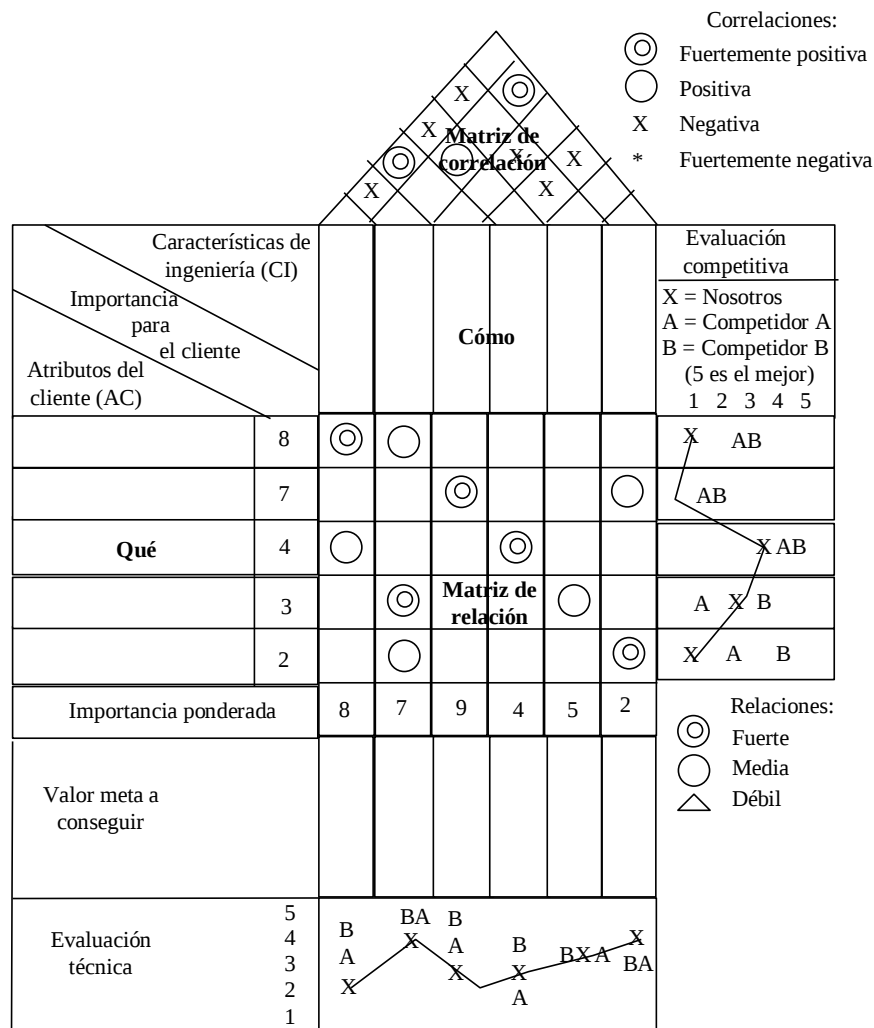


Figura 11: Casa de la calidad (Hauser y Clausing, 1988)

Posteriormente, hay que especificar en el lenguaje técnico interno las características de ingeniería (CI) necesarias para la fabricación del producto, y que afectan a uno o más atributos del cliente: indican «cómo» hay que hacer el producto. Las características de ingeniería se representan en la fila primera de la matriz y cada una de ellas ocupa una columna. Cualquier CI puede afectar a más de un AC. Las filas siguientes recogen la importancia ponderada de cada característica de ingeniería, el objetivo que se pretende alcanzar y, en la última fila, se comparan las características de ingeniería del producto con las características de ingeniería de los productos de la competencia. Se pueden añadir otras dos filas: una que recoja el grado de dificultad para mejorar cada característica de ingeniería y, otra más, para representar el coste de cada unidad de mejora de las características.

A continuación, un equipo interfuncional rellena el cuerpo de la casa indicando cómo afecta cada característica de ingeniería a cada atributo del cliente. Esta parte de la casa es conocida como «matriz de relaciones». Una señal es colocada en la celda localizada en la intersección de cada fila (atributo del cliente) con cada columna (características de ingeniería) que representa la intensidad de la relación entre ellas. Esas relaciones se

clasifican como fuertes o débiles, para lo cual se utilizan símbolos sencillos. La celda se deja en blanco si no existe relación.

El techo de la casa de la calidad ayuda a los ingenieros a especificar las características de ingeniería que tienen que mejorarse colateralmente. En ocasiones, la mejora de una característica implica la perturbación de tantas otras que lo más oportuno es no modificarla (Hauser y Clausing, 1988). Esta «matriz de correlaciones», en la que se comparan los *cómo* entre sí, facilita los intercambios de ingeniería necesarios y también permite localizar requerimientos conflictivos de diseño. Las correlaciones pueden ser positivas o negativas y se representan igualmente mediante símbolos sencillos.

En un equipo multifuncional los *cómo* de una casa se pueden convertir en los *qué* de otra. Por ejemplo, en una primera casa los *qué* pueden ser los atributos del cliente y los *cómo* las características de ingeniería. En la casa siguiente las características de ingeniería serán los *qué* y las características de las piezas serán los *cómo*. En una tercera serán los *qué* las características de las piezas y los *cómo* las operaciones de fabricación, y así sucesivamente (figura 12).

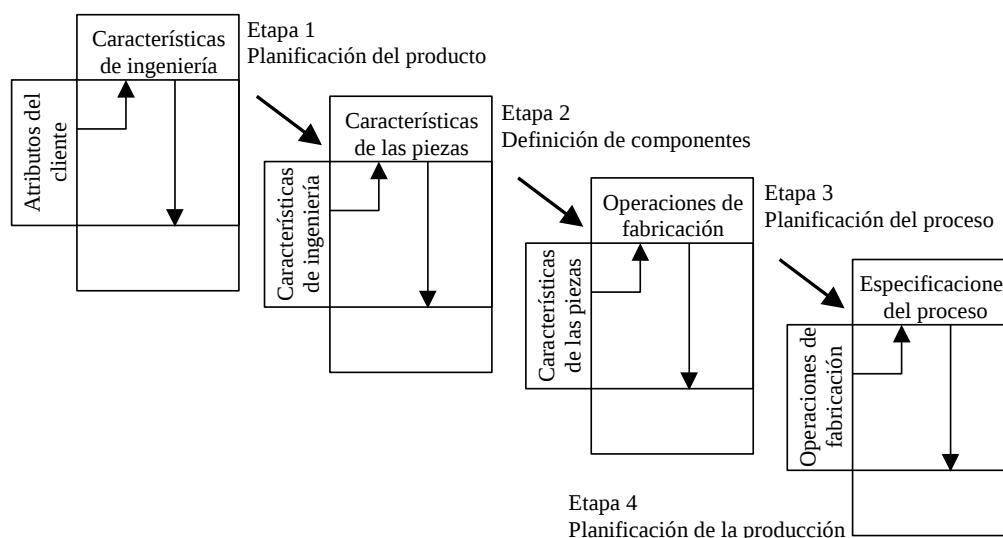


Figura 12: Despliegue de la calidad a lo largo del proceso de desarrollo de un producto

La técnica de despliegue de la calidad proporciona una serie de beneficios para la organización. Uno de los más importantes es que logra una mejora sustancial de la comunicación, tanto dentro de las funciones como entre las funciones. Esto lo consigue proporcionando un lenguaje común y un marco dentro del cual los ingenieros de diseño y los agentes comerciales pueden interactuar fructíferamente. El despliegue de la función de calidad mediante el uso juicioso de símbolos sencillos reúne gran cantidad de información en forma condensada y conveniente, y es útil para un enfoque sistemático (Juran, 1989). En definitiva, mejora el uso eficiente de la información interna y la vincula dentro del equipo. Otros beneficios importantes son el logro de un fuerte enfoque en el cliente y la mejora del trabajo en equipo.

Sin embargo, esta técnica presenta dificultades, entre las que cabe señalar la difícil traducción de las necesidades de calidad expresadas por el cliente en elementos o

características de ingeniería, así como el peligro que se corre en algunas ocasiones al construir gráficos demasiados complejos y poco manejables, con lo cual pierden su utilidad, ya que no facilitan la comprensión del proceso por parte de los trabajadores.

2.5. Proceso de evaluación

Una necesidad de mercado puede satisfacerse utilizando diferentes tecnologías. Por tanto, la empresa se enfrenta a una amplia gama de posibilidades tecnológicas, pero, al contar con recursos escasos, necesita seleccionar la más adecuada de acuerdo a los criterios previamente establecidos. Dos de los modelos de evaluación de innovaciones más socorridos son el proceso etapa–puerta y el embudo de las ideas.

Proceso etapa–puerta. Los sistemas de etapa–puerta (*stage–gate*) son una herramienta de dirección del proceso de desarrollo de la innovación para mejorar la posibilidad de lanzar nuevos productos de forma rápida y con éxito (Cooper, 1990).

Los sistemas de etapa–puerta consisten en dividir el proceso de innovación en un conjunto de etapas predeterminadas separadas por puntos de control o puertas. La comercialización es la última etapa. Cada etapa está formada por un conjunto de actividades prescritas. La entrada a cada etapa es una puerta. Cada puerta es un punto de evaluación donde personas con autoridad controlan la viabilidad del proyecto y su avance en el logro de los objetivos fijados por la dirección. Posteriormente, toman la decisión de desarrollar la etapa un poco más, seguir con el proyecto o abandonarlo.

Las puertas establecidas a lo largo del proceso y que preceden a cada una de las etapas cumplen las siguientes funciones (Cooper, 1996):

- Sirven como puntos de control de la calidad, de modo que, si el proyecto no reúne la calidad adecuada, no se permite que continúe adelante y, o bien se abandona, o se envía al comienzo de la etapa anterior para mejorarlo.
- Sirven como puntos de decisión sobre si seguir o abandonar un proyecto, de forma que los proyectos mediocres se abandonan, evitando, así, que se malgasten recursos que podrían emplearse en otros proyectos con mayor probabilidad de éxito.
- Son puntos donde se decide el camino hacia la siguiente etapa del proceso, así como los recursos que van a comprometerse.

Los sistemas de etapa–puerta son multifuncionales. Por un lado, las etapas son multifuncionales (cada una incluye actividades de muy diferentes funciones de la empresa), lo cual fomenta en estos sistemas el uso de equipos de proyectos multifuncionales. Estos equipos operan en paralelo, reduciendo el camino crítico, al coordinar mejor las últimas actividades con las más tempranas. Por otro lado, las puertas también son multifuncionales. En cada una de ellas, un grupo de directivos superiores de diferentes funciones, que poseen los recursos requeridos por el líder del proyecto y el equipo para la siguiente etapa, revisan conjuntamente el proyecto y deciden: a) proceder a la siguiente etapa, b) abandonar el proyecto o c) redirigirlo. Un comportamiento eficaz por parte de estos directivos es crítica para un efectivo sistema de etapa–puerta, especialmente en las etapas tempranas, en las que las ideas y los conceptos están desarrollándose y creciendo rápidamente. Estos directivos son decisivos a la hora de parar procesos que no están produciendo los resultados técnicos deseados o que no parece probable que vayan a proporcionar una ventaja competitiva (Cohen *et al.*, 1998).

La revisión del proyecto por el grupo de directivos en cada puerta se realiza de acuerdo con unos criterios acordados previamente, y que son cuantitativos –tasa de recuperación de la inversión y análisis de riesgo o sensibilidad, entre otros– y cualitativos –encaje estratégico, superioridad del producto o atractivo del mercado, por ejemplo–. Para lograr valor en el proyecto, el equipo de desarrollo y el grupo de directivos que controla cada puerta deben tener una visión clara y compartida de lo que será entregado en cada puerta y de los criterios de éxito asociados a estas entregas.

La efectividad de un proceso etapa–puerta no es mayor que la del grupo de directivos que flanquean cada puerta. Este grupo tiene que: a) tener experiencia en el desarrollo de productos, b) conocer en profundidad la disciplina requerida en cada una de las puertas (por ejemplo, ingeniería o investigación de mercados), c) controlar recursos económicos para ampliar la financiación si lo estiman necesario, d) conocer la relación de la tecnología y el producto con la estrategia de la empresa y d) ser objetivos en la aplicación de los criterios.

Los procesos de etapa–puerta suelen dividir el proceso de desarrollo de nuevos productos en un número de etapas, que oscila entre cuatro y ocho, dependiendo del tipo de producto y de la empresa. En la figura 13 se representa un proceso de etapa–puerta de cinco etapas.

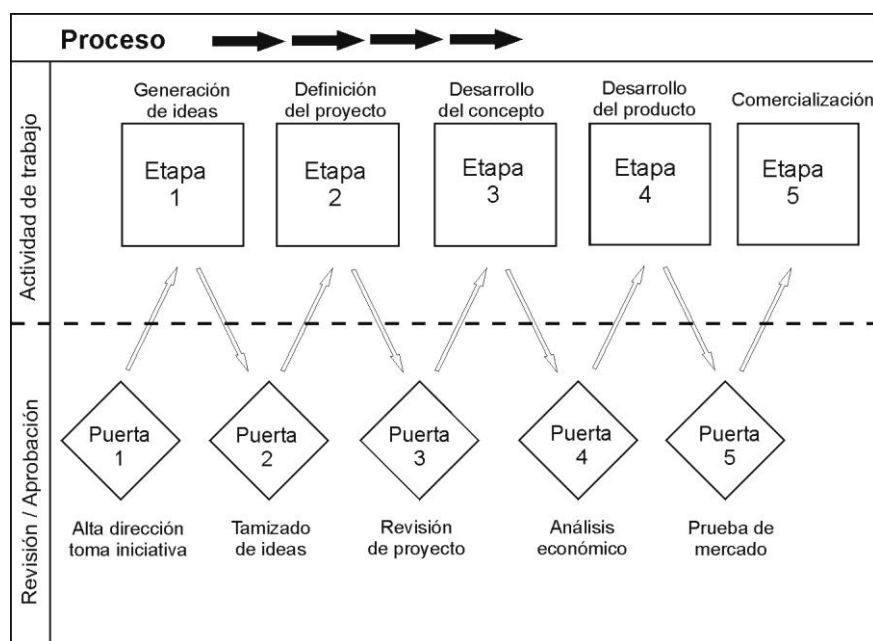


Figura 13: Proceso etapa–puerta

El primer filtro (o puerta 1) consiste en crear el equipo de desarrollo de una innovación y definir los objetivos que debe alcanzar. A continuación, el equipo se reúne para generar ideas. Esta etapa consiste en concebir ideas sobre innovaciones coherentes con los objetivos y estrategias de la empresa. Para ello, el equipo necesita toda la información de que pueda disponer, tanto interna como externa. En la puerta 2 se evalúan las ideas que les presenta el equipo de desarrollo. Conviene eliminar las ideas poco prometedoras lo más rápidamente posible, antes de que absorban recursos significativos. Saber más de una idea siempre implica incurrir en costes. Las

eliminaciones rápidas tienen la ventaja de dejar más recursos disponibles para las ideas que realmente valgan la pena.

Posteriormente, el equipo define uno o más proyectos de forma detallada. Cada proyecto debe contener la idea a desarrollar e indicar su forma, función, propósito y beneficios. En la puerta 3 se aprueba uno de los proyectos presentados. La selección de sólo una propuesta se justifica por las siguientes razones (Henriksen y Traynor, 1999): 1) los recursos a disposición de la empresa son, normalmente, inferiores a los necesarios para llevar a cabo los proyectos potenciales y 2) la organización debe guiar proactivamente los esfuerzos de investigación en una dirección consistente con su estrategia.

El desarrollo del concepto consiste en materializar la idea en un producto técnicamente viable y capaz de lanzarse al mercado en condiciones económicas. Esta etapa abarca las actividades de investigación y desarrollo, cuyo resultado es un prototipo. En la puerta 4 se evalúa el atractivo económico de la innovación en el caso de ser comercializada.

El desarrollo del producto consiste en perfeccionar técnicamente el prototipo, a la par que se busca su compatibilidad con el sistema productivo de la empresa y se diseñan las políticas de marketing. En la puerta 5 se realiza la prueba de mercado para tomar una de las siguientes actuaciones: a) introducir masivamente el producto en el mercado, b) modificar algunas de las características funcionales del producto o de la mezcla de marketing empleado, c) rediseñar con detalle el producto o d) rechazar el producto y decidir no comercializarlo.

La etapa de comercialización es la última del proceso y, previamente, requiere realizar cuantiosas inversiones en los procesos productivos para tener la capacidad suficiente que permita atender el mercado objetivo. También hay que desarrollar los canales de distribución oportunos, formar a los vendedores, comunicar el producto al mercado y fijar precios.

Desde su aparición en la literatura múltiples empresas han utilizado con éxito el proceso de etapa-puerta, que les han facilitado una rápida introducción de productos al mercado. Muchas empresas afirman que el sistema de etapa-puerta reduce, por un lado, el tiempo de desarrollo, al permitir la identificación de proyectos que deben abandonarse y, por otro, incrementa la ratio de proyectos desarrollados internamente que resultan en productos comerciales.

Indudablemente, un sistema de este tipo es mucho mejor que otro arbitrario. Así, el director del proyecto no tiene que andar continuamente negociando con la alta dirección para mantener vivo su proyecto y seguir avanzando. Como se ha señalado, la principal ventaja de este tipo de procesos es que evitan que se malgasten recursos en proyectos con escasas probabilidades de éxito, lo que permite que dichos recursos se asignen a otros proyectos, consiguiendo reducir el tiempo de desarrollo de los mismos. No obstante, cabe señalar otras ventajas (Tatikonda y Rosenthal, 2000). Un proceso de este tipo proporciona un sentido de estructura que permite reducir la ambigüedad del personal del proyecto. Por otro lado, permite detectar y resolver, lo antes posible en el proceso, problemas potenciales que pueden surgir en el diseño del producto, desarrollo y producción, reduciendo, así, el tiempo consumido y el esfuerzo realizado. Asimismo,

las revisiones pueden hacer más fácil ver dónde y cuándo es necesaria la asignación de recursos.

Sin embargo, los sistemas de etapa–puerta también presentan debilidades. La más relevante es que la revisión del proyecto en cada puerta consume demasiado tiempo. Además, si la aprobación para proceder a la siguiente etapa se demora, se detendrá todo el proceso, puesto que, sin dicha aprobación, no se puede comenzar la siguiente etapa. Para evitar este problema algunas empresas permiten que las decisiones de aprobación en cada puerta se tomen con información incompleta, de modo que se facilita el avance del proyecto a la siguiente etapa siempre y cuando los resultados sean positivos en la misma. Esto evita que el proyecto se detenga a la espera de la finalización de alguna actividad de la etapa previa (Cooper, 1996).

Otro inconveniente de estos sistemas es que la necesidad de superar el filtro de la puerta para poder continuar a la siguiente etapa puede implicar una falta de confianza en el equipo de desarrollo del producto. El equipo tiene que dedicar parte de su tiempo a prepararse para las revisiones periódicas, las cuales dirige un equipo de directivos superiores a los que les puede faltar el conocimiento técnico requerido (Anderson, 1996). La asignación a las distintas puertas de un equipo de directivos inadecuado puede hacer mucho daño al proceso. Son inadecuados los directivos que carecen de la experiencia o los conocimientos técnicos necesarios o desconocen la estrategia de la empresa.

También se critica que las reglas y revisiones pueden forzar la ejecución del proyecto de una manera predeterminada, más que permitir la adaptabilidad necesaria para hacer frente a las incertidumbres que surgen en los proyectos de desarrollo, cuando está disponible información de nuevos mercados o surgen problemas tecnológicos no anticipados (Tatikonda y Rosenthal, 2000).

Respecto a la implementación de los sistemas de etapa–puerta en la empresa, O’Connor (1994) observa que es complicada debido a las dinámicas de tiempo, a las estructuras organizativas cambiantes y a la transferencia de personal clave.

Finalmente, es importante subrayar que la aplicación de un sistema de etapa–puerta no es incompatible con el camino de la ingeniería concurrente. Ambos enfoques pueden considerarse conjuntamente y, en este sentido, las investigaciones sugieren que los beneficios de la ingeniería concurrente pueden lograrse sin tener que abandonar una estructura de desarrollo por etapas (Anderson, 1996).

Embudo de las ideas. El embudo del desarrollo, representado en la figura 14, proporciona una estructura gráfica para pensar sobre la generación e investigación de opciones alternativas de desarrollo, y combinar algunas de ellas en un concepto de producto. El proceso puede ser representado haciendo similitud con un embudo de boca muy ancha, donde las ideas entran por un lado y, por el otro, emergen los nuevos productos (Wheelwright y Clark, 1992).

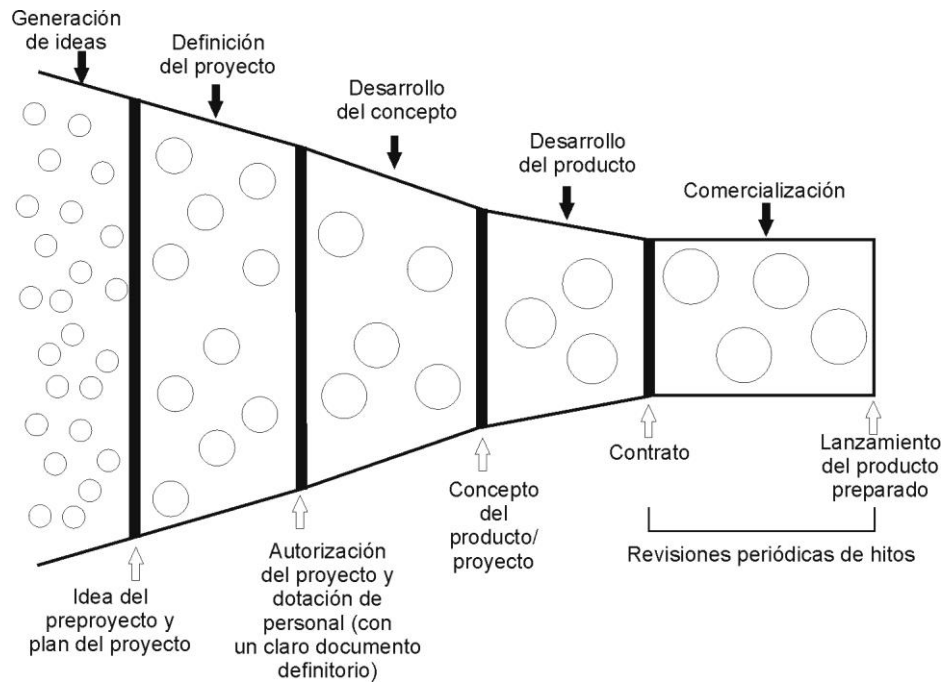


Figura 14: Embudo de las ideas

De acuerdo con Stevens y Burley (1997) se requieren 3000 ideas en bruto para producir un producto nuevo que sea un éxito comercial (figura 15). La industria farmacéutica es un ejemplo extremo: solo 1 de cada 10.000 compuestos tiene éxito como fármaco, con un tiempo global desde su descubrimiento hasta el lanzamiento al mercado de 12 años y un coste total de aproximadamente 350 millones de dólares (Schilling, 2008). De ahí que el proceso de innovación sea a menudo concebido como un embudo, con muchas potenciales ideas sobre nuevos productos que aparecen en el lado ancho, pero con muy pocas que logren avanzar a lo largo del proceso de desarrollo.

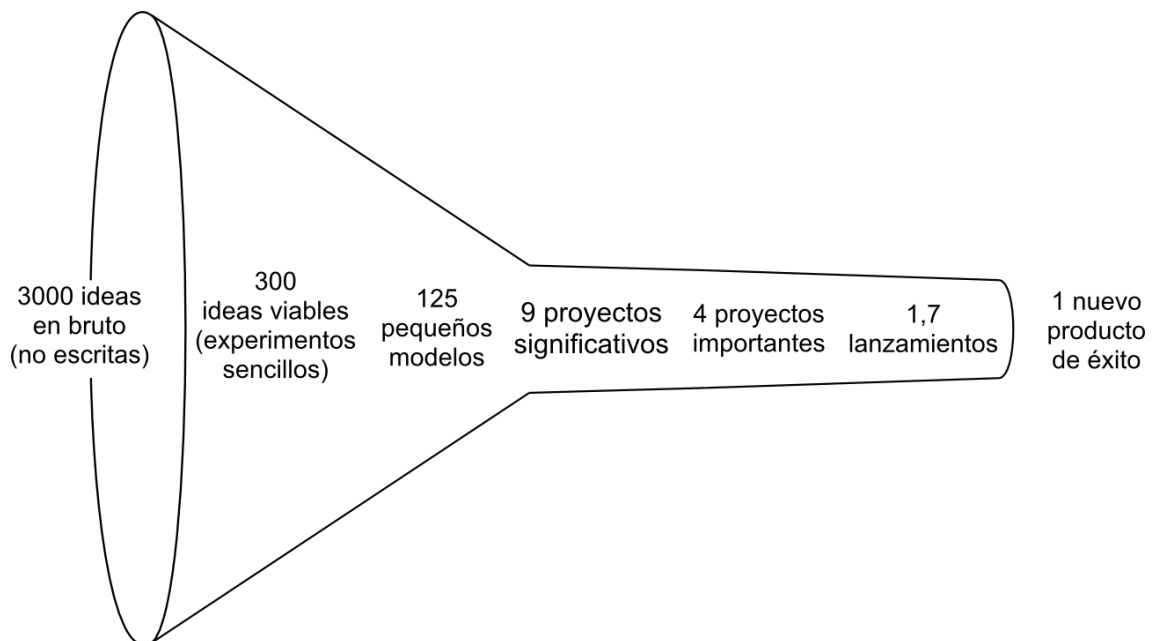


Figura 15: De la idea al producto (Stevens y Burley, 1997)

En suma, el embudo del desarrollo es una herramienta que permite seleccionar los proyectos adecuados. En la parte ancha del embudo entran una gran variedad de potenciales proyectos en forma de ideas para su investigación, pero, tras hacerles un examen cuidadoso, solamente quedan los que la empresa considera críticos. Para tomar esta decisión, los proyectos potenciales se someten a una serie de cribas antes de entrar en el cuello estrecho del embudo, con el fin de que los escasos recursos de desarrollo se apliquen solamente a aquellos proyectos que se alinean mejor con las capacidades de la empresa y los objetivos del negocio, pudiendo lograr, así, una cartera de proyectos de máximo valor para la empresa y equilibrada estratégicamente.

A medida que el embudo se va estrechando, los criterios para permanecer en el mismo cada vez resultan más rigurosos. El proceso requiere un período de reflexión, experimentación, investigación de mercado y desarrollo de un prototipo. Algunas ideas sobreviven a este proceso más tiempo que otras, pero sólo unas cuantas logran superar totalmente el paso hacia la comercialización.

Una empresa que deja que las ideas permanezcan en el embudo durante mucho tiempo tendrá menos probabilidades de desechar equivocadamente una buena idea. No obstante, sus costes serán más elevados que los de una empresa que deseche las ideas con gran rapidez. La empresa que tarda en rechazar las ideas también tardará más tiempo en lograr algún producto novedoso, es decir, cuanto más tiempo tarda en evaluar y desechar una idea, más tiempo permanecerán en la cola las ideas restantes. Esto provocará un aumento del ciclo total de tiempo de desarrollo de un producto.

Por el contrario, la empresa que desecha rápidamente las ideas, reducirá el tiempo y el coste de su ciclo total de desarrollo, pero puede provocar la eliminación accidental de una magnífica idea. Este error es muy fácil de cometer con las ideas más innovadoras, que requieren un mayor esfuerzo de evaluación. Un rechazo rápido de las ideas puede desmotivar y dejar insatisfechas a las personas encargadas de generar ideas.

Los criterios que utilizan las empresas para seleccionar su cartera de proyectos no deben ser únicamente financieros, sino que factores de tipo estratégico deben condicionar igualmente la selección de los proyectos.

2.6. Cuadro de mando integral

Un sistema de control ampliamente aceptado por la comunidad académica y por las empresas es el cuadro de mando integral, un conjunto de medidas capaz de visualizar una perspectiva global y rápida de la empresa (figura 16). Permite a los directivos contemplar la empresa desde cuatro ángulos importantes: a) ¿cómo nos ven los clientes? (perspectiva de los clientes); b) ¿en qué tenemos que destacar? (perspectiva interna); c) ¿podemos continuar mejorando y creando valor? (perspectiva de innovación y aprendizaje) y d) ¿qué les parecemos a los accionistas? (perspectiva financiera). Al mismo tiempo, minimiza la sobrecarga de información, al limitar el número de medidas utilizadas (Kaplan y Norton, 1992).

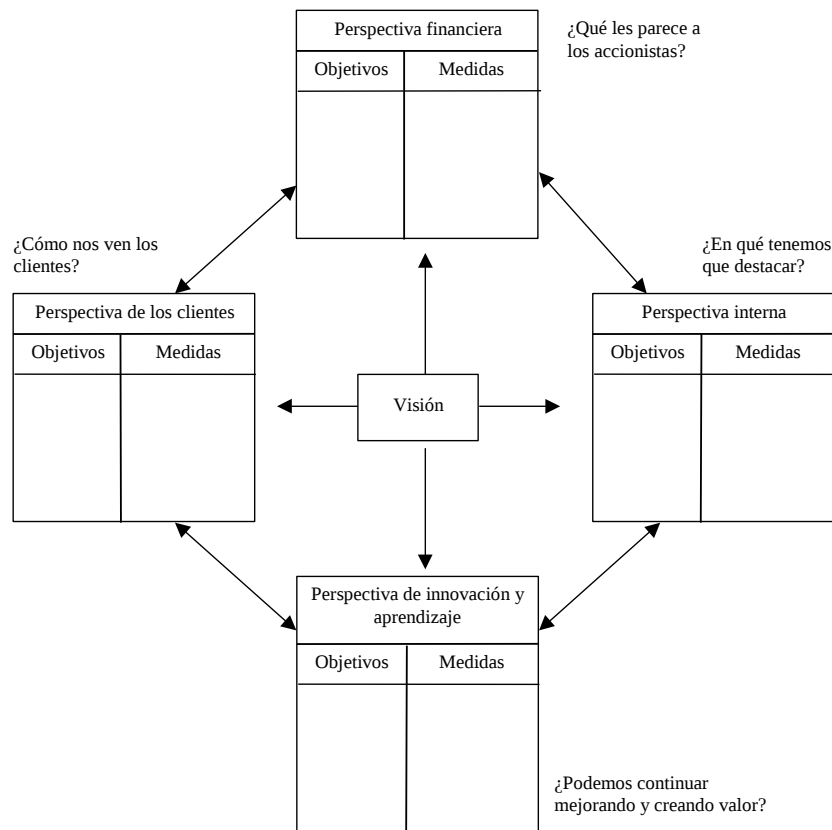


Figura 16: Traducción de la visión y de la estrategia: Cuatro perspectivas (Kaplan y Norton, 1992)

El hecho de depender de las valoraciones de los clientes para definir algunas de los objetivos de control obliga a la empresa a contemplar su rendimiento a través de las necesidades del mercado. A su vez, los gerentes necesitan concentrarse en las operaciones internas claves que mejoren la competitividad, especificando los procesos y las capacidades en los que desean sobresalir. Ahora bien, los objetivos para obtener éxito cambian continuamente. Por ello, la aptitud de una empresa para innovar y mejorar va unida directamente al valor de la empresa. Finalmente, las medidas financieras del rendimiento indican si la estrategia de la empresa y su ejecución están contribuyendo a la mejora del rendimiento neto.

El cuadro de mando integral reúne en un solo informe para la alta dirección los elementos necesarios para lograr una ventaja competitiva sostenible, vinculando los objetivos estratégicos a largo plazo con las acciones a corto plazo. También protege a la organización de la suboptimización. Al forzar a los altos directivos a considerar conjuntamente todas las medidas de resultados importantes, el cuadro de mando integral permite observar si la mejora en un área se ha conseguido a expensas de otras (Kaplan y Norton, 1992).

El cuadro de mando integral traduce la visión y estrategia de una empresa en un conjunto coherente de objetivos y medidas que se pretenden alcanzar en cada una de las cuatro perspectivas (clientes, interna, innovación y aprendizaje y financiera). Los objetivos deben ser consistentes y reforzarse mutuamente. Cada objetivo requiere una o

varias medidas apropiadas, a fin de que puedan gestionarse y validarse. Estas medidas de los objetivos tienden a ser efectos como la rentabilidad, cuota de mercado, satisfacción del cliente o capacidades de los empleados. Para cada medida se fija el valor a alcanzar y las iniciativas que hay que poner en marcha.

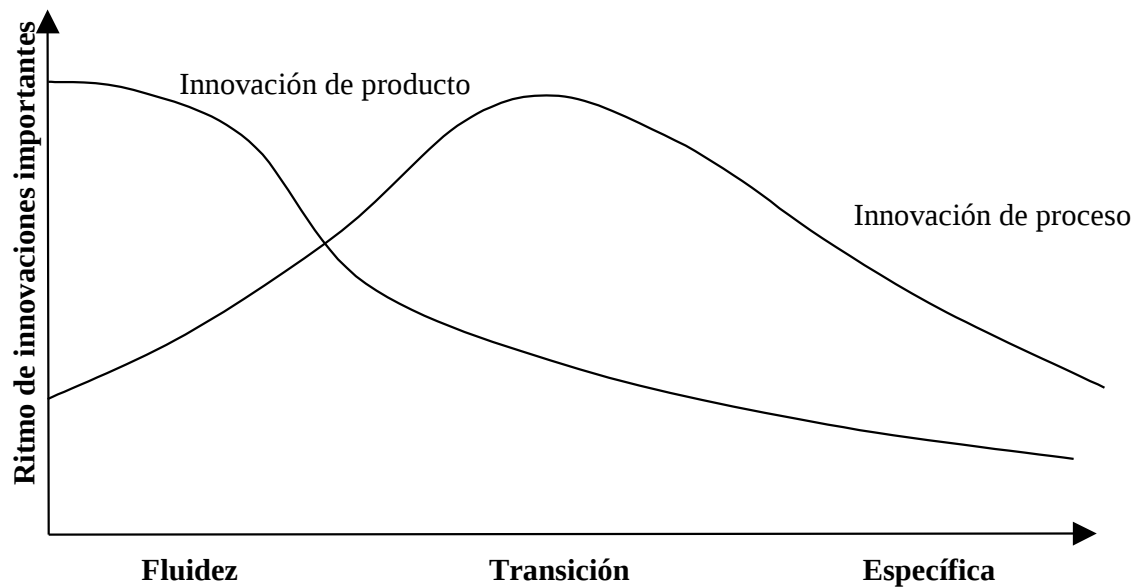
3. DINÁMICA DE LA INNOVACIÓN

Toda innovación evoluciona para satisfacer mejor las necesidades cambiantes del mercado. Esta evolución determina la estrategia a seguir y las capacidades que debe desarrollar la empresa.

3.1. Interacción producto/proceso

Los desarrollos del producto y del proceso constituyen un sistema integrado y su mutua dependencia generalmente se hace cada vez más fuerte con el transcurso del tiempo. En primer lugar, se desarrollan las características funcionales de los productos, mientras el proceso de producción se mantiene al margen. Después, al desarrollarse la industria, y una vez que se vislumbra un gran mercado potencial, la empresa se concentra en mejorar el proceso productivo con objeto de lograr las economías de escala que permitan fabricar el producto con un coste bajo. Finalmente, se estandariza el desarrollo tanto del producto como del proceso para, así, alcanzar la mayor eficiencia posible.

La figura 17 representa la evolución de las innovaciones en producto y en proceso de una empresa a lo largo de un período de tiempo. También puede interpretarse como el proceso evolutivo de la industria. La realidad es mucho más compleja, pero esta simplificación permite comprender mejor la evolución de la innovación.



Producto	Desde una gran variedad, a un diseño dominante, a innovación incremental de productos estandarizados.
Proceso	La fabricación progresa de una mano de obra cualificada y equipo de uso general a mano de obra poco cualificada y equipo de uso específico.
Organización	De una empresa con una estructura orgánica a una empresa mecánica, con tareas y procedimientos definidos y con escasas recompensas para la innovación radical.
Mercado	De mercados fragmentados e inestables con productos diferenciados y rápida retroalimentación a mercados con productos bastante indiferenciados.
Competidores	De muchas empresas pequeñas con productos diferentes (incompatibles) a un oligopolio de empresas con productos similares.

Figura 17: Dinámica de la innovación (Abernathy y Utterback, 1978; Utterback, 1994)

El modelo considera tres etapas (Abernathy y Utterback, 1978): a) la *etapa de fluidez* comprende el desarrollo del producto y la creación del mercado, b) la *etapa de transición* enfatiza las mejoras en el proceso productivo y c) la *etapa específica* concentra los esfuerzos en la reducción de costes y el control de calidad. En la figura 18 se recogen las características de cada una de estas tres etapas.

Etapa de fluidez. En esta etapa predominan las actividades de investigación dirigidas a obtener nuevos productos que satisfagan la demanda del mercado. Por tanto, hay una gran cantidad de incertidumbre tecnológica y de mercado. No está claro cuál es el mercado objetivo o qué características técnicas del producto servirán sus intereses. Las innovaciones pueden proceder de las fuentes más inesperadas, si bien la gran mayoría son externas a la industria. Existen diferentes diseños de productos incompatibles comercializados en su mayor parte por pequeñas empresas. Estos diseños, en cierto

modo, no son sino experimentos tecnológicos, que varían a medida que los fabricantes aprenden más sobre las necesidades del mercado y los clientes comprenden mejor el potencial de la tecnología en desarrollo. Por ejemplo, en la industria del automóvil no era en modo alguno obvio que el motor de combustión interna Otto de cuatro tiempos ganara a todos sus competidores. De hecho, en 1900 se fabricaron en Estados Unidos 4.192 coches. De estos, 1.681 tenían motor de vapor, 1.575 motor eléctrico y sólo 936 motor de gasolina (Basalla, 1988). El crecimiento de la demanda es lento debido a que el producto es nuevo y desconocido para los posibles compradores, pero también para la empresa que lo lanza al mercado. Cuando se crea una industria no está claro cuál es el mercado objetivo o qué características técnicas del producto servirán sus intereses, por lo que hay una gran incertidumbre tecnológica y de mercado. En el sistema productivo predominan las herramientas y las máquinas de uso general, con lo que el proceso estará desconectado y será bastante ineficiente. El personal a cargo de las máquinas está muy cualificado, pues, a menudo, consta de investigadores e ingenieros. Los cambios de proceso son frecuentes debido a la rápida evolución de la tecnología de producto. Los componentes utilizados en el proceso de transformación se compran en el mercado. Ello favorece las numerosas modificaciones que se deben realizar en el producto con objeto de mejorarlo técnicamente y corregir los fallos originales. La organización es orgánica, con una estructura flexible que favorece la descentralización de las decisiones y los cambios internos. El papel del marketing consiste en estimular la demanda y buscar nichos donde competir.

La ventaja competitiva del innovador se apoya en la originalidad y superioridad funcional del nuevo producto más que en su coste. En esta fase los costes serán altos, debido a que: 1) las series de producción son relativamente pequeñas, 2) los problemas tecnológicos no están totalmente resueltos, 3) los gastos de distribución y promoción son cuantiosos y 4) los canales de distribución están deficientemente desarrollados. Los productos tienen unos precios elevados, lo que limita el mercado potencial.

En la etapa de fluidez, ninguna empresa domina el mercado, son fabricantes desconocidos y los nombres comerciales importan poco. El producto necesita ser perfeccionado. Ninguna empresa tiene un proceso de fabricación eficiente ni ha conseguido un control inatacable de los canales de distribución. Los clientes no han desarrollado aún su propio sentido de diseño ideal del producto o lo que desean en cuanto a características técnicas o funcionales. Todo el mundo –productores y clientes– están aprendiendo sobre la marcha. Este ambiente conduce a la entrada en el mercado de muchas empresas, siempre que las barreras técnicas y de capital no sean demasiado elevadas (Utterback, 1994). Por ejemplo, Ford fue una de las 502 empresas creadas en Estados Unidos entre 1900 y 1908 para fabricar automóviles (Womack *et al.*, 1990). Las entrantes tienen orígenes muy diferentes. Algunas son empresas nuevas (entrantes de *novo*), mientras que otras son empresas establecidas que se diversificaron desde otros sectores (entrantes de *alio*). En esta etapa predominan las pequeñas empresas innovadoras (si bien algunas grandes empresas crean departamentos específicos para este tipo de actividades). Al final de la etapa deben optar entre retirarse del mercado o incrementar su capacidad para continuar comercializando el producto existente, lo que da paso a la etapa de transición.

Las barreras de entrada tienden a apoyarse en el acceso a los conocimientos tecnológicos. La rivalidad en esta etapa se apoya no tanto en los precios como en educar a los clientes, abrir canales de distribución y perfeccionar el diseño del producto. Tal rivalidad puede ser intensa, y la empresa que solucione primero los problemas de

diseño, a menudo, tiene la oportunidad de alcanzar una significativa posición en el mercado. Por ejemplo, en 1873, Remington and Sons (de Nueva York) fabricaron el primer modelo industrial de máquina de escribir dando origen a una nueva industria y también a nuevas profesiones, como la de mecanógrafo (en su origen este oficio era ocupado por hombres) (Hill y Jones, 2009).

	Fluidez	Transición	Específica
<i>Innovación</i>	Cambios frecuentes e importantes en el producto	Importantes cambios de proceso requeridos para una demanda creciente	Incremental para productos y con mejoras acumuladas en productividad y calidad
<i>Fuente de la innovación</i>	Pioneros en la industria; usuarios del producto	Usuarios, fabricantes	A menudo, proveedores
<i>Productos</i>	Diversidad de diseños, a menudo bajo pedido	Al menos un diseño de producto suficientemente estable para tener un volumen de producción significativo (diseño dominante)	Mayoritariamente indiferenciados, productos estándar
<i>Procesos productivos</i>	Flexibles e ineficientes, fácilmente adaptables a los cambios importantes en el producto	Comienzan a ser más rígidos produciéndose los cambios en grandes saltos	Eficientes, intensivos en capital y rígidos; altos costes de cambio
<i>Investigación y Desarrollo</i>	Poco específica, a causa del alto grado de incertidumbre técnica	Centrada en las características específicas del producto, una vez que aparecen los diseños dominantes	Centrada en tecnologías de producto incrementales; énfasis en tecnología de proceso
<i>Equipo</i>	De uso general, requiere trabajadores cualificados	Algunos subprocesos automatizados, creando islas de automatización	De uso específico, mayoritariamente automático, con mano de obra centrada en atender y controlar el equipo
<i>Planta</i>	Pequeña escala, localizada cerca del usuario o de la fuente de la innovación	De uso general con secciones especializadas	Gran escala, muy específica para productos concretos
<i>Coste de cambio de los procesos</i>	Bajo	Moderado	Alto
<i>Competidores</i>	Pocos, pero creciendo en número con cuotas de mercado muy fluctuantes	Muchos, pero disminuyendo en número tras la aparición del diseño dominante	Pocos, clásico oligopolio con estables cuotas de mercado
<i>Bases de la competencia</i>	Prestaciones funcionales del producto	Variaciones en el producto; adecuación al uso	Precio
<i>Control organizativo</i>	Informal y emprendedor	A través de proyectos y grupos de tareas	Estructura, normas y objetivos
<i>Vulnerabilidades de los líderes de la industria</i>	A los imitadores y copia de las patentes, al descubrimiento de productos radicales con éxito	A productores más eficientes que fabrican productos de mayor calidad	A las innovaciones tecnológicas que presentan los productos sustitutivos superiores

Figura 18: Características de las diferentes etapas del ciclo de vida producto–proceso (Utterback, 1994)

Cuando se percibe que el producto tiene un gran mercado potencial, será el momento de comenzar a realizar innovaciones en el proceso con objeto de reducir costes y aumentar el volumen de producción, ya que la competencia en precios resulta cada vez más intensa. Se inicia, pues, la *etapa de transición*.

Etapa de transición. En esta etapa el mercado está más definido y tiene lugar cierta estandarización de los componentes y una configuración definida del diseño del producto, es decir, surge un «diseño dominante» del producto, que supone una considerable reducción de la incertidumbre, la experimentación y cambios importantes en el diseño. Así pues, en el inicio de una industria existen diferentes variaciones del producto, una de las cuales triunfará y se convertirá en el diseño dominante.

Un diseño dominante² es aquél cuyos componentes principales y conceptos medulares básicos no varían considerablemente de un modelo de producto al otro, donde el diseño exige un alto porcentaje de participación en el mercado (Abernathy y Utterback, 1978).

La mayoría de los PC comparten un microprocesador Intel o compatible, memoria de acceso aleatorio (RAM), sistema operativo Microsoft, disco duro interno, una unidad de CD o DVD, un teclado, un monitor, un ratón, un modem, y más. Este conjunto común de características de un PC constituye un diseño dominante. Este diseño dominante derrotó a los diseños de las empresas Atari, Commodore y Apple, entre otras. La aparición del diseño dominante desplaza el énfasis competitivo a favor de aquellas empresas que son capaces de conseguir mayores capacidades en la innovación e integración de los procesos y que poseen unos conocimientos técnicos y de ingeniería más desarrollados. El campo de batalla pasa de la innovación de productos a la innovación de procesos (Utterback, 1994). Stobaugh (1988) examinando la experiencia de nueve empresas del sector petroquímico, encontró que el tiempo medio entre el momento de la innovación del producto y la innovación del proceso era de 5,7 años.

Un diseño dominante reduce drásticamente el número de requisitos de las prestaciones que debe cumplir un producto al hacer que muchas de estas exigencias estén implícitas en el propio diseño. No es necesariamente aquél que incorpora las prestaciones técnicas más avanzadas. Es aquél que satisface a muchos en cuanto a la interacción entre posibilidades técnicas y preferencias del mercado, en lugar de ser un diseño optimizado por unos pocos. Un diseño dominante es el resultado de una serie de decisiones técnicas acerca del producto, limitadas por elecciones técnicas previas y la evolución de las preferencias de los clientes (Utterback y Suárez, 1993). El diseño dominante incluye las necesidades de muchas clases de usuarios de un producto en especial, aun cuando no pueda satisfacer las necesidades de una clase particular tan bien como lo haría un diseño hecho específicamente para ella. Un diseño dominante es una trayectoria específica, a lo largo de la jerarquía de diseño de una industria, que establece un dominio entre las trayectorias de diseño que compiten entre sí. Anderson y Tushman (1990) encontraron que un diseño dominante siempre surge para conseguir una gran cuota de mercado a menos que la siguiente discontinuidad llegue demasiado pronto y rompa el ciclo, o varios productores patenten sus tecnologías propietarias y se nieguen a licenciarse entre sí. También observaron que el diseño dominante nunca tenía la misma forma que la discontinuidad original y que tampoco se encontraba en la vanguardia de la tecnología.

² El modelo de diseño dominante se adapta mejor a mercados de masas, en los que los gustos del consumidor son relativamente homogéneos. Además, valen solamente para industrias de fabricación de productos ensamblados (Teece, 1986b).

En lugar de maximizar el rendimiento o alguna dimensión individual de la tecnología, el diseño dominante tendía a aunar una combinación de características que satisfacen las demandas de la mayoría del mercado. Una trayectoria es el camino del progreso técnico establecido por la elección de un concepto tecnológico central en el comienzo. En la figura 19 se representan dos trayectorias.

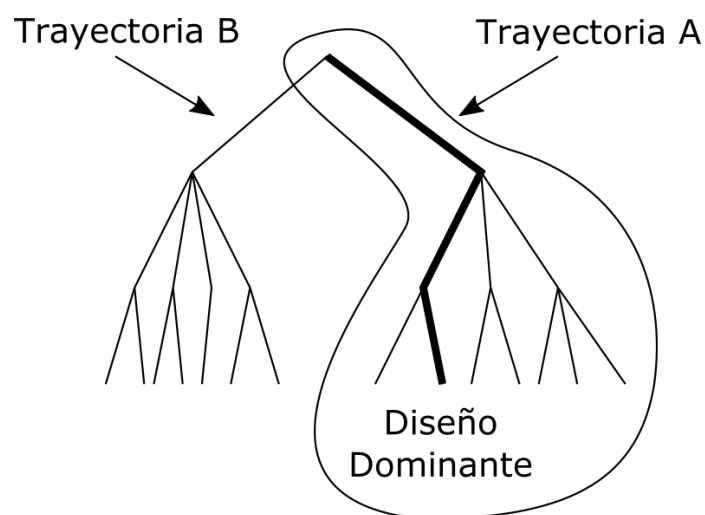


Figura 19: Diseño dominante (Utterback y Suárez, 1993)

El diseño dominante estabiliza la competencia entre empresas rivales, disminuye el número de competidores y cambia las bases competitivas desplazándose a refinamientos en las prestaciones del producto, la fiabilidad y el coste. Por tanto, se compite sobre la base de la diferenciación del producto. También se favorece la mecanización de los procesos y la expansión de los mercados a través de la extensión de la línea de productos. Las necesidades de los usuarios se entienden con mayor claridad. El centro de atención de las empresas se desplaza a la fábrica.

Antes de que finalmente un diseño se erija como dominante, los pioneros se centran en la mejora del producto, ya que, debido a la variedad de diseños incompatibles, existe una gran incertidumbre tecnológica (Abernathy y Utterback, 1978; Dosi, 1982). Una vez que surge el diseño dominante, el centro de atención de las empresas se desplaza a la fábrica (Suárez y Utterback, 1995). El campo de batalla pasa de la innovación del producto a la innovación del proceso productivo (Abernathy y Utterback, 1978). Por tanto, el diseño dominante proporciona una ventaja competitiva a las empresas capaces de conseguir mayores capacidades en la innovación e integración de los procesos productivos (Utterback, 1994). En este sentido, Klepper (1996) considera que el diseño dominante lo sustentan las economías de escala logradas por la innovación radical en proceso. Podría decirse que el diseño dominante es un modelo de negocio que facilita la producción en masa para atender el mercado mayoritario.

En la etapa de transición predominan las innovaciones radicales en maquinaria y/o de la organización del sistema productivo, surgiendo nuevas estructuras de producción. Como consecuencia, el sistema de producción tiende a ser más elaborado y fuertemente integrado, mediante una creciente automatización y control del proceso. Los materiales se hacen más especializados y costosos y se compite sobre la base de productos diferenciados. La organización tiende a burocratizarse, se comienzan a definir tareas, así

como a establecer controles rígidos y un reparto de funciones. En esta etapa predominan las grandes empresas. Sería de esperar que, con anterioridad a la aparición de un diseño dominante, surgiera una oleada de empresas entrantes con muchas versiones del producto. Posteriormente, una vez aceptado el diseño dominante, tendrá lugar una oleada de salidas de empresas competidoras (Christensen, 1997).

El surgimiento de un diseño dominante es importante por varias razones (Hamel y Prahalad, 1994; Hill y Jones, 2009). En primer lugar, la falta de un diseño dominante es un obstáculo para la difusión del nuevo producto, ya que los fabricantes de componentes y productos complementarios deben diseñar distintos modelos para las diferentes variedades de diseño, lo que se traduce en una subida de precios y un mercado mucho más lento en despegar. En segundo lugar, la existencia de una variedad de diseños incompatibles confunde a los clientes y los hace mostrarse menos proclives a comprar, pues prefieren esperar a que surja un claro vencedor. En tercer lugar, ayuda a reducir los costes de producción. Una vez que surge un diseño dominante, el producto y sus complementarios pueden fabricarse en masa, ya que a corto plazo no van a sufrir variaciones apreciables. En cuarto lugar, la existencia de un diseño dominante disminuye el riesgo de comprometer recursos en un producto que, al final, va a tener que abandonar el mercado por falta de demanda.

El diseño dominante redunda en la consecución de economías de escala en la producción, lo que, a su vez, facilita la reducción del coste en el avance desde el nicho hacia el mercado mayoritario. Este tránsito desde un nicho al mercado mayoritario, puede denominarse, utilizando términos de Moore (2002), ‘cruzar el abismo’, por lo peligroso que resulta.

La obligación de atender mercados cada vez más amplios requiere empresas grandes, con capacidad de obtener recursos para financiar su crecimiento y adquirir el tamaño adecuado para operar eficientemente. Los competidores invierten en capacidad para incrementar su participación en el mercado. Los costes unitarios bajan como consecuencia de un mayor volumen de producción, la acumulación de experiencia y la estandarización del producto. Los precios se mantienen al mismo nivel o se reducen ligeramente a medida que la demanda se incrementa. Los márgenes mejoran durante esta fase como consecuencia de la distribución de los costes fijos entre un mayor volumen de producción. El crecimiento y los beneficios durante esta fase atraerán competidores adicionales. Algunos competidores entran en el mercado con ‘fotocopias’ del producto original, mientras que otros introducen mejoras funcionales en productos y procesos. Esto fomenta la adaptación del producto a las necesidades específicas del mercado de masas, a cuyas exigencias no respondía con precisión la versión inicial. Con ello, el mercado se expande y el proceso de crecimiento se realimenta. Cada empresa tratará de diferenciar su producto y de ir creando una imagen de marca propia. De hecho, la entrada de nuevas empresas intensifica la actividad de marketing de la industria, lo que puede incidir favorablemente sobre las ventas. El mercado está más definido, hay una menor incertidumbre y los procesos de producción son más eficientes.

Etapa específica. En esta etapa proliferan los productos que se crean alrededor del diseño dominante. Las grandes empresas controlan el mercado con un producto mejorado funcionalmente sobre la idea original y compitiendo agresivamente en precios. También se aprovechan, por lo general, de su imagen corporativa, fruto de su fortaleza económica y de los muchos años de presencia en el mercado. En esta fase las

innovaciones en producto y en proceso alcanzan su menor grado de novedad, teniendo, por lo tanto, un carácter incremental. Los materiales son ahora muy especializados y el equipo de uso específico. Los productos están muy definidos y las diferencias con los productos de los competidores son menores que las similitudes.

A medida que el proceso está diseñado y organizado para fabricar productos específicos y que requiere inversiones elevadas, el perfeccionamiento selectivo de los elementos del proceso resulta cada vez más difícil (Abernathy y Utterback, 1978). El proceso llega a estar tan integrado que cualquier cambio resulta muy costoso, puesto que obliga a realizar modificaciones en los restantes elementos del proceso e, incluso, en el diseño del producto. El cambio del diseño del producto o del proceso se produce lentamente en esta fase, ya que ambas tecnologías están coespecializadas (Teece, 1986), lo que significa que en cierto sentido son únicas y complementarias y son más productivas cuando son utilizadas conjuntamente.

La base de la competencia se desplaza al precio, se reducen los márgenes de venta, el conjunto de los competidores tiende, en ocasiones, a convertirse en un oligopolio y en producción la atención se centra en la productividad y las economías de escala. A medida que la competencia en los precios se incrementa, los procesos de producción necesitan mayores inversiones de capital y cabe la posibilidad de que se decidan nuevos emplazamientos para obtener costes más bajos en los abastecimientos de las materias primas. Estas localizaciones pueden crear nuevas posibilidades en los mercados extranjeros.

Por último, cabe señalar que, en general, al analizar estáticamente el mercado, se observa que las empresas se encuentran en alguno de estos tres estados. No obstante, al estudiar los cambios en la empresa durante un largo período de tiempo, se puede percibir que algunas empresas han completado las tres etapas.

La implicación principal de este modelo es que, a medida que la tecnología evoluciona a través de las diferentes fases, una empresa necesita diferentes capacidades para mantenerse competitiva en el mercado. El patrón descrito se repite cuando llega al mercado una nueva tecnología con el potencial para desplazar la tecnología establecida. A menudo, un competidor ajeno a la industria introduce esta tecnología. Esto provoca una ruptura del modelo, lo que vuelve a precipitar el ciclo de innovación a la fase fluida, con otra oleada de empresas que entran en el mercado.

La experiencia con un cierto número de industrias respalda la opinión de que cada oleada de innovación repite el esquema de la interconexión de innovaciones de producto y de procesos, y la importancia del diseño dominante sobre el número de empresas que un sector industrial puede incluir en un momento dado. Cada nueva oleada de innovación tiene sus fases fluida, transición y carácter específico: el ritmo de innovaciones de producto alcanza un máximo más o menos pronto y experimenta un impulso de la innovación de procesos en el mismo momento en que la de producto disminuye. Cada uso se caracteriza por un máximo de empresas competidoras en algún momento alrededor de la aparición del diseño dominante, con una disminución a partir de ese punto. Ahora bien, merece la pena destacar el menor número total de empresas que participan en la segunda oleada de innovación. Así, mientras que había muchas empresas fabricando lámparas incandescentes (la primera oleada en iluminación

eléctrica), hubo menos intentando entrar en el negocio de las lámparas fluorescentes (la segunda oleada) (Utterback, 1994).

Este modelo asume varias conjeturas (Afuah, 1998). La primera es que la evolución de la innovación es lineal y previsible, al tiempo que se pueden distinguir las fronteras entre las diferentes etapas: fluida, transición y específica. La realidad nos dice que la linealidad y la duración de las etapas son muy cuestionables y las fronteras, en el mejor de los casos, borrosas. La segunda conjetura es que la innovación radical en proceso tiene lugar después de la innovación radical en el producto. Esto no ocurre en todos los casos: por ejemplo, Intel alterna las innovaciones radicales de producto con las de proceso. La tercera conjetura es que puede no surgir un diseño dominante y, en algunos casos, es difícil determinarlo. La cuarta es que no siempre resulta fácil separar la innovación en proceso de la innovación en producto.

3.2. Destrucción creativa

Cuando una empresa introduce una innovación radical, comienza en términos de Schumpeter (1934) un proceso de «destrucción creativa», que acabará transformando totalmente el mercado mayoritario que ataca.

Si bien se parte de la apreciación general de la tecnología como el conocimiento científico (u otro conocimiento organizado) aplicado a tareas prácticas, en este apartado se restringe su uso únicamente a una combinación schumpeteriana (plataforma o arquitectura); es decir, a un nuevo producto ensamblado formado por varios componentes. En el caso de una nueva combinación, con independencia de que se introduzca en un nicho o directamente en el mercado mayoritario, la denominación utilizada será tecnología emergente, mientras que, en el caso de una tecnología que haya triunfado en el mercado mayoritario, el calificativo será de dominante. Finalmente, se considera que una tecnología cuenta con varias dimensiones de valor (o métricas), que, para simplificar, resumiremos en una dimensión tecnológica primaria y otras secundarias (o auxiliares). La primaria representa la prestación tecnológica en la que se centran los clientes al comprar un producto. Las dimensiones secundarias caracterizan las otras prestaciones tecnológicas del producto esperadas por los clientes en un nivel satisfactorio; no contribuyen a mejorar la competitividad de la empresa, pero son unos requisitos tecnológicos básicos que, si el producto no los posee en los niveles adecuados, desembocaría en una desventaja competitiva. A partir de ahora, para simplificar la argumentación, solo hablaremos de una única dimensión secundaria (que contiene al conjunto de todas ellas).

El mercado principal es un mercado muy rentable en el que una tecnología dominante sustenta productos ensamblados con prestaciones muy definidas, de forma que las diferencias entre las distintas marcas competidoras son menores que las similitudes. No obstante, las empresas establecidas, para no verse rezagadas y perder cuota de mercado, continúan mejorando las prestaciones y la eficiencia de la tecnología dominante. Las exigencias del mercado principal respecto a la considerada dimensión tecnológica primaria aumentan con el tiempo, por lo que las empresas establecidas deben evolucionar las prestaciones del producto en esa dimensión. La competencia tecnológica entre las empresas es elevada, pues existe el temor de que, si un competidor consigue mejorar la dimensión primaria, captará clientes de las empresas establecidas que no lo hayan hecho (Langlois y Robertson, 1995).

El mercado principal puede ser atacado por una tecnología emergente de naturaleza discontinua o disruptiva. La tecnología discontinua ataca directamente el mercado principal –ataque directo–, mientras que la tecnología disruptiva utiliza previamente un nicho o segmento de mercado como cabeza de playa –ataque indirecto–. Ambas son innovaciones radicales que provocan la destrucción creativa schumpeteriana, pero presentando diferencias claras (figura 20).

Tipo de tecnología	Mercado atacado	Factor clave del éxito
Tecnología discontinua	Mercado mayoritario	Dimensión primaria
Tecnología disruptiva potencial ↓	Nicho	Dimensión primaria
	Segmento (bajo o alto)	
	Mercado extranjero	
Tecnología disruptiva	Mercado mayoritario	Dimensión secundaria

Figura 20: Características determinantes de las tecnologías disruptiva y discontinua

3.2.1. Tecnología discontinua

Una tecnología discontinua es aquella que ataca directamente el mercado principal, y supone una mejora radical en la dimensión primaria del desempeño tecnológico de la tecnología dominante. Si representamos la trayectoria tecnológica mediante una curva S (Rogers, 1962; Moore, 1991), la tecnología discontinua se escenifica mediante una pareja de curvas S que se cruzan en el primer cuadrante de un plano cartesiano (figura 21). El salto entre las dos curvas de la pareja representa una discontinuidad tecnológica: un cambio no evolutivo (Veryzer, 1998). Comparativamente, en el momento de introducir la tecnología discontinua en el mercado principal, la tecnología dominante tiene un mejor desempeño en la dimensión primaria, pero un potencial de mejora menor. Además, la tecnología discontinua tiene un límite técnico intrínsecamente superior al de la tecnología dominante (Tushman y Anderson, 1986). En algún momento, en lo que respecta a la dimensión primaria, el desempeño de la tecnología discontinua superará las prestaciones de la tecnología dominante, provocando que los clientes se desplacen a la nueva tecnología (Cooper y Schendel, 1976; Foster, 1986; Utterback, 1994).

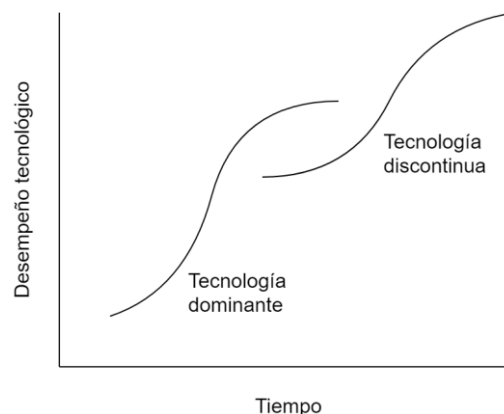


Figura 21: Discontinuidad tecnológica

El 13 de diciembre de 1907, el barco de vela *Thomas W. Lawson* se hundió frente a las islas Scilly en el canal de la Mancha. Todos sus tripulantes desaparecieron, excepto el capitán y un marinero. Podría haber sido otro desgraciado naufragio si el *Thomas W. Lawson* no hubiera sido un barco especial. Tenía siete mástiles y había sido diseñado para competir con los nuevos barcos de vapor, que, cada vez con más frecuencia, arrebatan a los veleros sus fletes. Era capaz de navegar a 22 nudos si soplaban un buen viento. Pero sus diseñadores habían tenido que sacrificar maniobrabilidad para ganar velocidad, por lo que resultaba un barco pesado y difícil de manejar. De hecho, era tan inestable que zozobró mientras estaba anclado durante una fuerte galerna. Nunca más se diseñó un velero mercante tan rápido. La era de los mercantes a vela se cerró con el *Thomas W. Lawson* y los barcos de vapor empezaron a dominar los mares (Foster, 1986). La discontinuidad es un cambio no evolutivo (figura 22).

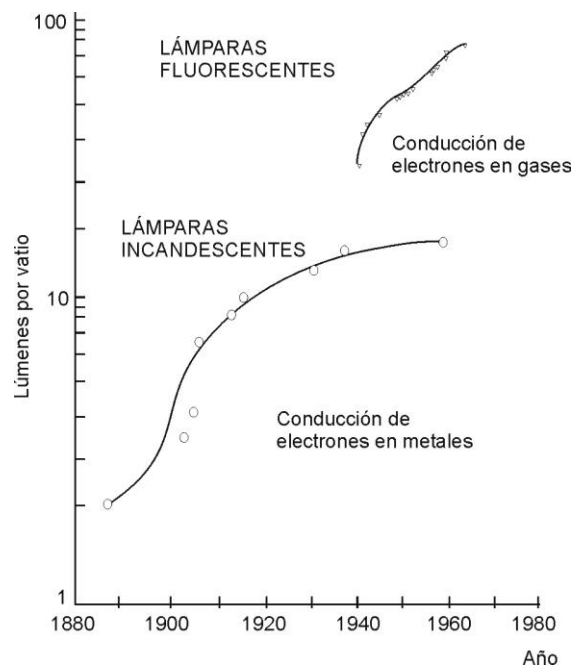


Figura 22: Curva S de la tecnología para el progreso tecnológico de las lámparas

La innovación discontinua satisface la necesidad del mercado principal, pero lo hace a partir de una base de conocimiento completamente nueva. Por ejemplo, el cambio de los discos de vinilo a los discos compactos o el cambio de los aviones de hélice a los reactores, son discontinuidades tecnológicas (Schilling, 2008). La tecnología discontinua es una innovación radical que emerge en el mercado de la tecnología dominante. Por tanto, el pionero ataca directamente la posición en el mercado de las empresas establecidas.

En el momento de la introducción de una tecnología discontinua (t_1), la tecnología dominante proporciona, generalmente, mejores prestaciones. No obstante, se produce entonces una ruptura en el desarrollo tecnológico, ya que empieza a evolucionar la nueva tecnología (figura 23), que no está sostenida por los mismos conocimientos en los que se apoyaba la dominante, sino que se desarrolla a partir de conocimientos de base completamente diferentes (Foster, 1986). Normalmente, la tecnología discontinua es muy tosca comparada con la dominante y requiere muchos perfeccionamientos y mejoras. A su vez, tampoco está clara la evolución del desarrollo tecnológico de la tecnología discontinua, puede alcanzar el valor A en la figura 23 pero también el B, desapareciendo su potencial de mejora sin poder superar a la tecnología dominante.

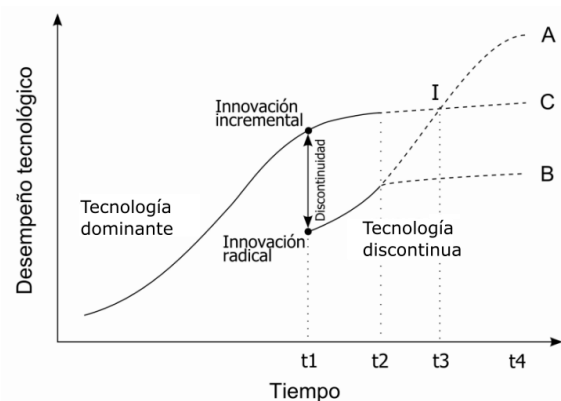


Figura 23: Discontinuidades tecnológicas

Sólo el tiempo (más bien las inversiones en tecnología) permite resolver los problemas asociados a la tecnología discontinua. Si evoluciona correctamente, llegará un momento en que ofrecerá unas ventajas similares a las de la dominante (t_3). Sin embargo, la discontinua al contar con un potencial de mejora mayor, desplazará del mercado a la dominante. Esto acontece cada vez con mayor rapidez en los mercados y conlleva un enorme riesgo para las empresas establecidas. La tecnología discontinua no desplaza la tecnología dominante en el inicio. Tiene que empezar captando clientes que por diversas razones la valoran más. En general, después de que la tecnología discontinua haya sido introducida, la tecnología dominante continúa durante un tiempo siendo mejorada (Cooper y Schendel, 1976; Utterback, 1994). El período de tiempo que transcurre desde la introducción de una tecnología discontinua hasta que sus ingresos superan los de la tecnología dominante varía entre 5 y 14 años (Cooper y Schendel, 1976).

La esencia de la dirección de la tecnología consiste en detectar la tecnología discontinua que puede superar el punto de inflexión (I) de la curva S (figura 23). El desafío consiste en cambiar de tecnología cuando se intersectan las curvas S de las tecnologías dominante y discontinua. La incapacidad de anticipar la amenaza de la tecnología discontinua y de pasarse a ella en el momento oportuno es la causa del fracaso de las empresas establecidas en el mercado principal y las líderes de la industria.

Conviene destacar que el inicio de la tecnología discontinua viene marcado por lo que Tushman y O'Reilly (1997) han denominado «período de fermentación tecnológica». Durante este período de fermentación, distintas variantes tecnológicas, cada una de ellas con distintos principios operativos, compiten por la aceptación del mercado (figura 24).

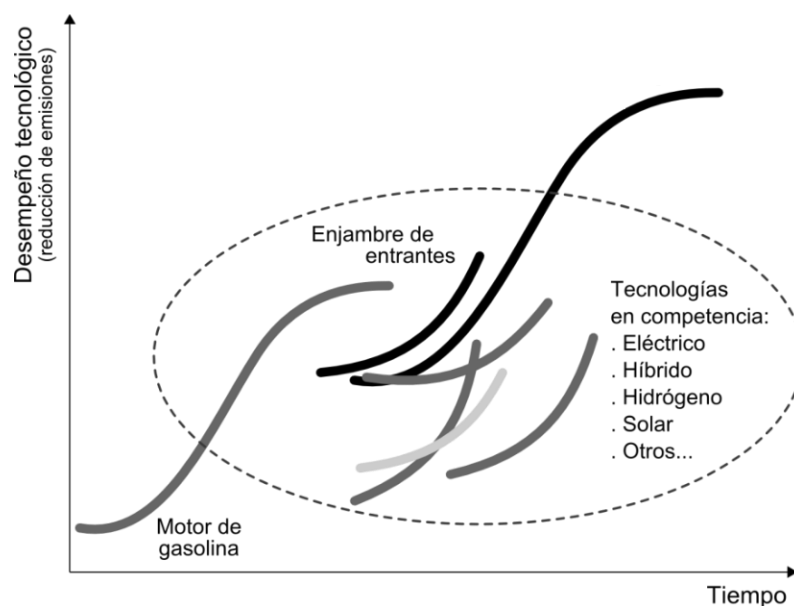


Figura 24: Enjambre de nuevos entrantes

En este sentido, Anderson y Tushman (1990) encontraron que la tecnología discontinua puede ofrecer importantes prestaciones, pero existe poco acuerdo sobre cuáles deberían ser los principales subsistemas de la tecnología o cómo deberían ser configurados en su conjunto. Así, mientras que la nueva tecnología desplaza a la vieja (período de *sustitución*), existe una considerable competencia conforme las empresas experimentan con diferentes diseños de la tecnología discontinua. El surgimiento de un diseño dominante señala la transición de la etapa de fermento a la de cambio incremental. En esta etapa, las empresas se centran en la eficiencia y la penetración de mercado. Las empresas tratan de conseguir una mayor segmentación de mercado ofreciendo diferentes modelos y precios. Pueden tratar de reducir los costes simplificando el diseño o mejorando el proceso de producción. Este período de acumulación de pequeñas mejoras puede suponer la mayor parte del progreso tecnológico en una industria y continúa hasta la siguiente discontinuidad tecnológica.

La competencia tiene lugar entre la tecnología discontinua y la dominante, así como entre las distintas variaciones de la tecnología discontinua. Ahora bien, la introducción de una tecnología discontinua no implica necesariamente una disminución inmediata de las ventajas de la tecnología dominante. Incluso muchas veces puede seguir expandiéndose durante un largo período de tiempo. Por ejemplo, uno de los primeros automóviles, introducido en 1771 por Nicolas Joseph Cugnot, nunca fue puesto en producción comercial debido a que era mucho más lento y difícil de manejar que un carruaje de caballos. Tenía tres ruedas, estaba impulsado por vapor y podía desplazarse a 3,5 kilómetros por hora. En el siglo XIX se introdujeron varios vehículos impulsados por vapor y gasolina, pero no fue hasta los inicios del siglo XX cuando los automóviles fueron producidos en masa (Schilling, 2008). En muchos casos, la tecnología dominante continúa siendo mejorada y alcanza un alto estado de desarrollo tecnológico después de que la tecnología discontinua haya sido introducida en el mercado. En este sentido, Strassmann (1959) señala que, para algunas de las más importantes innovaciones en producción de energía, metalurgia de hierro y otras industrias, la producción total con la antigua tecnología continuó creciendo en términos absolutos, mucho tiempo después de la introducción de la nueva tecnología. Los nuevos productos y las nuevas formas de hacer las cosas se van introduciendo paulatinamente en el mercado y durante largo tiempo coexisten con los antiguos. La tasa

de sustitución es lenta no solo porque los viejos productos pueden ser más eficientes que los nuevos, sino porque las empresas establecidas mejoran la calidad y las prestaciones de los mismos. También puede ocurrir que disminuya su precio, lo que retrasa su sustitución por los nuevos productos (Brozen, 1951). Las empresas establecidas que continúan comercializando la tecnología dominante se resisten a hacerla desaparecer del mercado, por lo que siguen perfeccionando sus productos y procesos con objeto de lograr una mayor eficiencia, que les permita reducir sus precios de venta o mejorar sus prestaciones.

Los recién llegados con innovaciones radicales deberían dar por sentado que las empresas establecidas responderán con inversiones en investigación para mejorar la tecnología dominante (figura 25). Con frecuencia las tecnologías discontinuas inducen respuestas vigorosas e imaginativas en las empresas establecidas a las que proporcionan un inmediato sustituto (Rosenberg, 1976). Por ejemplo, a medida que los fabricantes de hielo artificial ocupaban mayores espacios en el negocio del hielo, los comerciantes de hielo natural que los cosechaban en los grandes lagos de América del Norte continuaron con sus propias mejoras. Se utilizaron sierras circulares accionadas por vapor para la tarea de cortar bloques de hielo de ríos y lagunas. Simultáneamente, se instalaron cintas transportadoras mecánicas para proporcionar un transporte continuo de los trozos de hielo hasta las neveras, o a vagones de ferrocarril especialmente diseñados. De este modo, cuando la puerta de su industria se cerraba, los hombres del hielo del norte continuaron haciendo mejoras incrementales en los procesos de recogida, almacenamiento y entrega. Estas mejoras condujeron a mayores volúmenes y a costes unitarios inferiores. A pesar de la rápida expansión por el sur de EE.UU del hielo fabricado por máquinas, la recogida del hielo natural del año 1886 fue la mayor de toda su historia –25 millones de toneladas–, lo que indica cómo la muerte de una tecnología puede esconderse bajo un mercado creciente. En 1868 Nueva Orleans tenía su primera fábrica de hielo artificial. En 1879 había 30 plantas en los estados del sur y 5 en California. En 1889 había 222 plantas de hielo, la mayoría en el sur. Los cosechadores del hielo natural luchaban en una batalla perdida. A mediados de la década de 1920 la industria del hielo natural había desaparecido definitivamente, excepto en unas pocas zonas aisladas (Utterback, 1994). De hecho, cuando se producen discontinuidades el liderazgo cambia de manos en, aproximadamente, siete de cada diez casos (Foster, 1986).

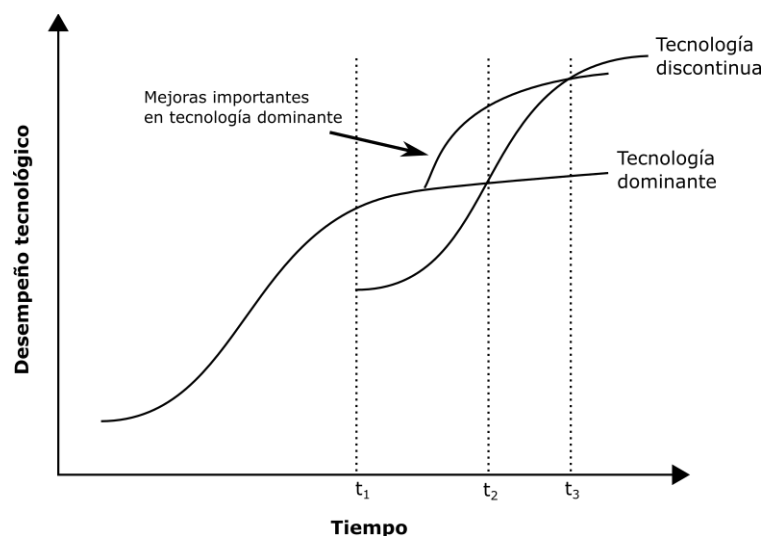


Figura 25: Prestaciones de un producto establecido y de un producto emergente (Utterback, 1994)

La curva *S* tiene algunas deficiencias (Afuah, 1998): a) los cambios inesperados en las necesidades de los usuarios o los avances en las tecnologías complementarias y de componentes pueden prolongar la vida de una tecnología que parecía estar cerca de su ocaso y la empresa tal vez descubra que cambió a la tecnología discontinua prematuramente; b) retrasar el cambio a la nueva tecnología hasta que ésta presente unos rendimientos superiores a los de la tecnología dominante puede privar a una empresa de una ventaja extra de la nueva tecnología (por ejemplo, la ventaja principal en pasar de las cajas registradoras electromecánicas a las electrónicas no es tanto la velocidad de los cálculos o del tamaño del producto como el que la electrónica puede dar a los clientes listas impresas de sus compras y, aún más importante, contribuir a controlar el inventario); c) llegar al límite físico tampoco es una condición suficiente para hacer el cambio, ya que la tecnología puede encontrar nuevas aplicaciones (por ejemplo, las empresas del gas no siguieron compitiendo en el mercado del alumbrado eléctrico, sino que encontraron nuevas aplicaciones para su tecnología –calentamientos de espacio y calor para procesos de fabricación– y abandonaron el mercado original del gas para la iluminación residencial y comercial) y d) no conviene olvidar que resulta difícil medir el progreso y que podemos estar cometiendo serios errores de apreciación.

En algunos casos también puede ocurrir que la tecnología discontinua, respecto a la tecnología dominante, tenga un mejor desempeño en la dimensión primaria desde el inicio (Sood y Tellis, 2005). Por tanto, un pionero discontinuo puede atacar a pequeña escala el mercado mayoritario, como ocurrió cuando Edson General Electric Company atacó a las empresas del gas en el muy rentable mercado del alumbrado, o también puede atacar a gran escala y abruptamente, como cuando IBM atacó el mercado de la máquina de escribir manual con la máquina eléctrica (Drucker, 1985). En ambos casos, la tecnología emergente tiene en la dimensión primaria un potencial de mejora (o una mejora real) muy superior a las prestaciones de la tecnología principal, representando un gran salto adelante, es decir, una discontinuidad (Veryzer, 1998). Asimismo, ninguno de los dos casos tiene una aplicación previa en un nicho o segmento.

Modelo de Foster. Foster (1986) elaboró un modelo para calcular el límite económico de una tecnología y determinar cuándo conviene abandonar las inversiones en investigación para el perfeccionamiento y mejora de la misma y dirigir las a una tecnología alternativa.

Toda inversión en tecnología genera un rendimiento para la empresa. Al principio, el progreso tecnológico es muy lento y los rendimientos escasos. Después, el ritmo se acelera, lo que también redundará positivamente sobre los resultados. En esta etapa, los resultados crecen, primero más que proporcionalmente a la inversión realizada en la mejora de la tecnología y, más adelante, menos que proporcionalmente. Con el tiempo, las inversiones para mejorar la tecnología dejan de producir resultados apreciables, ya que su nivel de desarrollo está rozando el límite económico factible. Llegará un momento en que, aunque se sigan realizando fuertes inversiones, cada vez resultará más difícil y costoso lograr progresos técnicos que proporcionen alguna ventaja competitiva adicional para la empresa. Hay un desfase entre las expectativas y las realizaciones. La mejora tecnológica resulta totalmente ineficaz. En este caso diremos que estamos alcanzando el límite de la tecnología, es decir, su maduración: no es el paso del tiempo lo que conduce al progreso,

sino la realización de un esfuerzo. Este límite económico se puede definir con relativa facilidad utilizando la relación rendimiento/inversión (Foster, 1986).

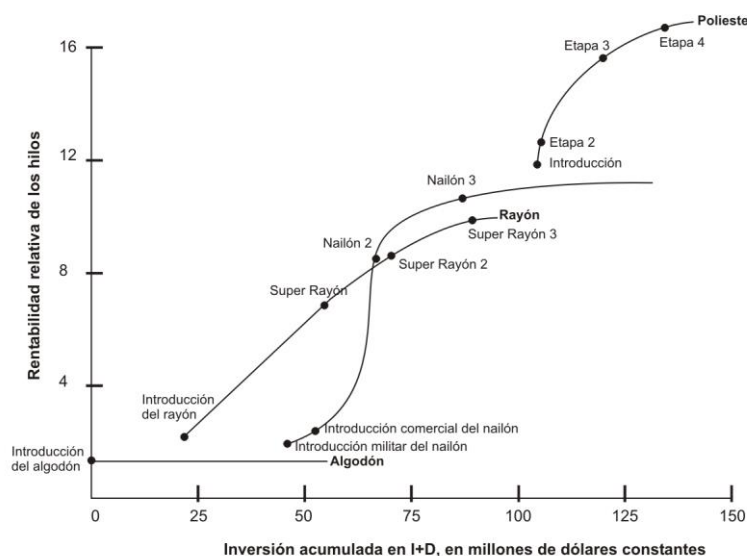


Figura 26: Curva en S de DuPont (Foster, 1986)

Las empresas no deberán invertir en una tecnología que no se puede mejorar, o cuyas mejoras no son perceptibles o aceptables por los clientes. Por ejemplo, en el caso de hilos para neumáticos, según Foster (1986), la empresa DuPont, no comprendiendo dónde se encontraba el *nylon* en la curva de rendimiento/inversión, ganó muy poco con los últimos 75 millones de dólares invertidos en I+D, mientras que su competidor Celanese progresó más deprisa con menos inversiones, porque invirtió en poliéster, una tecnología que se encontraba al principio de la curva (figura 26). Invertir en tecnologías que se encuentran en el límite tiene tan poco sentido como invertir en tecnologías que poseen un enorme potencial tecnológico intrínseco, pero poco interés para los clientes.

3.2.2. Tecnología disruptiva

El mercado principal también puede ser atacado por una tecnología disruptiva, término alrededor del cual existe mucha confusión. Christensen fue quien acuñó el calificativo disruptivo para la tecnología, si bien luego lo generalizó a una gran variedad de innovaciones, incluyendo las innovaciones en modelos de negocios, las innovaciones en servicios o las innovaciones sociales, con objeto de crear una teoría de la disrupción (Christensen *et al.*, 2018). De acuerdo con sus observaciones, en muchas ocasiones, mientras las empresas establecidas se centran en competir por mejorar la dimensión primaria de la tecnología dominante, puede emerger una nueva tecnología que logre proporcionar prestaciones superiores en la dimensión secundaria. Sin embargo, sus prestaciones en la dimensión primaria no satisfacen las necesidades del mercado principal, por lo que las empresas establecidas rechazan su comercialización. Ahora bien, el conjunto de prestaciones que proporciona la tecnología emergente puede satisfacer plenamente las necesidades del segmento del lado bajo del mercado mayoritario (Bower y Christensen, 1995; Christensen, 1997) o de un nicho no atendido (Christensen y Raynor, 2003). Con el tiempo, la tecnología emergente en el nicho o en el segmento, que al igual que Sood y Tellis (2011) denominamos «tecnología disruptiva potencial», irá mejorando sus prestaciones en las dimensiones primaria y secundaria.

Cuando la mejora en la dimensión primaria logra superar el nivel satisfactorio que exige el mercado principal, la tecnología disruptiva potencial lo invadirá, pasando en ese momento a denominarse tecnología disruptiva.

Los trabajos de Christensen recogen abundantes ejemplos de tecnologías disruptivas que inician su desarrollo en segmentos del lado bajo del mercado o en nichos no servidos. Uno muy ilustrativo es el de la excavadora hidráulica. A comienzos del siglo XX, el mercado mayoritario formado por grandes contratistas, los cuales movían ingentes cantidades de tierra, estaba atendido por la excavadora de cable que se desplazaba sobre una oruga. Una cuchara con un volumen superior a los 2 m^3 era lo mínimo aceptable por el mercado mayoritario (algunas excavadoras tenían una cuchara de más de 7 m^3). Una tecnología emergente surgió como una pala hidráulica montada en la parte trasera de un tractor con un volumen de cuchara de $0,2 \text{ m}^3$. En su origen, como tecnología disruptiva potencial, fue empleada por pequeños contratistas para sustituir a los obreros en las tareas de abrir drenajes y surcos. Ahora bien, cuando los pioneros mejoraron en la excavadora hidráulica el volumen de la cuchara por encima de los 2 m^3 , atrajo la atención de los grandes contratistas, que empezaron a valorar la precisión de la cuchara, la flexibilidad de los movimientos y la facilidad para la logística (el conjunto de todas ellas constituye la dimensión secundaria). En relación con esta dimensión secundaria, las excavadoras hidráulicas eran claramente superiores a las excavadoras de cable. En Estados Unidos, de las más de treinta empresas establecidas en los años cincuenta, solo cuatro estaban activas a finales de los años setenta, el resto tuvo que abandonar el mercado (Christensen, 1997).

Considerando el concepto original, la tecnología disruptiva potencial posee tres características importantes (Christensen, 1997; Christensen *et al.*, 2015): a) es más simple y más barata que la tecnología principal del mercado mayoritario, su rentabilidad es baja y su futuro es incierto; b) atiende las necesidades de un nicho no servido o de segmentos del lado bajo del mercado, y c) cambia la base de la competencia, desplazando la dimensión tecnológica en las que las empresas compiten. El planteamiento de Christensen también considera que las tecnologías disruptivas describen procesos por los cuales una empresa más pequeña y con menos recursos es capaz de desafiar con éxito a otras empresas más grandes y establecidas en el mercado mayoritario (Christensen *et al.*, 2018). Las empresas establecidas rechazan la tecnología disruptiva potencial porque sus clientes actuales no la valoran y, a su vez, tampoco está claro su potencial de mejora, pudiendo quedar estancada en una isla de aplicación (Adner y Levinthal, 2000).

Así pues, respecto a la dimensión tecnológica primaria, la tecnología dominante continúa manteniendo un desempeño superior a la tecnología disruptiva, aunque a veces sea a costa de provocar una sobreoferta tecnológica que el mercado principal no estará dispuesto a pagar (Christensen, 1997). En esa dimensión, la tecnología disruptiva mejora a un ritmo paralelo y por debajo de la tecnología dominante, por lo que sus trayectorias tecnológicas no se interceptan. Por tanto, no se trata de buscar la intersección de la dimensión primaria de las dos tecnologías, sino de determinar si se va a producir un cambio en las necesidades tecnológicas del mercado, desplazando la dimensión secundaria a la primaria como prioridad tecnológica. Ese cambio se producirá cuando la mejora de la dimensión primaria a lo largo de su trayectoria supere las necesidades del mercado mayoritario. La representación gráfica del modelo de Christensen en un plano cartesiano (figura 27) conlleva: a) dos curvas de desempeño

que no se cortan (para no complicar el gráfico, Christensen las representa mediante sendas rectas con pendientes positivas; en el eje de abscisas, el tiempo y, en el eje de ordenadas, la dimensión tecnológica primaria); y b) una recta que representa la evolución de la demanda tecnológica del mercado. Esta recta de demanda, con menor pendiente que la recta de la tecnología disruptiva, evoluciona siempre por debajo de la recta de la tecnología dominante y por encima de la recta de la tecnología disruptiva potencial. La tecnología disruptiva potencial en su inicio se introduce en un nicho o segmento de mercado, hasta que, a partir del punto de corte (punto 0, momento en que la tecnología disruptiva potencial pasa a ser disruptiva), el desempeño en la dimensión primaria satisface las necesidades tecnológicas del mercado principal.

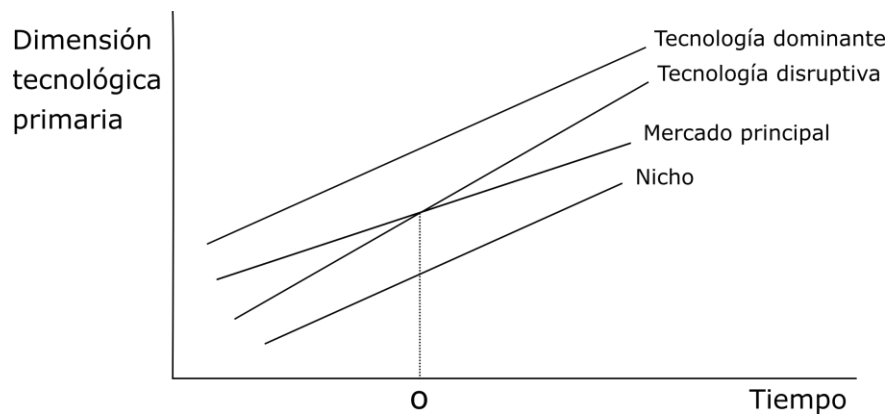


Figura 27: Disrupción tecnológica desde el lado bajo del mercado o de un nicho no atendido (Christensen, 1997)

Por ejemplo, en el mercado de las unidades de disco de los PC, Christensen (1997) observó cómo al comienzo, los clientes valoraban la capacidad, pero, una vez que vieron colmadas sus aspiraciones, trasladaron sus exigencias al tamaño físico, posteriormente a la fiabilidad y, finalmente, al coste. La unidad de disco de 5,25 pulgadas copaba la demanda del mercado de ordenadores personales. La unidad de disco de 3,5 pulgadas, que tenía una menor capacidad pero también un tamaño más reducido, en su origen fue utilizada para satisfacer el nicho de los ordenadores portátiles (en aquella época poco rentable y escasas ventas). Una vez que el disco de 3,5 pulgadas logra satisfacer la capacidad que demandaba el mercado de los ordenares personales, este mercado empezó a valorar su menor tamaño (figura 28). Generalmente, una vez logrado el nivel de desempeño requerido para la capacidad, los clientes emitían señales de no querer pagar un sobreprecio por una mejora incremental de ese criterio de desempeño. Por lo tanto, la sobreoferta de desempeño produce un desplazamiento en la base de la competencia y los criterios de señalización utilizados por los clientes se trasladan hacia cualidades todavía no proporcionadas por la tecnología dominante. En el momento en que existe el desplazamiento, en el caso de las unidades de disco, las empresas incumbentes pueden estar proporcionando unas elevadas prestaciones en capacidad a un precio relativamente bajo. No obstante, el mercado principal valora una tecnología disruptiva con un tamaño reducido que proporcione una capacidad suficiente a un precio adecuado. La tecnología dominante tiene unas prestaciones en capacidad muy altas, pero su mayor tamaño provoca el rechazo del mercado principal.

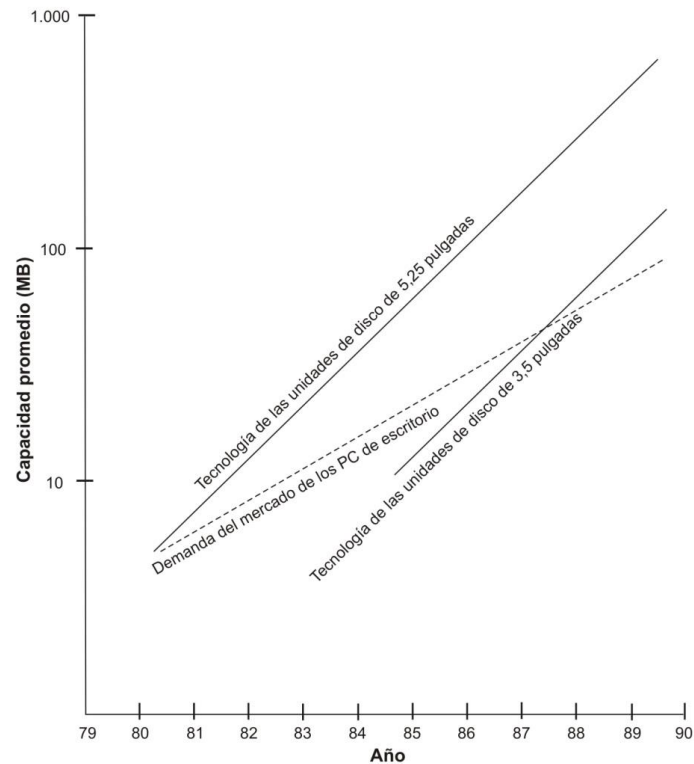


Figura 28: Invasión del mercado de los PC por las unidades de disco de 3,5 pulgadas (Christensen, 1997)

Si comparamos la tecnología disruptiva con la tecnología discontinua, la tecnología discontinua ataca a la tecnología principal mejorando de manera extraordinaria las prestaciones de la dimensión tecnológica primaria. En terminología de Grove (1996), en un factor 10X, es decir, multiplicando por 10 el desempeño de la dimensión primaria de la tecnología principal. Sin embargo, la tecnología disruptiva ataca a la principal con unas prestaciones menores en la dimensión primaria (aunque por encima de las necesidades del mercado principal) pero unas prestaciones muy superiores en la dimensión secundaria. El mercado principal, al tener satisfechas sus necesidades en relación con la dimensión primaria, comienza a valorar las prestaciones de la dimensión secundaria, llegando a convertirse esta en la nueva prioridad tecnológica. Cuando se produce un cambio en las preferencias de la dimensión tecnológica, es cuando surge la disrupción tecnológica.

Hasta ahora, parece claro, como también apunta Danneels (2004), que una tecnología disruptiva es una tecnología que cambia las bases de la competencia cambiando la métrica de rendimiento tecnológico con la que las empresas compiten. Sin embargo, si bien el cambio en la métrica de desempeño es una «condición necesaria» para que aflore la tecnología disruptiva, no es una condición suficiente, pues la tecnología puede cambiar la métrica y no ser disruptiva, como ocurrió cuando IBM introdujo la máquina de escribir eléctrica (a nuestro entender una tecnología discontinua) en el mercado de las máquinas mecánicas (tecnología dominante). En este sentido, Fernández et al (2019) consideran que, para ser disruptiva, se requiere como «condición suficiente» que, antes de atacar el mercado principal, la tecnología disruptiva potencial haya sido comercializada en algún otro mercado.

Esta delimitación de las condiciones necesaria y suficiente, además de permitir definir claramente lo que es una tecnología disruptiva, también permite ampliar el escenario abarcado por el concepto original de Christensen. Según la definición propuesta, la tecnología disruptiva, en su inicio como tecnología disruptiva potencial, puede aportar un valor inferior o superior a la tecnología dominante y, por tanto, tener un precio más bajo o más alto. Es decir, no solo puede comercializarse previamente en nichos o segmentos del lado bajo del mercado principal, sino también en segmentos de su lado alto o en cualquier otro mercado. Asimismo, en el contexto planteado, no existen impedimentos para que la tecnología disruptiva la puedan introducir no solo pequeñas empresas, como considera Christensen, sino todo tipo de empresas, pequeñas o grandes, nacionales o extranjeras, nuevas o establecidas en el mercado principal. De hecho, trabajos alternativos al de Christensen, ampliaron el origen de la tecnología disruptiva a segmentos del lado alto del mercado (elevadas prestaciones y altos precios), considerando igualmente que las grandes empresas establecidas en el mercado principal pueden introducir este tipo de tecnologías (Utterback y Acee, 2005; Govindarajan y Kopalle, 2006; Sood y Tellis, 2011). Asimismo, otros trabajos argumentan que puede ser disruptiva una tecnología inversa originada en las filiales de las multinacionales (Corsi y Di Minin, 2013).

Así pues, la tecnología disruptiva también se puede introducir desde el segmento de mayor valor añadido por una empresa establecida en el mercado principal, a pesar de que en un futuro pueda canibalizar la tecnología principal. En el inicio, Nespresso comercializó el café expreso como un producto gourmet, dirigido al segmento de profesionales entre 35 y 45 años y rentas altas. Finalmente, entró en el mercado principal de los hogares para sustituir al café molido y al café soluble (Brem *et al.*, 2016). Este ejemplo se caracteriza por ser una tecnología más compleja y cara que la tecnología dominante. Por tanto, con un precio más alto.

Cuando el ataque procede del lado alto del mercado, como en este último caso, la tecnología disruptiva, en relación con la tecnología dominante, tiene un precio muy alto que provoca el rechazo del mercado principal (en este caso, el precio actúa como dimensión primaria), y superioridad manifiesta en la dimensión secundaria (que viene representada por el desempeño tecnológico). Con el paso del tiempo, las mejoras (normalmente de proceso) en la tecnología disruptiva potencial logran reducir el coste/precio hasta un nivel aceptable para el mercado principal, momento en que lo invadirá con éxito cambiando la métrica del desempeño. Proporciona una prestación superior a la de la tecnología principal, a un precio aceptable. Si representamos este concepto gráficamente, obtenemos la figura 29. Se observa que la tecnología disruptiva potencial tiene un coste/precio superior a la principal hasta el punto 0. A partir de ese momento, el coste/precio de la tecnología disruptiva empieza a ser aceptable para el mercado principal, aunque se mantiene superior al de la tecnología dominante. No obstante, el mercado comienza a valorar las mejores prestaciones de la tecnología disruptiva, que finalmente desplazará a la principal.

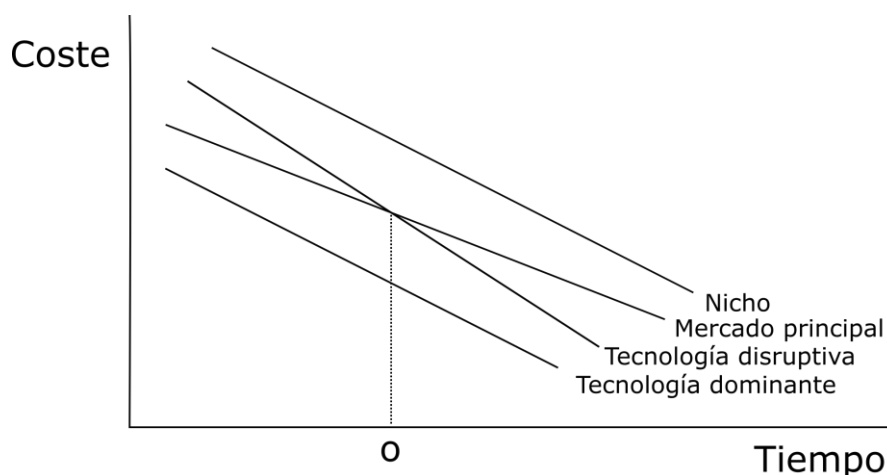


Figura 29. Disrupción tecnológica desde el lado alto del mercado

Existe evidencia empírica que muestra que cuando los pioneros introducen una tecnología disruptiva, normalmente logran provocar la salida de la mayoría de las empresas establecidas en el mercado (Christensen, 1997). Esta evidencia un tanto sorprendente podría ser explicada a partir de las siguientes consideraciones.

Las empresas establecidas pueden elegir entre invertir en la tecnología disruptiva potencial que satisface un nicho de mercado o en la tecnología dominante en el mercado principal. Pues bien, los directivos de estas empresas se comportan de acuerdo a lo que de ellos se espera, es decir, se podría decir que realizan una «gestión de libro», a saber (Christensen, 1997): escuchan cuidadosamente a los clientes, siguen de cerca las acciones de los competidores e invierten en I+D con objeto de mejorar las prestaciones de la tecnología dominante, mejorar la competitividad y lograr mayores beneficios. Así pues, las empresas establecidas atienden las necesidades del mercado principal mejorando las prestaciones de la tecnología dominante, ya que se esperan mayores beneficios y un consistente crecimiento. Por tanto, no suelen ser proclives a invertir en la tecnología disruptiva potencial, al no demandarla actualmente el mercado principal (Christensen, 1997). Estas empresas están, pues, sometidas a la «tiranía del mercado servido»: escucha solo a los clientes del mercado principal con el fin de satisfacer sus necesidades en relación con la tecnología dominante (Christensen, 1997).

En general, la inversión en la tecnología disruptiva potencial la desincentiva la rivalidad competitiva en el mercado principal. Cada empresa establecida, si no quiere verse rezagada y perder cuota de mercado, tiene que continuar mejorando las prestaciones y la eficiencia de la tecnología dominante. Provocando incluso una sobreoferta tecnológica que los clientes no están dispuestos a pagar. Aquellas empresas que decidan concentrar sus inversiones en la tecnología disruptiva potencial podrían ver deteriorada su competitividad si los rivales logran mejorar apreciablemente la tecnología dominante (Langlois y Robertson, 1995).

Por otra parte, las empresas establecidas pueden percibir la tecnología disruptiva potencial como una oportunidad de ganancias. Si tenemos en cuenta las investigaciones de Kahneman y Tversky (1979), la tristeza de perder algo es muy superior a la alegría de ganar ese algo, como media más del doble (Novemsky y Kahneman, 2005). Ello nos lleva a prestar más atención a lo que podemos perder que a lo que podemos ganar (Ariely, 2008). En suma, los individuos valoran más lo que poseen que lo que no

poseen, aunque esto se refiera a la percepción de unos ingresos que se verán reducidos por la pérdida de algunos clientes que se desplazan a los competidores por no haber mejorado la tecnología dominante en relación con la percepción del aumento de los ingresos que provienen de los nuevos clientes de la tecnología disruptiva potencial. De acuerdo con Thaler (1985) podemos considerar que, a pesar de que el dinero es fungible, las empresas perciben de forma independiente la disminución de ingresos de la tecnología dominante y el aumento de ingresos de la tecnología disruptiva potencial. Es esa percepción de pérdida la que concentra a los directivos en la defensa de los clientes del mercado principal (Christensen y Raynor, 2003), por lo que sobreinvierten en la tecnología dominante.

De igual forma, cuando las perspectivas de beneficios que aporta la tecnología disruptiva potencial no son claras o resultan menos atractivas que las de la tecnología dominante, a menudo es difícil justificar la inversión bajo el estricto criterio del rendimiento de la inversión (Christensen, 1997; Utterback, 1994). Existe un «efecto certeza» por el que los rendimientos de la inversión relativamente seguros derivados de las mejoras sostenidas de la tecnología dominante en el mercado principal se valoran más que las inversiones arriesgadas en la tecnología disruptiva potencial con un futuro incierto (Bercovitz, Figuereido y Teece, 1997).

La tecnología dominante proporciona respecto a la disruptiva potencial elevados beneficios y una apreciable tasa de crecimiento de las ventas, al menos en el corto plazo (Christensen, 1997). Por ejemplo, las empresas que fabricaban unidades de disco de 8 pulgadas para el mercado de los miniordenadores requerían márgenes brutos del 40 por ciento. Si se hubieran movido agresivamente mercado abajo, hacia las unidades de disco de 5,25 pulgadas que acababan de introducirse, se habrían encontrado con adversarios que habían diseñado sus estructuras de costes para obtener beneficios con márgenes brutos del 25 por ciento. Del otro lado, moverse mercado arriba, hacia las unidades de disco de 14 pulgadas, les permitiría, a la inversa, trasladar una estructura de costes relativamente baja hacia un mercado que se hallaba acostumbrado a otorgarles a sus proveedores márgenes brutos del 60 por ciento. ¿Cuál de las dos direcciones aparentaba, por lo tanto, ser la adecuada? Una asimetría similar encontraron los fabricantes de unidades de disco de 5,25 pulgadas en 1986, cuando debieron decidir si gastar sus recursos en alcanzar una posición de supremacía en el mercado emergente de las unidades de 3,5 pulgadas de los ordenadores portátiles o desplazarse hacia las empresas que adquirirían miniordenadores y mainframes con unidades de disco de 8 y 14 pulgadas (figura 30). La decisión de orientar recursos hacia el lanzamiento de tecnologías capaces de generar mayores márgenes brutos solía ofrecer, al mismo tiempo, mayores beneficios y ocasionaba menos inconvenientes. La migración de las empresas hacia el rincón superior derecho de la curva de trayectoria era, en ese sentido, muy racional, ya que, de esta forma, las empresas podían seguir creciendo en el mercado y obteniendo elevados beneficios (Christensen, 1997).

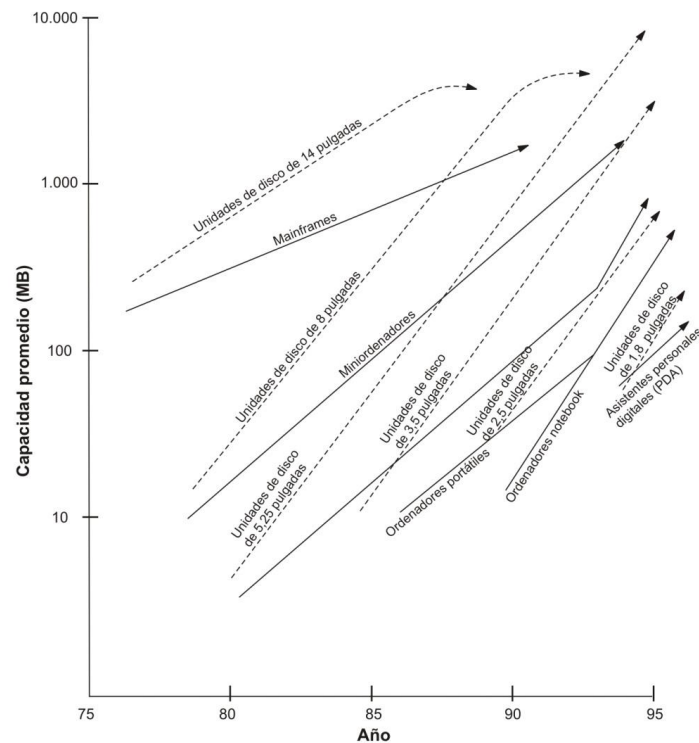


Figura 30 Incursión de las tecnologías en los mercados superiores (Christensen, 1997)

La perspectiva de elevados beneficios repercute favorablemente en la cotización de las acciones que, a su vez, permite a la empresa tener un mejor acceso a las fuentes de financiación y satisfacer las aspiraciones de los accionistas. Además los directivos con opciones de compra sobre acciones aumentarán sus ingresos significativamente. Asimismo, es de suponer que la tecnología dominante sustenta un período de crecimiento empresarial, lo que permitirá a los gerentes de alto desempeño promocionar a niveles superiores y evitar una rotación perniciosa. Cuando una empresa deja de crecer, muchos de los mejores directivos abandonan, ya que tienen pocas expectativas para promocionar (Milgrom y Roberts, 1992). Por otra parte, la alta dirección es consciente de que decantarse por una tecnología disruptiva potencial significaría, en la mayoría de las ocasiones, socavar los fundamentos económicos actuales de la empresa establecida, puesto que se reduciría el valor de sus activos especializados, al provocar que las bases de conocimiento y las capacidades actuales quedasen obsoletas (Leonard-Barton, 1995). En suma, las inversiones realizadas en la tecnología dominante, debido a su carácter coespecializado, son un coste hundido que la empresa tendría que asumir en caso de abandonarla, con el consiguiente deterioro a corto plazo de su valor en bolsa, lo que afectaría negativamente a las opciones sobre acciones de sus directivos (Henderson y Clark, 1990).

La falta de fondos y apoyo político también puede ser un factor clave en la explicación del fracaso de las empresas establecidas a la hora de reaccionar ante una tecnología disruptiva potencial. Ha sido verificado que, si el negocio principal de una empresa comienza a tener problemas y los directivos de mayor antigüedad están buscando formas de reducir costes o activos, la nueva tecnología se convierte en un objetivo fácil (Grove, 1996). El motivo es que los verdaderos réditos de esta nueva inversión no crecen sino hasta después de su jubilación, y ellos persiguen beneficios a corto plazo.

Además, a ningún alto directivo con una trayectoria profesional inmaculada le interesa apoyar una tecnología disruptiva potencial que pueda fracasar en el momento de su jubilación, empañando sus éxitos anteriores. En consecuencia, el apoyo político y los fondos necesarios para la inversión en la tecnología disruptiva potencial se evaporan (Day y Schoemaker, 2000). De igual forma, está claro que apostar por la tecnología disruptiva potencial podría suponer desplazar a la aristocracia de la empresa, siendo comprensible que los directivos sean reacios a perder poder (Leonard-Barton, 1995). Además, en general nadie en la empresa parece estar dispuesto a cuestionar a quienes han llevado la empresa a la hegemonía (Grove, 1996).

En algunas ocasiones, puede ocurrir que las empresas establecidas rechacen la tecnología disruptiva potencial debido a la existencia de una homogeneidad cultural en el mercado principal (Abrahamson, 1991). Así, por ejemplo, durante los años sesenta, los Tres Grandes fabricantes norteamericanos del sector del automóvil creían firmemente que los coches pequeños no eran rentables en su mercado, lo que abrió las puertas a los fabricantes japoneses y europeos (Womack *et al.*, 1990).

Por su parte, el comportamiento de los mandos intermedios también juega un papel importante. En muchas ocasiones, ellos son los que deciden en qué proyectos de investigación invertir (Bower, 1970). Si apoyan proyectos que fracasan, sus carreras profesionales pueden verse muy deterioradas, aunque no todos los fracasos suponen idéntica penalización. Aquellos proyectos que fracasan porque la tecnología no ha podido desarrollarse no suelen percibirse como fracasos, ya que se considera que el esfuerzo realizado permite obtener algunos conocimientos y que, además, el desarrollo tecnológico está sometido a prueba y error (Maidique y Hayes, 1984). Esto es lo que podría ocurrir con los proyectos de investigación relacionados con las tecnologías dominantes, que tendrían que enfrentarse tan solo a un riesgo tecnológico, puesto que el riesgo de mercado, debido a que las mejoras siempre van a satisfacer al mercado principal, sería casi inexistente. Sin embargo, los proyectos que fracasan porque no satisfacen las necesidades del mercado principal tienen consecuencias mucho más graves para los mandos intermedios, ya que la ortodoxia empresarial no acepta que desconozcan las necesidades de sus clientes. Las tecnologías disruptivas potenciales tienen un riesgo tecnológico pequeño, pues en el inicio son fácilmente mejorables, y, sin embargo, un riesgo de mercado muy alto, porque el mercado principal tiende a rechazar sus primeras mejoras, por muy avanzadas que sean técnicamente. En consecuencia, para evitar posibles penalizaciones, los mandos intermedios tienden a respaldar solo aquellos proyectos que satisfacen las necesidades del mercado principal. El «dilema del innovador» consiste en que las empresas que buscaron el crecimiento ingresando en nichos de mercado registraron veinte veces los ingresos de las que buscaron el crecimiento en el mercado principal (Christensen, 1997). Las empresas que cambiaron un «riesgo de mercado» o riesgo de que no se desarrollase finalmente un nuevo mercado para una nueva tecnología por un «riesgo tecnológico» o riesgo de no mejorar la tecnología dominante teniendo que enfrentarse a competidores ya consolidados hicieron un pésimo negocio.

A todo ello hay que añadir que la alta dirección sobreestima muchas veces su nivel de competencia. No le preocupa seguir centrándose en la tecnología dominante porque considera que, si algo nuevo tiene futuro, podrán imitarlo rápidamente (Grove, 1996). Esta seguridad quizá obedezca a que las empresas establecidas son muy hábiles en «acumular creatividad» (Pavitt, 1986) en el mercado principal, por lo que cosechan un

gran éxito durante la evolución de la industria. No obstante, se olvidan de la inercia organizativa.

La «inercia organizativa» (Ghemawat, 1991), es decir, la incapacidad de la empresa para adaptarse a su entorno, también resulta clave a la hora de desincentivar la inversión en una tecnología disruptiva potencial. A las empresas establecidas les resulta difícil romper con los hábitos, rutinas y procedimientos establecidos para explotar la tecnología dominante (Leonard-Barton, 1995), por lo que suelen hacerse complacientes y conservadoras (Grove, 1996). Un importante factor de inercia es la dependencia de la trayectoria (David, 1975), que asume que las decisiones y los enfoques de resolución de problemas utilizados en el pasado tiñen las decisiones que se toman con respecto al futuro. En este sentido, la investigación de Clark (1985) demuestra que las habilidades de diseño y el conocimiento tecnológico que poseen las empresas son el resultado de las tecnologías que abordaron en el pasado y de los conocimientos que las sustentan, y que han ido acumulando a lo largo del tiempo. Esta trayectoria tecnológica refleja la historia de la empresa, limita la gama de alternativas tecnológicas que pueden plantearse de forma realista (Schoemaker, 1992) y provocan un efecto *lock-in* (Arthur, 1989). Por tanto, la experiencia es un mal profesor (Levinthal y March, 1993), ya que restringe la racionalidad de la selección de oportunidades a los recursos que ha ido acumulando la empresa a lo largo del tiempo. Los éxitos pasados refuerzan el modo establecido de solucionar problemas y de tomar decisiones, haciendo que las empresas establecidas busquen en áreas que están estrechamente relacionadas con sus habilidades, formas de pensar, capacidades y tecnología actuales, restringiéndoles operar en nuevos campos de actividad (Amburgey, Kelly y Barnett, 1993). En el fondo, las empresas establecidas, cuando deciden seguir invirtiendo en la tecnología dominante por resultar su acumulación más rápida, menos costosa y más probable, están cayendo en la «trampa de su propio éxito» (o paradoja de Ícaro) (Leonard-Barton, 1995; Levinthal, 1995).

A su vez, las rutinas organizativas, que por un lado ayudan a las empresas establecidas a ser más eficientes en la satisfacción de sus clientes actuales (Tellis y Golder, 1996), por otro provocan que estas sean lentas a la hora de reaccionar ante la irrupción de las tecnologías disruptivas potenciales (Kimberly, 1976), fundamentalmente porque están basadas en conocimientos sustancialmente distintos (Henderson, 1993) y limitan el desarrollo de nuevo conocimiento (Levitt y March, 1988). La tecnología disruptiva potencial se apoya en nuevos conocimientos y capacidades, pudiendo ocurrir con gran probabilidad que las empresas establecidas no tengan las capacidades tecnológicas necesarias para desarrollarlos (Dewar y Dutton, 1986; Henderson, 1993). La adopción de una tecnología disruptiva potencial implicaría que muchas de las rutinas adquiridas hasta el momento quedasen obsoletas, teniendo que desarrollarse, en consecuencia, nuevas rutinas, lo cual es difícil, costoso y arriesgado (Nelson y Winter, 1982). Todo ello favorece el mantenimiento del statu quo (Samuelson y Zeckhause, 1988). Las rutinas impregnan la lógica dominante (Prahalad y Bettis, 1986) que, de manera eficiente, orienta las acciones de los directivos al tiempo que de forma implícita filtra las ideas y comportamientos que no encajan con ella, por lo que evita que la empresa emprenda nuevas iniciativas. La lógica dominante proporciona un marco cognitivo (Stiglitz y Greenwald, 2014) que afecta no solo al alcance de nuestro aprendizaje, sino a lo que aprendemos. Las empresas con un fuerte compromiso con la tecnología dominante se resistirán a la información que parece contradecirla, y que apoya la tecnología disruptiva potencial.

Del mismo modo, a las empresas establecidas también les puede faltar en muchas ocasiones perspectiva para evaluar la tecnología disruptiva potencial de manera completa (Day y Schoemaker, 2000). Ello es debido a que, en una tecnología dominante, la mayor parte de las mejoras gira en torno a los componentes, es decir, los avances tecnológicos pasan por la mejora radical o incremental de los componentes del producto, y no por la alteración de su arquitectura (Henderson y Clak, 1990). Por ello, lo más habitual en estos casos es que los investigadores estén agrupados en equipos especializados en el desarrollo de componentes. Por ejemplo, el departamento responsable de diseñar el mecanismo de dirección de un vehículo estará organizada en equipos que reflejan los componentes de dicho mecanismo: un grupo de columna de dirección, un grupo de piñón y corona, un grupo de varillaje, un grupo de bomba hidráulica, y así sucesivamente (Christensen, 1997). Todo ello provoca que las empresas establecidas centren sus esfuerzos en desarrollar conocimiento sobre los componentes que forman la arquitectura o tecnología dominante. Incluso es habitual que estas empresas externalicen muchos de esos componentes, encargando su diseño y fabricación a empresas independientes. En consecuencia, pocos investigadores y especialistas técnicos son capaces de tener una visión global de la tecnología disruptiva potencial, lo cual resulta imprescindible para poder desarrollarla.

Todos los hechos señalados provocan que las empresas establecidas a menudo subestimen la tecnología disruptiva potencial mientras se está desarrollando. Es decir, hasta que la disruptiva potencial, no irrumpe en el mercado principal como tecnología disruptiva, las empresas establecidas no perciben la amenaza ¿Qué ocurre entonces? Es de suponer que la tecnología disruptiva, cuando entra en el mercado principal, se encuentra en el inicio de la fase específica del modelo de Abernathy y Utterback (1978), fase de la que se derivan importantes economías de escala y un potente efecto experiencia. Es decir, el atacante disruptivo desarrolla en el nicho o segmento la innovación radical en producto y la posterior innovación radical en proceso, por lo que, al atacar el mercado principal, cuenta con un producto y un proceso coespecializados (Teece, 1986), lo que favorece una mayor productividad (Milgrom y Roberts, 1992) y representa una fuerza competitiva devastadora. Durante este tiempo, el pionero también aumenta su tamaño. Este acopio de recursos tecnológicos y económicos favorece la acumulación creativa (Pavitt, 1986) y las «eficiencias de masa» (Dierickx y Cool, 1989), lo cual enfatiza la mayor facilidad para acumular recursos cuanto mayor sea el nivel de partida. En el desarrollo del producto y del proceso, el atacante disruptivo puede que haya necesitado de fabricantes y proveedores de productos y activos complementarios con los que ha podido crear *joint ventures* o firmar acuerdos de cooperación exclusivos (Shapiro y Varian, 1999; Teece, 1986).

La tecnología dominante también se encuentra en la fase específica (al final) con un producto y un proceso igualmente coespecializados, pero se apoya en una trayectoria tecnológica diferente e incompatible con la tecnología disruptiva (Dosi, 1982; Markides y Geroski, 2005). En la coespecialización, cualquier cambio en la arquitectura del producto o en el proceso resulta lento y excesivamente costoso (Abernathy y Utterback, 1978). Además, las empresas establecidas suelen haber generado compromisos con sus proveedores y clientes, así como una alta dependencia de la trayectoria tecnológica desarrollada (Christensen y Rosenbloom, 1995). Todo ello dificulta que las empresas establecidas puedan imitar la tecnología disruptiva en un tiempo aceptable. Por otra parte, aunque la innovación de producto se puede replicar en un tiempo prudencial (Mansfield, 1986), no ocurre lo mismo con la innovación radical en el proceso, que,

debido a su complejidad, a apoyarse en gran medida en conocimiento tácito y a no ser visible en el mercado, resulta muy difícil de imitar (Levin et al., 1987; Utterback, 1994). Además, las empresas establecidas están afectadas por las «deseconomías de compresión del tiempo» (Dierickx y Cool, 1989), que sustentan que la tecnología disruptiva no puede imitarse de forma eficiente en un tiempo breve a pesar de que se intensifiquen las inversiones en su desarrollo.

El ataque disruptivo es muy rápido (Walsh et al., 2002) y la prioridad tecnológica cambia abruptamente. Cuando las empresas establecidas intentan reaccionar al ataque, puede que ni siquiera encuentren aliados o proveedores de los productos y activos complementarios, y desarrollarlos internamente requiere elevadas inversiones y, sobre todo, tiempo del que no disponen y conocimientos especializados que seguramente no poseen. Por tanto, la «interconexión de activos» tecnológicos y sociales, tangibles e intangibles, del atacante disruptivo juega en contra de la empresa establecida (Dierickx y Cool, 1989). También hay que tener en cuenta que los clientes son leales a la primera marca (Lieberman y Montgomery, 1988), y que el pionero ha podido lograr con la tecnología disruptiva potencial una buena reputación que contribuye a apalancar su posición en el mercado mayoritario. Todas estas circunstancias propician que las empresas establecidas tengan una escasa capacidad de reacción ante el ataque de la tecnología disruptiva.

Conviene asimismo destacar que el atacante disruptivo desempeña dos veces el rol de pionero: cuando introduce la tecnología disruptiva potencial en un nicho o segmento y, posteriormente, cuando introduce la tecnología disruptiva en el mercado mayoritario. En este segundo caso, lo hace con la experiencia de haber sido pionero en el nicho. El intervalo de tiempo entre los dos roles es un punto ciego para las empresas establecidas que, al rechazar la tecnología potencialmente disruptiva, propician que el pionero desarrolle en el nicho o segmento el diseño dominante que posteriormente introducirá como tecnología disruptiva en el mercado mayoritario.

3.2.3. Inercia estructural

La inercia surge cuando se deteriora la capacidad de la organización para adaptarse a su entorno. Normalmente la inercia es el resultado de una buena adaptación en el pasado. Puede ocurrir que esta adaptación inicial resulte tan eficiente que la empresa mantenga su ventaja competitiva durante bastante tiempo a pesar de utilizar una tecnología desfasada. Ello se debe a que las empresas que han adoptado una nueva tecnología, mucho más eficiente, están todavía inmersas en el proceso de expandir el mercado y perfeccionar la tecnología (Hanna y Freeman, 1989).

Sin embargo, cuando la inercia de las empresas establecidas retrasa la adopción de la tecnología emergente, serán las nuevas empresas las que la desarrollan, ya que no están atadas a compromisos previos ni tienen que sufrir la tiranía de los mercados servidos. A veces, la inercia se convierte en «inercia activa», cuando las empresas asentadas intensifican los métodos probados a pesar de ser ineficaces en el entorno actual. El aprendizaje es el antídoto contra la inercia porque permite que las organizaciones cambien su conducta y aumenten sus rutinas y capacidades (Langlois y Robertson, 1995).

La figura 31 representa a escala semilogarítmica las curvas de experiencia de dos tecnologías. Inicialmente, existen tres empresas que emplean la tecnología establecida:

un productor experto y con costes relativamente bajos, en A_1 y dos empresas intermedias, en C_1 y D_1 . Supongamos ahora que un nuevo productor B tiene dos opciones: a) entrar en el mercado actual comercializando la tecnología establecida con una experiencia B_1 y b) comercializar una tecnología emergente con una producción acumulada B_2 en una curva de experiencia de más pendiente. La empresa A tiene igualmente la posibilidad de colocarse en el mismo lugar de B en la tecnología emergente (A_2). La empresa B adoptará la tecnología emergente, pero la empresa A deberá hacer frente a la inercia, ya que está produciendo con costes más bajos en el mercado principal. La empresa A no sólo obtendría menores beneficios si adoptase la tecnología emergente, sino que podría ser expulsada del mercado actual si las empresas C y D decidieran competir agresivamente en precio. La empresa B, para sobrevivir con la tecnología emergente, deberá encontrar un nicho de mercado (Langlois y Robertson, 1995).

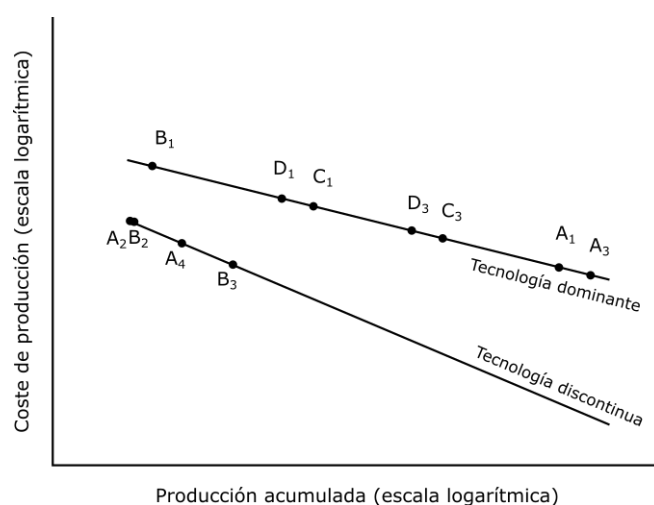


Figura 31: Liquidación y cambio de liderazgo entre empresas (Langlois y Robertson, 1995)

Después de un cierto tiempo, A se habrá situado en A_3 , y B en B_3 , el punto donde sus costes de producción se igualan. A continuación, B tomará la delantera, ya que sus costes disminuirán de manera más rápida. Si ahora la empresa A desea cambiarse a la tecnología emergente, tan sólo podrá aspirar a situarse en el punto A_4 , ya que no podrá acumular el conocimiento tácito y de propiedad de la empresa B. Paralelamente, las empresas C y D se encontrarán con las mismas desventajas que A. Así pues, las empresas que durante el dominio de la tecnología establecida fueron líderes pasarán ahora a ser imitadores, toda vez que la tecnología emergente se ha convertido en la dominante, ya que tendrán menos experiencia acumulada y, por consiguiente, costes relativamente más elevados. Bajo estas circunstancias, y si el período de transición es largo y el diferencial del coste inicial elevado, el curso lógico de las empresas A, C y D será el de liquidar sus inversiones, aunque, al final, esto signifique la transformación de las empresas líderes en empresas imitadoras, o quizá la expulsión de la industria una vez aparecida la tecnología emergente (Langlois y Robertson, 1995).

4. FACTORES ECONÓMICOS QUE INDUCEN LA INNOVACIÓN

En este apartado vamos a comentar una serie de hipótesis relacionadas con la innovación. Así, algunos investigadores consideran que el monopolio promueve las innovaciones. Otros enfatizan el tamaño como impulsores de la actividad innovadora. La integración vertical también se considera impulsora de la innovación. No son éstas todas las hipótesis relacionadas con la innovación, ni mucho menos, pero sí algunas de las más importantes y las que más investigación han generado.

Hipótesis del monopolio. Existen dos fuentes principales de interacción entre innovación y poder de monopolio. La primera es entre la innovación y la previsión de un poder de monopolio. El motivo para producir innovaciones es la obtención de un beneficio extraordinario. La segunda fuente de interacción se produce entre la innovación y la posesión de un poder de monopolio (Kamien y Schwartz, 1982). En este sentido, Demsetz (1969) sostiene que el monopolio no provoca desincentivos a la invención y que, en realidad, puede intensificarla. Hay varias razones que explican por qué el monopolio favorece la innovación mejor que la competencia.

Una empresa que tiene poder de monopolio respecto a los productos que comercializa puede extender ese poder a sus nuevos productos; por ejemplo, mediante el dominio de los canales de distribución o una fuerte imagen corporativa (Kamien y Schwartz, 1982).

Una empresa con beneficios extraordinarios se encuentra probablemente en mejores condiciones para asumir la financiación interna que la que sólo ha obtenido beneficios ordinarios. Los gastos en investigación y desarrollo suelen ser muy arriesgados y los bancos generalmente no están dispuestos a conceder préstamos para financiarlos. Si el proyecto de investigación fracasa o una empresa rival llega antes a la oficina de patentes, el prestatario puede no ser capaz de devolver el préstamo y el banco no podrá recuperar los fondos. Los inversores también suelen tener dificultades para valorar el potencial de un proyecto de investigación, pues los inventores siempre son optimistas respecto a sus ideas y se muestran reacios a revelar información relevante, por temor a que pueda llegar a los competidores. Los beneficios extraordinarios que se obtienen del monopolio permiten autofinanciar la innovación. En este sentido, Mowery (1986) observó que altos niveles de beneficios y *cash flow* están asociados con la intensidad de la I+D empresarial.

Proteger las innovaciones a través de los derechos de propiedad industrial es una actividad muy cara. Las empresas que obtienen beneficios de monopolio disponen de más fondos para proteger sus patentes que las empresas que están en industrias competitivas, por lo que pueden estar más dispuestas a realizar inversiones.

La empresa con beneficios de monopolio puede contratar a las personas más creativas, ya que les puede pagar sueldos más elevados. Ahora bien, no es fácil identificar la creatividad. Ésta generalmente se identifica más a posteriori que a priori y el conocimiento sobre quién la tiene no está tan extendido como para asegurar el perfecto funcionamiento de este mercado (Kamien y Schwartz, 1982).

Cuando existen muchas empresas en un mercado, el mercado se fragmenta y cada empresa, por tanto, produce menos. Esto significa que el beneficio que obtiene de reducir costes disminuye y, así, sus incentivos para innovar se ven atenuados (Stiglitz y Greenwald, 2014).

Por el contrario, Arrow (1962a) sostiene que los incentivos para innovar son mayores en las industrias competitivas que en las monopolistas, ya que es probable que el monopolio retrase el progreso técnico, además de implicar una pérdida estática del bienestar. La competencia imperfecta genera menos presión en las empresas para que introduzcan nuevas técnicas y productos, ya que en estos mercados hay menos competidores. Además, existen ventajas en tener un gran número de unidades tomadoras de decisiones independientes (competencia perfecta), pues es menos probable que un avance tecnológico importante sea trabado por la opinión interesada de unos pocos agentes.

La empresa que obtiene beneficios de monopolio puede estar menos motivada a buscar beneficios adicionales que la que sólo obtiene beneficios ordinarios. Existen varias razones que explican este hecho. La empresa con poder de monopolio puede comenzar a considerar que tener más tiempo para el ocio es más importante que obtener beneficios adicionales. Asimismo, puede llegar a preocuparse más de proteger su posición de monopolio que de desarrollar una nueva tecnología.

Una empresa con poder de monopolio, como señalan Baldwin y Childs (1969), debido a sus recursos, a su situación establecida y a sus canales de distribución, se encuentra en una situación ventajosa para imitar con rapidez cualquier innovación, por lo que puede permitirse esperar hasta que alguien desarrolle una innovación, e imitarla inmediatamente si considera que tiene potencial económico.

Una empresa que obtiene beneficios de monopolio de su producto o proceso actuales puede tardar más en sustituirlos por un producto o proceso superior que una empresa recién llegada. Es decir, la empresa que obtiene beneficios de monopolio calcula la ganancia que obtiene de la innovación como la diferencia entre sus beneficios corrientes y los que podría lograr con el nuevo producto, mientras que la empresa recién llegada considera los beneficios del nuevo producto como su única ganancia. Por ello, ésta última siempre tiene un incentivo mayor para innovar que una empresa similar que ya obtiene un beneficio de monopolio del producto existente. Este hecho fue puesto de relieve por Arrow (1962a) para las innovaciones de procesos y por Usher (1964) para las innovaciones de productos.

Las empresas en posiciones monopolistas consolidadas pueden afrontar elevados costes si la introducción de una innovación les obliga a reestructurar su industria. En este caso, se resistirán a introducir tales innovaciones. De acuerdo con este argumento, algunos recursos en las industrias monopolistas pueden desviarse hacia innovaciones secundarias y mejoras en el diseño (como, por ejemplo, la industria automovilística), en vez de concentrarse en investigaciones más radicales. Por razones similares, estos monopolistas consolidados podrían comprar y retener nuevas patentes, que favorecerían cambios rupturistas. Sin embargo, es probable que estos efectos sean importantes únicamente cuando hay barreras de entrada que protegen a los productores monopolísticos establecidos. En resumen, la información disponible no parece apoyar la hipótesis del monopolio. Sin embargo, esto no implica que la alternativa competitiva esté probada (Williamson, 1975).

Hipótesis del tamaño. Mucho se ha discutido acerca de si las empresas más innovadoras son las de mayor o las de menor tamaño. De acuerdo con Nutter (1956), al igual que la expectativa de lograr una posición de monopolio aumenta las posibilidades favorables a las innovaciones más arriesgadas, el gran tamaño posibilita las más costosas. La dimensión empresarial *per se* incentiva la investigación tecnológica. Por ejemplo, en los Estados

Unidos las empresas con más de 5.000 trabajadores agrupaban en 1972 el 53 por 100 de todo el empleo industrial y el 87 por 100 por ciento cuando se incluían las cifras de los proyectos de I+D de financiación pública (Scherer, 1980). Las multinacionales se cuentan entre los principales protagonistas de la carrera tecnológica: aproximadamente el 50 por 100 de las patentes registradas en Estados Unidos procede de menos de 700 grandes empresas y representa un porcentaje creciente en la I+D empresarial (Patel y Pavitt, 1991). Una serie de razones apoyan la relación entre tamaño e innovación.

Hay proyectos de investigación cuya magnitud y complejidad requieren recursos económicos extraordinarios. Por ejemplo, el análisis de la I+D espacial sugiere que puede ser necesario un gran tamaño para respaldar gran parte de esta investigación. Asimismo, en ocasiones la I+D implica tecnologías complejas de componentes múltiples o complicados procesos químicos para los cuales también es importante el gran tamaño (Nelson *et al.*, 1967). Así pues, en casos como éstos, la innovación puede ser factible únicamente para aquellas empresas que tengan una determinada masa crítica de recursos financieros, tecnológicos y/o humanos.

Dado que los proyectos de I+D comportan cierto riesgo (por cada proyecto con éxito muchos fracasan), las grandes empresas tendrán una ventaja sobre las pequeñas, al poder financiar varios proyectos de investigación paralelos, mientras que la pequeña empresa sólo puede financiar uno (Galbraith, 1952). Por lo tanto, una gran empresa puede disminuir el riesgo de sus inversiones en investigación al diversificar su cartera de proyectos, lo que no está al alcance de las empresas de menor tamaño (Braunstein *et al.*, 1980).

La gran empresa favorece las economías de escala en la investigación y desarrollo, por lo que cuenta con varias ventajas: los investigadores pueden intercambiar conocimiento con sus numerosos compañeros del departamento, facilitando el aprendizaje y los descubrimientos. Se puede agrupar a los investigadores por especialidades, lo que permite profundizar en cada una de las tecnologías que componen el producto. Al aumentar el número de investigadores, se favorece el fenómeno de *serendipity* (descubrimientos fortuitos). La productividad de la I+D, a su vez, puede incrementarse mediante la interacción de la I+D con otras funciones empresariales, como producción y marketing, que sólo están desarrolladas en las grandes empresas.

La gran empresa tiene ventajas en la explotación comercial de los inventos, al contar, generalmente, con una imagen de marca, una reputación, una capacidad de fabricación y unos canales de distribución que le permiten abastecer un gran mercado en mejores condiciones que una pequeña empresa, incrementando, de esta forma, su potencial para apropiarse de los beneficios de la innovación (Nelson, 1959).

A diferencia de las empresas de menor tamaño que se especializan en una línea de productos, la gran empresa suele estar diversificada en diferentes negocios. Una base tecnológica amplia asegura que, cualquiera que sea la dirección que pueda tomar la investigación, los resultados resultan valiosos para la empresa. Por esta razón, las empresas que apoyan la investigación aplicada son empresas con «huevos colocados en varias cestas». No es solo el tamaño de las empresas lo que favorece la investigación. Es más bien la amplia base tecnológica, el amplio conjunto de productos que fabrican o están dispuestas a fabricar si sus esfuerzos investigadores abren posibilidades, aspecto éste vinculado al tamaño (Nelson, 1959). Por otra parte, al poder utilizarse una misma tecnología en la fabricación de diferentes productos, la gran empresa diversificada podrá

obtener una mayor rentabilidad de sus tecnologías, al emplearlas en diversos negocios (Kamien y Schwartz, 1982).

La gran empresa puede tener ventajas en otros aspectos, tales como la destreza empresarial y los servicios legales. De esta forma, puede resolver con mayor eficacia los casos en los que existen problemas de apropiación de las rentas de monopolio originadas por la innovación y/o los costes de transacción relativamente elevados derivados de la cesión de derechos sobre su utilización.

Al ser la I+D un coste fijo, las grandes empresas lo pueden repartir entre un mayor número de productos. Cuando los costes fijos son elevados, las grandes empresas tienen unos costes medios más bajos que las pequeñas y disfrutan de una ventaja competitiva. No es sorprendente, pues, que la industria química, en la que la inversión en I+D es extraordinariamente importante, las empresas sean grandes.

Las empresas grandes tienen mayores incentivos para dedicarse a la investigación. Altas cuotas de mercado permiten una mayor apropiación de los beneficios (Mowery, 1986). Supongamos que una pequeña empresa produce cien mil bolígrafos al año. Si descubre una tecnología mejor, que reduce sus costes en 1 euro por bolígrafo, ahorra cien mil euros al año. Una gran empresa que realice ese mismo descubrimiento y produzca diez millones de bolígrafos al año ahorrará 10 millones. Por lo tanto, las grandes empresas tienen más incentivos para dedicarse a la investigación y el desarrollo, lo que les permite crecer más que las rivales más pequeñas.

En el lado contrario, Hamberg (1963) considera que las grandes empresas no han sido fuentes importantes de innovaciones radicales, por las siguientes razones: 1) prefieren proyectos de I+D que proporcionen rentabilidad a corto plazo, mientras que los proyectos de innovaciones radicales son siempre a largo plazo; 2) tienen intereses creados en la tecnología actual y pueden dudar en lo que respecta a perseguir innovaciones radicales que desplazarían los beneficios de la tecnología en uso; 3) científicos e inventores muy creativos consideran más atractivas las pequeñas empresas; 4) son aversas al riesgo; 5) las funciones de producción, ventas o finanzas asumen mayor peso que la actividad de investigación, por lo que acaparan más poder y presupuesto; 6) aparece una resistencia organizativa a los cambios internos; 7) prima la armonía sobre la creatividad y 8) se produce un fuerte control de la investigación por la alta dirección. También es posible que los laboratorios corporativos de una gran empresa se burocraticen, disminuyendo su creatividad. Por otra parte, los investigadores pueden estar menos motivados en una gran empresa que en una pequeña, ya que en esta última su remuneración puede estar relacionada de forma más directa con su rendimiento.

	<i>Pequeñas empresas</i>	<i>Grandes empresas</i>
<i>Marketing</i>	Capacidad para reaccionar con prontitud y mantenerse al día ante los cambios en el mercado. (Los costes de puesta en marcha en un mercado extranjero pueden ser prohibitivos)	Distribución extensa y facilidades de servicio. Alto grado de poder de mercado con los productos actuales.
<i>Dirección</i>	Ausencia de burocracia. Dinámicas; los empresarios reaccionan rápidamente para obtener ventajas de las nuevas oportunidades y están dispuestos a asumir riesgos.	Directivos capaces de controlar complejas organizaciones y de establecer estrategias corporativas. (Pueden sufrir un exceso de burocracia. A menudo son controlados por contables que pueden ser adversos al riesgo. Los directivos pueden llegar a ser meros gestores que carecen de dinamismo en relación a las nuevas oportunidades a largo plazo.)
Comunicación interna	Redes de comunicación interna eficientes e informales. Proporciona una rápida respuesta a la solución de los problemas internos; ofrece la capacidad de reorganizarse rápidamente para adaptarse a los cambios en el entorno externo.	(Las comunicaciones internas son, a menudo, muy voluminosas; esto puede llevar a una lenta reacción ante las amenazas y oportunidades externas)
Mano de obra técnica cualificada	(Con frecuencia carecen de suficientes especialistas técnicos y cualificados. A menudo son incapaces de sostener un esfuerzo formal en I+D a una escala apreciable.)	Capacidad de atraer a especialistas altamente cualificados. Puede sostener la financiación de un gran laboratorio I+D
Comunicación externa	(A menudo carecen del tiempo o de los recursos para identificar y utilizar importantes fuentes de conocimientos científicos y tecnológicos).	Son capaces de conectar con las fuentes externas de conocimientos científicos y tecnológicos. Pueden permitirse servicios de información y documentación. Pueden subcontratar I+D a los centros especializados. Son capaces de comprar crucial información científica y tecnológica.
<i>Financiación</i>	(Pueden experimentar grandes dificultades para atraer capital, especialmente, capital-riesgo. La innovación puede representar un riesgo financiero desproporcionadamente grande. Incapacidad para diversificar el riesgo mediante una cartera de proyectos.)	Capacidad para pedir prestado en el mercado de capitales. Capacidad para diversificar el riesgo mediante una cartera de proyectos. Mayor capacidad para diversificar riesgos en las nuevas tecnologías y en los nuevos mercados.
<i>Economías de escala y enfoque de sistemas</i>	(En algunas áreas las economías de escala crean sustanciales barreras de entrada a las pequeñas empresas. Incapacidad para ofrecer líneas de productos o sistemas integrados.)	Capacidad para conseguir economías de escala en I+D, producción y marketing. Capacidad para ofrecer una variedad de productos complementarios. Capacidad para ofrecer importantes proyectos llave en mano.
<i>Crecimiento</i>	(Pueden experimentar dificultades para adquirir el capital externo necesario para un rápido crecimiento. Algunas veces los empresarios son incapaces de enfrentarse con organizaciones cada vez más complejas).	Capacidad para financiar la expansión de los productos base. Capacidad para crecer vía diversificación y adquisición.
<i>Patentes</i>	(Pueden tener problemas al enfrentarse con el sistema de patentes. No pueden permitirse el tiempo ni los costes involucrados en un litigio de patentes.)	Capacidad para contratar especialistas en patentes. Pueden permitirse litigar para defender las patentes en caso de infracción.
<i>Regulaciones gubernamentales</i>	(No pueden hacer frente a las complejas regulaciones. Los costes unitarios de conformidad para las pequeñas empresas son, a menudo, altos.)	Pueden proveerse de servicios legales para enfrentarse con requerimientos regulatorios complejos. Pueden extender los costes de regulación. Pueden obtener la I+D necesaria por conformidad.

Nota: Las afirmaciones entre paréntesis representan áreas de potencial desventaja

Figura 32: Ventajas y desventajas de las pequeñas y grandes empresas en relación con la innovación (Rothwell y Zegveld, 1985)

Las pequeñas empresas parecen tener ciertas ventajas respecto a las grandes en el campo de la investigación y desarrollo (figura 32): a) tienen estructuras organizativas más planas, con líneas de comunicación cortas y directas entre los distintos niveles, por lo que es más improbable que los resultados inesperados de una investigación se pierdan, al contrario de lo que ocurre en las grandes empresas; b) las organizaciones son más flexibles, de manera

que desarrollan una mayor capacidad de asimilación y respuesta al cambio; c) la empresa más grande puede quedarse atrapada en la burocracia y en el papeleo, lo que hace que el personal investigador se desenvuelva en una atmósfera menos agradable para realizar aportaciones creativas, razón por la cual tiende a verse atraído por empresas más pequeñas en donde se tiene mayor libertad, más motivación y menos alienación (existe evidencia, referida a casos de investigadores de una gran empresa que la abandonan para explotar una innovación que aquélla no está dispuesta a patrocinar y que pone de manifiesto los efectos negativos que puede tener el tamaño de la empresa en la actividad innovadora) y d) cuanto más grande es la empresa, más dificultad entraña entender los problemas que necesitan solución (Kamien y Schwartz, 1982).

A las anteriores ventajas, Shimshoni (1970) añade las siguientes: los bajos costes y ahorros de tiempo en la labor de desarrollo (tienen una mayor conciencia respecto a los costes y toman decisiones con rapidez) y un mayor aprovechamiento de las economías externas bajo la forma de conocimientos tecnológicos traídos de otros lugares ajenos al sistema interno de I+D. Las pequeñas empresas, al contar con líneas de comunicación más cortas, facilitan que sus investigadores e ingenieros puedan conocer los problemas de producción y las demandas de productos de sus clientes con más precisión que los que trabajan en una gran empresa. Quizás sea ésta una de las razones por las que las pequeñas empresas emplean alrededor de las tres cuartas partes de las invenciones que patentan, mientras que las grandes empresas emplean sólo la mitad. De igual forma, el mayor contacto de los científicos e ingenieros de las pequeñas empresas con los problemas a los que se enfrentan puede llevar a que las empresas pequeñas y medianas gasten, como mucho, la mitad que las empresas grandes en investigación y desarrollo por invento patentado (Schmookler, 1968).

Con base en datos de la National Science Foundation, que considera tres clases de empresas según su tamaño (menos de 1.000 empleados, entre 1.000 y 5.000 empleados y más de 5.000 empleados), existe una clara tendencia a que las empresas más grandes gasten en I+D un mayor porcentaje sobre ventas que las pequeñas (Nelson *et al.*, 1967). Sin embargo, algunas investigaciones, que han sido llevadas a cabo principalmente en los Estados Unidos, sugieren que no existe una evidencia clara de que el esfuerzo o los resultados innovadores aumenten con el tamaño de la empresa. Una serie de autores han detectado la existencia de un «efecto umbral», con una intensidad investigadora creciente hasta cierto tamaño intermedio de las empresas, pero decreciente o constante por encima del mismo. Por ejemplo, Mansfield (1968b), al examinar los desembolsos en I+D (expresados como un porcentaje de las ventas) de las empresas muy grandes en relación con sus rivales grandes, en las industrias química, petrolera, farmacéutica, del acero y del vidrio, observó que, excepto para la industria de productos químicos, los resultados no presentan ninguna evidencia de que las empresas más grandes de estas industrias gasten más en I+D de lo que gastan las empresas algo más pequeñas. Scherer (1965) también observó que, entre una variedad más amplia de industrias, la intensidad en I+D de las empresas más grandes no excedía, por lo general, la de sus rivales de tamaño mediano.

Los resultados de la investigación de Mansfield (1963) muestran que las grandes empresas tienden a acaparar las innovaciones cuando su aplicación resulta costosa y requiere una comercialización a gran escala para resultar económica. En los mercados competitivos, cuando existe una escala eficiente mínima, las pequeñas y medianas empresas independientes pueden no ser capaces de emprender proyectos de I+D muy costosos. En algunas circunstancias, este problema se supera uniendo esfuerzos investigadores, pero

esta solución no tiene por qué ser siempre posible, dados los problemas financieros de las pequeñas empresas y los elevados costes de transacción que implica la organización de un programa de investigación conjunto. En esta corriente de opinión también participa Freeman (1974), al afirmar que existen diferencias significativas entre industrias en cuanto a las realizaciones relativas de empresas pequeñas y grandes. En la industria química, donde tanto la investigación como el desarrollo son casi siempre muy costosos, las grandes empresas predominan en la invención y en la innovación. En la industria mecánica, el ingenio barato puede jugar un papel importante, por lo que las empresas pequeñas o los inventores privados son más capaces de introducir innovaciones. En general, en las industrias intensivas en capital las innovaciones, tanto de procesos como de productos, han sido casi monopolizadas por grandes empresas. Las empresas pequeñas son más capaces de hacer contribuciones significativas en productos que les permitan dominar huecos o segmentos específicos, ignorados por los líderes del mercado, por ser demasiado pequeños o poco prometedores para merecer más que una ligera atención.

Pavitt y Wald (1971) encontraron que las oportunidades para las pequeñas empresas tienden a ser mayores en las primeras etapas del ciclo de vida del producto, cuando las economías de escala son relativamente poco importantes, las cuotas de mercado inestables y las tasas de entrada y de fracaso altas. Una entrada con éxito depende principalmente de la capacidad tecnológica en esta etapa. Así pues, los recursos modestos, con frecuencia, son suficientes para apoyar la invención y el trabajo de desarrollo de las primeras etapas (Jewkes *et al.*, 1969). El desarrollo de las etapas posteriores suele exigir mayores gastos (Scherer, 1980). A medida que maduran las tecnologías, la escala y la eficiencia en la producción se hacen más importantes y las oportunidades para las pequeñas empresas se reducen. Por lo tanto, los individuos y las empresas pequeñas que se dedican al trabajo de invención y de desarrollo de las primeras etapas pueden tener que enfrentarse a la posibilidad eventual de crecer y ser una gran empresa para poder comercializar la innovación con éxito. Probablemente las empresas que ya son grandes y poseen laboratorios de investigación e instalaciones de apoyo disfrutan de una ventaja en las etapas posteriores. De este modo, a menos que las empresas más pequeñas se aseguren de forma eficiente –por sí mismas o a través de los procesos del mercado– el acceso a una capacidad equivalente, pueden verse obligadas a salir del mercado (Williamson, 1975).

Varios estudios sobre el número de patentes, introducción de nuevos medicamentos e innovaciones tecnológicas han indicado que las pequeñas empresas obtienen mejores resultados que las grandes empresas. Por ejemplo, unos pocos estudios sobre patentes han concluido que las pequeñas empresas parecen gastar en I+D de manera más cuidadosa y son más eficientes, obteniendo un mayor número de patentes por dólar en I+D (Griliches, 1990). Un estudio de 116 empresas que desarrollaban productos business-to-business también encuentra que las empresas pequeñas (aquellas con ventas anuales menores a 100 millones de dólares) tienen ciclos de desarrollo significativamente más cortos que las grandes empresas (aquellas con más de 100 millones de dólares en ventas), incluso cuando se considera la magnitud relativa de la innovación (Griffin, 2002). Sin embargo, unos pocos estudios han indicado que las grandes empresas pueden obtener mejores resultados en innovaciones que las empresas pequeñas en algunas industrias (Cohen y Kleppler, 1996).

La gran empresa puede estar estructural y/o funcionalmente mal equipada para tolerar las invenciones radicales, ya que comportan un elevado riesgo. Además, una empresa no necesita ser el origen de invenciones para poder participar en el desarrollo comercial de las

mismas: puede adquirirlas a otras empresas o instituciones o incluso invertir en capital riesgo (Williamson, 1975). Los estudios y la evidencia sugieren que, mientras que las grandes empresas suelen ser mejores en avances incrementales y posiblemente más rápidas a la hora de adoptar la tecnología desarrollada fuera de la empresa, es más probable que la fuente de los mayores avances sean las empresas de pequeño y mediano tamaño. En su investigación, Grosvenor (1929) halló que sólo 12 de los 72 mayores inventos realizados desde 1889 se habían originado en laboratorios industriales. Jewkes *et al.* (1969) rastrearon los orígenes de 61 inventos mayores del siglo XX, la mayor parte de los cuales se efectuó entre 1930 y 1950. De éstos, 12 podían atribuirse a los laboratorios de las grandes empresas, 33 eran obra de inventores independientes, 5 provenían de los laboratorios de empresas pequeñas y el resto no podía clasificarse. Hamberg (1963) examinó 27 inventos realizados durante el período de 1946 a 1955: 7 fueron producto de grandes laboratorios industriales, 12 de inventores independientes y el resto provenía de empresas pequeñas, universidades y una estación de agricultura experimental. El estudio de Schmookler (1957) acerca de las estadísticas de patentes en el período 1950–57 reveló que entre el 50 y el 60 por 100 de los inventos recientes eran obra de inventores independientes, que realizaban la investigación por su cuenta, al margen de los laboratorios industriales. Incluso en una corporación tan grande como DuPont, de los inventos que puso en práctica desde 1920 a 1949, se originaron en los laboratorios de la empresa, en proporción, muchos más inventos para mejoras que inventos de nuevos productos (Mueller, 1962).

Hipótesis de la integración. Una empresa se halla integrada verticalmente en la medida en la que protagoniza las sucesivas etapas u operaciones productivas, técnicamente separables, que se necesitan para dar existencia a un producto y ponerlo en manos del usuario. A medida que crece el número de etapas asumidas por una sola empresa, lo hace el grado de integración vertical (Silver, 1984). En pocas palabras, la integración vertical es una cuestión de grado, no de sí o no.

Diferentes investigadores sustentan que la integración vertical favorece la innovación. Así, Wilson (1957) considera que integrar en una misma unidad de decisión los papeles del comerciante y del fabricante reduce “la fricción existente entre el primero, que es el único en saber qué está dispuesto a comprar el cliente, y el segundo, que a menudo está más interesado en convencer al comerciante de que acepte aquello que él ha hecho”. De igual forma, Florida y Kenney (1990) también consideran deseable un elevado nivel de integración vertical. Estos autores destacan la necesidad de coordinar la investigación aplicada con el desarrollo del producto y la fabricación a fin de obtener el máximo rendimiento de los avances científicos y de la ingeniería. Por otra parte, Lall (1978) sugiere que la integración vertical está relacionada con la aparición de productos muy específicos, que incorporan información nueva (por ejemplo, productos obtenidos con tecnología nueva).

En este sentido, Silver (1984) articula el esquema de una teoría de la integración vertical basada en los costes de información (coordinación). Para ello, considera que un empresario que intenta introducir algo novedoso desde el punto de vista cualitativo se encuentra, a menudo, con una fuerte resistencia. Dicha resistencia puede ser cultural y psicológica, como destacó Schumpeter (1934), aunque también informativa. El éxito de una innovación requiere frecuentemente la adaptación de las actividades complementarias: si la innovación es muy novedosa, muchas de las actividades necesarias para su puesta en práctica también lo serán. El problema para el empresario está en disponer de estas actividades especializadas. Hacerlo a través del mercado le

supondría informar y persuadir a los poseedores de las capacidades necesarias. Esto no será nada fácil, ya que la visión del empresario es nueva y virtualmente idiosincrásica por definición. Las potenciales partes contratantes tendrán que invertir en activos especializados y soportar el riesgo de una inversión irreversible que, bajo tales circunstancias, puede suponer un precio elevado. Por otro lado, al empresario esto puede suponerle un menor coste, al integrar las actividades especializadas y emplear a aquellas partes con capacidades relevantes, en vez de establecer contratos con ellas (Langlois y Robertson, 1995).

En la figura 33, las operaciones o etapas anteriores y posteriores que se necesitan para producir el bien X y ponerlo en manos del usuario se han dispuesto en orden creciente de desemejanza con respecto a la producción de X como tal³. Esto es, a medida que nos movemos hacia la derecha, las operaciones a las que se enfrenta el empresario tienen cada vez menos en común con la producción de X, en términos de conocimientos y experiencias de carácter tecnológico, de producción y de comercialización. Esta ordenación es la responsable de la pendiente positiva de la «curva M», que representa el incremento marginal en el coste de producción derivado de integrar una operación por encima del coste de producción de un productor independiente debidamente capacitado. El movimiento hacia la derecha a lo largo de la curva aumenta las deseconomías de alcance (Silver, 1984).

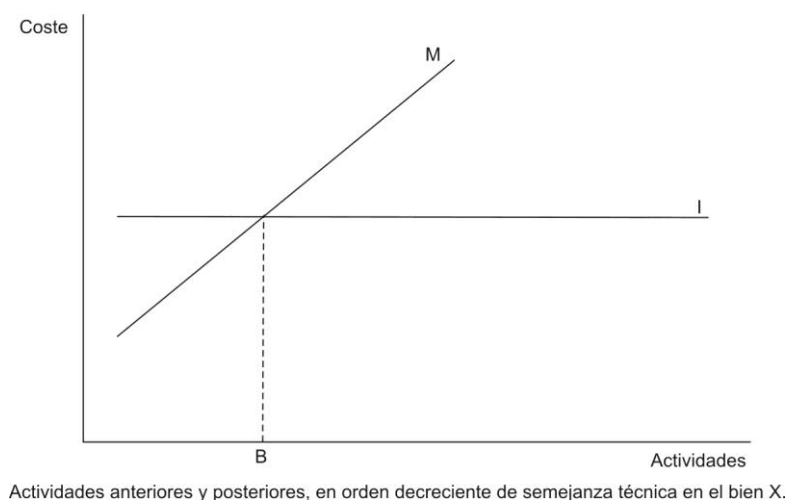


Figura 33: Nivel de integración de actividades que minimiza costes (Silver, 1984).

La pendiente plana de la «curva I», por su parte, refleja el supuesto simplificador de que la reducción marginal en los costes de transmisión de la información que resulta de integrar una operación, en lugar de comprársela a una empresa independiente, no cambia dependiendo de cual sea la operación en concreto integrada⁴. Hay un ahorro en

³ Suponemos que la diferencia de coste entre integrar una operación y adquirirla de una empresa independiente debidamente capacitada depende sólo de su grado de desemejanza con la producción de X. Se podrán quizá imaginar modelos en los cuales otras variables (por ejemplo, economías de escala) pudieran también desempeñar un papel a la hora de dar forma a este diferencial.

⁴ Tal vez la curva I podría también presentar una pendiente positiva o, de modo alternativo, quizá existan economías de escala de transmisión de información.

los costes de transmisión de información generado por la integración de las actividades correspondientes con respecto a la opción de encargar su realización a productores independientes.

Para simplificar la exposición, supondremos que una cantidad fija del esfuerzo potencial total de dirección del empresario se aplica a la producción de X como tal (y quizá también de otros productos). El esfuerzo de dirección residual queda disponible para integrar operaciones y para comunicarse con productores independientes. El punto en el cual la integración vertical minimiza los costes de la empresa innovadora es el punto B, en el que «M» es igual a «I».

Silver identifica los beneficios de este proceso, sobre todo en términos de costes de información: los innovadores pueden comunicar con mayor facilidad los procedimientos y rutinas a sus empleados, en vez de estipular contractualmente a productores independientes todas y cada una de las especificaciones del producto final. Existe también un coste en este proceso de organización interna: el innovador probablemente integrará aquellas actividades en las que sus propias capacidades son más aptas (Langlois y Robertson, 1995).

La integración vertical también puede corregir comportamientos oportunistas. Como señalamos, la innovación puede llevar a la sustitución de activos en más de una fase de la cadena de producción. Si la propiedad está descentralizada, puede que una innovación nunca llegue a producirse si alguno de los poseedores de activos o de los proveedores de bienes complementarios tiene el poder de bloquearla para proteger sus beneficios. Cuando una organización es al mismo tiempo la que desarrolla y la que utiliza una innovación, los beneficios se internalizan y el problema de la apropiabilidad desaparece (Langlois y Robertson, 1995).

Williamson (1975) sugiere que una empresa organiza sus actividades en función de sus costes de transacción. La transacción hace referencia a la transferencia de bienes y servicios entre unidades tecnológicamente separadas. Los costes de transacción son los derivados de identificar, explicar y mitigar todos los riesgos inherentes a una transacción. La teoría de los costes de transacción se apoya en el supuesto de que los humanos tienen una «racionalidad limitada» y, a veces, manifiestan un «comportamiento oportunista».

La racionalidad limitada considera que los agentes económicos están sujetos a restricciones derivadas de su capacidad limitada para formular y resolver problemas y para recibir, acumular, procesar y transmitir información (Simon, 1964b). De acuerdo con este supuesto a los agentes económicos les resulta difícil prever todas las contingencias de una transacción, por lo que redactar, supervisar y hacer cumplir los contratos suele ser muy costoso. La racionalidad limitada planteará un problema en aquellos entornos complejos caracterizados por una elevada incertidumbre y complejidad. Es demasiado costoso o imposible identificar todas las posibles situaciones que se puedan dar en un entorno incierto e incluirlas en el contrato. También es posible que una parte que participa en una transacción no pueda expresar inteligiblemente la información que posee, especialmente si la información es compleja⁵. E incluso si una de las partes pudiera expresarla, es posible que la otra parte

⁵ La información es compleja si, a causa de su carácter tácito o volumen total, es difícil de expresar, absorber o procesar, lo que hace que su transferencia sea muy costosa (von Hippel, 1994).

no sea capaz de comprenderla en su totalidad. Por tanto, para ciertos tipos de información, los costes de transacción pueden ser muy altos.

El oportunismo es la búsqueda del interés propio con astucia y la realización de afirmaciones que uno mismo no cree (por ejemplo, realizando falsas promesas de forma intencionada). El comportamiento oportunista lo pueden desarrollar determinadas personas algunas veces. Puede darse *ex ante*, cuando existe información asimétrica (una parte tiene más información que la otra, como ocurre cuando un vendedor de coches usados nos oculta un defecto) o *ex post* (después de la venta el vendedor no asume la promesa de la recompra si el cliente no está satisfecho).

El comportamiento oportunista *ex ante* conduce a la «selección adversa». Por ejemplo, si una aseguradora ofrece una póliza de seguro a una tarifa única de tipo medio, ésta sólo resultará atractiva para los enfermos de alto riesgo. El resultado será que la póliza acabará cubriendo a un conjunto de clientes en el que la parte de la población de alto riesgo estará sobre-representada. El oportunismo *ex post* se refiere a las acciones que las partes pueden llevar a cabo después de que se haya ejecutado la transacción. Esta asimetría informativa se conoce como «riesgo moral o acción oculta». Por ejemplo, a una empresa de seguros no le es posible conocer si el siniestro fue accidental o provocado. La información oculta es un concepto *ex ante*: se refiere a la información privada que existe antes de que las partes acuerden realizar una transacción. La acción oculta es un concepto *ex post*: se relaciona con la información privada que puede desarrollarse durante la ejecución de una transacción (Douma y Schreuder, 1991).

Las transacciones tienen tres dimensiones críticas: especificidad de los activos, incertidumbre/complejidad y frecuencia. Un activo es específico (o idiosincrásico) de una transacción si no puede ser reasignado a un uso alternativo sin una reducción significativa de su valor. Por otra parte, mientras más incertidumbre haya más espacio habrá para la asimetría de información y el oportunismo. La frecuencia de las relaciones se refiere a la cantidad de transacciones que llevan a cabo las partes. En suma, para aquellas transacciones con un alto grado de especificidad de activos, un alto grado de incertidumbre/complejidad y una alta frecuencia, los costes de las transacciones de mercado son extremadamente altos, mucho mayores que los costes de las transacciones internas. En consecuencia, tales transacciones tienden a ser realizadas dentro de las empresa (Williamson, 1975).

Las características de algún tipo de innovaciones, como, por ejemplo, las de carácter radical, aconsejan su internalización para minimizar los costes de transacción. La innovación radical es muy incierta, lo que hace muy difícil la especificación de los estados futuros del mundo y deja un gran espacio para el oportunismo. A causa de su naturaleza tácita o del volumen total de parte de la información, puede ser difícil especificar en contratos qué es lo que se espera de cada parte, ya que hay asimetría de información. Con frecuencia, este tipo de innovación requiere fábricas o individuos especializados, lo que de nuevo aumenta las probabilidades del oportunismo (Afuah, 1998).

Contratar la investigación con empresas independientes es una actividad no exenta de riesgo de comportamientos oportunistas por parte del vendedor, ya que esta actividad implica: a) especificaciones incompletas de los contratos, dada la incertidumbre sobre los resultados de la investigación; b) falta de protección adecuada de la información de

propiedad privada; c) posibilidad de fenómenos de «amarrar» (*lock-in*) con los oferentes de investigación, que, en consecuencia, pueden obtener rentas de esa ventaja asimétrica; d) pocos incentivos para minimizar costes y e) elevados costes de control (Dosi, 1988).

Además de las ventajas consideradas, la integración tendría la ventaja adicional de facilitar un mejor flujo de información desde el departamento de I+D hasta aquéllos que tienen que poner en práctica la nueva tecnología, y a la inversa. También serviría para limitar las filtraciones de información entre empresas. Por supuesto, se observan a menudo en el mercado transferencias de innovaciones y de competencias técnicas, tales como licencias y acuerdos de asistencia técnica. No obstante, estas formas de transferencia de tecnología entre empresas no constituyen un sustituto absoluto de la investigación dentro de la empresa. Se necesita tener suficiente capacidad interna como para reconocer, evaluar, negociar y, finalmente, adaptar la tecnología potencialmente disponible a otros (Dosi, 1988). En este sentido, los costes de transacción son una explicación insuficiente: incluso en las industrias en las que los contratos de I+D y las licencias de *know-how* técnico son una práctica común, éstos no representan casi nunca una alternativa a las actividades técnicas internas –incluyendo la I+D–, sino que son complementarias con ellas. De hecho, el conocimiento técnico rara vez puede ser obtenido de la estantería y casi siempre requiere el procesamiento y la modificación para ser utilizado efectivamente. Sin esta asimilación y mejora, no es factible que se produzca otra cosa que resultados ineficientes (Freeman, 1998).

La teoría de recursos ayuda a diferenciar las tecnologías que son relevantes para la empresa de aquellas otras que, aun siendo necesarias para la fabricación de los productos o la prestación de servicios, son de fácil acceso en el mercado.

Pero, ¿qué tecnologías debe desarrollar una empresa? La respuesta es las tecnologías relevantes para la empresa, denominadas «tecnologías nucleares», que han de ser idiosincrásicamente sinérgicas, inimitables e incontestables (Langlois y Robertson, 1995). Es decir, se trata de tecnologías específicas a la empresa que se combinan entre ellas para generar un resultado único y más valioso que el conjunto de resultados que se obtendrían de cada una de ellas si se utilizaran por separado. Por otra parte, se trata de tecnologías difíciles de imitar por los competidores. Finalmente, estas tecnologías las tiene que desarrollar internamente la empresa, ya que no se pueden comprar y vender en el mercado. Estas tecnologías trabajando juntas constituyen el potencial tecnológico de la empresa y la fuente de su ventaja competitiva.

Para saber qué tecnologías deben integrarse hay que basarse en la clasificación de Mahoney y Pandian (1992), que distinguen entre «sinergia contestable» y «sinergia idiosincrásica». Cuando existe contestabilidad, se da una combinación de tecnologías que crea valor competitivamente disponible, es decir, se obtienen sinergias a partir de tecnologías que se pueden adquirir en el mercado. De otra parte, la sinergia idiosincrásica se da cuando el incremento del resultado es el producto específico de la combinación de unas tecnologías únicas, no disponibles en el mercado. Esta última sinergia es posible alcanzarla mediante el potencial tecnológico. En suma, la empresa debería desarrollar internamente las tecnologías nucleares capaces de producir sinergias incontestables y, por tanto, ventajas competitivas.

El desarrollo de tecnologías incontestables se encuentra, a menudo, incorporado en rutinas organizativas. En realidad, los elementos heurísticos sobre cómo hacer las cosas y cómo mejorarlas suelen encontrarse incorporados en rutinas organizativas, que, a través de la práctica, la repetición y mejoras más o menos incrementales hacen que ciertas empresas sean buenas, explotando determinadas oportunidades técnicas y trasladándolas a productos de mercados específicos. En tales materias, existe un grado significativo de indivisibilidad organizativa, puesto que el aprendizaje organizativo puede no ser aditivo en el aprendizaje de los individuos o los grupos que componen la organización (Dosi, 1988).

Por su parte, Teece (1986b) justifica la importancia de la apropiabilidad y la especificidad de activos en su modelo de integración vertical (figura 34). La apropiabilidad es la capacidad de una de las partes (en este caso, el innovador) de capturar los beneficios de la innovación y dependerá del grado en el que se puede evitar la imitación. La habilidad del innovador a la hora de apropiarse de estas rentas determinará la extensión de la integración vertical. Por un lado, esta habilidad dependerá del grado de «complementariedad» y del «régimen de apropiabilidad». Éste último es, a su vez, la habilidad –legal y práctica– para fijar y hacer cumplir los derechos de propiedad industrial sobre la innovación (Teece, 1986b). Se observa que un innovador en busca de beneficios, con un régimen de propiedad débil y con la necesidad de recursos complementarios especializados, está obligado a ampliar sus actividades mediante la integración vertical. Cuando el conocimiento implicado está protegido con facilidad por patentes, como en el caso de los productos farmacéuticos, las licencias pueden hacer totalmente innecesaria la organización interna. Sin embargo, cuando éste no es el caso (como sucede quizá en algunas tecnologías de procesos), la organización interna puede ser la forma más eficiente para proteger los beneficios. El innovador tan sólo necesita tomar una posición –no de control– en los activos complementarios para obtener beneficios de su innovación. Por otro lado, cuando los activos implicados son especializados, los costes de transacción de mercado son elevados. En consecuencia, el innovador debe integrar la actividad para reducir los costes de transacción.

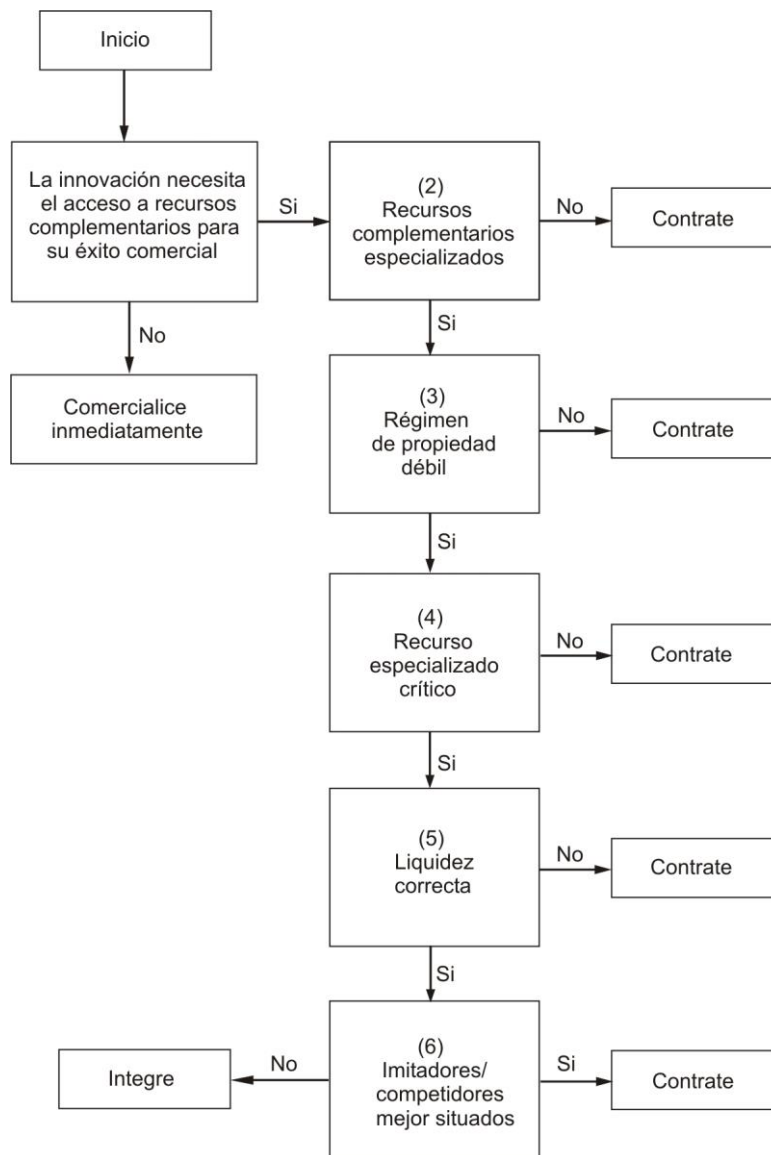


Figura 34: Diagrama de flujo de la integración frente a la contratación (Teece, 1986b).

El innovador también puede tener otros motivos para integrar activos no especializados. De lo contrario, los imitadores podrían entrar rápidamente y eliminar, así, los beneficios de la innovación.

Al objeto de distinguir entre estas dos variantes de la teoría de la integración, Langlois y Robertson (1995) denominan «apropiabilidad» a la versión de Teece y «empresarial» a la de Silver.

5. CLAVES DEL ÉXITO DE UNA INNOVACIÓN

Una innovación con éxito es la que otorga a la empresa una cuota de mercado que le permita recuperar la inversión realizada en su desarrollo y puesta en práctica; y obtener, cómo no, un beneficio que recompense el riesgo asumido. Un fracaso es una innovación que no logró una cuota de mercado aceptable y, por lo tanto, no proporciona beneficios, aun cuando haya constituido un gran adelanto sobre la tecnología en uso. Los beneficios no tienen por qué producirse de una manera inmediata.

El carácter aparentemente arbitrario del éxito de la innovación surge de la extrema complejidad en la determinación del avance tecnológico y de las imprevisibles alteraciones del mercado y de la competencia (Rothwell y Zegveld, 1981). Precisamente, la mayor fuente de dificultades para el éxito de la innovación reside en que el mercado y la tecnología están cambiando de forma continua. En consecuencia, los posibles fracasos empresariales surgen tanto de la incertidumbre técnica, inherente al propio proceso innovador, como de realizar juicios erróneos sobre las condiciones del mercado y la competencia.

El éxito técnico es más fácil de calcular que el éxito comercial. Un proyecto puede ser un resonante éxito técnico, pero estar muy lejos de ser un éxito comercial. A corto plazo es posible que un proyecto de poco riesgo logre los objetivos técnicos previstos, pero la innovación no tiene por qué convertirse en un éxito comercial, porque está expuesta a que salga un nuevo producto de la competencia, a que cambien las preferencias de los clientes o a que se produzca una alteración repentina en las prioridades de la empresa promotora (Mechlin y Berg, 1982). De hecho, la investigación realizada por Mansfield *et al.* (1971) concluye que la mayoría de los proyectos alcanzan los objetivos técnicos, pero sólo un poco más de la mitad se comercializa en el mercado. De éstos, aproximadamente un 60 por 100 no logran cubrir los costes de oportunidad de los recursos empleados en la investigación. La mayor parte de los riesgos de los proyectos de I+D son comerciales, no técnicos. El riesgo técnico suele ser bastante modesto, excepto en áreas donde se busca superar el estado del arte, como la industria espacial, la química o la farmacéutica (Schon, 1966). Ahora bien, Mansfield (1982) también observó que si los esfuerzos de la I+D se centran en proyectos ambiciosos, la probabilidad de éxito técnico disminuye, pero aumenta la del éxito comercial. En ningún caso la reducción de la probabilidad de éxito técnico está contrarrestada con un aumento de la de éxito comercial. Por otra parte, la innovación de producto comporta incertidumbre técnica y de mercado. La innovación de proceso puede comportar sólo incertidumbre técnica si se aplica dentro de la propia empresa y, de acuerdo con Hollander (1965), ésta puede ser mínima para las mejoras técnicas secundarias.

El compromiso de la alta dirección afecta positivamente al éxito de una innovación. Dicho compromiso debe concretarse en una dotación de recursos adecuada, seguimiento continuo del proyecto, objetividad en la evaluación de las posibilidades técnicas y comerciales del proyecto, respeto a la autonomía del equipo para la ejecución del trabajo y protección contra todo tipo de injerencias intencionadas.

La alta dirección estará más comprometida con la innovación si tiene cierta formación técnica. De hecho, las empresas que conceden menos importancia a la formación técnica de sus altos cargos, suelen encontrarse con problemas de falta de comunicación entre el personal económico y el técnico (Gluck y Foster, 1975); lo que debilita el proceso de innovación tecnológica y el éxito de nuevos productos en el mercado.

La actitud de la alta dirección puede asimismo ser la causante de algunos fracasos. Así, en ciertos casos, sobre todo en las PYMEs de reciente creación, la dirección tiende a sobrevalorar la importancia de la tecnología. Esto lleva, por ejemplo, a asignar de una forma desproporcionada un alto presupuesto a la investigación a costa de restringir el del resto de actividades funcionales (ingeniería, producción y marketing, fundamentalmente). El énfasis en la actividad investigadora provoca una cierta proliferación de nuevos productos con altas calidades técnicas, pero que no satisfacen las necesidades de los

clientes y, por tanto, constituyen fracasos rotundos en el mercado. En consecuencia, la empresa puede verse ahogada financieramente, ya que tiene muchos recursos inmovilizados en nuevos productos que carecen de demanda en el mercado. En general, las empresas cuya dirección está claramente orientada hacia la tecnología suelen tropezar con dos peligros (Frohman, 1982). El primero es el de sacrificar otros campos de actividad en aras de la técnica, con los inconvenientes que se acaban de exponer. El segundo es la excesiva participación de la dirección en las decisiones tecnológicas, que acaban entorpeciendo la labor de los ingenieros e investigadores.

De otra parte, Cooper (1979), al estudiar los factores relacionados con el éxito de 195 productos industriales, de los cuales 102 fueron considerados como éxitos y 93 como fracasos, observó la presencia de tres factores como determinantes del éxito del producto: 1) superioridad técnica del producto (mejores características funcionales o productos más innovadores) respecto a los productos de la competencia, 2) saber hacer en marketing de la empresa o la comprensión de las necesidades del mercado y la forma de satisfacerlas y 3) saber hacer tecnológico, que permite a la empresa fabricar buenos productos de una forma eficiente. Las probabilidades de éxito para cada uno de estos factores son del 82, 79 y 68 por 100, respectivamente. Si en un producto coinciden los tres factores, su probabilidad de éxito es del 90%.

Disponer de una tecnología superior es necesario, por supuesto, pero no es suficiente para garantizar el éxito en el mercado. La clave está en satisfacer las necesidades de los clientes. Se puede considerar que la característica fundamental de una innovación con éxito radica en el acoplamiento de las distintas fases del proceso tecnológico a las exigencias del mercado. En este sentido, Carter y Williams (1959) consideran que una buena comunicación con el mercado es uno de los factores más importantes para el éxito de una innovación. El producto o proceso tiene que ser diseñado, desarrollado y liberado de obstáculos para satisfacer las exigencias específicas de los futuros usuarios, con lo que la comprensión del mercado ha de estar presente desde las primeras etapas y a lo largo de todo el proceso tecnológico (Freeman, 1974).

La conexión entre los departamentos de I+D, producción y marketing representa una de las barreras más importantes para alcanzar el éxito de una innovación, aun cuando exista una necesidad en el mercado. Mansfield y Wagner (1975), después de entrevistar a los directivos de I+D, producción y marketing de 18 empresas norteamericanas, concluyeron que la probabilidad de éxito comercial habría sido un 50% superior (o sea 0,5 en lugar de 0,32) si los proyectos hubiesen sido fabricados y comercializados correctamente por los departamentos de producción y marketing. En general, la tecnología tendrá éxito si en su desarrollo y puesta en práctica existe una interrelación entre todos los departamentos funcionales implicados, así como entre éstos y el mercado. La mejor forma de integrar todas las funciones es creando un equipo para el desarrollo de un producto con un gerente de proyecto influyente.

Desarrollar competencias en innovación requiere que los gerentes tomen medidas proactivas para aprender de las experiencias en el desarrollo del producto e incorporar las lecciones de los éxitos y fracasos anteriores en los futuros procesos de desarrollo de nuevos productos (Edmondson, 2004).

LECTURA 4: RAZONES DE ÉXITO EN LAS INNOVACIONES
--

El estudio SAPPHO (Scientific Activity Predictor from Patterns with Heuristic Origins) de la Science Policy Research Unit de la Universidad de Sussex analiza cuarenta y tres pares de innovaciones en productos químicos y en instrumentos científicos, utilizando entrevistas extensas y análisis técnicos. Encontró las cinco categorías de variables más importantes para distinguir los éxitos de los fracasos:

1. Se vio que los innovadores que tuvieron éxito tenían un mejor conocimiento de las necesidades del usuario.
2. Los innovadores que tuvieron éxito habían puesto más atención en el marketing y la publicidad.
3. Los innovadores con éxito realizaban su trabajo de desarrollo con mayor eficiencia que los que fracasaban, pero no necesariamente con mayor rapidez.
4. Los innovadores con éxito hacían más uso de la tecnología exterior y del consejo científico, por lo menos en el área específica interesada, si ya no en general.
5. Los individuos responsables que había detrás de las innovaciones con éxito eran normalmente de mayor edad y gozaban de mayor autoridad que sus equivalentes en los proyectos sin éxito.

En la explicación del éxito o del fracaso aparecieron también algunos contrastes interesantes entre las industrias de productos químicos y las de instrumentos. Las innovaciones en productos químicos fueron predominantemente innovaciones en los procesos, desarrolladas por empresas muy grandes; la influencia de las empresas pequeñas fue menos importante, como también lo fue la comprensión de las necesidades del usuario, mientras que esto era importante en la industria de los instrumentos, caracterizada por empresas pequeñas y por las innovaciones del producto. El plazo de entrega, esto es, el hecho de ser el primero con un nuevo producto o un nuevo proceso fue de gran importancia en las innovaciones de la industria química, en la que, con frecuencia, hubo cambios completamente radicales, mientras que en la instrumentalización la primera empresa que introducía una innovación fallaba más a menudo de lo que lo hacía el segundo innovador.

Fuente: Rosenberg, N. (1982): *Inside the Black Box: Technology and Economics*, Cambridge University Press, Cambridge.

Es posible que un producto fracase en el mercado por una deficiente comercialización, a pesar de que técnicamente haya sido un éxito o también porque su potencial de ventas es escaso (Litvak y Maule, 1972). A veces, un enemigo de la innovación son los propios vendedores de la empresa, que no la promocionan correctamente, porque ello implica dejar obsoletos productos muy conocidos y rentables que les representan sustanciales comisiones y son fáciles de colocar en el mercado (Morita, 1986). Si se trata de bienes industriales, con frecuencia el producto físico es bueno, pero fracasa porque no se han desarrollado los servicios de apoyo necesarios, como el servicio de mantenimiento, la formación de operarios y asegurar el suministro de los componentes, entre otros. A la inversa, una buena comercialización puede neutralizar en gran medida las deficiencias de un producto mediocre.

Parece claro que un único factor no juega un papel importante en el éxito o fracaso de las innovaciones. En general, existe un cierto consenso acerca de la importancia de los siguientes factores en el éxito de una innovación: identificación de las necesidades del usuario, correcta formulación de los objetivos, políticas de marketing eficaces, comunicaciones abiertas tanto hacia el interior como el exterior, cooperación de todas las funciones desde el inicio, eficiencia del desarrollo técnico, buen departamento de investigación y desarrollo, cultura adaptativa, protección eficaz de la innovación y compromiso de la alta dirección. Todo lo anterior resultaría ineficaz si la empresa no tuviera un presupuesto de I+D equilibrado.

Presupuesto de I+D. La cantidad de recursos que una empresa puede dedicar a la investigación varía entre dos extremos. El extremo superior está determinado por factores como la disponibilidad de recursos financieros y de personal técnico, así como la capacidad de la empresa para producir y comercializar los nuevos productos que puedan

surgir del esfuerzo investigador. El extremo inferior, denominado «umbral de I+D» (o masa crítica), representa el nivel mínimo absoluto de recursos (no una ratio sobre ventas) que toda empresa debe asignar a I+D, con objeto de mantenerse al día en los cambios tecnológicos. Por debajo de este umbral será muy difícil proteger los antiguos productos o desarrollar otros nuevos en períodos suficientemente cortos para sobrevivir (Freeman, 1974).

El presupuesto total de investigación y desarrollo para un ejercicio dado se establece de varias formas, sin que ningún método presente unas bases totalmente satisfactorias. Naslund y Sellstedt (1974) encontraron que las empresas calculan el presupuesto de I+D de acuerdo con los siguientes métodos: porcentaje sobre las ventas previstas (12%), porcentaje sobre las ventas anteriores (5%), la misma cantidad que el competidor principal (1%), porcentaje incrementado sobre el año anterior (9%), análisis de proyectos individuales (62%), negociación (3%) y otros métodos (8%). El análisis de proyectos individuales es la mejor forma de asignar los fondos de investigación. El método del porcentaje sobre las ventas puede resultar muy ineficaz, ya que cuando las ventas varían el presupuesto también lo hace, y tales variaciones pueden socavar la estabilidad necesaria en el presupuesto de I+D. De hecho, una caída de las ventas provoca una reducción de los presupuestos de I+D (Terlecky, 1963); lo que introduce a la empresa en un círculo vicioso, puesto que esta actuación contribuirá a un bajo rendimiento a medio y largo plazo. No obstante, con independencia del método elegido, el juicio y la negociación jugarán inevitablemente un papel importante en la asignación de recursos a la I+D (Twiss, 1974). Existen múltiples fuerzas que actúan sobre la determinación del presupuesto de I+D, que pueden clasificarse en dos grandes grupos (Seiler, 1974):

- Factores no económicos, entre los que destacan el número de personas orientadas hacia la investigación en el consejo de administración, la edad del presidente ejecutivo y la presión de la opinión pública, entre otros.
- Factores económicos, como los beneficios de la empresa, las ventas periódicas, la inversión total, el porcentaje de utilización de la capacidad productiva, el presupuesto de I+D de la competencia y la cuota de participación en el mercado, entre otros.

La necesidad de mantener un grupo de personal investigador muy preparado hace que el esfuerzo de investigación y desarrollo represente un coste relativamente estable. No puede incrementarse y reducirse a corto plazo sin destruir en gran medida su eficacia. En algunas industrias existe una especie de código acerca de cuánto gastar en I+D: gastar menos es un signo de negligencia cara al futuro, gastar más se considera un desperdicio de recursos económicos (Naslund y Sellstedt, 1974). Muchas empresas suelen manifestar que aprueban el montante total del presupuesto estimado por el departamento de I+D. Sin embargo, análisis más detallados demuestran que los responsables de I+D ya calculan el presupuesto en función de lo que creen que es aceptable para la alta dirección (Seiler, 1974). Éste es un problema común en muchos procesos de elaboración de presupuestos.

En una industria en la que cada diez años surgen nuevos procesos o nuevas generaciones de productos, será necesario un nivel de intensidad investigadora relativamente alto para evitar la obsolescencia de la gama de productos o los excesivos costes. Un nivel muy bajo de actividad de I+D, o su ausencia absoluta, sería una estrategia viable en aquellas ramas de la industria donde la obsolescencia tecnológica no es un problema, o donde los cambios en la gama de productos obedecen fundamentalmente a la moda (Freeman, 1974).

A su vez, Comanor (1967) encontró que las inversiones en I+D eran relativamente más elevadas en el caso de productos complejos, cuya diferenciación en el mercado se basa fundamentalmente en sus características funcionales. En cambio, eran relativamente bajas en los bienes de producción estándar y en los bienes de consumo, cuya diferenciación puede ser ocasionada por la publicidad o la imagen más que por las diferencias físicas de las especificaciones del producto.

Mansfield *et al.* (1971) en un estudio realizado sobre una muestra de empresas de los sectores químico, electrónico y maquinaria, desglosan los costes como sigue: 10% para investigación básica, 7,5% para preparación de las especificaciones del producto, 30% para el prototipo o planta piloto, 35% para herramientas y proceso de fabricación, 10% para el inicio de la fabricación y 7,5% para el inicio del marketing. Existen diferencias significativas entre las distintas industrias: en química la investigación y las especificaciones del producto acaparan el 30% de los costes, mientras que en maquinaria y electrónica representan menos del 10%.

Por lo general, las empresas subestiman tanto el coste del proyecto como el tiempo necesario para completarlo, máxime si se trata de proyectos novedosos, que escapan a la experiencia previa de quien hace la estimación. Norris (1971) llegó a la conclusión de que los costes reales oscilaban de 0,97 a 1,51 veces de los costes previstos y que los tiempos de realización eran de 1,39 a 3,04 veces mayores que los estimados. Sin embargo, Marshall y Meckling (1962) llegan a unas conclusiones diferentes, ya que estiman que la ratio coste final de desarrollo/coste estimado fue entre 2,4 y 3 y el tiempo oscila entre 1/3 y 1/2 del previsto. En una muestra de empresas de aviones y misiles la tasa media actual fue de 3,2 en el coste y de 1,4 en el tiempo estimado (Mansfield, 1982). En los proyectos de grandes dimensiones las desviaciones en costes son superiores a las temporales. Sin embargo, en los proyectos pequeños ocurre lo contrario. Cualquier intento de acelerar el proyecto y abreviar el tiempo total hará aumentar los costes totales (Scherer, 1980). Por otra parte, los costes de I+D se pueden reducir si el personal de investigación ha tenido experiencia previa en el manejo de problemas similares (Mansfield y Rapoport, 1975).