



Universidad de Jaén

Facultad de Ciencias Sociales
y Jurídicas

Trabajo Fin de Grado

TEORÍA DE COLAS APLICADA A LA ADMINISTRACIÓN DE EMPRESAS

Alumna: ANA DEL RÍO HERVÁS

JUNIO, 2021

RESUMEN

La Teoría de Colas es el conjunto de herramientas matemáticas utilizadas para estudiar el comportamiento de las líneas de espera o colas que se producen porque, en un determinado momento, la demanda de un bien o servicio sobrepasa la capacidad que tiene el sistema para suministrar dicho bien o servicio.

El objetivo de este trabajo es realizar una introducción a la Teoría de Colas y estudiar algunos de los modelos más utilizados por las empresas. Para ello, se van a analizar las medidas de rendimiento de estado estable de esos modelos y se van a aplicar en algunos ejemplos, con el fin de aprender a interpretar tanto los datos como los resultados.

También se incluye el estudio económico de un sistema de colas que nos ayudará a hacer un balance entre el coste de ofrecer el servicio y el coste asociado a la espera de clientes para obtener ese servicio.

ABSTRACT

The Queueing Theory integrates mathematical tools that are used to study the behaviour of the waiting lines or queues that are produced because the demand for a good or service exceeds the capacity of the system to supply the good or service at a given time.

The objective of this work is to do an introduction to the Queueing Theory and to study some of the models most commonly used by companies. For this purpose, the steady-state performance measures of these models will be analyzed and applied to some examples in order to learn how to interpret both the data and the results.

In addition, this work includes the economic study of a queueing system that will help us to balance the cost of offering the service and the cost associated with customers waiting for that service.

ÍNDICE

1.- INTRODUCCIÓN.....	4
1.1.- MOTIVACIÓN.....	4
1.2.- OBJETIVOS.....	4
2.- ANTECEDENTES	5
2.1.- INTRODUCCIÓN A LA INVESTIGACIÓN OPERATIVA.....	5
2.2.- LA INVESTIGACIÓN OPERATIVA APLICADA A LA ADMINISTRACIÓN DE EMPRESAS	6
3.- METODOLOGÍA	8
3.1.- INTRODUCCIÓN A LA TEORÍA DE COLAS	8
3.2.- ELEMENTOS DE UNA LÍNEA DE ESPERA	9
3.2.1.- ESTRUCTURA BÁSICA DE LOS MODELOS DE COLAS	9
3.2.1.1.- PROCESO DE LLEGADAS	9
3.2.1.2.- LÍNEAS DE ESPERA O COLAS	11
3.2.1.3.- PROCESO DE SERVICIO	12
3.2.1.4.- PROCESO DE SALIDA.....	14
3.2.2.- NOTACIÓN DE KENDALL.....	14
3.3.- LA DISTRIBUCIÓN EXPONENCIAL Y EL PROCESO DE POISSON	15
3.4.- MODELO GENERAL DE COLAS DE POISSON	18
3.5.- MEDIDAS DE RENDIMIENTO DE ESTADO ESTABLE	21
3.6.- SISTEMAS DE COLAS	24
3.6.1.- SISTEMA DE COLAS M/M/1	24
3.6.2.- SISTEMA DE COLAS M/M/C	26
3.6.3.- SISTEMA DE COLAS M/M/1/K	29
3.6.4.- SISTEMA DE COLAS M/M/C/K	32
3.7.- ESTUDIO ECONÓMICO DE UN SISTEMA DE COLAS	35
4.- RESULTADOS.....	37
4.1.- DESCRIPCIÓN DEL SOFTWARE	37
4.2.- EJEMPLOS DE APLICACIÓN EN LA ADMINISTRACIÓN DE EMPRESAS	40
4.2.1.- EJEMPLO DEL MODELO M/M/1	40
4.2.2.- EJEMPLO DEL MODELO M/M/C.....	42
4.2.3.- EJEMPLO DEL MODELO M/M/1/K	44

4.2.4.- EJEMPLO DEL MODELO M/M/C/K.....	47
4.2.5.- COMPARACIÓN DE VARIOS MODELOS DE COLAS	49
5. CONCLUSIONES	53
6. BIBLIOGRAFÍA	54

ÍNDICE DE FIGURAS

FIGURA 1. PROCESO BÁSICO DE COLAS	9
FIGURA 2. SISTEMAS DE COLAS DE CANAL MÚLTIPLE CON UNA ÚNICA COLA	13
FIGURA 3. SISTEMAS DE COLAS DE CANAL MÚLTIPLE CON VARIAS COLAS	13
FIGURA 4. FUNCIÓN DE DENSIDAD DE PROBABILIDAD DE LA DISTRIBUCIÓN EXPONENCIAL	16
FIGURA 5. DIAGRAMA DE TRANSICIÓN EN COLAS DE POISSON	19
FIGURA 6. RELACIÓN ENTRE λ , λ_{ef} y $\lambda_{perdido}$	22
FIGURA 7. MODELO M/M/1	24
FIGURA 8. DIAGRAMA DE TRANSICIÓN DEL MODELO M/M/1.....	25
FIGURA 9. MODELO M/M/c.....	27
FIGURA 10. DIAGRAMA DE TRANSICIÓN DEL MODELO M/M/c.	28
FIGURA 11. MODELO M/M/1/K.....	29
FIGURA 12. DIAGRAMA DE TRANSICIÓN DEL MODELO M/M/1/K.....	30
FIGURA 13. MODELO M/M/C/K	33
FIGURA 14. DIAGRAMA DE TRANSICIÓN DEL MODELO M/M/C/K	34
FIGURA 15. MODELO DE DECISIÓN PARA LÍNEA DE ESPERA BASADO EN COSTO.	36
FIGURA 16. INTERFAZ DE R-PROJECT	38
FIGURA 17. INTERFAZ DE R-STUDIO	38

ÍNDICE DE TABLAS

TABLA 1. APLICACIONES REALES DE LA INVESTIGACIÓN OPERATIVA	7
TABLA 2. COMPARACIÓN DE LAS MEDIDAS DE EFECTIVIDAD.....	52
TABLA 3. COMPARACIÓN ENTRE LOS COSTES	53

1.- INTRODUCCIÓN

Este documento presenta el Trabajo de Fin de Grado (TFG) del Grado en Administración y Dirección de Empresas, cuyo título es “Teoría de Colas aplicada a la Administración de Empresas”. Este trabajo se ha desarrollado con la ayuda de la tutora Dña. María del Pilar Frías Bustamante, perteneciente al Departamento de Estadística e Investigación Operativa de la Universidad de Jaén. El proyecto consistirá en explicar y desarrollar la Teoría de Colas y en dar una solución a ciertos problemas presentes en las empresas, los cuales serán resueltos mediante un software.

A continuación, vamos a describir la motivación que nos ha llevado a realizar este proyecto y los objetivos de este.

1.1.- MOTIVACIÓN

Hoy en día, estamos acostumbrados a esperar y a introducirnos en unas largas colas para recibir un bien o un servicio. El hecho de esperar para pagar en una tienda, para entrar al médico, para cruzar un paso de peatones, entre otras muchas cosas, forma parte de nuestra vida cotidiana. Las colas se producen porque, en un determinado momento, la demanda de un bien o servicio sobrepasa la capacidad del sistema para suministrarlo. Estas colas nos hacen perder una gran parte de nuestro tiempo.

En el ámbito empresarial, el hecho de hacer esperar a los clientes puede ser bastante perjudicial para la empresa ya que genera ciertos costes derivados de la insatisfacción y pérdida de clientes. Por lo tanto, es necesario optimizar el rendimiento del sistema.

Estas reflexiones comentadas anteriormente fueron algunos de los motivos que ayudaron a elegir la Teoría de Colas como tema para este trabajo de fin de grado ya que esta teoría nos ayudará a solucionar el problema de las colas.

Además de esto, otro de los principales motivos que han llevado a escoger este tema ha sido que, durante el tercer curso del Grado, se ha estudiado la asignatura de Investigación Operativa y, en ella, no se incluyen contenidos sobre la Teoría de Colas. Por ello, se pensó que sería interesante desarrollar este tema ya que tener conocimientos sobre esto es de vital importancia para la toma de decisiones en una empresa.

1.2.- OBJETIVOS

Los objetivos de este proyecto son:

- Estudiar la Teoría de Colas, desarrollando principalmente los elementos de una línea de espera y las medidas de rendimiento de estado estable.
- Estudiar algunos de los modelos más comunes de la Teoría de Colas.
- Describir los softwares utilizados, que se trata de R-Proyect y R-Studio.
- Plantear ciertos problemas de los modelos de colas que surgen en las empresas, que serán resueltos con el software R-Studio, e interpretar sus resultados.
- Comparar varios modelos para saber cuál será el más adecuado para cierta empresa, es decir, cuál consigue optimizar el rendimiento de la empresa a un menor coste.
- Sacar conclusiones sobre la utilidad de esta teoría en la toma de decisiones de una empresa.

2.- ANTECEDENTES

2.1.- INTRODUCCIÓN A LA INVESTIGACIÓN OPERATIVA

La Investigación Operativa se ocupa de la aplicación del método científico para ayudar a la toma de decisiones óptima. Las fases de un estudio de Investigación Operativa son: la definición de un problema; la construcción del modelo, que consiste en ajustar el problema a un modelo matemático; la solución de ese modelo; la validación del modelo y, por último, su implementación.

Lawrence y Pasternak, 1998, definen la Investigación Operativa como, un enfoque científico para la toma de decisiones ejecutivas, que consiste en:

1. el arte de modelar situaciones complejas,
2. la ciencia de desarrollar técnicas de solución para resolver dichos modelos y
3. la capacidad de comunicar efectivamente los resultados.

Según la naturaleza de los datos los problemas de Investigación Operativa, se pueden clasificar en:

- Problemas determinísticos, en los cuales todos los datos son conocidos, por ejemplo, los problemas de transporte, asignación o redes.
- Problemas estocásticos, en los que los datos se consideran inciertos y pueden venir dados por una distribución de probabilidad, por ejemplo, la Teoría de Colas, Teoría de la Decisión o Modelos de Inventarios.

Atendiendo al objetivo del problema, los problemas de Investigación Operativa se pueden clasificar en:

- Modelos de optimización, cuyo objetivo es optimizar (maximizar o minimizar) cierta cantidad, por ejemplo, los problemas de asignación o de transporte.
- Modelos de predicción, cuyo objetivo es describir o predecir sucesos dadas ciertas condiciones, por ejemplo, los problemas de reemplazamiento, de inventario o de colas.

Ahora que tenemos una idea sobre lo que es la Investigación de Operaciones, veamos cuáles fueron sus inicios.

Muchos autores, como Taha, 2017, o Hillier y Lieberman, 2015, consideran que esta actividad nació en Inglaterra mientras tenía lugar la Segunda Guerra Mundial, ya que se encargó a un grupo de científicos ingleses unos servicios militares que consistían en la mejor asignación de recursos y materiales bélicos. Al acabar la guerra, muchos científicos llegaron a la conclusión de que las ideas que habían sido adoptadas en las actividades bélicas, podrían ser aplicadas en actividades diferentes a las militares. Por lo tanto, se empezó a utilizar la Investigación Operativa en organizaciones tanto del gobierno como de negocios o industriales. A partir de ahí, se han desarrollado rápidamente gracias al interés de muchos científicos a seguir investigando en esta área y al desarrollo de los computadores.

En cambio, otros autores, como Eiselt y Sandblom, 2012, creen que las bases de la Investigación de Operaciones se sentaron antes de la mano de ciertos ingenieros como, por ejemplo, Taylor y su famosa pregunta “¿cuál es la mejor forma de hacer un trabajo?”, que se puede decir que es el lema de esta actividad. Otros ejemplos serían Gantt, con sus gráficos de barras, o Agner Krarup Erlang, con la disciplina de hacer cola en 1909 cuando trabajaba en la central telefónica de Copenhague.

Nos centraremos en este último descubrimiento, ya que va a ser el tema que se va a desarrollar en este trabajo, la teoría de colas o líneas de espera.

2.2.- LA INVESTIGACIÓN OPERATIVA APLICADA A LA ADMINISTRACIÓN DE EMPRESAS

La Investigación de Operaciones es un instrumento indispensable e imprescindible en la Administración de Empresas y ha sido aplicada en empresas tan diversas como de transporte, de telecomunicaciones, de construcción, manufactureras e instituciones financieras, entre muchas otras. Así, la aplicación de la Investigación Operativa es bastante amplia.

Muchas empresas cuentan con un grupo de expertos para realizar distintos estudios de Investigación Operativa con el fin de resolver ciertos problemas que vayan surgiendo en la empresa. Este grupo de individuos debe tener experiencia y bastantes conocimientos en materias como economía, matemáticas, administración de empresas, estadística, ingeniería y otras técnicas especiales de Investigación de Operaciones.

Gracias a la Investigación de Operaciones, muchas empresas han mejorado su eficiencia, lo que ha dado lugar a un aumento de la productividad de la economía de muchos países.

La sociedad INFORMS (Institute for Operations Research and the Management Science) es una de las sociedades profesionales más importantes a nivel mundial para profesionales de Investigación de Operaciones y publica increíbles artículos sobre aplicaciones reales de esta disciplina.

En la Tabla 1 podemos observar algunas de estas aplicaciones reales. En las dos primeras columnas podemos ver la diversidad de empresas y de aplicaciones. En la tercera columna vemos el año en el que se llevaron a cabo. Con estas aplicaciones, las empresas han conseguido bastantes ahorros anuales.

En conclusión, podemos comentar que es muy importante ayudarnos de la Investigación Operativa para tomar las decisiones adecuadas en nuestra empresa para, así, conseguir que alcance el éxito.

TABLA 1. APLICACIONES REALES DE LA INVESTIGACIÓN OPERATIVA

Empresa	Aplicaciones	Año
TNT Express	Desarrollo de un programa de “Optimización global” que busca optimizar la red de transporte de la empresa mediante métodos avanzados. Resuelve problemas relacionados con las rutas óptimas de camiones, programación de personal y localización de almacenes	2012
HP	Gestión de su cartera de productos mediante la Investigación Operativa	2009
Zara	Mejora de su proceso de distribución mediante la Investigación Operativa	2009
Hoteles Marriott	Optimizar los precios	2009

Cocacola Enterprises	Programación de las rutas diarias de sus camiones	2007
-----------------------------	---	------

Fuente: Maroto, Alcaraz, Ginestar y Segura (2012)

3.- METODOLOGÍA

3.1.- INTRODUCCIÓN A LA TEORÍA DE COLAS

Como se ha comentado anteriormente, Erlang (1878-1929), un ingeniero, estadístico y matemático danés, trabajaba en la central telefónica de Copenhague y se dedicó durante un gran tiempo a estudiar el número exacto de circuitos telefónicos necesarios para proporcionar un buen servicio. Debido a esto, se considera que fue el creador de la teoría de colas. Gracias a este ingeniero y a muchos otros autores que siguieron con sus investigaciones, existe la teoría de colas, definida como el estudio de una secuencia de clientes o de otros elementos que demandan un servicio, teniendo que esperar si el servidor no se encuentra inmediatamente disponible, formándose así la línea de espera.

En general, la teoría de colas estudia *“el comportamiento de sistemas donde existe un conjunto limitado de recursos para atender las peticiones generadas por los usuarios, de tal manera que cuando un usuario envía una tarea al sistema, esta podrá tener que esperar para ser atendida por algún recurso, o, incluso, podrá ser rechazada si el sistema no tiene capacidad suficiente para almacenarla en espera de ser atendida”* (Pazos, Suárez y Díaz, 2003).

Esta teoría ha ayudado a tomar decisiones para resolver una gran variedad de problemas, ya que dedicamos una gran parte de nuestra vida a esperar. A continuación, vamos a ver algunos de estos problemas:

- Las personas esperamos para comer en un restaurante, para ser atendidos en el hospital, para pagar en una tienda, para montar en la atracción de un parque de atracciones y para sacar dinero en un cajero automático, entre otras muchas cosas.
- Los aviones esperan hasta que le dan la orden de que pueden despegar o aterrizar.
- Los vehículos esperan mientras está el semáforo en rojo o cuando hay atasco.
- Los trabajos esperan a ser procesados en una máquina.
- Las máquinas esperan para ser reparadas.

Es cierto que es muy difícil eliminar totalmente las líneas de espera de los casos anteriormente expuestos, a no ser que se incurra en unos elevados costes. Pero este sistema intenta reducir lo

máximo posible estas líneas de espera con unos costes moderados, es decir, hacer un balance entre el coste del servicio y el coste asociado a la espera por ese servicio.

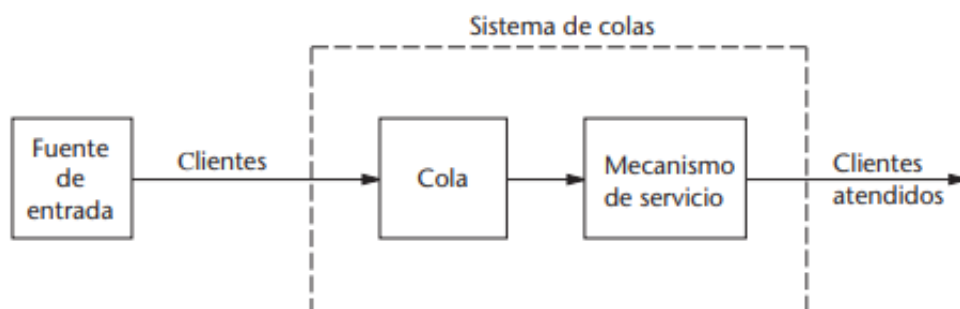
3.2.- ELEMENTOS DE UNA LÍNEA DE ESPERA

3.2.1.- ESTRUCTURA BÁSICA DE LOS MODELOS DE COLAS

En el proceso básico de la teoría de colas los actores principales son el cliente y el servidor. Este proceso se basa en que los “clientes”, cuando necesitan un servicio, se generan en el tiempo en una “fuente de entrada”. Después, llegan a la “instalación o estación de servicio”, y, si el servicio se encuentra disponible, lo recibirán inmediatamente; en cambio, si no lo está, los clientes forman una “cola” mientras esperan. Cuando el servicio esté disponible, se seleccionará a un cliente para proporcionarle el servicio mediante alguna regla conocida como “disciplina de cola”, y el servicio será proporcionado mediante un “mecanismo de servicios”. Por último, el cliente que ya ha sido atendido sale del sistema. Si la cola se queda vacía, el sistema se vuelve inactivo hasta la llegada de un nuevo cliente.

Este proceso podemos verlo y entenderlo mejor en la Figura 1. Sobre estos elementos del proceso del sistema de colas pueden hacerse muchas suposiciones.

FIGURA 1. PROCESO BÁSICO DE COLAS



Fuente: Hillier y Lieberman (2010)

A continuación, vamos a analizar los elementos comentados más detalladamente.

3.2.1.1.- PROCESO DE LLEGADAS

El proceso de llegadas está formado tanto la fuente de entrada y los clientes que se generan de ella como por el tamaño de esta y la forma en que llegan esos clientes.

El tamaño de la fuente de entrada es un aspecto muy característico. Podemos definirlo como el conjunto de todos aquellos clientes que podrían solicitar el servicio en un determinado momento, o lo que es lo mismo, el número total de clientes potenciales distintos.

Cuando hablamos de clientes, no solo nos referimos a los seres humanos, sino que nos podemos referir tanto a objetos como a tareas u otros elementos que pueden requerir un servicio en un determinado momento. Por ejemplo, en la fabricación de un pantalón, este es el que solicita el servicio de ser planchado y tiene que esperar hasta que alguna de las planchas se encuentre disponible.

Podemos suponer que el tamaño de la fuente de entrada es finita, limitando a los clientes que llegan al servicio, o infinita, como sería el caso de una central telefónica a la que pueden llegar infinitas llamadas. Esto sería lo mismo que decir que la fuente de entrada es limitada o ilimitada. Lo más común es suponer que el tamaño es infinito.

Hablaremos de cuatro características que inciden sobre la forma en que se producen las llegadas: estructura, tamaño, paciencia y distribución.

En primer lugar, refiriéndonos a la estructura de las llegadas, pueden ser “controlables” o “incontrolables”. Es cierto que las llegadas son más controlables de lo que se supone en un principio, hasta las llegadas que son incontrolables, como un servicio de urgencias en un hospital, pueden llegar a controlarse en cierto modo.

En segundo lugar, en cuanto al tamaño, nos encontramos las llegadas “único”, que son aquellas en las que los clientes acceden al sistema de forma unitaria, y “por lotes”, en las que acceden más de una persona al sistema.

En cuanto a la paciencia, se refiere a los distintos comportamientos que pueden adoptar los clientes. Destacan tres comportamientos principales: cuando los clientes se pasan de una línea a otra más corta intentando reducir la espera, lo que se denomina como “tramposo”; cuando renuncian a hacer cola porque es demasiado larga, lo que recibe el nombre de “rechazo”; y, en último lugar, el “abandono” que es cuando abandonan la cola porque llevan demasiado tiempo esperando.

Por último, el número de llegadas por unidad de tiempo o el tiempo entre dos llegadas consecutivas pueden seguir distintas distribuciones de probabilidad. Las más estudiadas son la distribución de Poisson o la distribución exponencial. Ambas distribuciones las analizaremos más adelante en el apartado 3.3.

Además de todo lo expuesto, el elemento más importante del proceso de llegadas es el tiempo entre llegadas consecutivas. Cuanto menor sea este tiempo, más cantidad de clientes llegarán por unidad de tiempo y más demanda de servidores habrá.

El tiempo entre llegadas consecutivas puede ser determinístico (fijo y conocido) o probabilístico (si el tiempo entre llegadas varía de forma aleatoria).

Por lo tanto, denominaremos T a la variable aleatoria “tiempo entre llegadas consecutivas”.

3.2.1.2.- LÍNEAS DE ESPERA O COLAS

La cola o línea de espera es el lugar donde permanecen los demandantes del servicio hasta recibirlo.

Existen, según su longitud, líneas de espera infinitas o finitas. En el primer caso, la cola tiene una gran capacidad, formada por un gran número de clientes. En cambio, en el segundo caso, hay un número limitado de clientes que pueden esperar. La mayor parte de los modelos suponen que existe una cola infinita, en parte porque resultan más fácil de calcular en la práctica. Pero hay que tener en cuenta que, en muchos casos reales, la capacidad de la cola es finita, pues llega un punto en el que esta llega a su límite y no puede admitir más clientes.

Por otro lado, en cuanto al número de colas, podemos hablar de sistemas de cola única, cuando todos los clientes esperan en la misma cola, o de colas múltiples, cuando los clientes pueden elegir entre varias filas. La desventaja de este último es que los clientes suelen cambiarse de cola frecuentemente si perciben que las demás van más rápido. Esta categoría coincidirá con el número de canales, de los que hablaremos en el siguiente apartado.

Finalmente, el último aspecto relevante en cuanto a las líneas de espera es la disciplina de cola, definida como el orden que siguen los clientes para recibir el servicio. Existen diferentes tipos de disciplinas de colas, entre los cuales destacan:

- FIFO (First in, first out), donde el primer cliente en llegar, es el primero en ser atendido y en salir. Esta es la disciplina que más se utiliza.
- LIFO (Last in, first out), donde el último cliente que llega, es al que primero atienden y el que primero sale.
- RSS (Ransom selection of service) o SIRO (Service in random order) que consiste en atender en orden aleatorio; por lo tanto, todos los clientes tienen la misma probabilidad de ser elegidos.
- PRI (Priority service), donde algunos clientes tienen cierta preferencia sobre otros y son atendidos antes que los que ya están esperando. Este sería el caso de las urgencias hospitalarias. Si hay una persona esperando para que le curen una pequeña herida y, con

posterioridad, llega una persona que ha tenido un infarto, esta última tiene prioridad en ser atendida.

3.2.1.3.- PROCESO DE SERVICIO

El proceso de servicio consiste en la manera en que los clientes o demandantes son atendidos por los diferentes servidores. Este proceso o mecanismo de servicio puede estar formado por una o más estaciones de servicio, las cuales pueden tener uno o más canales de servicio paralelos, denominados servidores.

Hay dos aspectos relevantes en este proceso, el número de clientes que pueden atenderse en una misma estación al mismo tiempo, y el tiempo de los servicios, o sea, el tiempo que transcurre desde el inicio del servicio hasta que termina, para un mismo cliente; cuanto más tiempo tarde el servidor en suministrar el servicio, más tiempo tendrán que esperar los demandantes en la cola.

Como hemos comentado en el primer párrafo de este apartado, las estaciones pueden estar compuestas por uno o varios servidores, y estos pueden atender a varios clientes. Por lo tanto, nos podemos encontrar sistemas con un “único servidor”, el cual podrá atender a los clientes de uno en uno o atender a varios a la vez. Por ejemplo, en una clínica de fisioterapia, los pacientes son atendidos de uno en uno por el servidor; pero, en una clase de ballet, los alumnos son atendidos todos a la vez por el mismo servidor. Además, un servidor no siempre es un solo individuo, podría ser un grupo de personas. Un ejemplo sería una cuadrilla de aceituneros que prestan el servicio de recoger la aceituna a un cliente. También existe el sistema “multiservidor”, que cuenta con “c” servidores iguales que atienden a “c” clientes a la vez. Y, por último, el sistema de “infinitos servidores”, donde cada cliente será atendido inmediatamente por el servidor.

En cuanto al tiempo de los servicios, puede ser de dos tipos: determinístico o probabilístico. El primero significa que el tiempo es fijo y conocido; en cambio, el segundo quiere decir que el tiempo de servicio varía de forma aleatoria. En este trabajo, consideraremos que el tiempo de servicio es probabilístico, es decir, es una variable aleatoria independiente e igualmente distribuida.

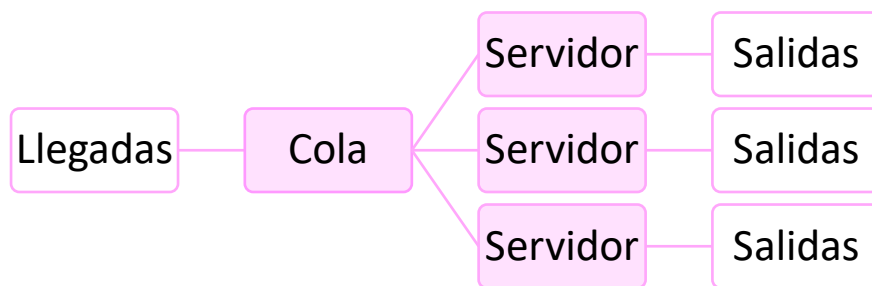
Por lo tanto, denominaremos S a la variable aleatoria “tiempo de servicio de cada cliente”.

Además de todo lo expuesto, la parte más importante del proceso de servicio son los canales ya que es el lugar del sistema donde se lleva a cabo el servicio.

Podemos encontrarnos sistemas con un único canal, en el que todos los clientes hacen la misma cola y son atendidos por el mismo servidor. Este sería el caso que vemos en la Figura 1. Un ejemplo de esto sería una frutería que cuenta con un único dependiente, entonces los clientes esperan por orden para ser atendidos por el frutero.

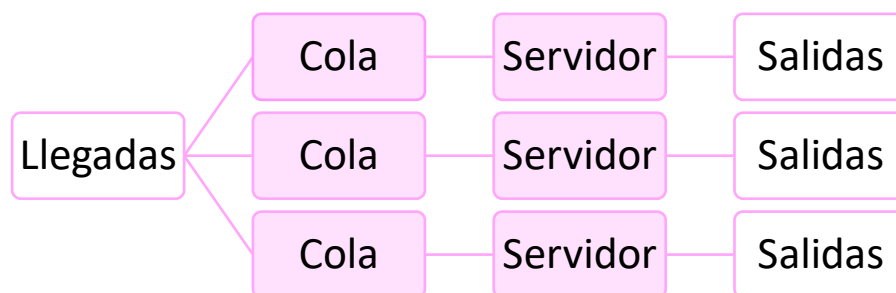
Otro caso son los sistemas en los que hay más de una canal, los llamados sistemas de colas de canal múltiple o multicanal. De estos sistemas de canal múltiple, nos encontramos dos tipos: en paralelo e independientes. Los canales de servicio en paralelo son aquellos en los que hay una única línea de espera con varios servidores. En la Figura 2 podemos verlo gráficamente. Un ejemplo sería Carrefour, donde los clientes esperan en la misma cola hasta que llega su turno y se les asigna uno de los servidores disponible. En cambio, los canales de servicio independientes tienen diferentes líneas de espera para los diferentes servidores, como vemos en la Figura 3. Este es el caso de Mercadona, donde cada cajero tiene una cola propia.

FIGURA 2. SISTEMAS DE COLAS DE CANAL MÚLTIPLE CON UNA ÚNICA COLA



Fuente: Elaboración propia

FIGURA 3. SISTEMAS DE COLAS DE CANAL MÚLTIPLE CON VARIAS COLAS



Fuente: Elaboración propia

3.2.1.4.- PROCESO DE SALIDA

Este proceso consiste en la manera en que abandonan el sistema los clientes. Podemos encontrarnos con dos tipos del proceso de salida: el sistema de colas de un paso y la red de colas.

En el sistema de colas de un paso, los clientes abandonan el sistema en cuanto han sido atendidos. Por ejemplo, en una tienda de ropa cualquiera donde el único servicio que recibes es el de ser atendido en la caja para pagar la ropa, una vez que hayas pagado, abandonas el sistema. Pero, en la red de colas, los clientes pueden, una vez adquirido el servicio en una estación, pasar a otra para recibir otro servicio diferente. Por ejemplo, si hablamos de un establecimiento como el Corte Inglés, en el que, además de vender ropa, vende comida, vende electrodomésticos y tiene restaurantes, entre muchas más cosas, el cliente puede pasar del servicio de comprar ropa a otro de los servicios existentes.

3.2.2.- NOTACIÓN DE KENDALL

Esta notación que vamos a ver, denominada así porque su creador fue David Kendall, se ha desarrollado para describir y escribir de forma resumida los modelos de colas y sus características, que son el proceso de llegada de los clientes, el proceso de servicio, la disciplina de cola, la capacidad del sistema y el número de canales. Por lo tanto, estos modelos se etiquetan de la forma siguiente:

$$A/B/c/K/Z$$

Donde:

- A es la distribución del tiempo entre llegadas.
- B es la distribución del tiempo de servicio.
- c es el número de servidores o canales. También suele denominarse con la letra s . Este símbolo tendrá un valor entre 0 y ∞ .
- K es la capacidad del sistema, es decir, el número máximo de clientes que puede soportar el sistema, y tendrá un valor entre 0 y ∞ , que será la suma de aquellos clientes que están siendo atendidos (" c ") y aquellos que esperan en la cola (" $k-c$ ").
- Z es la disciplina de cola, que puede ser FIFO, LIFO, SIRO o RSS, o PRI, como vimos en el apartado 3.2.1.2.

Tanto A como B pueden utilizar otros símbolos para indicar la distribución que siguen. Estas son las distribuciones que utilizaremos:

- M para indicar que siguen una distribución exponencial y los tiempos son aleatorios.
- D para indicar que siguen una distribución determinística y los tiempos son constantes.
- E_k para indicar que siguen una distribución Erlang de parámetro k y los tiempos son aleatorios.
- G para indicar que siguen una distribución general, es decir, una distribución distinta a las anteriores, y los tiempos son aleatorios.

A menudo, se utiliza la notación Kendall de forma abreviada, que sería:

$$A/B/c$$

Esta notación abreviada supone que la capacidad del sistema es infinita y que la disciplina seguida es FIFO (primero en entrar, primero en salir).

Por ejemplo, una tienda que cuenta con 5 dependientes, con tiempos exponenciales tanto entre llegadas como de servicio, capacidad del sistema infinita y donde el primero en llegar es el primero en ser atendido, se denotaría $M/M/5/\infty/FIFO$, que de forma abreviada quedaría $M/M/5$.

3.3.- LA DISTRIBUCIÓN EXPONENCIAL Y EL PROCESO DE POISSON

En la mayoría de los casos de colas, los clientes van llegando de forma aleatoria. Esto quiere decir que el hecho de que ocurra un evento, como que ha llegado un cliente o que ha terminado un servicio, no está influido por el tiempo que haya transcurrido desde que ocurrió el último evento.

Los tiempos aleatorios entre llegadas (T) y de servicio (S) de los sistemas de colas más estudiados se suelen representar mediante una distribución exponencial, o lo que es lo mismo, la razón de llegadas y de servicios tienen una distribución de Poisson. Estas distribuciones que se utilizan en la teoría de colas son complementarias, cuyo objetivo es la modelización de los tiempos entre llegadas y tiempos de servicio. Ambas son las distribuciones de probabilidad más importantes en esta teoría ya que son realistas y sencillas de manejar matemáticamente.

Por un lado, la variable aleatoria T posee una distribución exponencial si su función de densidad es:

$$f(t) = \lambda e^{-t\lambda}, \quad \lambda > 0, t > 0$$

Su función de distribución es:

$$F(t) = P[T \leq t] = 1 - e^{-t\lambda}$$

Y la esperanza (tiempo medio entre llegadas) es:

$$E[T] = 1/\lambda$$

donde λ es la media de llegadas en el tiempo t o razón de llegada.

Por otro lado, la variable aleatoria S posee una distribución exponencial si su función de densidad es:

$$f(s) = \mu e^{-s\mu}, \quad \mu > 0, s > 0$$

Su función de distribución es:

$$F(s) = P[S \leq s] = 1 - e^{-s\mu}$$

Y la esperanza (tiempo medio de servicio) es:

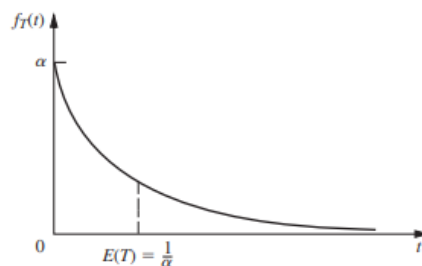
$$E[S] = 1/\mu$$

donde μ es el número medio de clientes atendidos por unidad de tiempo o razón de servicio.

A continuación, vamos a ver una serie de propiedades que hacen a la distribución exponencial adecuada para la modelización de los procesos de llegadas y de servicio en la teoría de colas:

Propiedad 1. La distribución exponencial es una función estrictamente decreciente. Relativamente hablando, es posible que T tome valores más cercanos a cero que a su media. En la figura 4 podemos ver representado lo que dice esta propiedad.

FIGURA 4. FUNCIÓN DE DENSIDAD DE PROBABILIDAD DE LA DISTRIBUCIÓN EXPONENCIAL



Fuente: Hillier y Lieberman (2010)

Propiedad 2. Esta es la propiedad del olvido o falta de memoria, lo que significa, que la distribución exponencial no considera lo que ha sucedido en las t unidades de tiempo anteriores al nuevo evento (llegada o fin del servicio). Se puede decir que el sistema “olvida” su pasado. Esta propiedad, la podemos expresar matemáticamente como:

$$P [T \leq t + h \mid T \geq t] = P [0 \leq T \leq h], h > 0$$

Según esto, como T es el tiempo entre llegadas consecutivas, si en t unidades de tiempo no se producen nuevas llegadas, la probabilidad de que un nuevo cliente llegue en las siguientes h unidades de tiempo es independiente del tiempo que haya pasado sin producirse nuevas llegadas, lo que significa que el proceso olvida que en t unidades de tiempo no se hayan producido llegadas y las siguientes llegadas dependen de h solamente.

Lo mismo pasaría con S , que es el tiempo de servicio de cada cliente. La probabilidad de que un servicio acabe en las siguientes h unidades de tiempo, es independiente del tiempo que lleve el cliente en el servicio.

Por ejemplo, si ahora son las 9:10 de la mañana y la última llegada fue a las 8:50, la probabilidad de que la siguiente llegada ocurra a las 9:15 es una función del intervalo de 9:10 a 9:15, y es totalmente independiente del tiempo transcurrido desde la ocurrencia del último evento, es decir, de 8:50 a 9:10.

Propiedad 3. El mínimo de variables aleatorias exponenciales independientes tiene una distribución exponencial. En concreto, sean $T_1, T_2, \dots, T_n$ variables aleatorias independientes que tienen distribuciones exponenciales con parámetros $\lambda_1, \lambda_2, \dots, \lambda_n$, respectivamente, la variable aleatoria $U = \min \{T_1, T_2, \dots, T_n\}$ sigue una distribución exponencial con parámetro $\lambda = \sum_{i=1}^n \lambda_i$.

Propiedad 4. Esta propiedad relaciona la distribución exponencial con la de Poisson.

La distribución de Poisson pretende modelizar la cantidad de hechos que se producen en un determinado intervalo de tiempo, de manera aleatoria y continua. Esta distribución cumple una serie de características:

1. La distribución de probabilidad del número de hechos es la misma cuando se trata de períodos de misma duración.
2. La probabilidad de que dos hechos se produzcan al mismo tiempo en un intervalo de tiempo pequeño es mínima. Por ello, se considera que las llegadas son individuales.
3. Si el tiempo entre eventos sucesivos sigue una distribución exponencial, el número de eventos en un determinado intervalo de tiempo sigue una distribución Poisson. Debido

a esto, en la teoría de colas se habla de tiempos entre llegadas o de servicios exponenciales o poissonianos indistintamente.

Por lo tanto, la probabilidad de que lleguen n clientes en un intervalo de tiempo t es:

$$P [X(t) = n] = \frac{(\lambda \cdot t)^n}{n!} \cdot e^{-\lambda t}$$

Siendo los tiempos entre dos llegadas consecutivas variables aleatorias independientes e igualmente distribuidas con distribución exponencial de media $1/\lambda$ y siguiendo la cantidad de llegadas una distribución de Poisson con razón de llegada λ .

De forma similar, la probabilidad de prestar n servicios en un intervalo de tiempo t es:

$$P [X(t) = n] = \frac{(\mu \cdot t)^n}{n!} \cdot e^{-\mu t}$$

Siendo los tiempos entre servicios variables aleatorias independientes e igualmente distribuidas con distribución exponencial de media $1/\mu$ y siguiendo la cantidad de servicios prestados una distribución de Poisson con razón de llegada μ .

3.4.- MODELO GENERAL DE COLAS DE POISSON

En este apartado vamos a analizar un modelo general de colas en el que se supone que las llegadas y las salidas de clientes se basan en un proceso de Poisson, y eso quiere decir que el tiempo entre dos llegadas consecutivas y entre dos servicios sigue una distribución de probabilidad exponencial.

Los sistemas de colas pueden alcanzar dos comportamientos, a largo plazo, también denominado estado estable o régimen estacionario, o a corto plazo, llamado estado transitorio. El estado transitorio se produce cuando al inicio del funcionamiento del sistema de colas, el estado del sistema se encuentra bastante afectado por el estado inicial y el tiempo que ha transcurrido desde el inicio. Una vez que ha pasado cierto tiempo desde el comienzo del sistema, el estado del sistema se vuelve totalmente independiente del estado inicial y del tiempo que ha pasado desde el inicio; por lo tanto, el sistema ha alcanzado su estado estable.

Este modelo supone que el comportamiento de la cola es a largo plazo o de estado estable, ya que el comportamiento a corto plazo o de estado transitorio es más complicado.

Además, este modelo supone que las tasas de llegada y de salida dependen del estado del sistema, es decir, del número de clientes en las instalaciones de servicio. Por ejemplo, en una

autopista de peaje, los empleados encargados de cobrar, cobrarán más rápido durante las horas pico.

Se definirá lo siguiente:

n = número de clientes en el sistema, tanto en la cola como en el servicio.

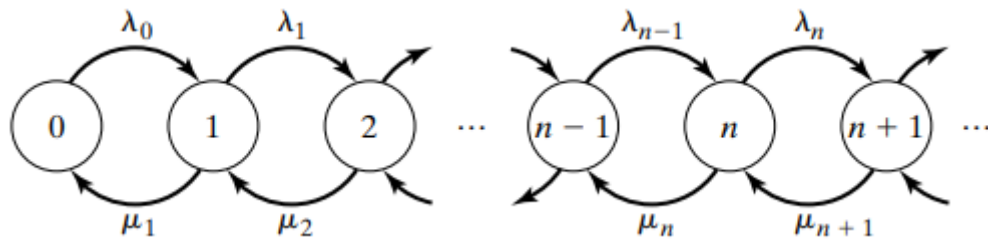
λ_n = Tasa de llegada dados n clientes en el sistema.

μ_n = Tasa de salida dados n clientes en el sistema.

p_n = Probabilidad de estado estable de que haya n clientes en el sistema.

La probabilidad p_n se determina usando el diagrama de transición, que se muestra en la Figura 5. Como vemos, “el sistema de colas está en el estado n cuando la cantidad de clientes en él es n . Para $n > 0$, el estado n solamente puede cambiar a dos estados: $n - 1$ cuando hay una salida con frecuencia μ_n y $n + 1$ cuando hay una llegada con frecuencia λ_n . El estado 0 solo puede cambiar al estado 1 cuando hay una llegada con la frecuencia λ_0 . Pero μ_0 no está definida porque no pueden haber salidas si el sistema está vacío” (Taha, 2004).

FIGURA 5. DIAGRAMA DE TRANSICIÓN EN COLAS DE POISSON



Fuente: Taha (2004)

Bajo las condiciones de estado estable, para $n > 0$, las tasas esperadas de flujo de entrada y de salida del estado deberán ser iguales. Por lo tanto, obtenemos que:

$$\text{Tasa esperada de flujo de entrada al estado } n = \lambda_{n-1} p_{n-1} + \mu_{n+1} p_{n+1}$$

$$\text{Tasa esperada de flujo de salida del estado } n = (\lambda_n + \mu_n) p_n$$

Al igualarlas, obtenemos una ecuación de balance:

$$\lambda_{n-1} p_{n-1} + \mu_{n+1} p_{n+1} = (\lambda_n + \mu_n) p_n$$

Cuando $n = 0$, la ecuación de balance sería:

$$\lambda_0 p_0 = \mu_1 p_1$$

La cual se resuelve, como todas las ecuaciones de balance, en función de p_0 :

$$p_1 = \left(\frac{\lambda_0}{\mu_1} \right) p_0$$

Cuando $n = 1$, las dos anteriores ecuaciones serían:

$$\lambda_0 p_0 + \mu_2 p_2 = (\lambda_1 + \mu_1) p_1 \rightarrow p_2 = \left(\frac{\lambda_1 \lambda_0}{\mu_2 \mu_1} \right) p_0$$

En general, podemos demostrar que:

$$p_n = \left(\frac{\lambda_{n-1} \lambda_{n-2} \dots \lambda_0}{\mu_n \mu_{n-1} \dots \mu_1} \right) p_0, \quad n = 1, 2, \dots$$

El valor de p_0 se suele obtener de la ecuación $\sum_{n=0}^{\infty} p_n = 1$.

Todo esto lo veremos aplicado en el siguiente ejemplo.

Ejemplo 1

En cierto lugar de “Carga y Descarga”, el estacionamiento para comerciantes está limitado a solo tres espacios. Los vehículos que ocupan estos espacios llegan siguiendo una distribución de Poisson, con una razón de llegada de cuatro vehículos por hora. El tiempo de estacionamiento se distribuye de manera exponencial con una media de 20 minutos. Los comerciantes que llegan y se encuentran todos los aparcamientos llenos, tienen que esperar en un espacio temporal hasta que un vehículo estacionado salga, el cual tiene una capacidad para solo dos autos. Todos los demás vehículos que no pueden estacionar en la zona de “Carga y Descarga” ni en el espacio temporal, deberán irse a otro lugar.

Sabiendo esto, debemos determinar la probabilidad de que n vehículos estén en el sistema.

En primer lugar, debemos tener claro que los aparcamientos serían los servidores y el espacio temporal en el que los vehículos esperan sería la cola. Por lo tanto, como el primero es 3 y el segundo es 2, la capacidad del sistema es de 5.

$\lambda_n = 4$ autos por hora, $n = 0, 1, 2, 3, 4, 5$.

$$\mu_n = \begin{cases} n \left(\frac{60}{20} \right) = 3n & \text{autos por hora, } n = 1, 2, 3. \\ 3 \left(\frac{60}{20} \right) = 9 & \text{autos por hora, } n = 4, 5. \end{cases}$$

En segundo y último lugar, vamos a hacer los cálculos necesarios para obtener la solución a nuestro problema.

Utilizando la ecuación de p_n definida anteriormente, obtenemos que:

$$p_1 = \left(\frac{\lambda_0}{\mu_1} \right) p_0 = \frac{4}{3 \cdot 1} p_0 = \frac{4^1}{1!3^1} p_0$$

$$p_2 = \left(\frac{\lambda_1 \lambda_0}{\mu_2 \mu_1} \right) p_0 = \frac{4 \cdot 4}{3 \cdot 2 \cdot 3 \cdot 1} p_0 = \frac{4^2}{2!3^2} p_0$$

$$p_3 = \left(\frac{\lambda_2 \lambda_1 \lambda_0}{\mu_3 \mu_2 \mu_1} \right) p_0 = \frac{4 \cdot 4 \cdot 4}{3 \cdot 3 \cdot 3 \cdot 2 \cdot 3 \cdot 1} p_0 = \frac{4^3}{3!3^3} p_0$$

$$p_4 = \left(\frac{\lambda_3 \lambda_2 \lambda_1 \lambda_0}{\mu_4 \mu_3 \mu_2 \mu_1} \right) p_0 = \frac{4 \cdot 4 \cdot 4 \cdot 4}{9 \cdot 3!3^3} p_0 = \frac{4^4}{3!3^5} p_0$$

$$p_5 = \left(\frac{\lambda_4 \lambda_3 \lambda_2 \lambda_1 \lambda_0}{\mu_5 \mu_4 \mu_3 \mu_2 \mu_1} \right) p_0 = \frac{4 \cdot 4 \cdot 4 \cdot 4 \cdot 4}{9 \cdot 3!3^5} p_0 = \frac{4^5}{3!3^7} p_0$$

El valor de p_0 lo vamos a determinar con la ecuación siguiente:

$$p_0 + p_0 \left\{ \frac{4^1}{1!3^1} + \frac{4^2}{2!3^2} + \frac{4^3}{3!3^3} + \frac{4^4}{3!3^5} + \frac{4^5}{3!3^7} \right\} = 1$$

De la cual obtenemos que:

$$p_0 \{3,87003826\} = 1 \rightarrow p_0 = 0,2583$$

Por lo tanto, teniendo el valor de p_0 podemos calcular las probabilidades anteriores:

$$p_1 = 0,3444$$

Existe una probabilidad del 34,44% de que haya un vehículo en el estacionamiento.

$$p_2 = 0,2296$$

Existe una probabilidad del 22,96% de que haya dos vehículos en el estacionamiento.

$$p_3 = 0,1020$$

Existe un 10,20% de probabilidad de que haya tres vehículos en el estacionamiento.

$$p_4 = 0,0453$$

Existe un 4,53% de probabilidad de que haya tres vehículos en el estacionamiento y otro vehículo esperando.

$$p_5 = 0,0201$$

Existe un 2,01% de probabilidad de que haya tres vehículos en el estacionamiento y dos vehículos esperando.

3.5.- MEDIDAS DE RENDIMIENTO DE ESTADO ESTABLE

Frecuentemente, el estudio del sistema de colas se centra en determinar una serie de características de este sistema, denominadas medidas de rendimiento, eficiencia o desempeño de una cola, las cuales cuantifican el rendimiento del sistema. Estas son:

L = Número medio de clientes en el sistema.

L_q = Número medio de clientes en la cola.

W = Tiempo medio de espera en el sistema.

W_q = Tiempo medio de espera en la cola.

$\hat{c}$ = Número medio de servidores ocupados.

ρ = Intensidad de tráfico en el sistema, que mide la relación entre la cantidad de clientes que llegan al sistema y la capacidad de atenderlos.

Cuando hablamos de sistema, este engloba a la cola y a la estación de servicio.

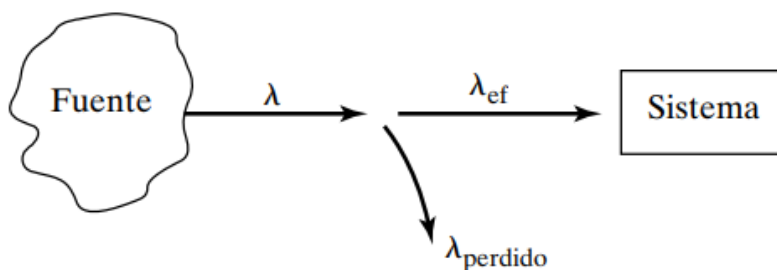
Bajo condiciones generales, la relación entre L y W , como entre L_q y W_q , es conocida como fórmula de Little, y está dada por:

$$L = \lambda_{ef} W$$

$$L_q = \lambda_{ef} W_q$$

Donde λ_{ef} es la tasa de llegada efectiva o razón de entrada al sistema. Será igual que λ (tasa de llegada) si todos los clientes que llegan entran al sistema. De otra forma, si algunos clientes no pueden entrar al sistema porque este está lleno, entonces $\lambda_{ef} < \lambda$ y $\lambda = \lambda_{ef} + \lambda_{perdido}$ (ver, Figura 6).

FIGURA 6. RELACIÓN ENTRE λ , λ_{ef} y $\lambda_{perdido}$.



Fuente: Taha (2004)

También existe una relación entre W y W_q , ya que el tiempo medio de espera en sistema (W) es igual a la suma del tiempo medio de espera en la cola (W_q) y el tiempo esperado de servicio ($1/\mu$). Esta fórmula quedaría como:

$$W = W_q + \frac{1}{\mu}$$

De forma similar, podemos relacionar L y L_q multiplicando ambos lados de la anterior ecuación por λ_{ef} y, junto con la fórmula de Little, obtenemos:

$$L = L_q + \frac{\lambda_{ef}}{\mu}$$

Por último, la diferencia entre el número medio de clientes en el sistema (L) y el número medio de clientes en la cola (L_q), es igual al número medio de servidores ocupados ($\hat{c}$). Esto quedaría como:

$$\hat{c} = L - L_q = \frac{\lambda_{ef}}{\mu}$$

Por lo tanto, la utilización de la estación de servicio será igual al número medio de servidores ocupados entre el número de servidores o canales ($\hat{c} / c$).

Ejemplo 2

Continuando con el ejemplo 1, vamos a determinar en este segundo ejemplo las medidas de rendimiento de estado estable que acabamos de definir.

1. Calcular la tasa efectiva de llegadas.

Como vimos en el anterior ejemplo, un vehículo no podrá acceder al estacionamiento si ya se encuentran 5 coches en él, tres estacionados y dos esperando. Por ello, la probabilidad de coches que no pueden entrar al estacionamiento es p_5 . Entonces, tenemos:

$\lambda_{perdido} = \lambda p_5 = 4 \cdot 0,0201 = 0,0804$ vehículos por hora abandonan el estacionamiento porque no pueden ni aparcar ni esperar.

$\lambda_{ef} = \lambda - \lambda_{perdido} = 4 - 0,0804 = 3,9196$ vehículos por hora pueden entrar al estacionamiento.

2. Calcular el número medio de vehículos en el estacionamiento, tanto los que están estacionados como los que esperan.

$L = 0 p_0 + 1 p_1 + \dots + 5 p_5 = 0 \cdot 0,2583 + 1 \cdot 0,3444 + 2 \cdot 0,2296 + 3 \cdot 0,1020 + 4 \cdot 0,0453 + 5 \cdot 0,0201 = 1,3913$ vehículos hay en el estacionamiento.

$L_q = L - \frac{\lambda_{ef}}{\mu} = 1,3913 - \frac{3,9196}{3} = 0,0848$ vehículos esperan en el espacio temporal.

3. Calcular el tiempo medio de espera de los vehículos en el sistema y en la cola.

$W = \frac{L}{\lambda_{ef}} = \frac{1,3913}{3,9196} = 0,35496$ horas es el tiempo medio de los vehículos en el estacionamiento.

$W_q = W - \frac{1}{\mu} = 0,35496 - \frac{1}{3} = 0,02163$ horas es el tiempo medio que esperan los vehículos en el espacio temporal.

4. Calcular el número medio de aparcamientos ocupados.

$\hat{c} = L - L_q = \frac{\lambda_{ef}}{\mu} = \frac{3,9196}{3} = 1,3065$ aparcamientos ocupados.

3.6.- SISTEMAS DE COLAS

El objetivo de este apartado es conocer algunos modelos del sistema de colas. Por lo tanto, aunque sea interesante conocer el origen de las fórmulas, nos centraremos principalmente en su resultado final. La deducción de estas fórmulas puede obtenerse en cualquier libro de teoría de colas, como los expuestos más adelante en la bibliografía.

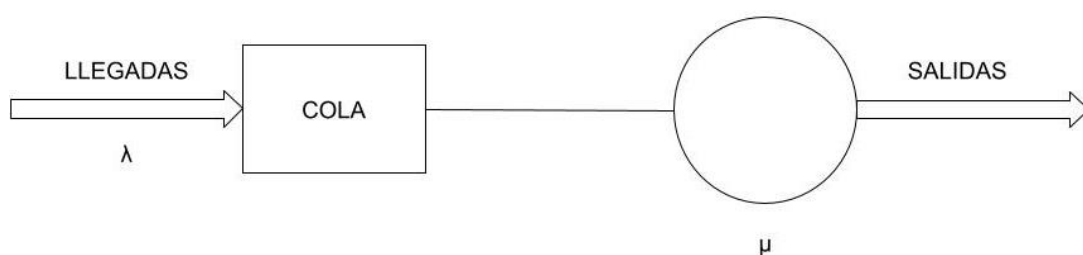
3.6.1.- SISTEMA DE COLAS M/M/1

El modelo M/M/1 es la abreviatura del modelo M/M/1/ ∞ /FIFO y se caracteriza por:

- Está formado por un único servidor.
- No existe límite en el tamaño del sistema. Por ello, la tasa de llegada efectiva es igual a la tasa de llegada ($\lambda_{ef} = \lambda$) y la tasa de pérdida de clientes es 0 ($\lambda_{per} = 0$).
- El orden que siguen los clientes para entrar al sistema, es decir, la disciplina de cola seguida es que el primero que entra es el primero que sale, como dice el método FIFO.
- Los tiempos aleatorios entre llegadas y de servicio siguen una distribución exponencial, es decir, los clientes llegan y son atendidos siguiendo una distribución de Poisson.

En la Figura 7 podemos ver un ejemplo visual de cómo es este modelo.

FIGURA 7. MODELO M/M/1



Fuente: Elaboración propia

Como definimos en apartados anteriores, λ hace referencia a la razón de llegada o a la media de llegadas en el tiempo t , donde $1/\lambda$ es el tiempo medio entre llegadas consecutivas, y μ hace referencia a la razón de servicio o a la media de clientes atendidos en el tiempo t , donde $1/\mu$ es el tiempo medio esperado de servicio.

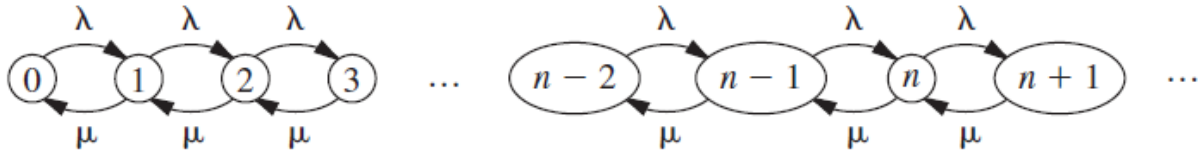
Tanto la tasa de llegada como la salida de salida son constantes, esto es:

$$\lambda = \lambda_n \quad n = 0, 1, 2, \dots$$

$$\mu = \mu_n \quad n = 0, 1, 2, \dots$$

En la Figura 8 podemos ver el diagrama de transición o de tasas de este modelo.

FIGURA 8. DIAGRAMA DE TRANSICIÓN DEL MODELO M/M/1.



Fuente: Hillier y Lieberman (2010)

Teniendo en cuenta el apartado 3.5. y las fórmulas de estado estable definidas en él, podemos obtener la distribución estacionaria de este modelo, y es:

$$p_n = \left(\frac{\lambda}{\mu}\right)^n \cdot p_0, \quad n \geq 1$$

A partir de esta, podemos deducir que la probabilidad de que el sistema esté vacío (p_0) es:

$$p_0 = 1 - \rho$$

donde la intensidad de tráfico en el sistema (ρ) es igual a λ/μ , coincidiendo en este caso con la probabilidad de que el sistema esté ocupado ($1 - \rho$), y, para que se den las condiciones de estado estable y para que el sistema no se sature, debe ser menor que 1. Por lo tanto, la tasa de llegadas debe ser menor que la tasa de servicio.

Por lo tanto, la distribución estacionaria resultante de ρ es:

$$p_n = (1 - \rho) \rho^n, \quad \text{para } n = 0, 1, 2, \dots$$

A continuación, vamos a describir las ecuaciones clave de este modelo.

1. Número de clientes en el sistema

$$L = \frac{\rho}{1 - \rho}$$

2. Tiempo medio de espera en el sistema

Como veíamos anteriormente, la fórmula de Little nos muestra que $L = \lambda_{ef}W$, y como en este modelo tenemos que $\lambda_{ef} = \lambda$, obtenemos que $W = L/\lambda$. Sustituyendo en esta última fórmula, tenemos que:

$$W = \frac{1}{\mu (1 - \rho)}$$

También podemos definir $W(t)$, que es la probabilidad de que un cliente permanezca más de t unidades de tiempo en el sistema, cuya fórmula sería:

$$W(t) = e^{-t/w}$$

3. Tiempo medio de espera en la cola

Como la relación entre W y W_q es $W = W_q + \frac{1}{\mu}$, obtenemos que:

$$W_q = \frac{\rho}{\mu (1 - \rho)}$$

También podemos definir $W_q(t)$, que es la probabilidad de que un cliente permanezca más de t unidades de tiempo en la línea de espera, cuya fórmula sería:

$$W_q(t) = \rho e^{-t/w}$$

4. Número de clientes en la cola

Utilizando tanto la fórmula de Little como la fórmula de L , obtenemos que:

$$L_q = \frac{\rho^2}{1 - \rho}$$

3.6.2.- SISTEMA DE COLAS M/M/C

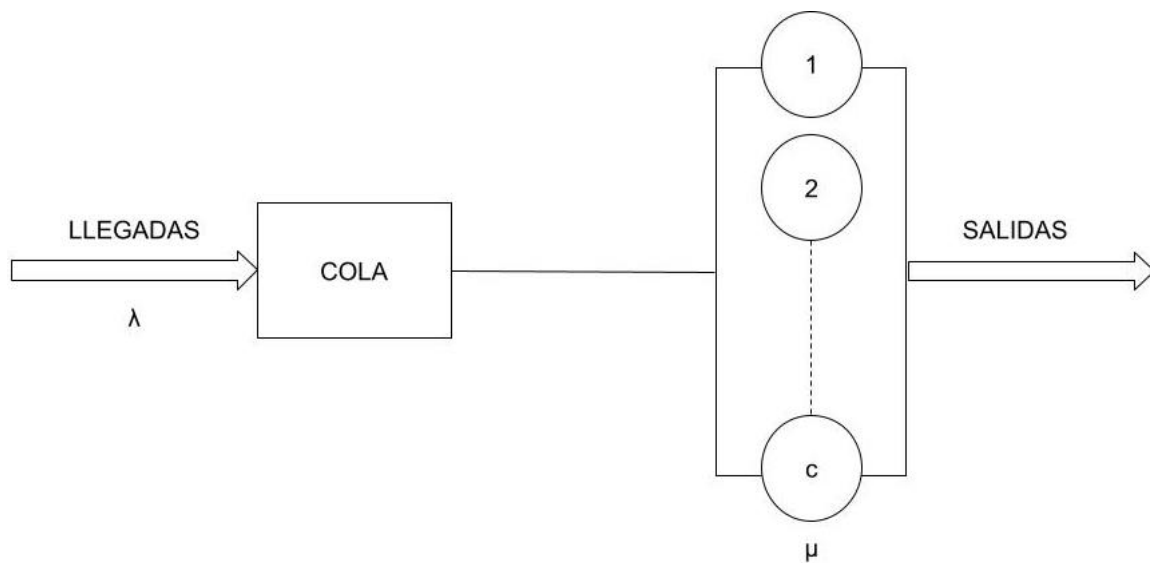
El modelo M/M/C es la abreviatura del modelo M/M/c/∞/FIFO y se caracteriza por:

- Está formado por c servidores, los cuales son idénticos y paralelos, lo que significa que realizan la misma función con el mismo desempeño o eficiencia. Por ello, los tiempos de servicio son variables aleatorias e igualmente distribuidas con distribución exponencial.
- No existe límite en el tamaño del sistema. Por ello, la tasa de llegada efectiva es igual a la tasa de llegada ($\lambda_{ef} = \lambda$) y la tasa de pérdida de clientes es 0 ($\lambda_{per} = 0$).

- El orden que siguen los clientes para entrar al sistema, es decir, la disciplina de cola seguida es que el primero que entra es el primero que sale, como dice el método FIFO.
- Por un lado, los tiempos aleatorios entre llegadas son exponenciales, es decir, los clientes llegan siguiendo una distribución de Poisson. Por otro lado, cada uno de los servidores tiene una distribución exponencial independiente e idéntica de servicio.

En la Figura 9 podemos ver un ejemplo visual de cómo es este modelo.

FIGURA 9. MODELO M/M/c



Fuente: Elaboración propia

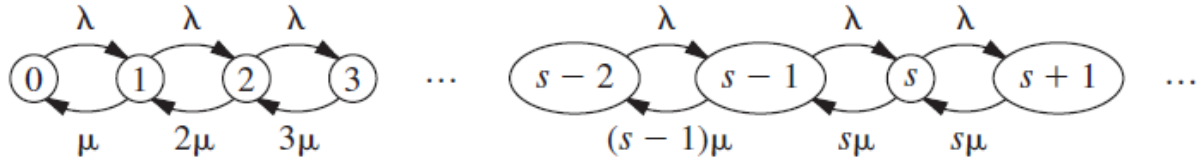
Al igual que en el modelo anterior, usaremos λ para referirnos a la razón de llegada o a la media de llegadas en el tiempo t , donde $1/\lambda$ es el tiempo medio entre llegadas consecutivas, y usaremos μ para referirnos a la razón de servicio o a la media de clientes atendidos en el tiempo t , donde $1/\mu$ es el tiempo medio esperado de servicio.

La tasa de llegada sigue siendo constante como en el modelo anterior. En cambio, la tasa de salida variará según el número de clientes que haya en el sistema. Si hay n clientes en el sistema y n es menor que c , la tasa de salida será $n \cdot \mu$, ya que hay c servidores, suficientes para servir a los n clientes; pero, si n es mayor que c , hay más clientes que servidores, entonces la tasa de salida será $c \cdot \mu$. Por tanto:

$\lambda = \lambda_n \quad n = 0, 1, 2, \dots$	$\mu_n = \begin{cases} n \cdot \mu & n = 1, \dots, c-1 \\ c \cdot \mu & n = c, c+1, \dots \end{cases}$
--	--

En la Figura 10 podemos ver el diagrama de transición o de tasas de este modelo.

FIGURA 10. DIAGRAMA DE TRANSICIÓN DEL MODELO M/M/c.



Fuente: Hillier y Lieberman (2010)

Usando las fórmulas de estado estable definidas en el apartado 3.5. y las tasas de llegada y de salida anteriores, podemos obtener la distribución estacionaria de este modelo, y es:

$$p_n = \begin{cases} p_0 \frac{r^n}{n!}, & \text{si } 1 \leq n \leq c \\ p_0 \frac{r^n}{c!c^{n-c}}, & \text{si } n > c \end{cases}$$

Siendo la probabilidad de que el sistema esté vacío (p_0):

$$p_0 = \left[\sum_{n=0}^{c-1} \frac{r^n}{n!} + \frac{r^c}{c! \left(1 - \frac{r}{c}\right)} \right]^{-1}$$

donde $r = \frac{\lambda}{\mu}$. Además, para que se den las condiciones de estado estable, ρ debe ser menor que

1 y, en este modelo, $\rho = \frac{\lambda}{c\mu}$. Por lo tanto, la tasa de llegadas debe ser menor que el producto de

la tasa de servicio y el número de servidores para que el sistema no se sature.

A continuación, vamos a describir las ecuaciones clave de este modelo.

1. Número de clientes en el sistema

$$L = \frac{\lambda}{\mu} + L_q$$

2. Tiempo medio de espera en el sistema

$$W = W_q + \frac{1}{\mu}$$

3. Tiempo medio de espera en la cola

$$W_q = \frac{L_q}{\lambda} = p_0 \frac{\left(\frac{\lambda}{\mu}\right)^s \mu}{(s-1)!(s\mu - \lambda)^2}$$

4. Número de clientes en la cola

$$L_q = p_0 \frac{\left(\frac{\lambda}{\mu}\right)^s \lambda \mu}{(s-1)!(s\mu - \lambda)^2}$$

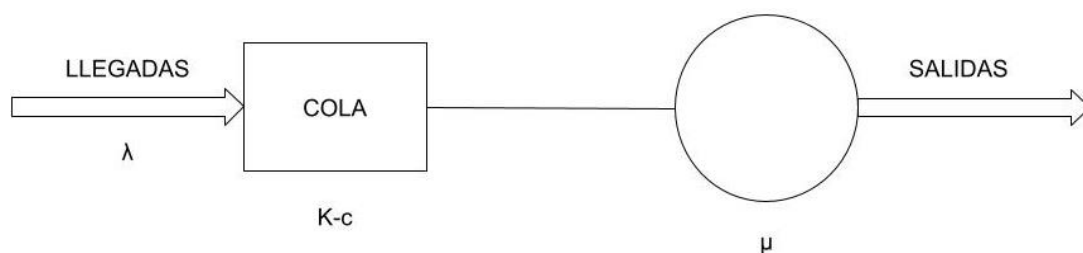
3.6.3.- SISTEMA DE COLAS M/M/1/K

El modelo M/M/1/K es la abreviatura del modelo M/M/1/K/FIFO y se caracteriza por:

- Está formado por un único servidor.
- La capacidad del sistema es finita. Esto quiere decir que el número de clientes en el sistema no puede superar un límite, que en este caso es K; por lo tanto, la capacidad de la cola es “K-c”. Si se da el caso de que llegue un cliente cuando la cola está llena, este no podrá entrar al sistema y lo dejará para siempre.
- El orden que siguen los clientes para entrar al sistema, es decir, la disciplina de cola seguida es que el primero que entra es el primero que sale, como dice el método FIFO.
- Los tiempos aleatorios entre llegadas y de servicio siguen una distribución exponencial, es decir, los clientes llegan y son atendidos siguiendo una distribución de Poisson.

En la Figura 11 podemos ver un ejemplo visual de cómo es este modelo.

FIGURA 11. MODELO M/M/1/K



Fuente: Elaboración propia

Al igual que en los modelos anteriores, usaremos λ para referirnos a la razón de llegada o a la media de llegadas en el tiempo t , donde $1/\lambda$ es el tiempo medio entre llegadas consecutivas, y usaremos μ para referirnos a la razón de servicio o a la media de clientes atendidos en el tiempo t , donde $1/\mu$ es el tiempo medio esperado de servicio.

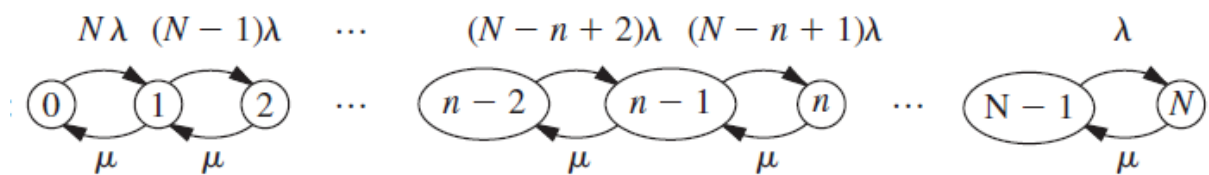
En este caso, la tasa de salida es constante. En cambio, la tasa de llegada variará según el número de clientes que haya en el sistema. Si hay menos de K clientes en el sistema, la tasa de entrada será λ ; pero, si hay más de K clientes, entonces la tasa de entrada será 0, porque se supera la capacidad del sistema. Por tanto:

$$\mu = \mu_n \quad n = 0, 1, 2, 3, \dots$$

$$\lambda_n = \begin{cases} \lambda & n = 0, 1, \dots, K-1 \\ 0 & n = K, K+1, \dots \end{cases}$$

En la Figura 12 podemos ver el diagrama de transición o de tasas de este modelo.

FIGURA 12. DIAGRAMA DE TRANSICIÓN DEL MODELO M/M/1/K



Fuente: Hillier y Lieberman (2010)

Usando las fórmulas de estado estable definidas en el apartado 3.5., las tasas de llegada y de salida anteriores y sabiendo que $\rho = \frac{\lambda}{\mu}$, podemos obtener la distribución estacionaria de este modelo, y es:

$$p_n = \rho^n \cdot p_0, \quad n = 0, \dots, K$$

Donde p_0 es:

$$p_0 = \begin{cases} \frac{1}{K+1}, & \rho = 1 \\ \frac{1-\rho}{1-\rho^{K+1}}, & \rho \neq 1 \end{cases}$$

Por lo tanto, si sustituimos p_0 en la fórmula de la distribución estacionaria, obtenemos:

$$p_n = \begin{cases} \frac{1}{K+1}, & \rho = 1 \\ \frac{(1-\rho)\rho^n}{1-\rho^{K+1}}, & \rho \neq 1 \end{cases} \quad n = 0, 1, \dots, K$$

Podemos decir que la distribución estacionaria existe siempre en este modelo ya que, como la capacidad del sistema está limitada, la cola no puede crecer indefinidamente y el sistema se estabiliza.

A continuación, vamos a describir las ecuaciones clave de este modelo.

1. Tasa de llegada efectiva o razón de entrada al sistema

Como ya hemos comentado, este modelo tiene una capacidad del sistema limitada. Por ello, la tasa de llegada es distinta a la tasa de llegada efectiva o razón de entrada. Mientras que la tasa de llegada está compuesta por todos los clientes que llegan al sistema, la tasa de llegada efectiva está compuesta solo por aquellos clientes que consiguen entrar al sistema. Por lo tanto, aunque la tasa de llegada siga una distribución de Poisson, la tasa de llegada efectiva no la sigue.

La probabilidad de que el sistema no esté completo sería $1 - p_K$, y es igual a la probabilidad de que un cliente pueda acceder al sistema. Sabiendo esto, podemos decir que la tasa de llegada efectiva es:

$$\lambda_{ef} = \lambda \cdot (1 - p_K)$$

2. Tasa de pérdidas

Esta tasa hace referencia al número medio de clientes que llegan al sistema en un periodo determinado de tiempo, pero que no pueden entrar porque el sistema se encuentra lleno. Es decir, sería el resultado de restar la tasa de llegada efectiva a la tasa de llegada. Esta tasa se calcula como:

$$\lambda_{perdido} = \lambda - \lambda_{ef} = \lambda p_K$$

3. Número de clientes en el sistema

$$L = \begin{cases} \frac{K}{2}, & \rho = 1 \\ \frac{\rho \{1 - (K + 1) \rho^K + K \rho^{K+1}\}}{(1 - \rho)(1 - \rho^{K+1})}, & \rho \neq 1 \end{cases}$$

4. Tiempo medio de espera en el sistema

A partir de la fórmula de Little, podemos obtener que:

$$W = \frac{L}{\lambda_a}$$

5. Tiempo medio de espera en la cola

$$W_q = W - \frac{1}{\mu}$$

6. Número de clientes en la cola

$$L_q = W_q \cdot \lambda_a = L - c \cdot \rho$$

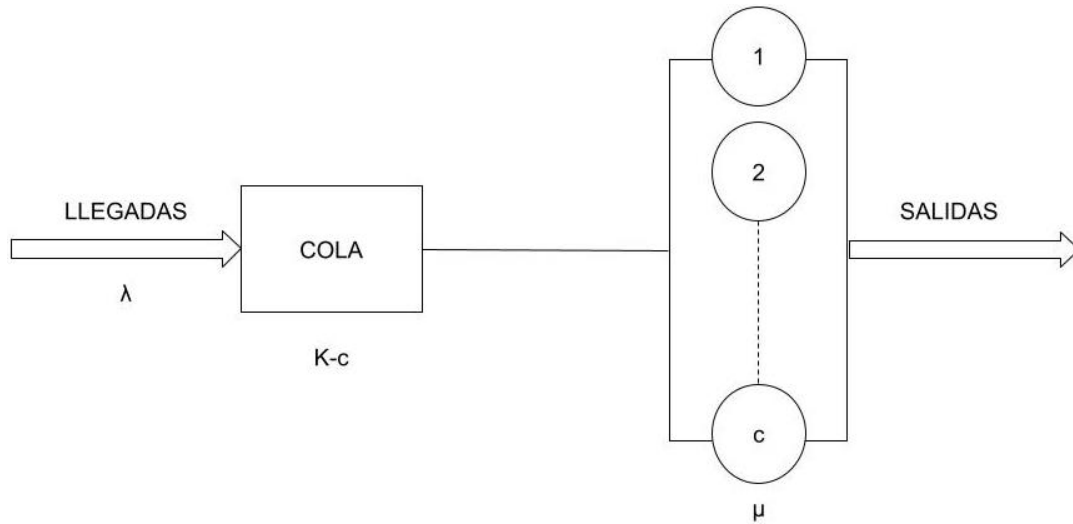
3.6.4.- SISTEMA DE COLAS M/M/C/K

El modelo M/M/C/K es la abreviatura del modelo M/M/C/K/FIFO y se caracteriza por:

- Está formado por c servidores, los cuales son idénticos y paralelos, lo que significa que realizan la misma función con el mismo desempeño o eficiencia. Por ello, los tiempos de servicio son variables aleatorias e igualmente distribuidas con distribución exponencial.
- Tiene una cola finita. Esto quiere decir que el número de clientes en el sistema no puede superar un límite, que en este caso es K; por lo tanto, la capacidad de la cola es “K-c”. Si se da el caso de que llegue un cliente cuando la cola está llena, este no podrá entrar al sistema y lo dejará para siempre.
- El orden que siguen los clientes para entrar al sistema, es decir, la disciplina de cola seguida es que el primero que entra es el primero que sale, como dice el método FIFO.
- Modelo con llegadas Poissonianas.

En la Figura 13 podemos ver un ejemplo visual de cómo es este modelo.

FIGURA 13. MODELO M/M/C/K



Fuente: Elaboración propia

En este caso, ni la tasa de llegada ni la tasa de salida son constantes, ambas variarán según el número de clientes que haya en el sistema.

Por una lado, la tasa de llegada será λ , si hay menos de K clientes en el sistema; pero, si hay más de K clientes, entonces la tasa de entrada será 0, porque se supera la capacidad del sistema. Por otro lado, la tasa de salida será $n \cdot \mu$ si hay n clientes en el sistema y n es menor que c , ya que hay c servidores, suficientes para servir a los n clientes; pero, si n se encuentra entre c y K , hay más clientes que servidores sin superar la capacidad del sistema, entonces la tasa de salida será $c \cdot \mu$.

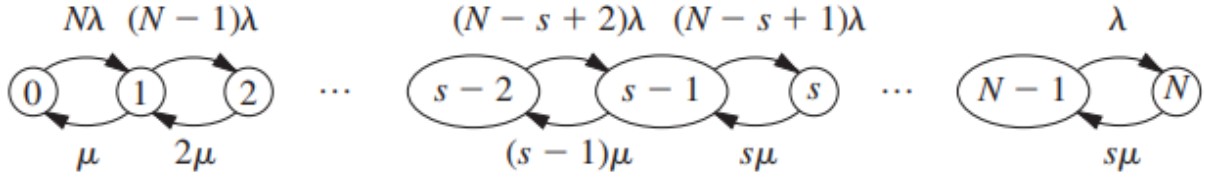
Por lo tanto, las fórmulas de estas tasas serían:

$$\lambda_n = \begin{cases} \lambda & 0 \leq n < K \\ 0 & n \geq K \end{cases}$$

$$\mu_n = \begin{cases} n \cdot \mu & 0 \leq n < c \\ c \cdot \mu & c \leq n \leq K \end{cases}$$

En la Figura 14 podemos ver el diagrama de transición o de tasas de este modelo.

FIGURA 14. DIAGRAMA DE TRANSICIÓN DEL MODELO M/M/C/K



Fuente: Hillier y Lieberman (2010)

Usando las fórmulas de estado estable definidas en el apartado 3.5., las tasas de llegada y de salida anteriores y sabiendo que $r = \frac{\lambda}{\mu}$, podemos obtener la distribución estacionaria de este modelo, y es:

$$p_n = \begin{cases} \frac{r^n}{n!} p_0, & 0 \leq n < c \\ \frac{r^n}{c! r^{n-c}} p_0, & c \leq n \leq K \end{cases}$$

Donde p_0 sería:

$$P_0 = \begin{cases} \left[\sum_{n=0}^{c-1} \frac{r^n}{n!} + \frac{r^c (1 - (\frac{r}{c})^{K-c+1})}{c! (1 - \frac{r}{c})} \right]^{-1}, & \frac{r}{c} \neq 1 \\ \left[\sum_{n=0}^{c-1} \frac{r^n}{n!} + \frac{r^c}{c!} (K - c + 1) \right]^{-1}, & \frac{r}{c} = 1 \end{cases}$$

A continuación, vamos a describir las ecuaciones clave de este modelo.

1. Tasa de llegada efectiva o razón de entrada al sistema

Como el modelo anterior, este modelo también tiene una capacidad del sistema limitada. Así pues, la tasa de llegada es distinta a la tasa de llegada efectiva o razón de entrada, y esta última no sigue una distribución de Poisson.

Por lo tanto, la tasa de llegada efectiva es:

$$\lambda_{ef} = \lambda \cdot (1 - p_K)$$

Donde λ es la tasa de llegada y $1 - p_k$ es la probabilidad de que un cliente pueda acceder al sistema.

2. Tasa de pérdidas

El número medio de clientes que llegan al sistema en un periodo determinado de tiempo, pero que no pueden entrar porque el sistema se encuentra lleno, se calcularía como:

$$\lambda_{perdido} = \lambda - \lambda_{ef} = \lambda p_K$$

3. Número de clientes en el sistema

$$L_q = \begin{cases} \frac{r^{c+1}}{(c-1)!(c-r)^2} \left[1 - \left(\frac{r}{c}\right)^{K-c+1} - (K-c+1) \left(1 - \frac{r}{c}\right) \left(\frac{r}{c}\right)^{K-c} \right] p_0, & \frac{r}{c} \neq 1 \\ \frac{r^c (K-c)(K-c+1)}{2c!} p_0, & \frac{r}{c} = 1 \end{cases}$$

4. Tiempo medio de espera en el sistema

Despejando de la fórmula de Little, tenemos que:

$$W = \frac{L}{\lambda_a}$$

5. Tiempo medio de espera en la cola

$$W_q = W - \frac{1}{\mu}$$

6. Número de clientes en la cola

$$L = \lambda_a \cdot W = \lambda_a \left(W_q + \frac{1}{\mu} \right) = L_q + \frac{\lambda_a}{\mu} = L_q + c \cdot \rho$$

3.7.- ESTUDIO ECONÓMICO DE UN SISTEMA DE COLAS

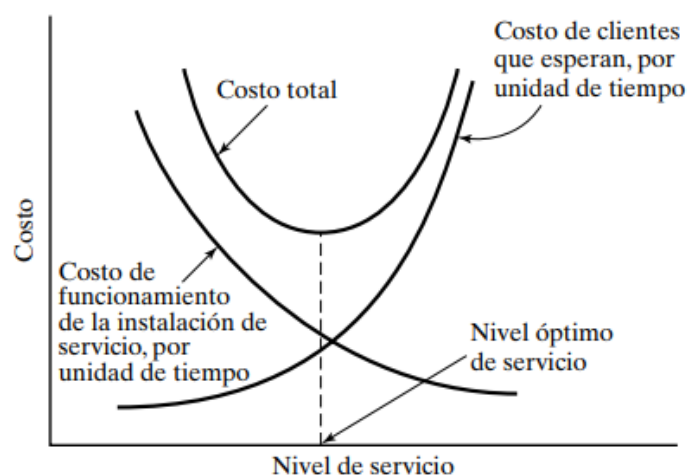
La teoría de colas puede utilizarse tanto para calcular las medidas de rendimiento como para la optimización. Por ejemplo, el dueño de una tienda tiene que decidir cuántos empleados contratar. Evidentemente, al aumentar el número de empleados aumenta el coste. Pero, a su

vez, supondrá menos tiempo de espera para los clientes, lo que puede traducirse en la reducción del coste asociado al abandono de clientes.

Así pues, podemos decir que uno de los principales objetivos de los sistemas de colas, como veíamos en la introducción, es hacer un balance entre dos costos opuestos: el coste de ofrecer el servicio y el coste asociado al tiempo de espera del cliente por ese servicio, es decir, el coste de retraso en la oferta del servicio.

Estos dos costes son contrarios el uno del otro ya que, como podemos observar en la Figura 15, al aumentar uno disminuye el otro de forma automática, y viceversa.

FIGURA 15. MODELO DE DECISIÓN PARA LÍNEA DE ESPERA BASADO EN COSTO.



Fuente: Taha (2004)

Entonces, la fórmula del modelo de costos sería:

$$\boxed{\begin{array}{c} \text{Coste total} \\ \text{esperado por} \\ \text{unidad de} \\ \text{tiempo} \end{array}} = \boxed{\begin{array}{c} \text{Coste del} \\ \text{servicio} \end{array}} + \boxed{\begin{array}{c} \text{Coste asociado} \\ \text{a la espera de} \\ \text{los clientes} \end{array}}$$

Por un lado, el coste de servicio está formado por costes como las nóminas, los seguros, los costes de las instalaciones de servicios, etc. Por lo tanto, engloba tanto al coste asociado con el tiempo que el servidor esté ocupado como al coste asociado al tiempo que el servidor esté vacío. Sabiendo esto, podemos decir que los costes de servicio serán mayores cuantos más servidores haya en el sistema.

La forma más sencilla de este tipo de coste es la siguiente función lineal:

$$\text{Coste de servicio } (x) = C_1 x$$

Donde C_1 es el coste por unidad de x por unidad de tiempo.

Por otro lado, el coste asociado a la espera de los clientes puede referirse a la pérdida de clientes debido a la larga espera, al tiempo perdido o malgastado, a la gasolina malgastada, etc. Por lo tanto, engloba tanto al coste asociado al tiempo que esperan los clientes en la cola como al coste asociado al tiempo que pasan los clientes en la estación de servicio.

Sabiendo esto, podemos decir que los costes asociados a la espera de los clientes serán menores cuantos más servidores haya en el sistema, debido a que habrá un mejor servicio.

La forma más sencilla de este tipo de coste es la siguiente función lineal:

$$\text{Coste asociado a la espera de los clientes}(x) = C_2 L$$

Donde C_2 es el coste de la espera por unidad de tiempo por cada cliente que espera.

4.- RESULTADOS

4.1.- DESCRIPCIÓN DEL SOFTWARE

En este trabajo vamos a utilizar el software *R-Project* y el software *R-Studio*.

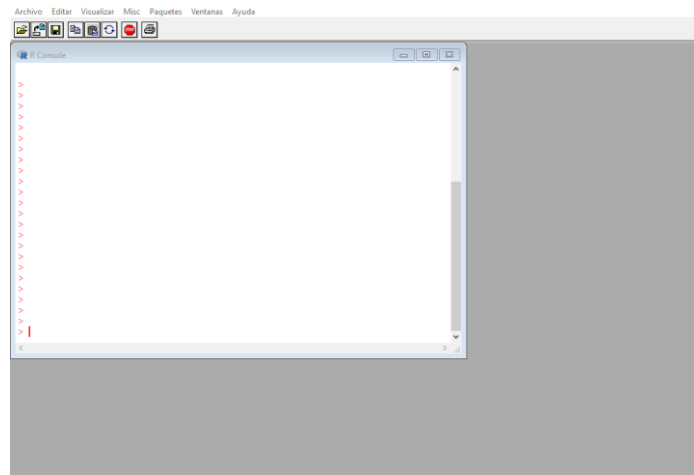
Por un lado, *R* es un lenguaje de programación, de alto nivel, que se utiliza para realizar tareas estadísticas, tanto las más básicas como más complicadas. Está formado por una serie de programas para manipular datos, realizar cálculos y hacer gráficos. Además, es un software gratuito y fácil de descargar e instalar. Algunas de sus características son:

- Puede almacenar y manipular datos.
- Contiene una gran cantidad de herramientas para analizar datos, coherentes e integradas.
- Posibilidad de realizar gráficos.
- Dispone de un lenguaje de programación simple, efectivo y bien desarrollado.

R puede parecer bastante complejo a priori, pero, en realidad, es fácil de utilizar gracias a su simplicidad y su flexibilidad.

En la Figura 16 podemos ver el interfaz de este software.

FIGURA 16. INTERFAZ DE R-PROJECT¹



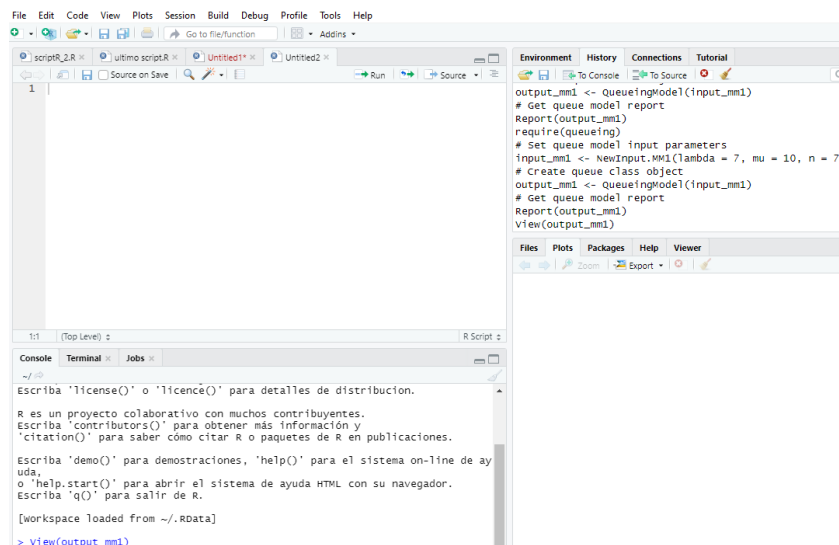
(2021)

Por otro lado, *R-Studio* es un entorno de desarrollo integrado para *R*. También es gratuito y fácil de descargar e instalar. Entre otras cosas, *R-Studio* nos permite:

- Abrir, a la vez, varios scripts.
- Ver el historial.
- Ver la consola donde aparecen los resultados, a la vez que vemos el script.

En la Figura 17 podemos ver cómo es el programa y sus cuatro panales de trabajo.

FIGURA 17. INTERFAZ DE R-STUDIO



(2021)

¹ Tanto la Figura 16 como la Figura 17 se han obtenido de los propios programas (*R-Project* y *R-Studio*), tal y como son en la versión descargada en abril de 2021.

En el siguiente apartado vamos a ver cinco problemas: para el modelo M/M/1, para el modelo M/M/C, para el modelo M/M/1/K, para el modelo M/M/C/K y un último problema en el que se compararán varios modelos.

Estos problemas los vamos a resolver con el programa *R-Studio*, que funciona de forma similar al programa *R-Project*. Para ello, vamos a utilizar el paquete “queueing”.

Una vez instalado el paquete, se introducirán los parámetros de entrada, los cuales se introducirán de forma diferente, dependiendo del modelo.

- Para el modelo M/M/1 se utiliza la función “NewInput.MM1”.
- Para el modelo M/M/C se utiliza la función “NewInput.MMC”.
- Para el modelo M/M/1/K se utiliza la función “NewInput.MM1K”.
- Para el modelo M/M/C/K se utiliza la función “NewInput.MMCK”.

Los componentes de la función “NewInput.MM1” son:

- λ : tasa de llegada
- μ : tasa de salida
- n: número de clientes en el sistema

Los componentes de la función “NewInput.MMC” son:

- λ : tasa de llegada
- μ : tasa de salidas
- n: número de clientes en el sistema
- c: número de servidores

Los componentes de la función “NewInput.MM1K” son:

- λ : tasa de llegada
- μ : tasa de salidas
- k: capacidad del sistema

Los componentes de la función “NewInput.MMCK” son:

- λ : tasa de llegada
- μ : tasa de salidas
- c: número de servidores
- k: capacidad del sistema

En el ejemplo del modelo M/M/1, se deberán introducir los datos con el siguiente código:

```
Input_mm1 <- NewInput.MM1 (lambda=7, mu=10, n=7)
```

Para construir el modelo de este sistema de colas y calcular las medidas de efectividad, utilizaremos la función “Queueingmodel”, que se introducirá de la siguiente manera:

Output_mm1 <- Queueingmodel (Input_mm1)

El inicio de estos dos códigos (Input_mm1 u Output_mm1) es el nombre que nosotros queramos darle.

Para obtener todas las medidas de efectividad de forma detallada, introduciremos:

Report(Output_mm1)

Y si lo que queremos es obtener estas medidas de una forma resumida, introduciremos:

Summary(Output_mm1)

Por último, para el quinto ejemplo, que realizaremos una comparación de varios sistemas de colas, deberemos utilizar el código:

“CompareQueueingModels(modelo1,modelo2,modelo3)”

Donde modelo1, modelo2 y modelo3 son los nombres que le hemos ido dando a cada uno de los modelos usados en el ejemplo.

4.2.- EJEMPLOS DE APLICACIÓN EN LA ADMINISTRACIÓN DE EMPRESAS

4.2.1.- EJEMPLO DEL MODELO M/M/1

Una empresa dedicada a la fabricación de cajas de cartón, realiza un control de calidad para asegurarse de que la caja final cumple con todos los requisitos necesarios. Aquellas cajas que tengan defectos, pasarán a una estación de trabajo donde serán analizadas y reparadas. Se sabe que las cajas defectuosas llegan a la estación de trabajo siguiendo una distribución de Poisson, con una media de 4 cajas por hora, y que el tiempo que se tarda en reparar cada caja sigue una distribución exponencial, con una media de 12 minutos por caja.

Para este enunciado, vamos a calcular las medidas de efectividad.

Se trata de un modelo M/M/1 con parámetros $\lambda = 4$ cajas por hora y $\mu = 60/12 = 5$ cajas por hora.

Para que se den las condiciones de estado estable, ρ debe ser menor que 1 y, en este modelo,

$\rho = \frac{\lambda}{\mu}$. Por lo tanto, se cumplen las condiciones de estado estable ya que la tasa de llegadas es menor que la tasa de salidas.

FORMULACIÓN Y RESULTADOS OBTENIDOS

En el siguiente cuadro se muestra el código *R* (o script de *R*) utilizado para la resolución del ejemplo.

```
require(queueing)
# Set queue model input parameters
input_mm1 <- NewInput.MM1 (lambda = 4, mu = 5, n = 5)
# Create queue class object
output_mm1 <- QueueingModel(input_mm1)
# Get queue model report
Report(output_mm1)
The inputs of the M/M/1 model are:
Lambda: 4 mu: 5, n: 5
The outputs of the M/M/1 model are:
The probability (p0, p1, ..., pn) of the n = 5 clients in the system
are: 0.2 0.16 0.128 0.1024 0.08192 0.065536
The traffic intensity is: 0.8
The server use is: 0.8
The mean number of clients in the system is: 4
The mean number of clients in the queue is: 3.2
The mean number of clients in the server is: 0.8
The mean time spend in the system is: 1
The mean time spend in the queue is: 0.8
The mean time spend in the server is: 0.2
The mean time spend in the queue when there is queue is: 1
The throughput is: 4
```

INTERPRETACIÓN DE LOS RESULTADOS

La probabilidad de que la estación de trabajo dedicada a reparar las cajas defectuosas esté vacía (p_0) es del 20%. Sabiendo esto, podemos decir que la probabilidad de que una caja llegue y sea reparada inmediatamente es del 20%.

La probabilidad de que haya una caja es del 16%, la probabilidad de que haya dos cajas es del 12,8%, la probabilidad de que haya tres cajas es del 10,24% y, así, sucesivamente. Por lo tanto, podemos decir que hay menor probabilidad de que haya muchas cajas para ser reparadas en esta estación de trabajo.

El número medio de cajas en la estación de trabajo (L) es de 4 cajas, permaneciendo en término medio cada caja en el sistema (W) una hora.

El número medio de cajas en la cola (L_q), esperando a ser reparadas, es de 3,2 cajas. Cada caja esperará en término medio en la cola (W_q) unos 48 minutos ($0,8 \cdot 60 = 48$).

El número medio de cajas siendo reparadas 0,2, tardando en arreglarse cada una unos 12 minutos ($0,2 \cdot 60 = 12$) en término medio.

En consecuencia de lo anterior, podemos decir que las cajas permanecen en término medio más tiempo esperando que siendo reparadas.

Tanto el uso del servidor ($\frac{\lambda}{c\mu}$) como la intensidad de tráfico en el sistema ($\rho = \frac{\lambda}{\mu}$) es de 0,8. Por ello, como hemos comentado antes, se dan las condiciones de estado estable.

Por último, la productividad de la estación de trabajo es 4.

Como conclusión, podemos decir que, como las cajas permanecen más tiempo en la cola que siendo reparadas, quizá la empresa debería ser más cuidadosa en el proceso de fabricación para que lleguen menos cajas a la estación de reparación.

4.2.2.- EJEMPLO DEL MODELO M/M/C

El proceso productivo de una empresa dedicada a la fabricación y venta de relojes cuenta con una fase en la que se comprueba si los relojes son resistentes al agua. Para ello, se dispone de cinco máquinas encargadas de realizar esta comprobación. Los relojes llegan a esta fase siguiendo una distribución de Poisson, con una media de 50 relojes por hora. El tiempo que se tarda en comprobar cada reloj sigue una distribución exponencial, con una media de 5 minutos por reloj.

Para este enunciado, vamos a calcular las medidas de efectividad.

Como vemos, es un modelo M/M/5 con parámetros $\lambda = 50$ relojes por hora y $\mu = 60/5 = 12$ relojes por hora.

Para que se den las condiciones de estado estable, ρ debe ser menor que 1 y, en este modelo,

$\rho = \frac{\lambda}{c\mu}$. Por lo tanto, se cumplen las condiciones de estado estable ya que $\rho = \frac{50}{5 \cdot 12} = 0,83$.

FORMULACIÓN Y RESULTADOS OBTENIDOS

En el siguiente cuadro se muestra el código *R* (o script de *R*) utilizado para la resolución del ejemplo.

```
require(queueing)
# Set queue model input parameters
input_mmc <- NewInput.MMC(lambda = 50, mu = 12, c = 5, n = 5, method
= 0)
# Create queue class object
output_mmc <- QueueingModel(input_mmc)
# Get queue model report
Report(output_mmc)
The inputs of the model M/M/C are:
Lambda: 50, mu: 12, c: 5, n: 5, method: Exact
The outputs of the model M/M/C are:
The probability (p0, p1, ..., pn) of the n = 5 clients in the system
are:  0.009875998  0.04114999  0.08572915  0.1190683  0.1240294
0.1033579
The traffic intensity is: 4.166666666666667
The server use is: 0.833333333333333
The mean number of clients in the system is: 7.26740254024984
The mean number of clients in the queue is: 3.10073587358317
The mean number of clients in the server is: 4.166666666666667
The mean time spend in the system is: 0.145348050804997
The mean time spend in the queue is: 0.0620147174716634
The mean time spend in the server is: 0.083333333333334
The mean time spend in the queue when there is queue is: 0.1
The throughput is: 50
```

INTERPRETACIÓN DE LOS RESULTADOS

La probabilidad de que esta estación de trabajo esté vacía (p_0), estando las máquinas desocupadas, es del 0,98%. Sabiendo esto, podemos decir que la probabilidad de que un reloj llegue y se realice la comprobación inmediatamente es bastante baja.

La probabilidad de que haya un cliente es del 4,11%, la probabilidad de que haya dos clientes es del 8,57%, la probabilidad de que haya tres clientes es del 11,9% y, así, sucesivamente. Observando los datos, podemos decir que la probabilidad va aumentando a medida que aumenta

el número de relojes hasta llegar a 4 relojes; a partir de ahí, la probabilidad empieza a disminuir conforme aumentan los clientes. Por ello, lo más probable es que haya 4 relojes en la estación de trabajo.

El número medio de relojes a los que hay que realizar la comprobación de si son resistentes al agua (L) es de 7,26 clientes, permaneciendo en término medio cada reloj en esta estación de trabajo (W) casi 9 minutos ($1,26249 \cdot 60 = 8,718$).

El número medio de relojes en la cola esperando a ser comprobados (L_q) es de 3,1. Cada reloj permanecerá en término medio en la cola (W_q) casi 4 minutos ($0,083 \cdot 60 = 3,72$).

El número medio de relojes siendo comprobados por las máquinas es de 4,16, tardando en ser comprobada la resistencia al agua de cada reloj unos 5 minutos ($0,083 \cdot 60 = 4,98$) en término medio.

En consecuencia de lo anterior, podemos decir que los relojes permanecen en término medio casi el mismo tiempo en la esperando que siendo comprobados por las máquinas.

En este problema, el uso del servidor ($\frac{\lambda}{c\mu}$) es de 0,83. En cambio, la intensidad de tráfico en el sistema ($\frac{\lambda}{\mu}$) es de 4,167.

Por último, la productividad de las máquinas es 50.

4.2.3.- EJEMPLO DEL MODELO M/M/1/K

Una clínica dental privada, en la que solo hay un dentista, ofrece a sus pacientes revisiones gratis un día a la semana. Los pacientes van llegando siguiendo una distribución de Poisson, con una media de 8 personas por hora. Pero, en la sala de espera, solo hay 6 asientos, por lo que los clientes se van si no encuentran ningún sitio libre. El tiempo que tarda el dentista en realizar cada revisión sigue una distribución exponencial, con una media de 15 minutos por revisión.

Para este enunciado, vamos a calcular las medidas de efectividad.

Como vemos, es un modelo M/M/1/7 con parámetros $\lambda = 8$ pacientes por hora y $\mu = 60/15 = 4$ revisiones por hora.

Podemos decir que la condición de estado estable se da siempre en este modelo ya que, como la capacidad del sistema está limitada, la cola no puede crecer indefinidamente y el sistema se estabiliza.

FORMULACIÓN Y RESULTADOS OBTENIDOS

En el siguiente cuadro se muestra el código *R* (o script de *R*) utilizado para la resolución del ejemplo.

```
require(queueing)
# Set queue model input parameters
input_mm1k <- NewInput.MM1K(lambda = 8, mu = 4, k = 7)
# Create queue class object
output_mm1k <- QueueingModel(input_mm1k)
# Get queue model report
Report(output_mm1k)
The inputs of the model M/M/1/K are:
Lambda: 8, mu: 4, k: 7
The outputs of the model M/M/1/K are:
The probability (p0, p1, ..., pk) of the clients in the system are:
0.003921569 0.007843137 0.01568627 0.03137255 0.0627451 0.1254902
0.2509804 0.5019608
The probability (q0, q1, ..., qk-1) that a client that enters meets
n clients in the system are: 0.007874016 0.01574803 0.03149606
0.06299213 0.1259843 0.2519685 0.503937
The traffic intensity is: 0.996078431372549
The server use is: 0.996078431372549
The mean number of clients in the system is: 6.03137254901961
The mean number of clients in the queue is: 5.03529411764706
The mean number of clients in the server is: 0.996078431372549
The mean time spend in the system is: 1.51377952755906
The mean time spend in the queue is: 1.26377952755906
The mean time spend in the server is: 0.25
The mean time spend in the queue when there is queue is:
1.27380952380952
The throughput is: 3.9843137254902
```

INTERPRETACIÓN DE LOS RESULTADOS

La probabilidad de que la clínica esté vacía (p_0) es del 3,92%. Sabiendo esto, podemos decir que la probabilidad de que un paciente llegue y se le preste el servicio inmediatamente es bastante baja.

La probabilidad de que haya una persona es del 0,78%, la probabilidad de que haya dos personas es del 1,56%, la probabilidad de que haya tres personas es del 3,13% y, así, sucesivamente. Observando los datos, podemos decir que la probabilidad va aumentando a medida que aumentan los pacientes. Por ello, es más probable que en la clínica haya 6 pacientes esperando para realizarse la revisión, es decir, es más probable que el sistema esté lleno.

El número medio de personas que acuden a la clínica (L) es de 6,031 personas, permaneciendo en término medio cada persona en el sistema (W) una hora y media aproximadamente ($1,513 \cdot 60 = 90,78$).

El número medio de personas esperando en la sala de espera de la clínica (L_q) es de 5,035. Cada persona permanecerá en término medio en la cola (W_q) una hora y cuarto ($1,263 \cdot 60 = 75,78$).

El número medio de pacientes atendidos por el dentista es de 0,99, tardando cada persona en realizarse la revisión unos 15 minutos ($0,25 \cdot 60 = 15$) en término medio.

En consecuencia de lo anterior, podemos decir que los personas pasan en término medio más tiempo en la cola que siendo revisados por el dentista.

En este problema, el uso del servidor ($\frac{\lambda_{ef}}{\mu}$) es de 0,996, que coincide con la intensidad de tráfico en el sistema ($\frac{\lambda_{ef}}{\mu}$).

Por último, la productividad del trabajador que se encuentra en la ventanilla es 3,98.

Como el programa no nos muestra el valor de la tasa de llegada efectiva o razón de entrada al sistema, ni la tasa de pérdida, vamos a calcularlas manualmente.

$$\lambda_{ef} = \lambda \cdot (1 - p_K) = 8 \cdot (1 - 0,5019608) = 3,9843136$$

$$\lambda_{perdido} = \lambda - \lambda_{ef} = 8 - 3,9843136 = 4,0156864$$

Viendo estos resultados, podemos decir que aproximadamente consiguen entrar al sistema 4 personas pero, las 4 restantes, deben abandonar la clínica al encontrarse la sala de espera llena. Además, la probabilidad de que el sistema no esté completo ($1 - p_K$), que es igual a la probabilidad de que una persona pueda acceder al sistema, es de un 49,8%.

Como recomendación, cabe decir que quizá la clínica debería contratar a otro dentista para que los clientes no tengan que esperar tanto tiempo y para reducir el número de clientes perdidos, debido a que se encuentran la sala de espera llena.

4.2.4.- EJEMPLO DEL MODELO M/M/C/K

Una empresa aérea cuenta con una centralita telefónica para que los clientes puedan realizar sus reservas por teléfono. Esta centralita está compuesta por 3 líneas, cada una de ellas con un agente, por lo que hay 3 telefonistas. Cuando los telefonistas están ocupados, las llamadas que llegan quedan en espera con una musiquita, pudiendo quedar en espera hasta 6 llamadas. Estas serán atendidas por orden de llegada cuando los agentes estén disponibles. En cambio, las llamadas que lleguen habiendo ya 6 llamadas en espera, recibirán el tono de “línea ocupada” y el cliente abandonará el sistema. Las llamadas llegan a la centralita siguiendo una distribución de Poisson, con una media de 22 llamadas por hora. El tiempo que tarda cada telefonista en atender cada llamada sigue una distribución exponencial, con una media de 5 minutos por llamada.

Para este enunciado, vamos a calcular las medidas de efectividad.

Como vemos, es un modelo M/M/3/9 con parámetros $\lambda = 22$ llamadas por hora y $\mu = 60/5 = 12$ llamadas por hora.

Podemos decir que, al igual que en el modelo anterior, la condición de estado estable se da siempre ya que, como la capacidad del sistema está limitada, la cola no puede crecer indefinidamente y el sistema se estabiliza.

FORMULACIÓN Y RESULTADOS OBTENIDOS

En el siguiente cuadro se muestra el código *R* (o script de *R*) utilizado para la resolución del ejemplo.

```

require(queueing)
# Set queue model input parameters
input_mmck <- NewInput.MMCK(lambda = 22, mu = 12, c = 3, k = 9)
# Create queue class object
output_mmck <- QueueingModel(input_mmck)
# Get queue model report
Report(output_mmck)
The inputs of the model M/M/C/K are:
Lambda: 22, mu: 12, c: 3, k: 9
The outputs of the model M/M/C/K are:
The probability (p0, p1, ..., pk) of the clients in the system are:
0.1414287 0.2592859 0.2376787 0.1452481 0.08876274 0.0542439
0.03314925 0.02025775 0.01237974 0.007565395
The traffic intensity is: 1.81946344279175
The server use is: 0.606487814263918
The mean number of clients in the system is: 2.30448317916684
The mean number of clients in the queue is: 0.48501973637509
The mean number of clients in the server is: 1.81946344279175
The mean time spend in the system is: 0.105547745788125
The mean time spend in the queue is: 0.0222144124547914
The mean time spend in the server is: 0.08333333333333
The mean time spend in the queue when there is queue is:
0.0622705660033522
The throughput is: 21.8335613135011

```

INTERPRETACIÓN DE LOS RESULTADOS

La probabilidad de que no llegue ninguna llamada a la centralita (p_0) es del 14,14%. Sabiendo esto, podemos decir que la probabilidad de que una llamada sea atendida inmediatamente es del 14,14%.

La probabilidad de que llegue una llamada es del 25,92%, la probabilidad de que lleguen dos llamadas es del 23,76%, la probabilidad de que lleguen tres llamadas es del 14,52% y, así, sucesivamente. Observando los datos, podemos decir que la probabilidad va disminuyendo a medida que aumenta el número de llamadas. Por ello, es más probable que en la centralita haya solo 1 llamada.

El número medio de llamadas telefónicas que llegan a la centralita para reservar un billete de avión (L) es de 2,304 llamadas, permaneciendo en término medio cada llamada en el sistema (W) unos 6 minutos ($0,1055 \cdot 60 = 6,33$).

El número medio de llamadas en espera para hablar con los agentes (L_q) es de 0,485. Cada llamada permanecerá en término medio en espera (W_q) poco más de un minuto ($0,0222 \cdot 60 = 1,332$).

El número medio de clientes hablando con los telefonistas o de llamadas siendo atendidas es de 1,8194, tardando cada llamada en finalizar casi 5 minutos ($0,0833 \cdot 60 = 4,998$) en término medio.

En consecuencia de lo anterior, podemos decir que las llamadas pasan en término medio más tiempo siendo atendidas que en espera.

En este problema, el uso del servidor ($\frac{\lambda_{ef}}{c\mu}$) es de 0,60. En cambio, la intensidad de tráfico en el sistema ($\frac{\lambda_{ef}}{c\mu}$) es de 1,81.

Por último, la productividad de los telefonistas que reciben las llamadas es 21,83.

Como el programa no nos muestra el valor de la tasa de llegada efectiva o razón de entrada al sistema, ni la tasa de pérdida, vamos a calcularlas manualmente.

$$\lambda_{ef} = \lambda \cdot (1 - p_K) = 22 \cdot (1 - 0,007565395) = 21,83$$

$$\lambda_{perdido} = \lambda - \lambda_{ef} = 22 - 21,83 = 0,16$$

Viendo estos resultados, podemos decir que aproximadamente consiguen ser atendidas casi las 22 llamadas prácticamente, ya que el valor de $\lambda_{perdido}$ es bastante bajo. Además, la probabilidad de que el sistema no esté completo ($1 - p_K$), que es igual a la probabilidad de que una llamada pueda entrar en el sistema, es de un 99,24%.

4.2.5.- COMPARACIÓN DE VARIOS MODELOS DE COLAS

En el departamento de contabilidad de cierta empresa hay 4 fotocopadoras que son usadas por los empleados. Pero, debido a las quejas constantes de los empleados de este departamento sobre la cantidad de tiempo que pierden esperando que alguna fotocopadora quede libre para imprimir o fotocopiar los documentos necesarios, el jefe se ha planteado aumentar el número de fotocopadoras. Se sabe que los empleados llegan a la fotocopadora siguiendo un proceso de Poisson, con una tasa de llegada de 30 empleados por hora. Además, se supone que el tiempo que tarda cada empleado en usar una fotocopadora sigue una distribución

exponencial, con una media de 5 minutos por persona. A parte de esto, el coste medio de la productividad perdida debido al tiempo que pierde cada empleado en la zona de fotocopiado es de 25€ por hora (coste asociado a la espera) y el coste de cada fotocopiadora es de 2€ por hora (coste del servicio).

Sabiendo esto, vamos a calcular el número de fotocopiadoras necesarias para cumplir los dos objetivos siguientes:

- Los empleados no tengan que esperar más de un minuto para imprimir o fotocopiar los documentos necesarios.
- Se minimice el coste total por hora.

En primer lugar, tenemos un modelo M/M/4 con parámetros $\lambda = 30$ empleados por hora y $\mu = 60/5 = 12$ empleados por hora. Este es el sistema con el que cuenta inicialmente la empresa. Por lo tanto, vamos a comparar este modelo con otros modelos M/M/C, aumentando el número de servidores, hasta que encontremos el que cumple nuestros objetivos.

Para que se den las condiciones de estado estable, ρ debe ser menor que 1 y, en este modelo, $\rho = \frac{\lambda}{c\mu}$. Por lo tanto, se cumplen las condiciones de estado estable ya que $\rho = \frac{30}{4 \cdot 12} = 0,625$.

FORMULACIÓN Y RESULTADOS OBTENIDOS

En el siguiente cuadro se muestra el código R (o script de R) utilizado para la resolución del ejemplo.

```

require(queueing)
# Set queue model input parameters
entrada_1<- NewInput.MMC(lambda = 30, mu = 12, c = 4, n = 5, method
= 0)
# Create queue class object
alternativa_1 <- QueueingModel(entrada_1)
require(queueing)
# Set queue model input parameters
entrada_2<- NewInput.MMC(lambda = 30, mu = 12, c = 4, n = 5, method
= 0)
# Create queue class object
alternativa_2 <- QueueingModel(entrada_2)
require(queueing)
# Set queue model input parameters
entrada_3<- NewInput.MMC(lambda = 30, mu = 12, c = 4, n = 5, method
= 0)
# Create queue class object
alternativa_3 <- QueueingModel(entrada_3)
require(queueing)
# Set queue model input parameters
entrada_4<- NewInput.MMC(lambda = 30, mu = 12, c = 4, n = 5, method
= 0)
# Create queue class object
alternativa_4 <- QueueingModel(entrada_4)
CompareQueueingModels(alternativa_1, alternativa_2, alternativa_3,
alternativa_4)

```

	Lambda	mu	c	k	m	ro	p0
1	30	12	4	NA	NA	0.625	0.07369498
2	30	12	5	NA	NA	0.500	0.08010013
3	30	12	6	NA	NA	0.4166667	0.08162026
4	30	12	7	NA	NA	0.3571429	0.08198040
	Lq		Wq		L		W
1	0.53309451		0.0177698169		3.033095		0.10110315
2	0.13037130		0.0043457099		2.630371		0.08767904
3	0.03388915		0.0011296384		2.533889		0.08446297
4	0.00857971		0.0002859903		2.508580		0.08361932

Como el programa no calcula el coste total esperado por unidad de tiempo, vamos a calcularlo manualmente, para cada una de las alternativas.

Alternativa 1

Coste Total = $C_1 x + C_2 L = 2 \cdot 4 + 25 \cdot 3,033095 = 83,82\text{€}$ por hora es el coste para 4 fotocopadoras.

Alternativa 2

Coste Total = $C_1 x + C_2 L = 2 \cdot 5 + 25 \cdot 2,630371 = 75,75\text{€}$ por hora es el coste total para 5 fotocopadoras.

Alternativa 3

Coste Total = $C_1 x + C_2 L = 2 \cdot 6 + 25 \cdot 2,533889 = 75,34\text{€}$ por hora es el coste total para las 6 fotocopadoras.

Alternativa 4

Coste Total = $C_1 x + C_2 L = 2 \cdot 7 + 25 \cdot 2,508580 = 76,71\text{€}$ por hora es el coste total para las 7 fotocopadoras.

INTERPRETACIÓN DE LOS RESULTADOS

En la tabla 2, vamos a recoger de forma resumida las medidas de efectividad obtenidas con cada uno de los modelos, así nos resultará más fácil compararlos y ver cuál cumple con nuestro primer objetivo, que los clientes no tengan que esperar más de un minuto en la cola.

TABLA 2. COMPARACIÓN DE LAS MEDIDAS DE EFECTIVIDAD

	4 fotocopadoras	5 fotocopadoras	6 fotocopadoras	7 fotocopadoras
ρ	0,625	0,50	0,416	0,357
p_0	0,07369498	0,08010013	0,08162026	0,08198040
L	3,033095	2,630371	2,533889	2,508580
W	$0,10110315 \cdot 60$ = 6,06 min	$0,08767904 \cdot 60$ = 5,26 min	$0,08446297 \cdot 60$ = 5,06 min	$0,08361932 \cdot 60$ = 5,01 min
L_q	0,53309451	0,13037130	0,03388915	0,00857971
W_q	$0,0177698169 \cdot 60$ = 1,06 min	$0,0043457099 \cdot 60$ = 0,2607 min	$0,0011296384 \cdot 60$ = 0,067 min	$0,0002859903 \cdot 60$ = 0,017 min

Fuente: Elaboración propia

Como podemos observar, a medida que se aumenta el número de fotocopadoras, que serían los servidores en este ejemplo, aumenta la probabilidad de que la zona de fotocopiado esté vacía y de que un empleado pueda utilizar alguna de las fotocopadoras inmediatamente, sin tener que esperar. También, a medida que aumenta el número de fotocopadoras, disminuye tanto el número medio de empleados en el sistema y en la cola como el tiempo que permanece cada empleado en el sistema y en la cola en término medio.

Una vez comentado esto, podemos decir que para cumplir el objetivo de que los empleados no tengan que estar más de un minuto esperando en la cola, la empresa debe contar con 5 o más fotocopadoras.

Para decantarnos por establecer 5, 6 o 7 fotocopadoras, vamos a analizar en la tabla 3 qué opción nos ayudaría a minimizar el coste total por hora.

TABLA 3. COMPARACIÓN ENTRE LOS COSTES

	4 fotocopadoras	5 fotocopadoras	6 fotocopadoras	7 fotocopadoras
Coste del servicio	8 €/h	10 €/h	12 €/h	14 €/h
Coste de la espera	75,82 €/h	65,75 €/h	63,34 €/h	62,71 €/h
Coste Total	83,82 €/h	75,75 €/h	75,34 €/h	76,71 €/h

Fuente: Elaboración Propia

Como podemos ver, los costes van bajando a medida que aumenta el número de fotocopadoras. Pero, al pasar de 6 a 7, los costes aumentan. Por ello, la decisión del jefe debe ser comprar 2 fotocopadoras más y, así, contar con 6 fotocopadoras ya que con 6 fotocopadoras se cumplen los dos objetivos planteados al principio: que los empleados no tengan que esperar más de un minuto para imprimir o fotocopiar los documentos necesarios y que se minimice el coste total por hora.

En conclusión, el sistema de colas M/M/6 sería el más eficiente y rentable para esta empresa.

5. CONCLUSIONES

En este TFG se ha estudiado la Teoría de Colas y algunos de sus modelos más utilizados en el mundo empresarial. Para ello, hemos analizado las medidas de efectividad y las hemos aplicado

en algunos ejemplos con el objetivo de asimilar las explicaciones teóricas y saber interpretar los datos y resultados obtenidos. Además, se ha analizado el estudio económico de un sistema de colas con el fin de conseguir un nivel óptimo de servicio, es decir, conseguir un equilibrio entre el coste de ofrecer el servicio y el coste asociado a la espera de los clientes para obtener ese servicio.

La Teoría de Colas nos ha permitido estudiar los fenómenos de espera a los que nos enfrentamos día a día en cualquier lugar y la forma de actuar de las empresas para poner solución a estas colas y satisfacer a sus clientes.

En este trabajo se ha realizado una introducción a la Teoría de Colas, además se han presentado algunos ejemplos reales de aplicación de esta teoría en empresas, comprendiendo los distintos resultados y estudiando la importancia que tiene en el ámbito empresarial.

Las Aplicaciones a la Administración de Empresas propuestas se han resuelto con el programa *R-Studio*. Los resultados obtenidos para las medidas calculadas son valores que sirven de ayuda a las empresas para la toma de decisiones óptima. Las empresas deben, en primer lugar, recoger datos sobre la tasa de llegadas y la tasa de salidas para, después, realizar un análisis que les permitirá saber cuánto tiempo y cuántos clientes hay en su empresa por unidad de tiempo, entre otras cosas. Una vez realizado este estudio, deberán tomar decisiones como, por ejemplo, ampliar o reducir el número de servidores con el fin que de los clientes no tengan que esperar demasiado y mejorar así su satisfacción o, todo lo contrario, con el fin de que los empleados no estén demasiado tiempo desocupados.

Además, el ejemplo de mayor importancia es aquel en el que se han comparado varios modelos de colas ya que nos ha ayudado a deducir que no hay un sistema de colas adecuado para todas las empresas, sino que, dependiendo de la empresa y de los objetivos que esta se plantee, será más recomendable un sistema u otro. Por ello, en este ejemplo donde los objetivos eran que los empleados no esperasen más de un minuto en la cola y minimizar el coste total, hemos ido aumentando el número de servidores hasta que hemos conseguido estos dos objetivos.

6. BIBLIOGRAFÍA

Aplicaciones de la Investigación de Operaciones. (s.f.). Investigación de Operaciones. Disponible on line:

https://www.investigaciondeoperaciones.net/aplicaciones_de_la_investigacion_de_operaciones.html

Eiselt, H. A. y Sandblom, C. L. (2012), *Operations Research*, Estados Unidos: Springer.

Eppen, G. D., Gould, F. J., Schmidt, C. P., Moore, J. H. y Weaterford, L. R. (2000), *Investigación de Operaciones en la Ciencia Administrativa*, México: Pearson Prentice Hall.

Hillier, F. S. y Lieberman, G. J. (2010), *Introducción a la Investigación Operativa*, México: McGraw-Hill.

Llano Monelos, P. (1997), *Aplicabilidad de la Teoría de Colas al Fenómeno Hospitalario* (Tesis Doctoral). Disponible on line: <http://hdl.handle.net/2183/18274>

Maroto, C., Alcaraz, J., Ginestar, C. y Segura, M. (2012), *Investigación Operativa en Administración y dirección de Empresas*, Valencia: Editorial Universitat Politecnica.

Martín Martín, Q. (2003), *Investigación Operativa*, Madrid: Pearson Prentice Hall.

Pazos Arias, J.J., Suárez González, A. y Díaz Redondo, R.P. (2003), *Teoría de Colas y Simulación de Eventos Discretos*, Madrid: Pearson Prentice Hall.

R Core Team. (2000), *Introducción a R. Notas Sobre R: Un Entorno de Programación Para análisis de Datos Y Gráficos*.

Taha, H. A. (2004), *Investigación de Operaciones*, México: Pearson Prentice Hall.

Teoría de Colas: procesos de nacimiento y muerte (s.f.). Portal Estadística Aplicada. Disponible on line: <https://www.estadistica.net/INVESTIGACION/TEORIA-COLAS.pdf>

Van Der Loo, M. P. J., y De Jonge, E. (2012), *Learning RStudio for R Statistical Computing*, Birmingham, UK: Packt publishing.