

Previsión del suministro usando modelos estadísticos de supervivencia

Díaz Martínez, Zuleyka (zuleyka@ccee.ucm.es)

Departamento de Economía Financiera y Actuarial y Estadística

Facultad de Ciencias Económicas - UCM

Fernández Menéndez, José (jfernán@ccee.ucm.es)

Departamento de Organización de Empresas y Marketing

Facultad de Ciencias Económicas - UCM

Minguela Rata, Beatriz (minguela@ccee.ucm.es)

Departamento de Organización de Empresas y Marketing

Facultad de Ciencias Económicas - UCM

RESUMEN

La previsión de la demanda en las organizaciones tiene un gran interés y abundan los modelos estadísticos para estimar la distribución de probabilidad de sus valores. Por el contrario, la previsión de las cantidades que se podrán obtener de los proveedores no ha recibido apenas atención. Las organizaciones hacen una previsión de la demanda a medio plazo para negociar con los proveedores las condiciones de suministro. Las inevitables variaciones que se produzcan en la práctica con respecto a las condiciones acordadas se gestionan sobre la marcha y no suelen ser suficientemente relevantes como para modelizarlas. Sin embargo, no siempre es así. Por ejemplo, para los distribuidores de productos farmacéuticos existen restricciones legales o contractuales que les obligan a garantizar el suministro a sus clientes y mercados. En este caso es frecuente que se produzca racionamiento de la oferta por parte de los laboratorios suministradores, lo que hace aconsejable modelizar estadísticamente tanto la demanda como la oferta para garantizar que la probabilidad de desabastecimiento no rebase unos límites aceptables. En este trabajo se

presentan modelos estadísticos de supervivencia o fiabilidad, cuyo rasgo distintivo es la presencia de observaciones censuradas, para modelizar adecuadamente la oferta por parte de los proveedores.

Palabras clave: Previsión, modelos de supervivencia, distribución de Weibull.

Área temática: A5 - Aspectos Cuantitativos de Problemas Económicos y Empresariales.

ABSTRACT

Demand forecasting in organizations is a very noteworthy issue, and many statistical models have been developed to estimate the probability distribution of its values. However, the forecast of the amounts of goods that can be gotten from the suppliers has not received almost any attention. Usually, firms make a mid term forecast of the demand in order to bargain with suppliers the details of the procurement. The unavoidable variations that spring over in practice regarding to the agreed conditions, are managed on the fly, and typically they are not relevant enough in order to justify their statistical modeling. But that's not always the case. For example, there are legal and contractual requirements for pharmaceutical distributors that force them to provide a given service level to their customers. In this situation, there is often some degree of supply rationing from the pharma laboratories, and this makes it advisable to model both demand and supply in order to guarantee that the medication shortage probability is kept under an acceptable level. In this work we present some survival, or reliability, statistical models, whose distinctive feature is the presence of censored observations, to properly model the supply from the providers.

1. INTRODUCCIÓN

Disponer de previsiones razonablemente válidas de la demanda futura es una cuestión de gran importancia en las organizaciones. Estas previsiones resultan esenciales para planificar el proceso productivo, determinar las necesidades de materias primas y personal y gestionar las contrataciones, gestionar los inventarios, ajustar la capacidad productiva, negociar las compras con los proveedores, etc. (Heizer & Render, 2007; Hopp & Spearman, 2011; Thomopoulos, 2015). Para realizar esas previsiones de la demanda se han desarrollado, y se utilizan de forma rutinaria, una

gran cantidad de técnicas, que se suelen diferenciar entre cualitativas y cuantitativas (Hill, 2012). Las cualitativas son técnicas basadas de forma esencial en opiniones de expertos, como ocurre por ejemplo con el método Delphi. Las cuantitativas por su parte van desde sencillas heurísticas y aproximaciones basadas en la experiencia pasada, hasta complejos métodos estadísticos y econométricos, y se suelen clasificar en métodos basados en series temporales, por un lado, y modelos causales, que son básicamente modelos de regresión, por otro lado. En cualquier caso la evidencia parece indicar que con herramientas de complejidad moderada se pueden conseguir resultados satisfactorios en la práctica (Silver, Pyke & Thomas, 2017). En la gran mayoría de manuales y tratados sobre Dirección de Operaciones, por ejemplo en Heizer y Render (2007), en Thomopoulos (2016), o en Hopp y Spearman (2011), y también en manuales de Gestión de Inventarios, como Axsäter (2006) o Silver, Pyke y Thomas (2017), se pueden encontrar excelentes, y muy comprensibles, descripciones de estos métodos.

En los entornos empresariales, para la previsión de la demanda son especialmente usados en la práctica, por su sencillez y por la amplia disponibilidad de software que los implementa, los modelos de series temporales. Pero las preferencias se orientan no tanto hacia los modelos de tipo ARIMA y sus desarrollos como ARCH, GARCH, etc., originados en la obra de Box y Jenkins (1970), que son más habituales en la econometría y las finanzas, como hacia los modelos basados en medias ponderadas, particularmente el alisado exponencial, denominado así porque las ponderaciones decrecen exponencialmente para las observaciones más antiguas. El uso del alisado exponencial para realizar previsiones en las empresas y en la administración surge a mediados del siglo pasado como consecuencia del trabajo de Robert G. Brown (Brown, 1959, 1963), que luego fue extendido por Holt (2004, edición original de 1957) y Winters (1960), para dar lugar al popular método de Holt-Winters, que amplía la versión básica del alisado exponencial para incluir términos de tendencia

y estacionalidad.

Posteriores desarrollos han permitido implementar estos métodos de series temporales basados en el alisado exponencial en términos de modelos de espacio de estados (Hyndman et al., 2008; Hyndman et al., 2002). Los modelos de espacio de estados surgieron inicialmente en el área de la ingeniería, en automática y procesamiento de señales, pero también han demostrado ser muy útiles en el campo de las series temporales.

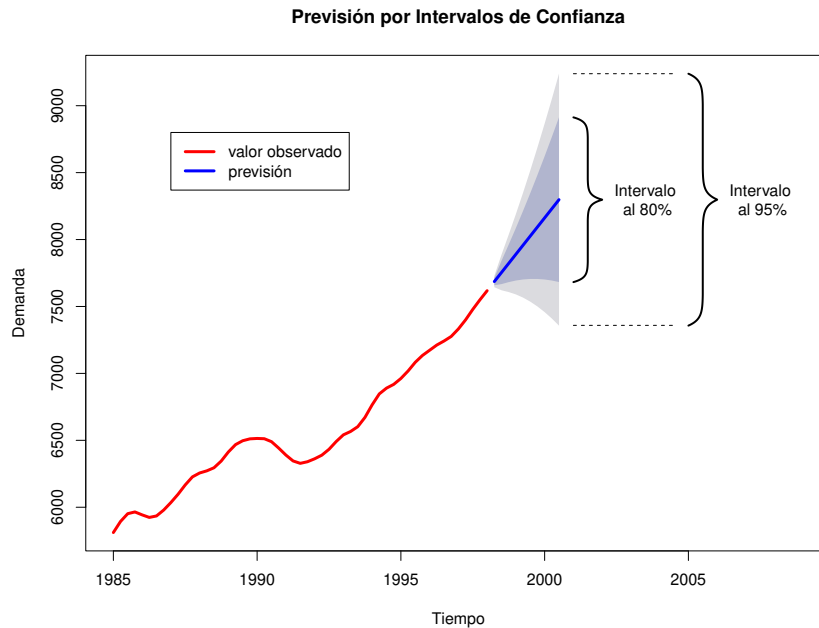
De acuerdo con la formulación de espacio de estados un modelo de alisado exponencial se representa por medio de dos ecuaciones. Una de ellas expresa la variable de interés sobre la que se quieren hacer predicciones (y_t) como función lineal de un vector de variables de estado para el periodo anterior que incluye los efectos de la tendencia y la estacionalidad (x_{t-1}) y de un término de error (ϵ_t):

$$y_t = w'x_{t-1} + \epsilon_t \quad (1)$$

La segunda ecuación muestra cómo el vector de estado en un instante dado depende linealmente (a través de la matriz de coeficientes F) del correspondiente vector en el instante anterior:

$$x_t = Fx_{t-1} + g\epsilon_t \quad (2)$$

El término de error ϵ_t es, en principio, la diferencia entre los valores observados y los predichos: $\epsilon_t = y_t - \hat{y}_t$. Asumiendo que esta diferencia sigue una determinada distribución de probabilidad, que puede ser estimada a partir de los datos históricos, se pueden dar valores a los términos de error ϵ_t para obtener la distribución de probabilidad de las predicciones y_t . De esta forma resulta posible responder a cuestiones como: ¿cuál es la probabilidad de que la demanda no supere las mil unidades?, o ¿cuál es el valor de la demanda que no se superará el 90% de las ocasiones? Esto permite también obtener intervalos de valores dentro de los que la demanda prevista se encontrará con un determinado nivel de confianza, como se representa en la siguiente figura:



La estimación de este tipo de modelos es relativamente sencilla y está disponible en la mayoría del software estadístico de uso habitual, por ejemplo, en el paquete *forecast* de R (Hyndman & Athanasopoulos, 2014; Hyndman et al., 2020; Hyndman & Khandakar, 2008)

Así como la previsión de la demanda ha sido objeto de una gran atención, la previsión de la oferta por parte de los proveedores, es decir, la modelización de las cantidades que se pueden conseguir de los proveedores cuando se les hacen pedidos, ha sido completamente ignorada por parte de las organizaciones. En general, lo que se hace habitualmente es realizar previsiones de la demanda con un horizonte temporal típicamente de entre uno y dos años y, sobre la base de esas previsiones, se negocia el suministro por parte de los proveedores, se planifica la adquisición de maquinaria y la contratación de mano de obra, etc. Las habituales desviaciones en las cantidades recibidas con respecto a las pedidas a los proveedores, retrasos, etc., se gestionan sobre la marcha y no suelen tener la suficiente entidad como para causar problemas relevantes ni para que esté justificado el esfuerzo que supondría su

modelización. Sin embargo, no siempre es así. Por ejemplo, en el caso de los distribuidores de productos farmacéuticos existen habitualmente restricciones legales que les obligan a garantizar el suministro de determinados productos en los mercados en los que operan. Esto significa que si es previsible un racionamiento de la oferta por parte de sus propios proveedores, los laboratorios que les suministran las especialidades farmacéuticas, los distribuidores deben concentrar la oferta en sus mercados habituales y restringir las exportaciones hacia mercados secundarios, incluso aunque estos fuesen más rentables. Estos distribuidores de productos farmacéuticos tienen, por lo tanto, un gran interés en estimar tanto la distribución de la demanda futura como la de las cantidades que podrán obtener de sus proveedores, para garantizar en la medida de lo posible que la probabilidad de desabastecimiento se mantenga dentro de unos límites que se consideren aceptables. Se trataría de una situación en la cual el distribuidor de productos farmacéuticos, ante la posibilidad de exportar una determinada cantidad de un producto, digamos 100 unidades, evalúe la demanda previsible y la oferta previsible por parte del laboratorio, y concluya que para mantener la probabilidad de desabastecimiento en su mercado local por debajo de, por ejemplo, el 2 %, lo máximo que puede destinar a la exportación son 40 unidades. Para poder llevar a cabo este tipo de análisis es necesario disponer de modelos que permitan analizar y estimar las cantidades que se pueden obtener por parte de los proveedores y cómo es la distribución de probabilidad de estas cantidades. Lo que se busca con este trabajo es desarrollar algunos modelos que permitan realizar de forma razonablemente sencilla este tipo de análisis, es decir, se trata de presentar modelos estadísticos que permitan realizar para la oferta el mismo tipo de previsiones que se llevan a cabo de forma rutinaria para la demanda.

2. MODELOS DE PREVISIÓN DE LA OFERTA

Consideremos una situación de compra de bienes por parte de una empresa a un proveedor que, pese a su sencillez y notable grado de idealización, sí que recoge los aspectos básicos de dicho proceso que son relevantes de cara a la modelización que aquí se plantea. Vamos a denominar a este modelo **Modelo I**.

Supongamos que la empresa desea adquirir una determinada cantidad de unidades, Q_0 , y que realiza un pedido al proveedor por esta cantidad. Habitualmente el proveedor le suministrará la cantidad demandada, pero en ocasiones le enviará una cantidad menor. El objetivo es modelizar este fenómeno, es decir, modelizar la probabilidad de obtener una u otra cantidad, basándose en la experiencia pasada y, eventualmente, en determinadas variables que puedan influir sobre esta cantidad, como determinadas cotizaciones en los mercados, época del año, precios de materias primas o productos intermedios, etc.

Denominemos Q_m a la cantidad máxima que el proveedor está dispuesto a suministrarnos en cada transacción. Ésta será la variable aleatoria que intentaremos modelizar utilizando la información disponible. Cuando la cantidad deseada sea menor o igual que este valor máximo que el proveedor está dispuesto a suministrar se obtendrá la cantidad deseada, pero si es mayor se obtendrá el máximo Q_m . Si denominamos Q a la cantidad que se obtiene resulta que $Q = \min(Q_0, Q_m)$. Q_0 es una cantidad conocida en cada caso y Q es una cantidad observada, pero el valor real de Q_m sólo se observa en algunas ocasiones, cuando se obtiene menos de lo que se desea, es decir, cuando $Q_0 > Q_m$, en cuyo caso $Q = \min(Q_0, Q_m) = Q_m$. Sin embargo, si el proveedor está dispuesto a proporcionarnos más unidades de las que deseamos, es decir, si $Q_0 < Q_m$, entonces lo que se obtiene es justamente lo que deseamos, $Q = \min(Q_0, Q_m) = Q_0$, y no se observa ni se conoce el valor real de Q_m . Ésta es una situación totalmente análoga a la que se produce en el caso de modelos

de supervivencia, o de fiabilidad, en los que se tienen observaciones con censura a la derecha (*right censoring*), es decir, observaciones en las que se sabe que el tiempo hasta que ocurre un determinado evento (curación, fallecimiento, fallo, relapso, o el que sea en cada caso dependiendo de la situación que se esté analizando) es simplemente superior a un determinado valor, típicamente denominado *follow-up time*, o tiempo de seguimiento, pero no se conoce su valor real. Esto sugiere la conveniencia de usar modelos de supervivencia para la modelización de la variable Q_m .

En un modelo de supervivencia (Cox & Oakes, 1984; Kalbfleisch & Prentice, 2002; Lawless, 2003) existe una variable aleatoria que se desea modelizar y que será típicamente un tiempo, T_m , hasta que ocurre un evento determinado. Para ello se dispondrá de una serie de observaciones de ese tiempo de supervivencia, pero algunas de ellas estarán censuradas (normalmente por la derecha), de manera que en algunos casos (los de observaciones sin censurar) lo que se observa es realmente el tiempo de interés, T_m , mientras que para otras observaciones (las censuradas), lo que se observa es el tiempo de seguimiento, T_0 , pasado el cual concluye la observación y lo único que se sabe respecto al tiempo real en el que ocurre el evento es que $T_m > T_0$. De esta manera, si T es el tiempo observado, lo que se tiene es que $T = \min(T_0, T_m)$. Un aspecto muy relevante de cara a la estimación de este tipo de modelos es que siempre se sabe si el valor observado se corresponde con el valor real de la variable analizada o es un valor censurado. Esto se suele recoger en un indicador δ , al que podemos llamar «indicador de evento», que adopta el valor $\delta = 1$ cuando el tiempo que se observa es el de ocurrencia del evento, es decir, cuando $T = T_m$, y adopta el valor $\delta = 0$, cuando el tiempo observado es el de censura, $T = T_0$. De esta forma, para cada observación la información de la que se dispone es la proporcionada por la variable aleatoria conjunta (T, δ) . En nuestro Modelo I de previsión de la oferta lo que se observa en cada caso es (Q, δ) .

Como el Modelo I puede ser expresado completamente en términos de un mo-

delo de supervivencia, no habrá ninguna dificultad, como se indicará en el siguiente apartado, a la hora de estimarlo utilizando de forma inmediata todas las herramientas disponibles en los paquetes de software habituales en estadística para tratar este tipo de modelos.

Una cuestión que conviene resaltar con respecto al Modelo I es que requiere conocer los valores de Q_0 , las cantidades que se *desean* obtener realmente del proveedor, y en muchas ocasiones lo que existe disponible en los registros son las cantidades pedidas, es decir, los volúmenes de los pedidos realizados. En muchas ocasiones estas cantidades registradas coincidirán con las efectivamente deseadas, pero en otros casos, especialmente cuando existan, o se anticipen, problemas de racionamiento por parte de los proveedores, la práctica habitual por parte de los responsables de compras será la de hinchar o incrementar, de forma incluso desmesurada, las cantidades pedidas con el propósito de conseguir «por lo menos algo», y luego cancelar pedidos adicionales. De esta manera las cantidades registradas como pedidas no reflejarían las necesidades reales, que son las que se tienen en cuenta en el Modelo I. Se deben tener entonces registros de las cantidades deseadas, Q_0 , y asumir que el departamento de compras, para intentar obtener estas cantidades, gestionará las compras de la manera más adecuada posible en cada caso, lo que podrá suponer realizar pedidos por volúmenes que no tengan nada que ver con las cantidades realmente deseadas. Esos volúmenes pedidos no se tienen en cuenta para estimar el Modelo I.

El Modelo I es sencillo, fácil de estimar, como se indica más adelante, y puede ser suficientemente realista en muchas ocasiones, pero conviene ampliarlo para tener en cuenta situaciones como la siguiente: es frecuente que al realizar un pedido, incluso aunque no exista racionamiento de la oferta por parte del proveedor, las cantidades recibidas no coincidan exactamente con las solicitadas, y que se reciba algo menos de lo pedido, o incluso algo más. Esto responde a razones operativas, como errores,

descuidos, prisas, etc., a la hora de preparar el pedido (*picking*), tramitarlo, hacer el recuento de unidades, por comodidad o conveniencia, para completar o deshacerse de determinados lotes de productos, alcanzar unos objetivos de ventas modificando al alza el tamaño de algunos pedidos, y un sinnúmero de circunstancias que se pueden observar en la práctica. Para tener en cuenta todas estas eventualidades se considera un **Modelo II**.

En el Modelo II la cantidad que se recibe del proveedor en ausencia de racionamiento de la oferta es $Q_0 + \epsilon$, siendo ϵ una variable aleatoria, que por simplicidad supondremos normalmente distribuida con media 0 y desviación típica ς . De esta manera se tiene en cuenta las posibles diferencias entre la cantidad deseada y la obtenida debido a problemas operativos. Existirá además, como en el Modelo I, una variable aleatoria Q_m que indica la cantidad máxima que el proveedor está dispuesto a suministrar debido a que pueda haber un racionamiento de la oferta, y que es independiente de los aspectos operativos que determinan las características de ϵ . De esta manera la cantidad recibida será $Q = \min(Q_0 + \epsilon, Q_m)$. Habrá entonces dos fenómenos aleatorios independientes actuando en paralelo, por una parte la alta dirección del proveedor fijará un máximo que puede ser suministrado, Q_m , debido a la existencia de un racionamiento en la oferta por los motivos que sean en cada caso, y, por otro lado, debido a una serie de circunstancias relacionadas con la operativa del día a día, la cantidad que se podría recibir ($Q_0 + \epsilon$) no coincidirá exactamente con la deseada (Q_0).

Una diferencia esencial de este Modelo II con respecto al I, y que afecta a la manera de estimarlo, es que en el Modelo I se sabe cuándo una cantidad observada está censurada o no, y se refleja en el valor del indicador de evento δ , mientras que en el Modelo II esta información no está en general disponible. Cuando se recibe una determinada cantidad Q del proveedor como respuesta a un pedido con el que se desea obtener Q_0 no se sabrá en general, al menos de forma suficientemente fiable,

si el valor de la cantidad recibida es el que es por un racionamiento de la oferta o simplemente debido a la «modulación» impuesta sobre las cantidades deseadas por parte de las fluctuaciones derivadas de la operativa del día a día.

Tanto en el Modelo I como en el II el hecho de que una observación esté censurada significa que la cantidad que se recibe del proveedor es, al menos aproximadamente, la que se pretendía obtener. Es decir, la censura significa que no se tiene realmente información sobre la variable de interés, Q_m , la cantidad máxima que el proveedor está dispuesto a ofrecer, ya que la cantidad pedida está por debajo de ese valor. Esto significa que en circunstancias «normales», de ausencia de racionamiento de la oferta, se obtendrá del proveedor la cantidad deseada y las observaciones estarán censuradas, no habrá información sobre la cantidad máxima que el proveedor estaría dispuesto a proporcionar, más allá del hecho de que sea mayor de lo que se necesita.

3. ESTIMACIÓN DE LOS MODELOS

Como se ha indicado, el Modelo I supone que estamos interesados en la variable aleatoria Q_m , cantidad máxima que los proveedores están dispuestos a enviarnos. Esta variable aleatoria seguirá una determinada distribución de probabilidad que asumiremos que es una Weibull, ya que es una distribución ampliamente usada en modelos de supervivencia (aunque nada impide usar en la práctica otra distribución que resulte ser más adecuada). Para cada observación se conoce el valor de la cantidad deseada, Q_0 , y de la cantidad recibida, $Q = \min(Q_0, Q_m)$, lo que permite conocer si lo que se observa es Q_0 o Q_m , es decir, si el indicador de evento es $\delta = 0$ o $\delta = 1$, respectivamente. En resumidas cuentas, para cada observación i , se conoce el valor del par (Q_i, δ_i) .

Esta información se puede usar como *input* para un modelo paramétrico de

supervivencia que se puede estimar con cualquier software estadístico. Por ejemplo, usando R (R Core Team, 2020), en concreto el paquete de R *survival* (Therneau, 2020; Therneau & Grambsch, 2000), está disponible la función *survreg*, que permite estimar este tipo de modelos paramétricos de supervivencia. Para ello es necesario especificar la distribución de probabilidad de la variable medida (que asumiremos Weibull) y las eventuales variables explicativas, que de existir entrarían a formar parte de un predictor lineal de la forma $\mu = \beta_0 + \beta_1 x_1 + \dots + \beta_k x_k$.

La distribución de Weibull está caracterizada por dos parámetros, que en la parametrización usada por R-base son un parámetro de forma, α , y otro de escala, σ . Sin embargo, en el paquete *survival* se utiliza una parametrización alternativa, basada en el hecho de que el logaritmo de una distribución de Weibull es una distribución de Gumbel, y que utiliza un parámetro de escala que será $\frac{1}{\alpha}$, y un parámetro de localización $\mu = \log \sigma$, que es justamente el que se modeliza con el predictor lineal indicado más arriba. Se tiene que $\frac{Q}{\sigma} = Q^*$ sigue una distribución Weibull que no depende de σ (su parámetro de escala es 1), y cuyo logaritmo sigue una distribución de Gumbel de parámetro de localización 0 y parámetro de escala $\frac{1}{\alpha}$. La función *survreg* utiliza este hecho para estimar un modelo que puede ser puesto en la forma $\log Q = \log \sigma + \log Q^* = \mu + \log Q^*$, donde μ es el predictor lineal que nos proporciona el valor del parámetro de escala σ de la distribución de Weibull, y $\log Q^*$ sería un «término de error» cuyo parámetro de escala $\frac{1}{\alpha}$ también se estima y que nos proporciona el valor del parámetro de forma de la Weibull.

Podemos simular 100 observaciones de Q , parte de ellas censuradas y parte no y generadas asumiendo que Q_m sigue una distribución de Weibull con parámetro de forma $\alpha = 6$ y parámetro de escala $\sigma = 1000$. Estimando con *survreg* estos dos parámetros a partir de las observaciones simuladas se obtienen los resultados presentados en la Tabla 1, de la que se ha eliminado alguna información superflua generada por el modelo. Se puede ver que $\mu = 6,94$, que es el valor del intercepto y

que en este caso coincide con el predictor lineal al no haber variables explicativas. De acuerdo con esto el valor estimado del parámetro de escala será $\hat{\sigma} = e^{\mu} = \exp(6,94) = 1032,77$, muy próximo al valor real de 1000. El valor reportado como $\text{Log}(\text{scale})$ por *survreg* es el logaritmo de $\frac{1}{\alpha}$, con lo que en este caso tendremos que $\alpha = \exp(1,84) = 6,29$, también en excelente acuerdo con el valor real, $\alpha = 6$, del parámetro de forma de Q_m .

Modelo I (survreg)	
(Intercept)	6,94 (0,02)***
Log(scale)	-1,84 (0,09)***
Log Likelihood	-508.18
Num. obs.	100

*** $p < 0,001$, ** $p < 0,01$, * $p < 0,05$

Tabla 1: Estimación Modelo I con *survreg*

Se puede llevar a cabo la estimación de una forma más flexible y general utilizando alguna herramienta de optimización numérica para maximizar la función de verosimilitud del modelo y obtener así los parámetros de la distribución de Q_m . Para ello es necesario establecer primero cuál es la forma de la función de verosimilitud correspondiente. Al tratarse de un modelo de supervivencia convencional esto no presenta ninguna dificultad. En general la función de verosimilitud es el producto de las funciones de densidad de probabilidad para las observaciones, expresada como función de los parámetros que se desean estimar. Las observaciones son de la forma (Q_i, δ_i) y para establecer su función de densidad de probabilidad es necesario conocer la distribución de los tiempos de censura en un modelo de supervivencia, o, en nuestro caso, de las cantidades que se desean y que serán los valores en los que eventualmente se produce la censura. Opciones habituales son (Lawless, 2003; Moore, 2016) la censura aleatoria (*random censoring*), en la que los valores en los

que se censura se distribuyen de forma aleatoria e independiente de la variable analizada, o la censura de tipo I, que es un caso particular de la aleatoria en la que los valores en los que se censura están prefijados. Es razonable suponer en nuestro caso censura de tipo I ya que los valores en los que se censura, las cantidades deseadas Q_0 , se pueden considerar prefijados.

Denominaremos $f(q)$, $F(q)$ y $S(q) = 1 - F(q)$, respectivamente, a las funciones de densidad de probabilidad, de distribución y de supervivencia de una distribución Weibull. Entonces, en el caso de una observación sin censura, es decir, cuando $\delta_i = 1$ y se observa el valor de Q_m , se tiene que $Prob(Q_i = q_i, \delta_i = 1) = Prob(Q_m = q_i) = f(q_i)$, mientras que si la observación está censurada, es decir, si $\delta_i = 0$ y se observa el valor de Q_0 , se tiene que $Prob(Q_i = q_i, \delta_i = 0) = Prob(Q_m > q_i) = S(q_i)$. De acuerdo con esto la función de verosimilitud será un producto de términos, uno por cada observación, y tendrá una forma como la siguiente: $L(\alpha, \sigma) = \prod_i f^{\delta_i}(q_i) S^{1-\delta_i}(q_i)$, y tomando su logaritmo, para obtener la función *log-likelihood*, el producto se convierte en una suma, que es la expresión que habitualmente se maximiza:

$$\log L(\alpha, \sigma) = \sum_i f^{\delta_i}(q_i) S^{1-\delta_i}(q_i) \quad (3)$$

Esta función se puede maximizar con facilidad usando una función estándar de optimización, como la función *optim* de R. Un ejemplo de cómo hacerlo se encuentra en Moore (2016, p. 141). Estimando de esta manera los parámetros de la distribución Weibull del ejemplo anterior se obtienen unos valores de $\hat{\alpha} = 6,27$ y $\hat{\sigma} = 1032,55$, prácticamente idénticos a los obtenidos con *survreg* y en excelente acuerdo con los valores reales usados para simular las observaciones.

La estimación del Modelo II es más compleja y no se puede hacer con la función *survreg*, sino que es necesario maximizar directamente la función de verosimilitud. La razón de esta mayor complejidad estriba en que en el Modelo II se observa el valor $Q = \min(Q_0 + \epsilon, Q_m)$, pero no existe un indicador δ que informe sobre si hay o no censura, es decir, no se sabe cuál de estas dos cantidades, $Q_0 + \epsilon$ o Q_m , se está

observando. En consecuencia, a la hora de construir la función de verosimilitud es necesario, para cada observación, considerar ambas posibilidades, que corresponda a una observación censurada o que no.

La función de densidad de probabilidad para una observación $Q = q_i$, para la que la cantidad que se desea obtener es q_{0i} , será $Prob(Q = q_i) = Prob(Q = q_i, evento) + Prob(Q = q_i, censura)$, o lo que es lo mismo, $Prob(Q = q_i) = Prob(Q_m = q_i, q_{0i} + \epsilon > q_i) + Prob(Q_m > q_i, q_{0i} + \epsilon = q_i)$. En definitiva, como Q_m y ϵ son independientes, se tiene que:

$$Prob(Q = q_i) = Prob(Q_m = q_i) \cdot Prob(q_{0i} + \epsilon > q_i) + Prob(Q_m > q_i) \cdot Prob(q_{0i} + \epsilon = q_i) \quad (4)$$

Lógicamente, habrá que tener en cuenta que si ϵ es una variable aleatoria normal de media 0 entonces $q_{0i} + \epsilon$ será normal de media q_{0i} .

A un resultado idéntico se puede llegar a partir de la distribución de probabilidad de la variable aleatoria $W = \min(X, Y)$, donde en nuestro caso X sería Q_m , que seguirá una distribución de Weibull, e Y sería $Q_0 + \epsilon$, con una distribución normal de media Q_0 que dependería de cada observación, pero conocida en todos los casos, y desviación típica ς . Se tiene que la función de distribución de W es $F_W(w) = Prob(\min(X, Y)) = Prob(X \leq w \cup Y \leq w) = 1 - Prob(X > w, Y > w)$. Finalmente, al ser X e Y independientes resulta que $F_W(w) = 1 - Prob(X > w) \cdot Prob(Y > w) = 1 - S_X(w) S_Y(w)$. Derivando para obtener la función de densidad se tiene el mismo resultado que antes:

$$f_W(w) = f_X(w) S_Y(w) + S_X(w) f_Y(w) \quad (5)$$

De acuerdo con esto cada observación contribuye a la función de verosimilitud con un factor que será la suma de dos elementos, uno que corresponde al caso de que la observación esté censurada y otro al de que no lo esté. Tomando logaritmos

para trabajar con la función *log-likelihood* y convertir los productos en sumas, el problema se convierte en el de estimar los parámetros de la distribución de Weibull y de la normal maximizando la siguiente función:

$$\log L(\alpha, \sigma, \varsigma) = \sum_i \log(f_{weib}(q_i) S_{norm}(q_i) + S_{weib}(q_i) f_{norm}(q_i)) \quad (6)$$

Para la maximización se puede usar algún algoritmo estándar, como los que proporciona, una vez más, la función *optim* de R, pero teniendo cuidado de usar algún método, como L-BFGS-B (Byrd et al., 1995), que permite establecer restricciones sobre los valores máximos y mínimos que pueden adoptar los parámetros sobre los que se optimiza (*box constraints*), ya que los tres parámetros, α , σ y ς deben ser mayores que cero. El método L-BFGS-B es un clásico método numérico de optimización de tipo quasi-newton que lleva a cabo una optimización local utilizando el gradiente de la función objetivo. Adicionalmente y para comprobar la validez de la solución obtenida, se han utilizado otros dos métodos de optimización global, el Recocido simulado (*Simulated annealing*), implementado en el paquete *GenSA* de R (Xiang et al., 2013) y la Evolución diferencial (*Differential evolution*), disponible en el paquete *DEoptim* (Mullen et al., 2011). Ambos métodos utilizan sofisticadas técnicas de búsqueda del óptimo global dentro de la región indicada del espacio de parámetros.

Para comprobar la eficacia de los procedimientos de estimación se ha creado un conjunto de 100 observaciones simuladas, algunas censuradas y otras no, con parámetros $\alpha = 6$ y $\sigma = 1000$ para la distribución de Weibull que sigue la variable Q_m y parámetro $\varsigma = 40$ para la desviación típica de ϵ , y se han estimado estos parámetros a partir de los valores observados. Para ello se han utilizado los tres procedimientos de optimización indicados, L-BFGS-B, Recocido simulado y Evolución diferencial. Los valores estimados aparecen recogidos en la Tabla 2, donde se ve que son virtualmente idénticos con los tres métodos y muestran un excelente acuerdo

con los valores reales.

Estimación con	$\alpha = 6$	$\sigma = 1000$	$\varsigma = 40$
L-BFGS-B	6,111775	1016,318904	42,110373
Rec. Simulado	6,111775	1016,318862	42,110373
Ev. Diferencial	6,111775	1016,318857	42,110373

Tabla 2: Parámetros estimados Modelo II

4. CONCLUSIONES

La previsión de la demanda es ampliamente usada en las organizaciones debido a las grandes ventajas que aporta de cara a planificar sus actividades con la suficiente antelación. Existe por ello un gran número de herramientas de software y modelos analíticos disponibles que abarcan todo el espectro imaginable en cuanto a niveles de complejidad y sofisticación, y que son además utilizados de una forma rutinaria.

Por el contrario, la previsión de la oferta por parte de los proveedores no ha recibido virtualmente ninguna atención debido a que la incertidumbre asociada a ella se considera manejable en la práctica sin que sea necesaria su modelización. Sin embargo, no siempre es así, y existen situaciones en las que la previsión adecuada de esa oferta es importante de cara a garantizar un suministro adecuado en el mercado.

En este trabajo se proponen dos modelos para estimar la oferta que son sencillos pero suficientemente realistas y que tienen en cuenta lo que se puede considerar que es el rasgo distintivo fundamental de la oferta en situaciones de racionamiento, a saber, que existe un umbral, que es el que interesa modelizar, que marca el límite superior de la cantidad que se puede obtener, de manera que los pedidos que superen el umbral quedarán recortados en ese valor.

La existencia de ese «efecto umbral» hace que los modelos propuestos puedan

ser estimados de forma similar a los típicos modelos de supervivencia con datos censurados. Se describe además en este trabajo cómo se puede llevar a cabo esa estimación utilizando herramientas convencionales disponibles en paquetes de software estadístico y numérico como R, y se presentan algunos ejemplos de dicha estimación que muestran cómo se puede realizar de forma rápida y eficiente.

5. REFERENCIAS BIBLIOGRÁFICAS

- Axsäter, S. (2006). *Inventory Control*. Springer.
- Box, G. E. P. & Jenkins, G. M. (1970). *Time Series Analysis, Forecasting, and Control*. Holden-Day.
- Brown, R. G. (1959). *Statistical Forecasting for Inventory Control*. McGraw-Hill.
- Brown, R. G. (1963). *Smoothing, Forecasting and Prediction of Discrete Time Series*. Prentice-Hall.
- Byrd, R. H., Lu, P., Nocedal, J. & Zhu, C. (1995). A Limited Memory Algorithm for Bound Constrained Optimization. *SIAM Journal on Scientific Computing*, 16(5), 1190-1208.
- Cox, D. R. & Oakes, D. (1984). *Analysis of Survival Data*. Chapman & Hall/CRC.
- Heizer, J. & Render, B. (2007). *Dirección de la Producción y de Operaciones* (8^a ed). Pearson Education.
- Hill, A. V. (2012). *The Encyclopedia of Operations Management*. Pearson Education.
- Holt, C. C. (2004). Forecasting seasonals and trends by exponentially weighted moving averages. *International Journal of Forecasting*, 20(1), 5-10.
- Hopp, W. J. & Spearman, M. L. (2011). *Factory Physics* (3^a ed). Waveland Press.
- Hyndman, R. J. & Athanasopoulos, G. (2014). *Forecasting. Principles and Practice*. OTEXTS.COM.

- Hyndman, R. J., Athanasopoulos, G., Bergmeir, C., Caceres, G., Chhay, L., O'Hara-Wild, M., Petropoulos, F., Razbash, S., Wang, E. & Yasmeeen, F. (2020). *forecast: Forecasting functions for time series and linear models* (R package version 8.11).
- Hyndman, R. J. & Khandakar, Y. (2008). Automatic time series forecasting: the forecast package for R. *Journal of Statistical Software*, 26(3), 1-22.
- Hyndman, R. J., Koehler, A. B., Ord, J. K. & Snyder, R. D. (2008). *Forecasting with Exponential Smoothing. The State Space Approach*. Springer.
- Hyndman, R. J., Koehler, A. B., Snyder, R. D. & Grose, S. (2002). A state space framework for automatic forecasting using exponential smoothing methods. *International Journal of Forecasting*, 18(3), 439-454.
- Kalbfleisch, J. D. & Prentice, R. L. (2002). *The Statistical Analysis of Failure Time Data* (2nd ed). Wiley-Interscience.
- Lawless, J. F. (2003). *Statistical Models and Methods for Lifetime Data* (2nd ed). Wiley-Interscience.
- Moore, D. F. (2016). *Applied Survival Analysis Using R*. Springer.
- Mullen, K., Ardia, D., Gil, D., Windover, D. & Cline, J. (2011). DEoptim: An R Package for Global Optimization by Differential Evolution. *Journal of Statistical Software*, 40(6), 1-26.
- R Core Team. (2020). *R: A Language and Environment for Statistical Computing*. R Foundation for Statistical Computing. Vienna, Austria. <https://www.R-project.org/>
- Silver, E. A., Pyke, D. F. & Thomas, D. J. (2017). *Inventory and Production Management in Supply Chains* (4^a ed). CRC Press.
- Therneau, T. M. (2020). *A Package for Survival Analysis in R* [R package version 3.1-11].

- Therneau, T. M. & Grambsch, P. M. (2000). *Modeling Survival Data: Extending the Cox Model*. Springer.
- Thomopoulos, N. T. (2015). *Demand Forecasting for Inventory Control*. Springer.
- Thomopoulos, N. T. (2016). *Elements of Manufacturing, Distribution and Logistics*. Springer.
- Winters, P. R. (1960). Forecasting Sales by Exponentially Weighted Moving Averages. *Management Science*, 6(3), 324-342.
- Xiang, Y., Gubian, S., Suomela, B. & Hoeng, J. (2013). Generalized Simulated Annealing for Efficient Global Optimization: the GenSA Package for R. *The R Journal*, 5(1).