

**UNIVERSIDAD AUTÓNOMA DE BAJA CALIFORNIA
FACULTAD DE INGENIERÍA
MAESTRÍA Y DOCTORADO EN CIENCIAS E INGENIERÍA**



ÁREA DE CONOCIMIENTO: COMPUTACIÓN

**PLANIFICACIÓN DE LA PRODUCCIÓN EN UN TALLER DE
FLUJO HÍBRIDO CON RESULTADOS ESTOCÁSTICOS DE
UNA OPERACIÓN INTERMEDIA**

TESIS

**para obtener el grado de
MAESTRO EN CIENCIAS
presenta**

SALVADOR RAMÍREZ RODRÍGUEZ

DIRECTOR

Dra. Larysa Burtseva

RESUMEN de la Tesis de Salvador Ramírez Rodríguez, presentada como requisito parcial para la obtención de grado de MAESTRO EN CIENCIAS DE LA COMPUTACION, Mexicali, Baja California, ... de 2012.

PLANIFICACIÓN DE LA PRODUCCIÓN EN UN TALLER DE FLUJO HÍBRIDO CON RESULTADOS ESTOCÁSTICOS DE UNA OPERACIÓN INTERMEDIA

Resumen aprobado por:

Dra. Larysa Burtseva
Director de Tesis

En esta tesis se investiga el proceso de manufactura de interruptores eléctricos (*Reed Switch*) en una planta, con el propósito de asegurar el cumplimiento de los pedidos de consumidores en las fechas de entrega, mientras que los resultados de una operación intermedia en la ruta tecnológica son estocásticos lo que dificulta a la planificación y la programación de la producción. Más aun, estos resultados son sujetos al tipo de material utilizado en lote, y desechos generados en las etapas previas. Como consecuencia, la cantidad necesaria de lotes y su orden de procesamiento son desconocidos, produciendo un número incontrolable de inventario, y la eficiencia en la producción es afectada.

La empresa necesita herramientas que apoyen la planificación y programación de la producción, ayudando a minimizar la cantidad necesaria de lotes. Se toman en cuenta los resultados estocásticos de una operación en la ruta tecnológica, el inventario generado y un nivel fijo de desechos. El manejo de material se realiza bajo las características del sistema de control de inventario de la empresa, así que toda solución se debe adaptar a las condiciones establecidas bajo este sistema.

ABSTRACT of the Thesis, presented by Salvador Ramirez Rodríguez, in order to obtain the MASTER OF SCIENCE DEGREE in COMPUTER SCIENCE. Mexicali, Baja California, Mexico. June 2012.

PRODUCTION PLANNING IN A HYBRID FLOW SHOP WITH STOCHASTIC RESULTS OF AN INTERMEDIATE OPERATION

Abstract approved by:

Dr. Larysa Burtseva
Thesis Advisor

This thesis investigates the process of manufacture of electric switches (Reed Switch) in a plant, for the purpose of ensuring compliance with orders from consumers at delivery dates, while the stochastic results of an interim operation on the route complicated technological planning and production scheduling. Moreover, these results are subject to the type of material used in batch. As a result, the necessary amount and processing order of lots are unknown, generating uncontrolled growth of inventory, thereby affecting the efficiency in production.

The firm needs tools that support the planning and production scheduling, helping to minimize the required number of batches. They take into account the stochastic results of an operation in the technological route the inventory generated and a fixed level of waste. Material management is performed under the system inventory control features of the company, so any solution must adapt to the conditions set out under this system.

ÍNDICE

1 INTRODUCCIÓN

- 1.1 Planteamiento del problema
- 1.2 Modelo de la producción de interruptores eléctricos
- 1.3 Objetivos
- 1.4 Metodología
- 1.5 Esquema general de la tesis

2 BASES TEÓRICOS DE LA PLANIFICACIÓN EN UN TALLER DE FLUJO HÍBRIDO CON PROCESAMIENTO POR LOTES

- 2.1 Talleres de flujo híbrido
 - 2.1.1 Definición
 - 2.1.2 Fórmula
 - 2.1.3 Flexibilidad de un taller de flujo
 - 2.1.4 Otras definiciones
- 2.2 Particularidades de la planificación de la producción con procesamiento por lotes
- 2.3 Manejo de inventarios en manufactura serial
- 2.4 Conclusiones del capítulo

3. ANÁLISIS ESTADÍSTICO Y PROBABILÍSTICO DE DATOS

- 3.1 Recopilación y análisis de datos reales
- 3.2 Reclasificación
- 3.3 Ajuste a la distribución normal
- 3.4 Enfoque probabilístico
- 3.5 Conclusiones del capítulo

4 PLANIFICACIÓN DE LA PRODUCCIÓN EN LA OPERACIÓN ESTOCÁSTICA DE CLASIFICADO

- 4.1 Estimación de cantidad de lotes
 - 4.1.1 Notaciones
 - 4.1.2 Modelo de la Programación Lineal

- 4.1.3 Orden de procesamiento de pedidos
 - 4.1.4 Resolución del modelo
- 4.2 Ajuste de la aproximación al proceso estocástico
 - 4.2.1 Análisis de desviaciones
 - 4.2.2 Criterio de ajuste
- 4.3 Algoritmo
 - 4.3.1 Estructura del algoritmo
 - 4.3.2 Etapa de estimación
 - 4.2.3 Etapa de ajuste
- 4.4 Conclusiones del capítulo

5 VERIFICACIÓN DEL ALGORITMO EN EL EXPERIMENTO COMPUTACIONAL

- 5.1 Generación de las muestras
 - 5.1.1 Diseño del banco de pruebas
 - 5.1.2 Complejidad computacional
 - 5.1.3 Algoritmo de generación de muestras
- 5.2 Diseño de experimento
- 5.3 Resultados
- 5.4 Conclusiones del capítulo

6 CONCLUSIONES

7 FUTURO TRABAJO

REFERENCIAS

APÉNDICES

Apéndice A. Las muestras G1

Apéndice B. Las muestras G2

INDICE DE LAS FIGURAS

Fig. 1.1. Interruptor eléctrico (*Reed Switch*)

Fig. 1.2. Modelo de producción de interruptores eléctricos

Fig. 1.3. Operación de clasificado: a) M máquinas idénticas con R repositorios, o búferes limitados; b) Búferes ilimitados (WIP).

Fig. 1.4. Los repositorios de una máquina

Fig. 2.1. Un ejemplo que ilustra un modelo de un TFH con tres etapas

Fig. 3.1. Distribución inicial de un lote entre 25 repositorios de la máquina

Fig. 3.2. Región de aceptación de la hipótesis nula para las medias iguales

Fig. 3.3. Un ejemplo de distribuciones reales de tres lotes entre los repositorios.

Fig. 3.4. Las gráficas de residuos de la muestra global $G1$ contra diferentes distribuciones: a) Normal, b) Lognormal, c) Weibull, d) Rayleigh, e) Extreme Value y f) Exponencial

Fig. 3.5. Distribución de datos reales de un lote entre 3 niveles CR y 35 bins, recibida por el algoritmo de reclasificación

Fig. 3.6. El diagrama de validación de una pieza por niveles CR y rangos unitarios de operación

Fig. 3.7. Comparación entre distribución de la empresa y distribución reclasificada. a) Empresa. b) Reclasificada.

Fig. 3.8. Modelo de obtención de probabilidades.

Fig. 4.1. Representación gráfica de la región factible de la solución para el ejemplo.

Fig. 4.2. Comparación de la distribución esperada con la distribución de una muestra.

Fig. 4.3. Diagrama de casos.

Fig. 4.4. Fases de la etapa de estimación del algoritmo.

Fig. 4.5. Fases de la etapa de ajuste del algoritmo.

Fig. 5.1. Un ejemplo de muestras semejantes adquiridas a partir del algoritmo iterativo enumerativo para el primer tipo de lote.

Fig. 5.2 Un ejemplo de muestras variables adquiridas a partir del algoritmo SFSA para el primer tipo de lote.

INDICE DE LAS TABLAS

Tabla 3.1. Prueba de bondad de ajuste.

Tabla 3.2. Resultados de la prueba de bondad de ajuste a la distribución normal de los lotes reclasificados utilizando la prueba χ -cuadrada para tres niveles de CR: N_1 , N_2 y N_3 .

Tabla 3.3. Resultados de la prueba de bondad de ajuste a la distribución normal de los lotes reclasificados utilizando la prueba Jarque-Bera.

Tabla 3.4. Las probabilidades por los repositorio.

Tabla 3.5. Cantidad de piezas en % por repositorios para tipos de vidrio G1 y G2.

Tabla 5.1. Representación de un conjunto de demandas utilizadas en el experimento.

Tabla 5.2. Historial de un experimento.

Tabla 5.3. Escenario 1: Variables estáticas.

Tabla 5.4. Escenario 2: Tamaño de demanda variable.

Tabla 5.5. Escenario 3: Porcentaje de piezas por repositorio.

Tabla 5.6. Escenario 4: Simulación con variación en todos los parámetros.

Tabla 5.7. Escenario 5: Simulación con 5% de nivel de porcentaje.

1 INTRODUCCIÓN

1.1 Planteamiento del problema

El problema del cumplimiento de pedidos en la fecha de entrega es crítico en cualquier empresa manufacturera y en muchos casos está relacionado con el uso eficiente de materia prima. El poder utilizar la mínima cantidad de material para el mismo número de pedidos es indispensable tanto para ahorro de material como para la eficiencia de la producción en general.

En muchos problemas reales de manufactura, la información involucrada se caracteriza como determinística, es decir, es conocida de antemano y no varía con respecto a cualquier entrada, por ejemplo, recursos de maquinaria en la planta, rutas tecnológicas, tiempos de procesamiento de piezas. Sin embargo, los datos acerca de muchos ambientes de producción no pueden ser descritos como determinísticos debido su naturaleza estocástica. La incertidumbre que se provoca es causada por eventos inesperados, así como fallas de máquinas, cambio en fechas límites de entrega de productos al consumidor, cancelación de demandas, etc. (vea la lista completa en el artículo de Vieira et al. 2003). Más aún, ciertos valores y resultados de algunas operaciones siempre son aleatorios, por ejemplo, resultados de mediciones físicas, el componente de piezas defectuosas (*scrap*) en una cantidad total de productos, etc.

Resultados aleatorios de una o más variables en el proceso de manufactura hacen difícil la toma de decisiones acerca de las necesidades en materia prima e implican tamaños incontrollables del almacenamiento. La información que acompaña tal producción, con frecuencia tiene fallas y errores, y causa la debilidad de los resultados. Los riesgos de la planificación en tales casos incluyen tanto incumplimiento, como sobreproducción. El incumplimiento de pedidos de consumidores implica crecimiento de penalidades, mientras que empresas, que mantienen reservas grandes en inventarios, se consideran como inatractivas en mercados de hoy (Tsourveloudis, 2010).

El cumplimiento de pedidos de consumidores es una condición clave para que manufactureros se posicionen encabezando la competencia. Una de las estrategias reconocidas para abastecer consumidores con un servicio satisfactorio es el aumentar continuamente la capacidad de manufactura y servicios a responder los cambios, lo que se realiza a través del control eficiente y efectivo del flujo de material e incluye (Qiu, 2005):

- el monitoreo del nivel de inventario en cada buffer y cada máquina;
- el rastreo oportuno del estado de material en el piso de taller;
- el responder pedidos del piso de taller en tiempo real.

Las decisiones oportunas y bien justificadas se obtienen cuando estas son basadas en el análisis estadístico y probabilístico, así como en el modelado de la producción.

En esta tesis se investiga el proceso de manufactura de interruptores eléctricos (*Reed Switch*) en una planta, con el propósito de asegurar el cumplimiento de los pedidos de consumidores en las fechas de entrega, mientras que los resultados de una operación intermedia en la ruta tecnológica son estocásticos lo que dificulta a la planificación y la programación de la producción. Más aun, estos resultados son sujetos al tipo de material utilizado en lote y desechos generados en las etapas previas. Como consecuencia, la cantidad necesaria de lotes y su orden de procesamiento son desconocidos, se produce incontrolable inventario, y la eficiencia de la producción es afectada.

La empresa necesita políticas que apoyen a la planificación y la programación de la producción en el minimizar la cantidad necesaria de lotes de material tomando en cuenta resultados estocásticos de una operación en la ruta tecnológica, el inventario generado, y un nivel fijo de desechos.

1.2 Modelo de la producción de interruptores eléctricos

El interruptor eléctrico es un ensamblado que contiene un par de cuchillas ferro magnéticas de contacto (abiertas o cerradas), herméticamente selladas en una envoltura de vidrio, llenado de gas inerte (Figura 1.1). La medida en que se evalúa un interruptor es el valor de operación ó “vueltas por ampere” (*amper turns*), que equivale a las vueltas que da un ampere sobre un embobinado. Cada modelo de interruptor tiene un rango aceptable del valor de operación, llamado “rango de operación”, que junto con cierta forma de cuchillas se identifica como un “número de parte”. La fabricación de interruptores se efectúa de acuerdo a los pedidos sometidos por los consumidores. Un pedido incluye un número de parte, una cantidad de piezas y una fecha de entrega.



Figura 1.1. Interruptor eléctrico (*Reed Switch*)

La fabricación de interruptores se realiza de acuerdo a un proceso tecnológico compuesto por ocho operaciones (Figura 1.2). El proceso tecnológico está dividido de una manera natural en dos partes por una operación que se realiza en exteriores de la planta, de tal manera que las dos últimas operaciones, el clasificado y forma de cuchillas, son separadas de la primera parte de la ruta tecnológica. La operación de clasificado, en la cual se enfoca este proyecto, es la penúltima en la ruta (Figura 1.3, a). La última operación (forma de cuchillos) realizada en cuatro líneas paralelas, así como las primeras, son sujetas a los resultados de clasificado que retroalimenta las operaciones anteriores con la información actual.

La clasificación de interruptores se ejecuta por una de las máquinas paralelas, cuya función es aplicar dos mediciones a cada interruptor: la liberación del contacto de cuchillos (*blade contact release*, CR) y el valor de operación (*operate value*), y depositarlo en uno de 25 repositorios de la máquina. Esta operación se realiza por lotes de producción. En resultado de procesar un lote, las piezas se distribuyen entre los repositorios. Inicialmente, un lote contiene 5500 piezas con envolturas hechas del mismo tipo de vidrio (lo que se refiere a continuación como “el tipo de lote”). En la planta se utilizan, generalmente, dos tipos de vidrio, es decir, dos tipos de lote. La cantidad de piezas que llega a la máquina, varía por razón de desechos, generados en las etapas anteriores. La etapa de clasificado también genera desechos. El concepto de lote desaparece cuando su contenido se distribuye entre los repositorios, y desde entonces el producto se maneja en piezas.

Tres niveles de CR se consideran en la planta. A cada repositorio de una máquina, corresponden un nivel de CR y un rango de valores de operación, tal que estos valores en general no se intersecan. Contenido de un repositorio completo se empaca, se etiqueta y se almacena esperando que un pedido lo reclame. Las cajas etiquetadas forman una parte del trabajo-en-proceso (*work-in-process*, WIP).

Después de clasificar un lote, los interruptores se mandan las de líneas de forma de cuchillas, de acuerdo a la ruta tecnológica asociada con el número de parte de la demanda. Mientras tanto El material excedente se acumula para el cumplimiento de próximas demandas (Figura 1.3, b).

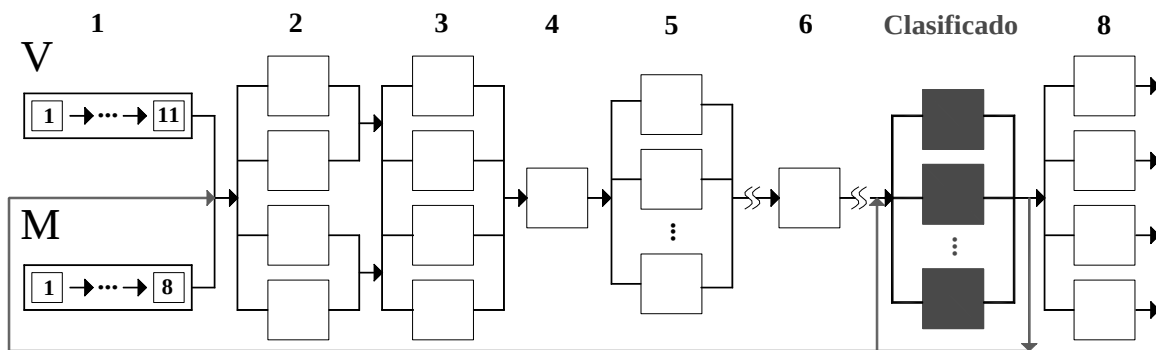


Figura 1.2. Modelo de producción de interruptores eléctricos

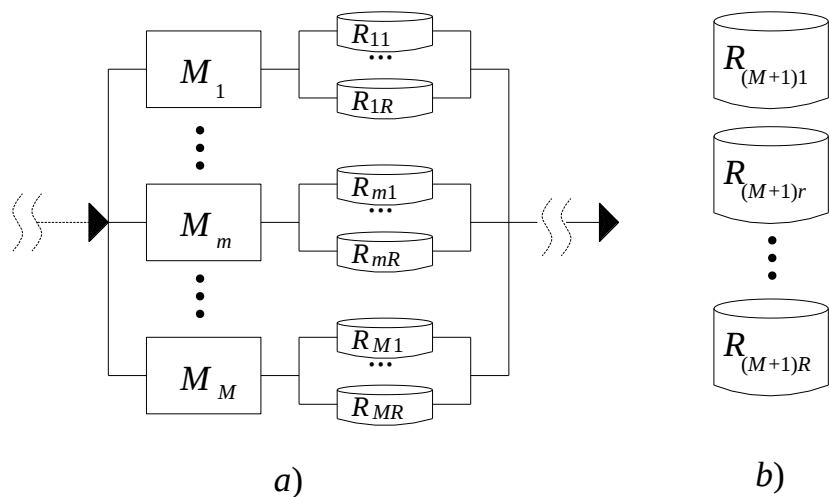


Figura 1.3. Operación de clasificado: a) M máquinas idénticas con R repositorios, o búferes limitados; b) Búferes ilimitados (WIP).

En la foto (Figura 1.4), se observan repositorios reales de una máquina. En la etiqueta de un repositorio se indican el nivel de CR y el rango de operación de piezas que este contiene.



Figura 1.4. Los repositorios de una máquina.

La operación de clasificado provoca un disturbio en la planificación de la producción, puesto que sus resultados son aleatorios. Debido a la complejidad de información, las decisiones en la planta se toman de una manera empírica, basándose en la experiencia previa, y por tanto la compañía sufre de los resultados inesperados cuando la producción se finaliza. Las consecuencias de la operación de clasificado afectan directamente a la eficiencia de la producción, la generación de la sobreproducción y por tanto a la cantidad necesaria de lotes de producción y el tiempo de ejecución de demandas, puesto que cada lote se distribuye entre los repositorios estocásticamente y solo después de esta operación se aclara si el cumplimiento de un pedido es posible en el periodo planeado.

De la experiencia se observó que el tipo de vidrio, del cual son hechos los tubos de interruptores en un lote, influye a la distribución de piezas entre los repositorios, por lo tanto el

conocimiento acerca de las tendencias de distribución de piezas para cada tipo de lote es de suma importancia para la planificación de la producción en la planta.

1.3 Objetivos

Objetivo general:

Optimizar la planificación de las necesidades en lotes de producción de diferentes tipos para cumplir con una lista de pedidos de consumidores, considerando una producción en serie con resultados estocásticos de una operación en la ruta tecnológica y el inventario generado en esta operación.

Objetivos particulares:

- Investigar las tendencias en la distribución de piezas para lotes de producción de diferentes tipos, en una operación, cuyos resultados son estocásticos.
- Desarrollar políticas que apoyen a la planificación y la programación de la producción en el minimizar la cantidad necesaria de lotes considerando el inventario generado.
- Diseñar un procedimiento para minimizar la cantidad necesaria de lotes aprovechando las diferencias en la distribución de dos tipos de lotes.

1.4 Metodología

En desarrollo de esta investigación se utilizan herramientas y métodos de la programación matemática (*mathematical programming*) y la planificación de la producción (*scheduling*). Basándose en el análisis del tipo de taller de flujo (*flow shop*), se estudia la literatura acerca de la específica de la producción con el procesamiento por lotes de producción, los sistemas de manufactura, y manejo de inventario en industrias con la producción en serie. A partir de este estudio, se hace el análisis estadístico de datos reales de manufactura y el análisis probabilístico de las distribuciones para poder predecir los resultados aleatorios de clasificado.

La metodología principal empleada en este trabajo se basa en un método conocido, que se utiliza para abordar un problema estocástico de optimización proveniente de la manufactura:

modificarlo a un problema determinístico sustituto bajo un incertidumbre estocástico. El procedimiento básico contiene dos pasos iterativos (Marti, 2008):

1. *Aproximación.* Los parámetros desconocidos, o variados estocásticamente, en el modelo del proceso se reemplazan por sus valores estimados, nominales, o aproximados. Después se aplica una decisión con respecto al problema determinístico sustituto.
2. *Ajuste.* La desviación del desempeño real del proceso con respecto a los valores prescritos se compensa entonces con la evaluación (medición) directa y las acciones de corrección.

La aproximación, o la solución inicial, que representa la cantidad mínima de lotes de cada tipo, se realiza a partir de un modelo matemático basado en los datos estadísticos acumulados. Esta solución se analiza después del procesamiento de siguientes lotes y en caso de la necesidad se ajusta al resto de pedidos, de acuerdo a un criterio.

El procedimiento diseñado se verifica en un experimento computacional.

1.5 Esquema general de la tesis

En el capítulo 2 se exponen los tópicos de teorías de la planificación y programación de la producción relacionados con el modelo de producción considerado, y algunos conceptos de procesamiento por lotes que se utilizan en la tesis. Un análisis de sistemas de inventarios en diferentes sistemas de la producción finaliza el capítulo.

El capítulo 3 se dedica a la recopilación y al análisis estadístico de la información real, el enfoque probabilístico, la obtención de las distribuciones iniciales a partir de reclasificación de los datos y aplicación de una serie de pruebas estadísticas. El proceso de la investigación se describe como una serie de experimentos con los datos.

En el capítulo 4 se describe la solución al problema: Se proponen el modelo de la programación lineal y su solución para la aproximación de cantidad de lotes necesarios para el cumplimiento de pedidos, así como el criterio y el algoritmo de su ajuste.

Un experimento computacional en el que se verifica el procedimiento diseñado, se describe en el capítulo 5. Se analiza el efecto producido por variables a la solución del problema.

Las conclusiones de la investigación y futuro trabajo descritos en los capítulos 6 y 7, respectivamente, finalizan la tesis.

2 BASES TEÓRICAS DE LA PLANIFICACIÓN EN UN TALLER DE FLUJO HÍBRIDO CON PROCESAMIENTO POR LOTES

La producción por lotes es típica en los sistemas de manufactura modernos debido al nivel alto de estandarización en el diseño de productos y volúmenes grandes de los pedidos. En este capítulo se describen las bases teóricas en la planificación de la producción aun cuando esta se realiza por lotes.

2.1 Talleres de flujo híbrido

2.1.1 Definición

En la planificación de la producción, un proceso multi-etapa con la propiedad de que todos los productos tienen que pasar por un número de etapas en el mismo orden, se clasifica como un taller de flujo (*flow shop*). En un taller de flujo simple (*simple flow shop*), cada etapa consiste de una sola máquina, la cual procesa a lo máximo, una operación en cada momento. Es más realístico suponer que en cada etapa son disponibles un número de las máquinas que operan en paralelo. Tal modelo es conocido como un taller de flujo híbrido (TFH), o en inglés *a hybrid flow shop* (HFS) (Ruiz & Vázquez-Rodríguez, 2010). Algunas etapas tienen una sola máquina, pero el modelo, para ser clasificado como un TFH, en al menos una etapa debe disponer varias máquinas en paralelo. Estas máquinas pueden ser idénticas o tener capacidades diferentes. Cada trabajo se procesa por no más que una máquina en cada etapa. El flujo de trabajos en el modelo es unidireccional; cada trabajo se procesa por una sola máquina en cada etapa.

Los modelos de TFH son comunes en las industrias que tienen la misma ruta tecnológica para todos los productos como una secuencia de etapas, donde algunas etapas tienen un grupo de máquinas capaces de realizar la misma operación. Varios procesos industriales, tales como químicas, metalúrgicas, semiconductores, tableros de circuito impreso (*printed circuit board*, PCB), farmacéuticas, petroleros, alimenticios, manufactura de automóviles, entre otros, pueden ser modelados como un TFH (Ruiz et al., 2008). En tales industrias, las instalaciones de equipo son duplicados en paralelo para incrementar la capacidad total, o para balancear las capacidades de las etapas, eliminar o reducir el impacto de etapas con cuello de botella en capacidades del piso de taller.

Un TFH tiene las siguientes características:

- Existen k etapas de procesamiento que ocurren en un orden lineal: 1, 2, ..., k , para procesar n trabajos. Cada etapa tiene un número predeterminado m_k de máquinas en paralelo, y el número de máquinas varía de etapa a etapa, $m_k = 1, \dots, M_k$.
- Cada de n trabajos visita las etapas en aquel orden, aunque no todas los trabajos necesitan visitar todas las etapas. Algunas etapas pueden ser saltadas para un trabajo particular, pero el flujo de trabajos por etapas es el mismo.
- El tiempo de procesamiento de cada trabajo en cada máquina que visite es conocido de antemano y es constante.
- Un trabajo no se divide entre las máquinas, y una máquina procesa no más que un trabajo en el mismo tiempo. Cuando una operación se inicia en una máquina, no se interrumpe hasta terminar.
- Típicamente, el orden de procesamiento de trabajos no tiene una relación de precedencia.
- Entre las etapas se disponen buffers para guardar productos intermedios.

2.1.2 Fórmula

Para describir todas las detalles de un TFH considerado, se utiliza la notación $\alpha|\beta|\gamma$ de tres campos, llamada la tripleta de Graham (Pinedo, 2008). El campo α denota la configuración del taller, incluyendo el tipo de taller y el ambiente de las máquinas por etapas. El campo α se descompone en cuatro parámetros, es decir α_1 , α_2 , α_3 , y α_4 , posiciones como $\alpha_1\alpha_2$ ($\alpha_3\alpha_4^{(1)}$, $\alpha_3\alpha_4^{(2)}$, ..., $\alpha_3\alpha_4^{(\alpha_2)}$). Aquí el parámetro α_1 indica el taller considerado, el parámetro α_2 muestra el número de etapas. Para la notación de TFH, en la posición α_1 se denota FH, y el valor del parámetro α_2 debe ser mayor a uno. Para cada etapa, los parámetros α_3 y α_4 indican el ambiente del conjunto de las máquinas. Más específico, α_3 denota información acerca de tipo de máquinas, mientras α_4 indica el número de máquinas por las etapas.

Los posibles desarrollos del conjunto de las máquinas en la etapa k de un TFH son los siguientes:

- *Una máquina (1)*: un caso especial; algunas etapas (no todas) en un TFH tienen una sola máquina;
- *Máquinas idénticas en paralelo (Pm_k)*: trabajo j puede ser procesado en cualquiera de m_i máquinas;

- *Máquinas uniformes en paralelo* (Qm_i): las m_k máquinas en el conjunto tienen diferentes velocidades; cualquier máquina del conjunto es capaz de procesar un trabajo j , mientras el tiempo de procesamiento es proporcional a la velocidad de la máquina;
- *Máquinas no relacionadas en paralelo* (Rm_i): existe un conjunto de m_k máquina diferentes en paralelo. El tiempo que un trabajo se queda en la máquina depende tanto del trabajo como de la máquina;
- *Máquinas dedicadas en paralelo* (Dm_i): para cada máquina se conoce de antemano un subconjunto de los trabajos que procesa.

Cuando existen varias etapas consecutivas con el mismo ambiente de las máquinas en diferentes etapas, los parámetros α_3 y α_4 se agrupan como $((\alpha_3\alpha_4)^k)_{k=s}^i$, donde s e i son los índices de la primera y la última etapas consecutivas, correspondientemente. Por ejemplo, la notación $FH4, (1, P2^{(i)})_{i=2}^3, Q3^{(4)}$ refiere a una configuración de TFH con cuatro etapas, donde una sola máquina se encuentra en la primera etapa, dos máquinas paralelas idénticas en la segunda y tercera etapas, y tres máquinas no relacionadas en la cuarta etapa.

El campo β incluye las propiedades del taller, los detalles y condiciones del procesamiento, también puede ser vacío si estos no existen. A continuación se describen algunas propiedades asociadas frecuentemente con un problema de planificación en un TFH (el campo β):

batch(b) Procesamiento por batches. Una máquina es capaz procesar hasta b trabajos continuamente, sin ningún ajuste.

Brkdown Interrupción (*breakdown*) de una máquina. Implica que una máquina puede no ser disponible continuamente.

M_{jk} Restricciones de elegibilidad de la máquina. El procesamiento del trabajo j es restringido con un conjunto M_j de las máquinas en la etapa k .

S_{sd} En taller existen tiempos de ajuste de las máquinas dependientes de la secuencia de trabajos.

w_j El factor de prioridad denota el peso, o importancia de un trabajo j con relación a otros trabajos en el sistema.

El campo γ estabiliza objetivo a minimizar. Las funciones objetivo más comunes en un problema de planificación en un TFH son (el campo γ):

$C_{\max}$ – el tiempo máximo de cumplimiento (*makespan*);

$F_{\max}$ - el tiempo máximo de flujo de un trabajo (*flow time*);

$L_{\max}$ - el retraso (*lateness*) máximo;

$T_{\max}$ – la tardanza (*tardiness*) máxima;

$E_{\max}$ – el anticipado (*earliness*) máximo;

El criterio más utilizado para un problema de planificación en un TFH, es el tiempo de cumplimiento, que es cuando el último trabajo sale del sistema, que se refiere como un *makespan*, o $C_{\max}$.

Un problema estándar de planificación en un TFH con k etapas y un numero de máquinas en cada etapa se denota como $FHk, ((PM^{(k)})_{k=1}^K) || C_{\max}$. En este caso, la fórmula define un TFH con k etapas, $|M^{(k)}|$ máquinas idénticas en paralelo en la etapa k , $k = 1, \dots, K$; el campo β es vacío, y el objetivo es el mínimo de makespan.

La Figura 2.1 ilustra la relación física entre máquinas y etapas, que corresponde a la notación $FH3, (1, P3^{(2)}, R2^{(3)}) | M_{j3}, S_{sd} | T_{\max}$, refiriendo a un TFH de tres etapas. La etapa uno tiene una sola máquina, la etapa dos tiene tres máquinas idénticas en paralelo, y la etapa tres tiene dos máquinas no relacionadas; M_{j3} y S_{sd} indican que en la etapa tres existen las restricciones de la elegibilidad de las máquinas y tiempos de ajuste dependientes de la secuencia de trabajos. El objetivo es minimizar la máxima tardanza. Más aun, la Figura 2.1 muestra que existen buffers ilimitados entre etapas para el almacenamiento del WIP.

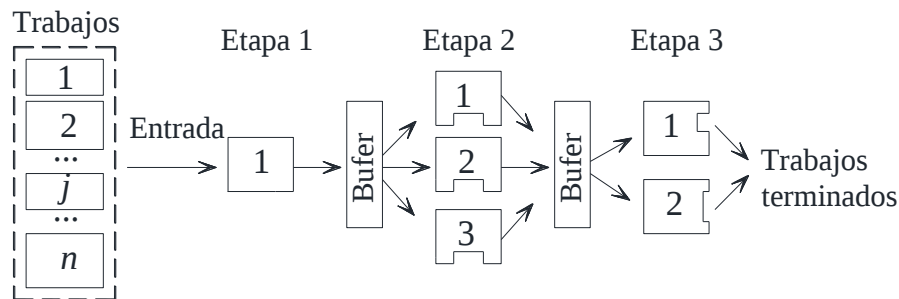


Figura 2.1. Un ejemplo que ilustra un modelo de un TFH con tres etapas.

2.1.3 Flexibilidad de un taller de flujo

Un sistema de producción para clasificarse como un TFH debe ser flexible. Es importante de conocer la diferencia entre un sistema flexible de la producción y un tradicional; qué menciona exactamente el concepto de la flexibilidad y justifica el uso de modelos especiales para sistemas flexibles de la producción. Los sistemas automatizados de la manufactura muestran la flexibilidad en modos múltiples y entrecruzados, pertinentes al equipo, procesos, productos, volúmenes de producción, etc. Entre los más importantes conceptos son los siguientes (Crama, 1997), (Vairaktarakis, 2004):

- *flexibilidad de máquinas (machine flexibility)*, la cual es la habilidad de las máquinas a realizar varios tipos de operaciones sin requerir un costoso esfuerzo para el cambio de una operación a la otra;
- *flexibilidad en manejo de material (material flexibility)*, el cual es la habilidad del sistema de manejo de material para mover eficientemente diferentes partes de material para la colocación propia y el procesamiento a través de facilidades de manufactura;
- *flexibilidad de operaciones (operation flexibility)*, la cual es la habilidad de las mismas a realizarse en diferentes modos;
- *flexibilidad del procesamiento (processing flexibility)*, la cual menciona que trabajos pueden saltar etapas, o existe un conjunto de tipos de partes los cuales pueden producirse en el sistema sin mayores ajustes;
- *flexibilidad en asignaciones de ruta (routing flexibility)*, la cual es la habilidad de sistemas manufactureros de fabricar una parte de productos en el sistema utilizando rutas alternativas.

Un modelo de la planificación de la producción con conjuntos de máquinas en paralelo debe cumplir con uno de estos conceptos para ser clasificado como un taller de flujo flexible, tomando en cuenta que el flujo de productos en la planta es unidireccional. La *hibridación* ocurre cuando algunos productos requieren algunas condiciones especiales para su fabricación, por ejemplo, diferentes calidades o capacidades de máquinas en la misma etapa, asignación de algunos trabajos a ciertas máquinas, u otras condiciones especiales.

2.1.4 Otras definiciones

Mientras el modelo de TFH ha sido bien investigado a partir de los 70^s, los investigadores prestan mucha atención a este modelo, y varios nuevos diseños han sido descubiertos en años recientes. Este hecho, probablemente, implica confusiones en terminología y notaciones. Actualmente, no existe en la literatura una clasificación convencional de esta clase de talleres de flujo. Una variedad de modelos conocidos se recomienda interpretar como un TFH o su caso especial. Estos son:

Taller de flujo flexible (TFF), o Flexible flow shop (FFS); es un TFH en un desarrollo de máquinas paralelas cuando las máquinas en cada conjunto son idénticas (flexibilidad de procesamiento en cada etapa de producción, la cual esta derivada desde habilidad de procesar un trabajo en cualquier máquina de etapa). Algunos autores reconocidos, como, por ejemplo, Pinedo (2008), Jungwattanakit et al. (2009), no utilizan la notación de TFH, y describen las más complejas configuraciones como un TFF con las máquinas no idénticas al menos en una etapa. Más aun, una variedad de autores no difieren entre términos de TFF y TFH, refiriendo este modelo como un *Taller de flujo flexible (hibrido)* (Allahverdi et al., 2008); o utilizan el término de TFH en un desarrollo con máquinas paralelas idénticas (Naderi et al., 2009).

Línea de flujo flexible, o Taller de flujo con multiples procesadores (refiriendo en inglés como *Flexible flow line*, FFL y *Flow shop with multiple processors*, FSMP, MPFS) son equivalentes a un TFF (Lin & Liao, 2003). Zandieh et al. (2006) consideran que un modelo de TFH es conocido como una línea de flujo flexible, por la razón que el flujo de trabajos en tal sistema es unidireccional.

Taller de flujo hibrido flexible o Línea de flujo hibrido flexible (Hybrid flexible flow shop or Flexible hybrid flow line (HFFL); este modelo es equivalente a un TFH donde trabajos pueden saltar etapas (*flexibilidad de procesamiento* atravesando etapas) (Ruiz & Vazquez-Rodriguez, 2010), (Allahverdi et al., 2008).

El Sistema de TFHs paralelos (TFHP) (*Parallel HFS system*, PHFS) representa un TFH descompuesto en un número de menores TFHs, sub-diseños operados en paralelo. Mas especifico, un TFHP es compuesto de un número de sub-diseños independientes, cada uno de los

cuales es un TFH con una ruta unidireccional (*flexibilidad en asignaciones de ruta*) (Vairaktarakis, 2004).

Un TFH restringido con dos etapas de procesamiento, hasta el caso más simple, cuando una etapa contiene dos máquinas paralelas idénticas y una sola máquina en la segunda etapa, ya es NP-duro (NP-*hard*) de acuerdo a (Gupta, 1988). Más aún, un caso particular cuando existe una máquina en la etapa, conocido como un taller de flujo (*simple flow shop*), y el caso más simple, cuando existe una sola etapa con varias máquinas, conocido como un desarrollo con máquinas paralelas idénticas, también son NP-duros (Glover & Laguna., 1997). La complejidad de problemas de la planificación en un TFH es esencialmente más alta cuando se trata de procesos estocásticos.

2.2. Particularidades de la planificación de la producción con procesamiento por lotes

Los problemas de planificación en sistemas de manufactura, donde los componentes del producto se procesan por lotes (paletas, contenedores, boxes) de muchas idénticas piezas, atrae recientemente la atención de investigadores debido a su importancia práctica en la industria y servicio modernos. Ejemplos más comunes de la organización de la producción por lotes se encuentran en la producción en serie como manufactura de dispositivos digitales y sus componentes, líneas de ensamblaje, facilidades de servicio de telecomunicaciones. El procesamiento por lotes y el uso de máquinas de procesamiento por *batches* para ejecución paralela de una o más operaciones tecnológicas son típicos en sistemas modernos de manufactura.

En un problema tradicional de planificación, un trabajo es indivisible, y no puede ser transferido a la siguiente máquina antes de terminar el proceso (Marimuthu et. al 2009; Potts & Baker 1989). Al contrario, la división de trabajos en sublotos en muchas situaciones es permisible y deseable para acelerar la fabricación o con respeto de sistemas Justo-En-Tiempo (*Just-In-Time*, JIT) a través de procesamiento paralelo. Un TFH con procesamiento por lotes tiene dificultades adicionales tales que diferencias en finalización de sublotos en la máquina, tamaños desconocidos de lotes (sublotos), tratamiento de tiempos de ajuste de las máquinas a través de secuenciación de lotes, la no-proporcionalidad entre tamaños de lotes y pedidos de

consumidores, la cual conduce a la división de pedidos y problemas de asignación (*allocation problems*), etc. Esta clase de problemas posee una complejidad computacional alta; por lo tanto, la mayoría de resultados publicados de investigaciones son dedicados a los talleres de flujo simple o TFH de dos etapas, además, en un ambiente estático.

La planificación de la producción por lotes tiene diferencias en comparación con la planificación tradicional de trabajos. La notación del trabajo depende de las suposiciones del problema y requiere una interpretación explícita. Tanto un lote, como un pedido pueden ser considerados como un trabajo. El interés teórico, así como el práctico, representan modelos donde se permite la división de trabajos (*job splitting*), es decir, es permitido dividir un trabajo en sublotes para procesar en paralelo en la misma o varias máquinas. En muchos modelos, el tamaño de lote no es evidente y debe ser optimizado. Cuando un sublote requiere un ajuste considerable, la complejidad del problema se incrementa. El agrupamiento de sublotes es el segundo aspecto importante de procesamiento por lotes.

El enfoque estocástico en problemas de manufactura serial introduce cambios en criterios tradicionales. Los métodos estándares de la planificación solamente consideran eficiencia de taller utilizando objetivos basadas en tiempo, tales como makespan, tiempo máximo de flujo, etc. Una planificación dinámica se desvía significativamente de la pre-planificación (*prescheduling*) y lleva a los nuevos actividades y criterios basados en la gestión de materiales, empleo de recursos humanos, etc. (Kopanos et al., 2008).

Las suposiciones típicas para un problema de la planificación en un TFH con procesamiento por lotes son las siguientes:

- Sean k etapas de producción por lotes que ocurren en un orden lineal y al menos una etapa tiene máquinas en paralelo. Un trabajo representa el procesamiento de un entero lote de piezas idénticas, o un pedido a producir una cierta cantidad de piezas idénticas, los cuales se procesan por lotes.
- Cada trabajo visita etapas en un orden unidireccional. Algunas etapas pueden ser saltadas para un trabajo en particular, pero el flujo del proceso para todos los trabajos es el mismo.
- El tamaño de lote y tiempo de procesamiento de una pieza son conocidos de antemano en cada máquina y son constantes. Los sublotes que pertenecen al mismo lote pueden tener diferentes tiempos de procesamiento, a una consecuencia trivial de diferentes tamaños.

- Sublotes que no tienen un tiempo de ajuste entre sí, forman un *batch* de producción y se procesan juntos sin requerir cualquier ajuste de la máquina. Sublotes, que pertenecen a diferentes batches, pueden necesitar un tiempo de ajuste esencial. Una máquina de procesamiento por batches es capaz de procesar varios sublotes simultáneamente, hasta que la capacidad de la máquina no excede.
- Todos los sublotes incluidos en el mismo batch, inician y terminan juntos, al mismo tiempo. Interrupciones en procesamiento de sublotes no se permitan.
- Para almacenar productos intermedios se colocan los buffers entre etapas y son mencionados cuando son limitados.

La planificación en un TFH atrae interés de la comunidad científica, sin embargo, el asunto de procesamiento por lotes en este tipo de talleres todavía no consiguió un interés suficiente.

2.3 Manejo de inventarios en manufactura serial

Los sistemas de inventario incluyen tradicionalmente tres componentes:

- 1) Materia prima (*raw material*);
- 2) WIP;
- 3) Productos terminados (*finished goods*).

Comúnmente, los productos manufacturados completamente, los cuales son listos para la venta y entrega a consumidores, son llamados “productos terminados”, o *finished goods*, mientras el WIP incluye las partes no terminados que se encuentran pasando el sistema de manufactura, y la materia prima es un sujeto del labor. Realmente, los tres componentes son mutuamente dependientes (Tsourveloudis, 2010).

Un aspecto fundamental a nivel operacional en los sistemas de producción es la obtención de la cantidad deseada del producto con la mínima cantidad de inventario en proceso. Conway et al. (1988) observan que en líneas de producción serial, el propósito de WIP es dar alguna independencia operacional a cada etapa del sistema de producción, así, el WIP se recomienda utilizar en cada planta que tiene una línea de producción. En manufacturas de ensamblaje, WIP se genera continuamente a consecuencia de estocastibilidad de resultados. Demasiado WIP prolonga el tiempo desde el inicio hasta que productos sean listos para abandonar instalaciones

de la planta, para seguir la cadena de la producción y distribución o llegar al usuario terminal. Con un WIP demasiado pequeño, cuando hay variaciones en tiempos o cantidades de producción, se provocan más “hambriento” y bloqueo en la producción que en caso con un WIP más grande (Pettersen, 2009).

En últimos años, el problema de optimización y control de inventario atrae mucha atención de investigadores. Hay varios artículos enfocados en el problema de control de WIP, donde se proponen diferentes métodos para mantener WIP en bajo nivel, sin embargo, un valor exacto y óptimo de WIP no puede ser determinado para manufacturas reales (Gershwin, 2000). Yang and Posner (2010) subrayan la importancia del conocimiento sobre la locación propia y el tamaño del WIP necesario en un sistema de producción. Qiu (2005) presenta una solución potencial práctica para un sistema manejable y distribuido de control de WIP dirigiendo a los asuntos tales como tiempo real de procesamiento, escalabilidad y reconfigurabilidad. También se propone la Línea de Producción Virtual (*Virtual Production Line*, VLP) como un método para mantener un cierto nivel de WIP.

Para intentar a optimizar el WIP, se propusieron diferentes métodos. Lin (2009) utiliza redes neuronales artificiales con el objetivo de identificar un nivel óptimo del WIP y maximizar el índice de rendimiento; simultáneamente, se minimizan la media y desviación estándar del tiempo de ciclo. El inventario de WIP se mide por el número de partes no terminados en los buffers intermedios del sistema de manufactura, y debe permanecer tan reducido como es posible (Tsourveloudis, 2010). Algunos investigadores intentan controlar el WIP utilizando algoritmos heurísticos donde el objetivo es encontrar una asignación óptima que minimice la media del inventario WIP; se provee que el material involucrado no excede un nivel dado (Papadopolus 2001), otros intentan optimizar el inventario WIP empleando la programación lineal donde un modelo introducido es desarrollado utilizando un marco general de la teoría de la producción dinámica (Kim 1994).

En artículos de (Sivakumar & Arivarignan, 2009) y (Kilic & Tarim, 2010) se investigan sistemas estocásticas, donde llegada de pedidos en proceso de manufactura es incontrolable. En (Sivakumar & Arivarignan, 2009), el sistema de inventario es modelado a través de productos perecederos. Las llegadas de consumidores regulares y negativos son descritas como un Proceso de Llegadas Markoviano (*Markovian Arrival Process*, MAP). En el artículo está dado el índice

del costo óptimo. En (Kilic & Tarim, 2010), una demanda representa una variable discreta aleatoria que varía de periodo a periodo. Se presentan las bases para la medición del nerviosismo (*nervousness*) de un sistema en ambientes no-estacionarios de demandas, se evalúa el comportamiento de costos de pólizas (R, S) y (s, S), tomando en cuenta la estocasticidad de demanda.

Existen distintos sistemas de control de la producción implementados en la industria (Hopp, 1998), (Pettersen, 2009). Estos sistemas se suelen denominar con carácter general por el efecto ejercido sobre el flujo de materiales: *push* - efecto de empujar -, *pull* - efecto de tirar -, e híbridos si se dan ambos efectos. Push sistema se define como “hazlo todo solo en caso” (“*make all we can just in case*”), lo que quiere decir que es categorizado por la aproximación de producción, los usos previstos, los lotes grandes, los inventarios altos, pérdidas, la gestión por la lucha contra incendios, la falta de comunicación. Al contrario, pull sistema se define como “haz lo que necesitas cuando es necesario” (“*make what’s needed when we need it*”) y es categorizado con la producción exacta, consumo actual, lotes pequeños, reducción de pérdidas, gestión por la vista, una mejor comunicación (Pettersen, 2009). Las ventajas de pull sistemas sobre push son las siguientes (Hopp, 1998):

- 1) Observabilidad: el WIP es directamente observable, mientras capacidades de las máquinas con respecto a fechas de entrada a las cuales estas deben ser ajustadas, no las son.
- 2) Eficiencia: Pull sistemas alcanzan el índice de material involucrado igual que un push sistema con el menor nivel de WIP en medio.
- 3) Variabilidad: Tiempos de flujo son menos variables en pull sistemas que en push, puesto que en pull sistemas se regula la fluctuación del nivel de WIP, mientras en push sistemas no es así.
- 4) Robustness: Pull sistemas son menos sensibles a los errores en el WIP comparando con la sensibilidad de push sistemas a errores en los momentos de llegada.

Es común de pensar que pull sistemas son mejores que push, sin embargo, Bonney (1999) muestra que ambos sistemas tienen resultados parecidos en la representación de sistemas.

El modelo JIT apareció en base de pull sistema. Su objetivo es simple: producir la producción de la calidad y cantidad deseadas en un periodo requerido (Azadeh et al., 2005). En

manufactura, el modelo JIT está bien acreditado con muchos beneficios holísticos. Estos incluyen: niveles reducidos del inventario, invitaciones reducidas en el inventario, la calidad aumentada de entrante material, y productos de la calidad alta constantemente (Yasin, 2003). De acuerdo a la filosofía JIT, el WIP no debe generarse especialmente. El inventario de la producción, el cual es considerado en este trabajo, tiene características de sistema pull con el modelo de la producción JIT, lo que implica producir piezas solo para cumplimiento de los pedidos y no para hacer reservas, así, la generación de WIP debe ser evitada.

2.4 CONCLUSIONES DEL CAPÍTULO

Problemas de planificación y programación de la producción en TFH atraen mucho interés de investigadores, el cual se refleja en la literatura especializada. Sin embargo, las condiciones de la manufactura moderna proporcionan nuevos modelos y problemas no solucionados hasta el momento. El enfoque de procesamiento por lotes en este tipo de taller no obtuvo una atención suficiente.

El modelo considerado representa un TFH con ocho etapas. La investigación se concentra en la séptima – el clasificado, con varias máquinas paralelas idénticas. Los criterios clásicos a minimizar, como mínimo de tiempo total de procesamiento, la tardanza o el anticipado de trabajos, son los más investigados. Sin embargo, los criterios que reflejan el balanceo de tiempos de procesamiento por las máquinas, eficiente uso de material, minimización de inventarios tienen mucha importancia por las condiciones de las manufacturas en serie modernos. Estos criterios son basados en la minimización del WIP.

El proceso de ensamble de interruptores sigue el sistema de inventario *pull*, con modelo de la producción *JIT*. El conocimiento sobre requerimientos de JIT es necesario para optimizar un sistema de inventario que existe actualmente en la producción de interruptores.

El análisis de la literatura muestra que el problema de la planificación en un TFH, donde se busca satisfacer los pedidos con un número mínimo de lotes, tomando en cuenta un proceso estocástico de la producción y diferentes tipos de material, es nuevo y no ha sido solucionado hasta el momento.

3. ANÁLISIS ESTADÍSTICO Y PROBABILÍSTICO DE DATOS

3.1 Recopilación y análisis de datos reales

En la empresa de ensamblado de interruptores eléctricos se manejan diversos tipos de vidrio con el cual esta formado el tubo del interruptor. Se recolectaron alrededor de 50 lotes ensamblados con tubos de dos tipos de vidrio. Los lotes fueron procesados en la etapa de clasificado en diferentes días y con diferentes máquinas (todas las máquinas se consideran como similares).

La información, que acompaña cada pieza en un lote, incluye:

- la identificación del lote,
- valor de operación,
- valor de calidad (CR, por siglas en ingles),
- numero de repositorio, y
- el resultado de la prueba de paso (*pass-test*).

Los datos obtenidos en la empresa inicialmente no muestran ningún tipo de relación entre las cantidades de piezas depositadas en los repositorios. Lo primero que se realizó fue separar los lotes por tipos de vidrio, para después analizar la distribución de cada uno (cantidad de piezas contenidas en los repositorios de las máquinas) y encontrar similitudes. La Figura 3.1 muestra el diagrama de la distribución de un lote entre 25 repositorios en una máquina de clasificado.

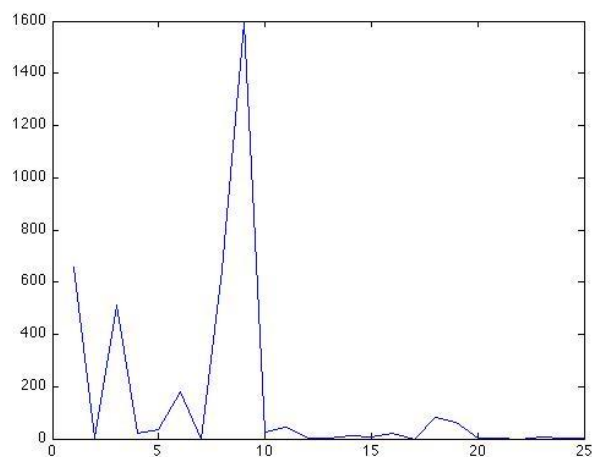


Figura 3.1. Distribución inicial de un lote entre 25 repositorios de la máquina.

Se sabe que la empresa maneja cuatro tipos de vidrio que se diferencian entre sí por el grosor del tubo del interruptor. Dado que a cada lote corresponde una distribución, es necesario comprobar de una manera experimental si existe diferencia entre ensamblar lote con un tipo de vidrio u otro.

El primer experimento se realizó en la planta con el propósito de obtener la información real sobre el proceso de clasificado. Para la recopilación de datos solo se consideraron 2 tipos de vidrio, ya que estos son los más utilizados para la fabricación de piezas en la empresa. Los tipos de vidrio seleccionados son nombrados como G1 y G2, cada uno con el mismo grosor.

Se define *muestra* como la representación de un lote que ha sido procesado por la máquina clasificadora, y colocado en repositorios de la máquina. Es importante considerar las muestras de la empresa para el experimento, ya que reflejan los resultados reales de la etapa de clasificación.

En el primer paso, se obtuvo un conjunto de muestras para cada tipo de lote. Se seleccionaron 18 muestras con tipo de vidrio G1 y 25 muestras con tipo G2. Los datos sobre los lotes se presentan en los Apéndices A y B. La información mostró dos diferencias: variación en el tamaño del lote y variación en la cantidad de piezas depositadas en los repositorios. Se calculó la media y desviación estándar de la *muestra global* (representa la población de piezas para el tipo de vidrio seleccionado). Esta muestra es el acumulado de piezas de los lotes reales que fueron ensamblados con el mismo tipo de vidrio. Los cálculos son basados únicamente en el *valor de operación* de las piezas.

Al realizar los cálculos para cada tipo de vidrio, se observó que tanto la media como la desviación estándar son diferentes entre si. Para poder establecer si son estadísticamente diferentes se realizó una prueba estadística que consistió en lo siguiente:

Se planteó la hipótesis del experimento como:

H0: Las medias μ_{G1} , μ_{G2} son iguales, por lo tanto son estadísticamente iguales.

H1: Las medias μ_{G1} , μ_{G2} son diferentes, por lo tanto son estadísticamente diferentes o existe una diferencia significativa.

Para comprobar estas hipótesis, es necesario descubrir si las medias son estadísticamente iguales o no; demostrando la hipótesis nula se pueden sacar conclusiones.

Considerando las características del proceso, se define que los lotes se producen en manera

independiente, y las condiciones de la producción no afectan a la clasificación de las piezas. La variable a estudiar es el valor de operación de la pieza, con el cual se realiza la clasificación. Debido a estas características el estudio se considera *seccional cruzado* (Rosner, 2010).

El segundo experimento consiste en encontrar un método estadístico para comprobar la hipótesis nula. Sabiendo que nuestro problema es de tipo *seccional*, se manejan varias opciones de comprobación, una de ellas es la *prueba t de Student* para dos muestras independientes. El primer paso se define encontrando si ambos conjuntos de datos tienen varianzas iguales (Rosner, 2010). La comprobación consiste en calcular la varianza para las muestras globales con tipo de vidrio G1 y G2, y aplicar la prueba de igualdad de dos varianzas. A partir del experimento se obtienen los valores de varianza, la *prueba F de Fisher*, y la probabilidad:

$$s_{G_1}^2 = 12.8016 \quad (3.1)$$

$$s_{G_2}^2 = 11.9323 \quad (3.2)$$

$$F = \frac{s_{G_1}^2}{s_{G_2}^2} = 1.0729 \quad (3.3)$$

$$p = 2 * 1 - P F_{n_1-1, n_2-1} > F = 0 \quad (3.4)$$

Como el valor de p es significativo, puesto que es menor a 0.05, se concluye que las varianzas son diferentes.

Al saber que las muestras tienen varianzas distintas, ahora debemos comprobar la hipótesis nula. Para esto se aplica la prueba t para dos muestras independientes con varianzas diferentes utilizando el *método de Satterthwaite* (Rosner, 2010).

El método de Satterthwaite, consiste en tres pasos:

1) Calcular

$$t = \frac{x_1 - x_2}{\sqrt{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}}} = -89.31 \quad (3.5)$$

2) Calcular el grado de libertad aproximado

$$d' = \frac{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}}{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}} = 101026.9 \quad (3.6)$$

3) Redondeando d' al entero inferior, se obtiene $d''=101026$.

Ahora, si $t > t_{d'', 1-\frac{\alpha}{2}}$ o $t < -t_{d'', 1-\frac{\alpha}{2}}$, entonces la hipótesis nula es rechazada, o en otras palabras, las medias son diferentes. Al contrario, si $-t_{d'', 1-\frac{\alpha}{2}} \leq t \leq t_{d'', 1-\frac{\alpha}{2}}$, entonces la hipótesis nula es aceptada, lo que significa que las medias son iguales. La Figura 3.2 muestra de manera gráfica la colocación de las regiones de aceptación y negación. El valor obtenido de t , se encuentra en la región de rechazo, puesto que es menor que 0. Por lo tanto, se niega la hipótesis nula, y se acepta la hipótesis alternativa, que indica que las medias pertenecientes a los lotes del tipo G1 y G2 son estadísticamente diferentes entre sí. Al encontrar que los lotes de tipo de vidrio 3 y G2 son diferentes estadísticamente, para el resto de la investigación todos los cálculos tomarán en consideración esta diferencia.

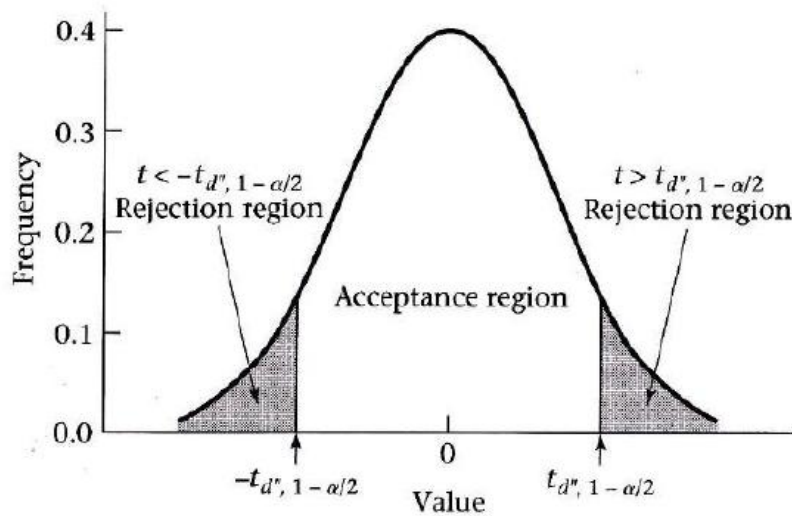


Figura 3.2. Región de aceptación de la hipótesis nula para las medias iguales (Rosner, 2010)

Para comparar un lote frente a otro, se necesitan los datos de la distribución de las piezas entre los repositorios (Figura. 3.3).

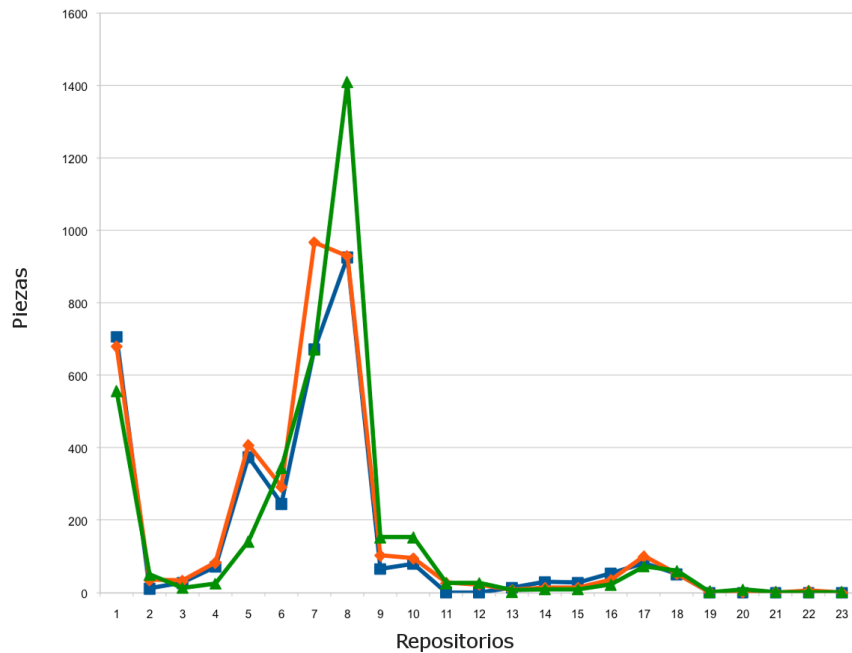
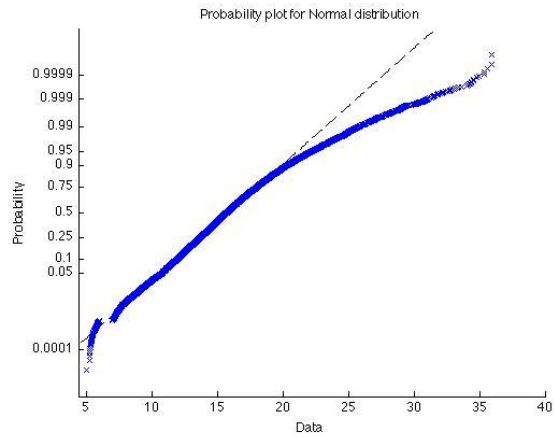


Figura 3.3. Un ejemplo de distribuciones reales de tres lotes entre los repositorios.

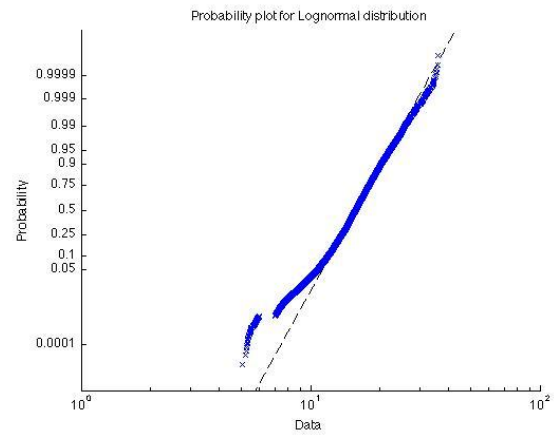
Para estimar la cantidad de piezas depositadas en los repositorios, cada lote fue analizado por separado, se realizó un análisis en los lotes buscando la distribución que le correspondía a proceso de clasificación de piezas. Para definir si la distribución de un lote entre los repositorios tenía semejanza con alguna distribución conocida, se realizó un experimento.

El tercer experimento consistió en comparar mediante una gráfica de residuos el acumulado de piezas ensambladas con un tipo de vidrio en específico, en contra de varias distribuciones conocidas, con el propósito de validar si se encontraba la semejanza.

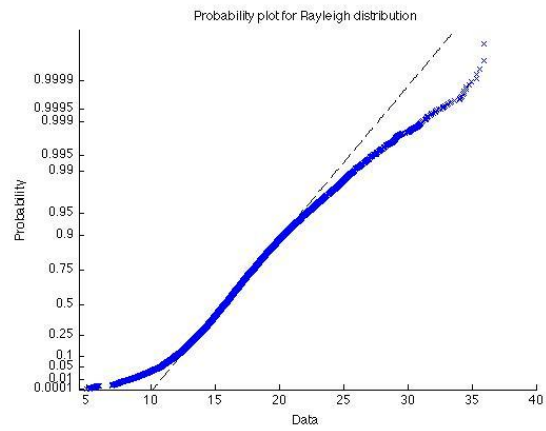
Como lo muestra la Figura 3.4, los lotes no se ajustan completamente a las siguientes distribuciones: *Normal*, *Lognormal*, *Weibull*, *Rayleigh*, *Extreme Value* y *Exponencial*, por lo tanto se consideró a la distribución de piezas entre los repositorios, como una distribución desconocida.



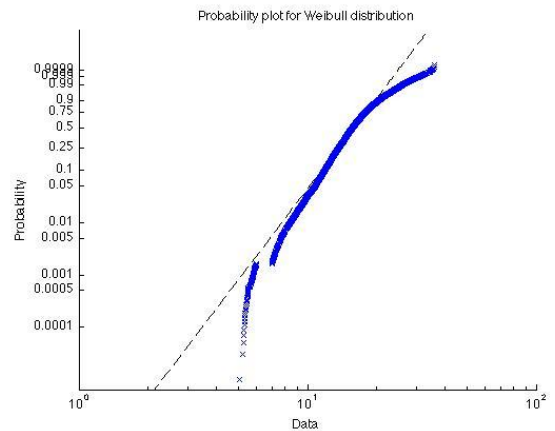
a)



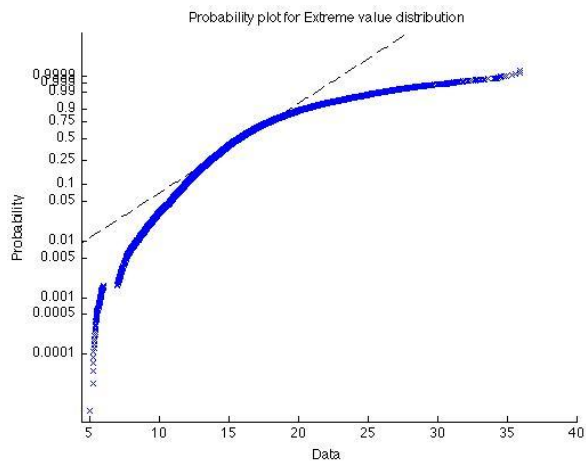
b)



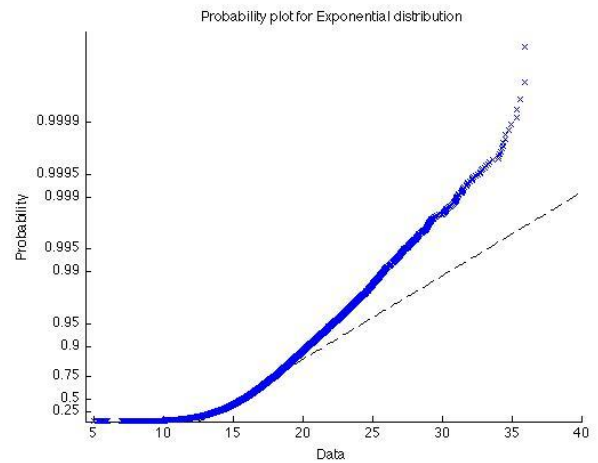
c)



d)



e)



f)

Figura 3.4. Las gráficas de residuos de la muestra global G1 contra diferentes distribuciones: a) Normal, b) Lognormal, c) Weibull, d) Rayleigh, e) Extreme Value y f) Exponencial.

Al no obtener resultados satisfactorios en definir una distribución, se optó por realizar una prueba de normalidad que indicará si la distribución de piezas entre repositorios corresponde a una distribución normal. La razón de usar la distribución normal como base es por el Teorema del Límite Central, indicando que toda variable aleatoria con un gran número de muestras, con su media y varianza se aproxima a una distribución normal (Monthgomery & Runger, 2001).

El cuarto experimento fue aplicar una *prueba de bondad de ajuste (chi-square goodness of fit)* (Rosner, 2010), en lotes con ambos tipos de vidrio obteniendo los siguientes resultados (Tabla 3.1):

Tabla 3.1. Prueba de bondad de ajuste

Lotes Reales	ChiStat	Valor P
Lote 1	132.47	6.73E-22
Lote 2	52.37	1.16E-06
Lote 3	99.89	4.97E-15
Lote 4	10576.00	0
Lote 5	677.71	3.03E-132
Lote 6	131.48	2.76E-19
Lote 7	290.34	2.93E-49
Lote 8	87.83	1.30E-13
Lote 9	31.20	0.0018
Lote 10	61.47	2.95E-07
Lote 11	282.19	1.23E-50
Lote 12	546.36	5.28E-106
Lote 13	193.45	2.21E-30
Lote 14	811.72	2.68E-163

En la tabla se observa que los valores p se acercan al cero, por lo tanto los datos no se ajustan a una distribución normal. En esta investigación la distribución de los datos entre los repositorios es crucial, ya que es importante saber qué distribución siguen los datos, y así utilizar esta información como un medio de estimación, conocer la manera en la que los datos se comportan y calcular las distribuciones de futuros lotes. Por esta razón es necesario encontrar la distribución a

la cual se ajustan los datos reales.

Al analizar el comportamiento de los datos, se encontraron dos razones que provocan la problemática en el ajuste de los lotes reales a una distribución normal:

- 1) El rango amplio del valor de operación en repositorios.
- 2) Dos mediciones involucradas en la clasificación.

La primera razón existe por los diferentes rangos del valor de operación en los repositorios. La consecuencia de la existencia de rangos amplios del valor de operación en los repositorios es la acumulación dispareja de piezas entre los repositorios. Para homogeneizar la distribución de piezas se introduce el término de *bin*. Así un *repositorio* se define como un conjunto de 35 bins, donde a cada bin le corresponde un valor de operación de rango unitario, es decir, el rango de una unidad del valor de operación. Por ejemplo, el repositorio *A* tiene un rango de operación de 15 a 20, así que este tiene cinco unidades (o bins) del valor de operación: 15, 16, 17, 18, y 19.

La segunda razón está relacionada con los niveles de la calidad de piezas. La medida del valor de calidad está dada por la medida del CR (*contact release*), tomando en cuenta este valor para la clasificación. Por ejemplo, repositorio *A* opera el rango de operación de 15 a 20 con el nivel CR igual a 1, repositorio *B* tiene el mismo rango de operación pero con el nivel 2, y el repositorio *C* opera rango de 30 a 35 con los tres niveles. Así, la distribución de piezas entre los repositorios no muestra un real comportamiento de datos.

3.2 Reclasificación

Para encontrar una distribución propia de la probabilidad, se ha creado un algoritmo que proporciona una solución para estos datos problemáticos. El algoritmo reclasifica los datos de la empresa utilizando el valor de operación y el valor de calidad CR, de tal manera, que en lugar de dar una distribución entre los repositorios, da la distribución entre los bins separados en niveles de calidad. Así, cada lote se separa en niveles y bins, para conocer el comportamiento real de los datos (Figura 3.5).

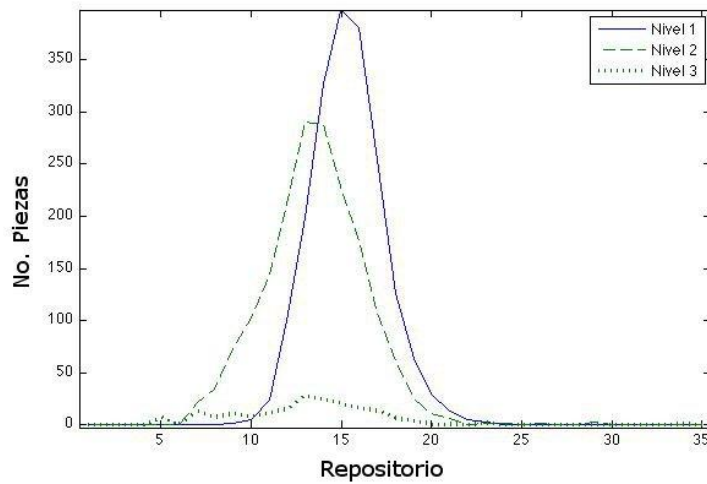


Figura 3.5. Distribución de datos reales de un lote entre 3 niveles CR y 35 bins, recibida por el algoritmo de reclasificación.

El algoritmo consiste en la simulación del comportamiento de una máquina clasificadora a partir de la información obtenida en el primer experimento. Al llegar un lote de tamaño N , entonces en cada pieza se mide el valor de operación y el valor de calidad CR. Al contar con estos valores, primero se valida el nivel de calidad de la pieza, después se compara el valor de operación de la pieza con los rangos de operación de los bins, y se coloca al que corresponda. Así se define que cada nivel contiene la misma cantidad de repositorios con un rango de operación unitario (Figura 3.6).

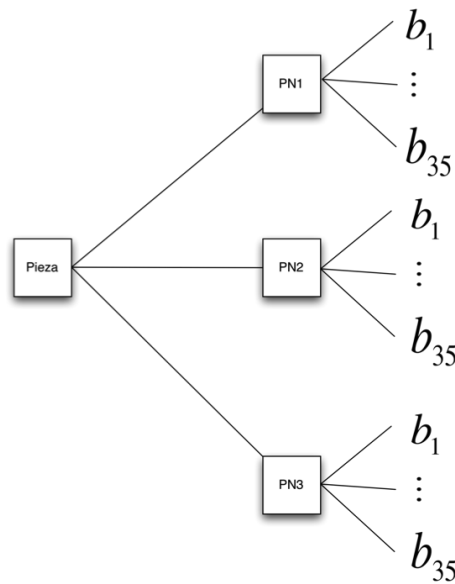


Figura 3.6. El diagrama de validación de una pieza por niveles de calidad y rangos unitarios de operación.

Como resultado de la reclasificación son obtenidas tres distribuciones con rangos de operación unitarios, donde cada distribución corresponde a un nivel de calidad. Teniendo estas distribuciones, se calculan la media y la desviación estándar para cada nivel.

3.3 Ajuste a la distribución normal

El siguiente paso, fue evaluar el acumulado de piezas por niveles, donde se encontró que cada nivel tiene su propias estadísticas, por lo tanto su propia distribución, así que se deben considerar como independientes. Debido a esta característica, se concluyó que cada nivel debe tomarse en cuenta por separado para encontrar la distribución que corresponde a un lote real. La comprobación de que la distribución de piezas de un lote separadas por niveles de calidad corresponde a una distribución normal, se basa en el Teorema de Limite Central (Monthgomery. & Runger, 2001).

En el quinto experimento, se utilizaron dos pruebas estadísticas para validar si los datos reclasificados se ajustan a una distribución normal. La primera prueba fue la de *bondad de ajuste* (*chi square goodness of fit test*) realizada con 95% de confianza. Consistió en la comparación de distribuciones para cada nivel de calidad con una distribución normal, y la comprobación si el modelo se ajusta o no a una distribución normal. Se utilizó el siguiente procedimiento (Rosner, 2010):

1. Los datos crudos de la empresa se dividen en g grupos.
2. Estimar los k parámetros del modelo de probabilidad a partir de los datos, donde k es el número de bins.
3. Utilizar las estimaciones recibidas en el paso 2 para calcular la probabilidad p y obtener un valor dentro de un grupo particular, la frecuencia esperada np correspondiente dentro de este grupo, donde n es el número total de puntos de datos.
4. Calcular

$$\chi^2 = (O_1 - E_1)^2 / E_1 + \dots + (O_i - E_i)^2 / E_i + \dots + (O_g - E_g)^2 / E_g, \quad (3.7)$$

Donde O_i y E_i son respectivamente los números observados y esperados de unidades dentro del

grupo i , $i = 1, \dots, g$ y g es el número de grupos.

5. Para una prueba con el nivel de significancia χ se rechaza la hipótesis nula H_0 si $\chi^2 \leq \chi_{g-k-1, 1-\alpha}$. Se acepta la hipótesis nula H_0 , si $\chi^2 > \chi_{g-k-1, 1-\alpha}$.

La segunda prueba es de *Jarque-Bera* (*Jarque-Bera test*), que se utiliza para muestras grandes. La prueba Jarque-Bera es una verificación de bondad de ajuste entre dos distribuciones (*two-sided-goodness-of-fit test*) apropiada cuando la distribución inicial (*a fully-specified null distribution*) es desconocida y sus parámetros deben ser estimados (Jarque &, 1987, pp. 163–172). La estadística de la prueba es:

$$JB = \frac{n}{6} \left(s^2 + \frac{k-3}{4} \right), \quad (3.8)$$

donde n es el tamaño de la muestra, s es la asimetría (*skewness*) y k es curtosis (*kurtosis*) de la muestra. Para tamaños grandes de la muestra, la estadística de la prueba tiene la distribución χ -cuadrada con dos grados de libertad. El experimento consistió en evaluar los niveles de calidad de manera independiente en cada lote reclasificado. Para un mejor ajuste se eliminaron valores extremos en los datos (Figura 3.7).

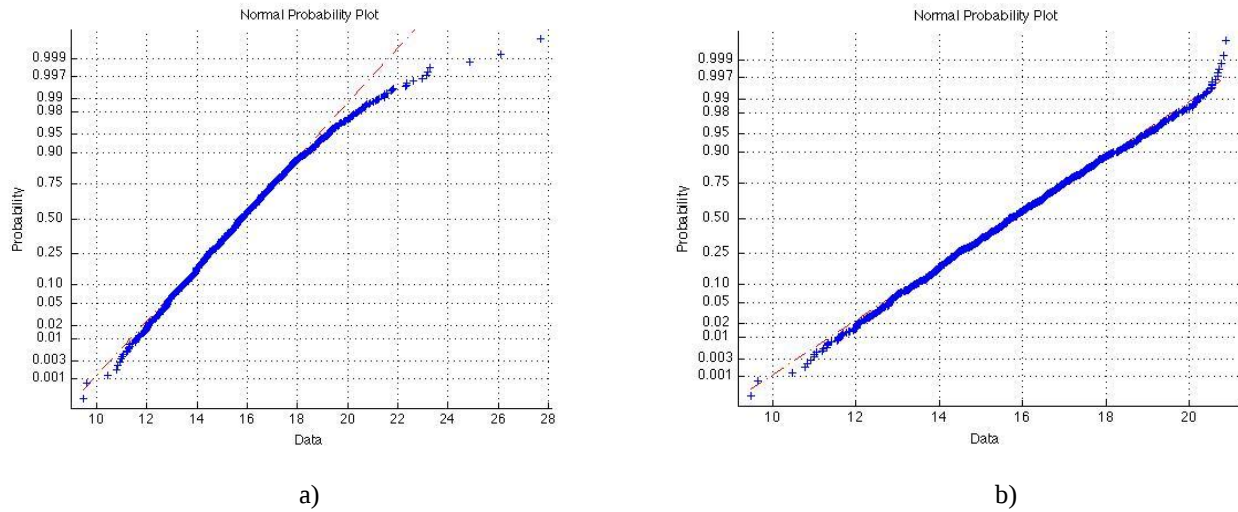


Figura 3.7. Comparación entre distribución de la empresa y distribución reclasificada. a) Empresa. b) Reclasificada.

Los resultados a estas pruebas para los 14 lotes y tres niveles de calidad (N_1 , N_2 y N_3) se muestran en las siguientes Tablas 3.2 y 3.3, las conclusiones fueron que en su mayoría los lotes siguen una distribución normal, considerando niveles de calidad y rangos unitarios de operación.

Tabla 3.2. Resultados de la prueba de bondad de ajuste a la distribución normal de los lotes reclasificados utilizando la prueba χ -cuadrada para tres niveles de CR: N_1 , N_2 y N_3 .

χ -cuadrada	N_1	N_2	N_3
Lote 1	0.2653	0.0771	0.1929
Lote 2	0.001	0.001	0.05
Lote 3	0.0022	0.228	0.2963
Lote 4	5.44E-19	7.46E-04	0.5327
Lote 5	0.025	1.05E-18	2.13E-09
Lote 6	0.33	4.17E-08	0.0271
Lote 7	5.89E-54	1.31E-15	0.3984
Lote 8	0.0154	0.4574	0.0451
Lote 9	0.2631	0.2797	0.0078
Lote 10	4.04E-07	0.1127	0.001
Lote 11	2.96E-11	8.73E-34	0.3189
Lote 12	1.74E-04	0.3784	0.0146
Lote 13	1.18E-35	4.57E-23	0.1026
Lote 14	1.07E-14	0.0955	0.5385

Tabla 3.3. Resultados de la prueba de bondad de ajuste a la distribución normal de los lotes reclasificados utilizando la prueba Jarque-Bera

Jarque-Bera	N_1	N_2	N_3
Lote 1	0.2653	0.0771	0.05
Lote 2	1.00E-03	0.0024	0.3803
Lote 3	0.0544	1.00E-03	0.2867
Lote 4	1.00E-03	0.0015	0.3752
Lote 5	0.5	1.00E-03	1.27E-02
Lote 6	0.0688	0.0852	0.3478
Lote 7	1.00E-03	1.00E-03	0.5
Lote 8	0.1709	1.00E-03	0.0255
Lote 9	0.0524	0.1748	0.4816
Lote 10	1.00E-03	0.318	0.2106
Lote 11	1.00E-03	1.00E-03	0.1465
Lote 12	1.00E-03	0.5	0.5
Lote 13	1.00E-03	1.00E-03	0.1808
Lote 14	1.00E-03	0.2333	0.2102

Los valores mostrados en las tablas, que son mayores a 0.05 (pintados de color), se consideran como no significativos, por lo tanto los lotes correspondientes se ajustan a la normal. Como se observa, la mayoría de distribuciones se ajustan a la distribución normal. Los lotes que no

pasaron la prueba, se le atribuye a dos razones: los valores atípicos (*outliers*) y que el lote sufre pérdidas por piezas defectuosas. La conclusión de este experimento se enfoca en saber que la mayoría de las distribuciones se ajustan a una distribución normal, por lo tanto, se toma esta como base para la estimación de piezas en los repositorios. El conocimiento si los datos se ajustan a una Gaussiana, permitió obtener la función de la densidad de probabilidad para cada nivel y así obtener los resultados del clasificado. Con una distribución correcta es posible estimar estrechamente el comportamiento de los datos con respecto a los repositorios.

3.4 Enfoque probabilístico

Al comprobarse que las distribuciones de un lote por niveles y bins, se ajustan a una distribución normal, el siguiente paso es la estimación de la cantidad de piezas en un repositorio para cualquier tipo de vidrio.

El evento que una pieza es depositada en un repositorio, es la información estadística que conocemos. Falta encontrar la probabilidad con que ocurra este evento. Se sabe que la cantidad de piezas que se depositen en un repositorio depende de tres factores:

- ✓ El tipo de vidrio utilizado.
- ✓ La cantidad de desechos en el lote.
- ✓ La probabilidad de que la pieza corresponda a ese repositorio.

Los dos primeros factores, son conocidos por la empresa previamente a la etapa de clasificación, pero el tercer factor, en ningún momento es tomado en cuenta para la planeación de lotes y menos para el cumplimiento de pedidos, ya que es un factor desconocido y que en investigaciones anteriores (Romero & Burtseva, 2010) no se ha tomado en cuenta. Para descubrir ese factor, es necesario entender el comportamiento de piezas en la clasificación desde un punto de vista probabilístico.

Debido a la reclasificación, cada nivel cuenta con sus propias estadísticas, es decir, existe una media y desviación estándar para la distribución de cada nivel, y dado que estas distribuciones se ajustan a una Gaussiana, se expresa de la siguiente manera:

$$F_{N1} \sim N_{x_1, s_3^2} = P(X_{N1} \leq x) \quad (3.9)$$

$$F_{N2} \sim N(x_2, s_3^2) = P(X_{N2} \leq x) \quad (3.10)$$

$$F_{N3} \sim N(x_3, s_3^2) = P(X_{N3} \leq x) \quad (3.11)$$

donde F_1 se define como la función de densidad de probabilidad para la distribución normal del CR nivel 1, así como F_{N2} y F_{N3} para los niveles CR 2 y 3, respectivamente.

Ya que se conocen los parámetros x_1 y s_1^2 , que son valores obtenidos del calculo de media y desviación estándar de los lotes para cada uno de los niveles, podemos definir un medio de estimación para las piezas contenidas en cualquiera de los niveles de calidad.

El siguiente paso es considerar los eventos que intervienen al momento de realizar la clasificación de piezas.

El primer evento que debemos identificar, se define como la pieza que se deposita en el bin b . Como la pieza se coloca únicamente en un repositorio, se dice que es mutuamente excluyente y se expresa de la siguiente manera:

$$e_{b_1} \cap e_{b_2} = \emptyset, \forall b, \quad (3.12)$$

donde e_{b_1} y e_{b_2} representan los eventos de que una pieza pertenece al bin 1 y bin 2 respectivamente.

Dado que a los bins corresponden los valores de operación de una unidad, una pieza no puede tener dos valores de operación; en otras palabras, la misma pieza no puede ser depositada en dos bins distintos.

La probabilidad de que una pieza sea depositada en un bin, se obtiene de la función de densidad de probabilidad F_{N1} , F_{N2} o F_{N3} dependiendo del nivel de calidad de la pieza, ya que existe una equivalencia entre bin y el valor de operación de la pieza.

Al ser una muestra grande, la probabilidad de que una pieza corresponda a un nivel, se calcula con la frecuencia relativa de ese nivel. Es decir,

$$P_{N1} = \frac{\text{Cantidad de piezas del Nivel 1}}{\text{Cantidad total de piezas}}; \quad (3.13)$$

$$P_{N2} = \frac{\text{Cantidad de piezas del Nivel 2}}{\text{Cantidad total de piezas}}; \quad (3.14)$$

$$P_{N_3} = \frac{\text{Cantidad de piezas del Nivel 3}}{\text{Cantidad total de piezas}}. \quad (3.15)$$

Ahora definimos el evento que una pieza está asignada a un bin. Para encontrar la probabilidad de que una pieza sea depositada en un bin, primero se debe saber que a un bin corresponde uno de los niveles. Sabiendo que cada nivel sigue una distribución normal con propia media y desviación estándar, se define que un bin es equivalente a un punto de la función de densidad normal del nivel a que pertenece. Conociendo el nivel y el bin al que le corresponde la pieza, es posible calcular la probabilidad de que esa pieza se deposite en ese nivel y ese bin.

El siguiente diagrama (Fig. 3.8) muestra el modelo de obtención de probabilidades para cualquier bin de cualquier nivel.

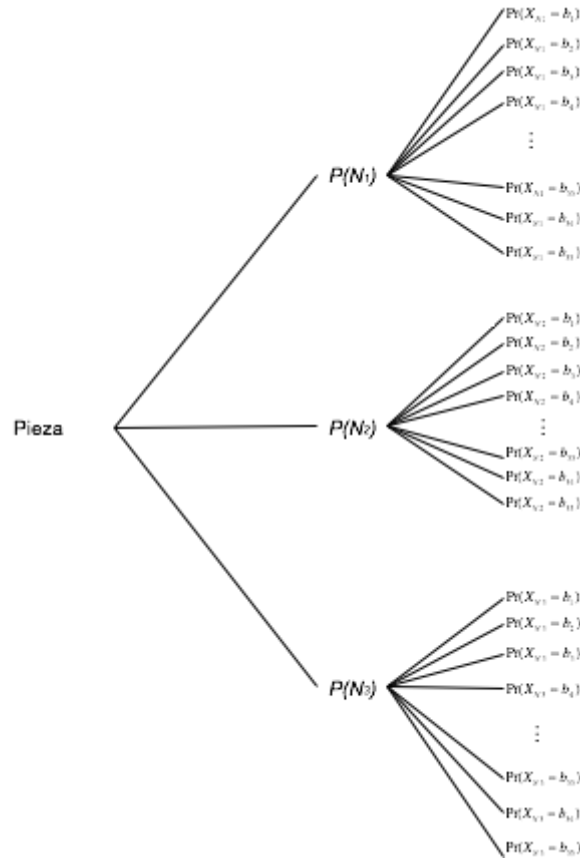


Fig. 3.8. Modelo de obtención de probabilidades.

La probabilidad de que la pieza se deposite en un repositorio, se calcula a partir de la función de densidad de probabilidad para la distribución normal correspondiente. Utilizando las funciones de densidad de probabilidad se calculan las probabilidades por repositorios que estiman la cantidad de piezas depositados en cada repositorio de la máquina considerando el rango de operación y nivel de calidad CR (Tabla 3.4).

Tabla 3.4. Las probabilidades por los repositorios

Repositorio	Bins	Nivel	Probabilidad
1	10 – 14	1	$P_{N_1} (0 < x < 15) \bar{P}(N_1)$
2	23 – 24	1	$P_{N_1} (3 < x < 25) \bar{P}(N_1)$
3	14 – 15	2	$P_{N_2} (4 < x < 16) \bar{P}(N_2)$
4	7	1,2	$P_{N_1} (x = 7) \bar{P}(N_1) + P_{N_2} (x = 7) \bar{P}(N_2)$
5	8	1,2	$P_{N_1} (x = 8) \bar{P}(N_1) + P_{N_2} (x = 8) \bar{P}(N_2)$
6	9 – 10	1,2	$P_{N_1} (x = 9) \bar{P}(N_1) + P_{N_2} (0 < x < 11) \bar{P}(N_2)$
8	11 – 13	2	$P_{N_2} (1 < x < 14) \bar{P}(N_2)$
9	15 – 19	1,2	$P_{N_1} (5 < x < 20) \bar{P}(N_1) + P_{N_2} (6 < x < 20) \bar{P}(N_2)$
10	21 – 22	1,2	$P_{N_1} (1 < x < 23) \bar{P}(N_1) + P_{N_2} (1 < x < 23) \bar{P}(N_2)$
11	20, 23 , 24	1,2	$P_{N_1} (x = 20) \bar{P}(N_1) + P_{L2} (3 < x < 25) \bar{P}(N_2) + P_{N_2} (x = 20) \bar{P}(N_2)$
12	26 – 28	1,2	$P_{N_1} (6 < x < 29) \bar{P}(N_1) + P_{N_2} (6 < x < 29) \bar{P}(N_2)$
13	25	1,2	$P_{N_1} (x = 25) \bar{P}(N_1) + P_{N_2} (x = 25) \bar{P}(N_2)$
14	7	3	$P_{N_3} (x = 7) \bar{P}(N_3)$
15	8	3	$P_{N_3} (x = 8) \bar{P}(N_3)$
16	9 – 10	3	$P_{N_3} (0 < x < 11) \bar{P}(N_3)$
18	11 – 14	3	$P_{N_3} (1 < x < 15) \bar{P}(N_3)$
19	15 – 19	3	$P_{N_3} (5 < x < 20) \bar{P}(N_3)$
20	21 – 22	3	$P_{N_3} (1 < x < 23) \bar{P}(N_3)$
21	20, 23 – 24	3	$[P_{N_3} (x = 20) + P_{N_3} (3 < x < 25)] \bar{P}(N_3)$
22	25 – 29	3	$P_{N_3} (5 < x < 30) \bar{P}(N_3)$
23	5	3	$P_{N_1} (x = 5) \bar{P}(N_1) + P_{N_2} (x = 5) \bar{P}(N_2) + P_{N_3} (x = 5) \bar{P}(N_3)$

24	30 – 33	1,2,3	$P_{N_1} \left(\left(0 < x < 34 \right) \left(\left(V_1 \right) \right) \right) + P_{N_2} \left(\left(0 < x < 34 \right) \left(\left(V_2 \right) \right) \right) + N_{N_3} \left(\left(0 < x < 34 \right) \left(\left(V_3 \right) \right) \right)$
25	34 – 35	1,2,3	$P_{N_1} \left(\left(4 < x < 36 \right) \left(\left(V_1 \right) \right) \right) + P_{N_2} \left(\left(4 < x < 36 \right) \left(\left(V_2 \right) \right) \right) + P_{N_3} \left(\left(4 < x < 36 \right) \left(\left(V_3 \right) \right) \right)$

Con los resultados obtenidos a partir de función de densidad de probabilidad, se construyen las distribuciones por repositorios para cada tipo de lote. Esto significa que la cantidad de piezas en repositorios es diferente entre lotes que fueron ensamblados con un tipos de vidrio debido a que los tipos de vidrio entre sí son independientes (propia media y desviación estándar). Con este hecho, se obtuvo la posibilidad de estimación de la cantidad necesaria de lotes a través de evaluación de los resultados ajustados de clasificado y la cantidad de piezas en repositorios. Por lo tanto, cuando se obtiene la minimización de la desviación entre lo calculado y lo obtenido, se facilita la reducción de los niveles de inventario, y así el WIP se mantiene en niveles bajos.

En la Tabla 3.5 se muestran los porcentajes de cada tipo de vidrio estudiado, obtenidos a partir de la función de densidad para los 25 repositorios.

Tabla 3.5. Cantidad de piezas en % por repositorios para tipos de vidrio G1 y G2.

Repositorio	G1, %	G2, %
R1	24.168	13.258
R2	0.026	2.990
R3	12.44	8.76
R4	0.732	0.441
R5	1.568	0.845
R6	5.465	3.258
R7	0.000	0.000
R8	14.786	9.619
R9	32.48	44.99
R10	0.704	4.146
R11	1.714	5.883
R12	0.000	0.162
R13	0.010	0.269
R14	0.039	0.094
R15	0.083	0.140
R16	0.422	0.457
R17	0.000	0.000
R18	2.253	1.619
R19	2.028	1.932
R20	0.058	0.255

R21	0.095	0.317
R22	0.001	0.048
R23	0.440	0.409
R24	0.000	0.003
R25	0.000	0.000

Obteniendo estas estadísticas, el siguiente paso es encontrar la manera de aproximar la cantidad de futuros lotes, basándose en esta función de densidad de probabilidad para cada tipo de vidrio.

3.5 Conclusiones del capítulo

Se realizaron 5 experimentos:

1. Se recopilaron muestras reales de la empresa de 18 lotes de tipo G1 y 20 lotes de tipo G2, considerando 2 tipos de vidrio más frecuentes en la empresa y 25 repositorios para cada máquina. Se formaron dos muestras globales, G1 y G2.

2. Se compararon dos muestras globales. Se formuló la hipótesis nula: medias iguales. Se aplicó primero la prueba F de Fisher la cual es la prueba de igualdad de dos varianzas. Se concluyó que las varianzas son estadísticamente distintas. Utilizando el método de Satterthwaite para aplicar la prueba t de Student para dos muestras independientes, se rechazó la hipótesis nula, es decir, las medias de G1 y G2 son diferentes, y los lotes pueden ser analizados por separado.

3. Se compararon las distribuciones de G1 y G2 contra la Normal, Lognormal, Weibull, Rayleigh, Extreme Value y Exponencial mediante las gráficas de residuos, y se concluyó que los datos reales no se ajustan a ninguna distribución conocida y por tanto se definen como una distribución desconocida. Para comprobar los resultados obtenidos con las gráficas de residuos se realizó una prueba de bondad de ajuste que confirmo la conclusión sobre una distribución desconocida.

4. Se concluyó que las distribuciones de piezas entre repositorios no reflejan correctamente los resultados de mediciones en la operación de clasificado. Se formaron virtualmente los bins, cada uno con un rango de operación de una unidad en lugar de repositorios con rangos amplios y

diferentes. Se reclasificaron las piezas de cada lote entre los bins de acuerdo al nivel de calidad CR y el valor de operación, utilizando la información que acompaña cada lote. Se desarrolló un algoritmo de reclasificación.

5. Se utilizaron dos pruebas estadísticas para validar si los datos para cada nivel CR se ajustan a una distribución normal. La primera prueba fue χ^2 cuadrada bondad de ajuste realizada con 95% de confianza para distribuciones reclasificadas. La segunda prueba fue de Jarque-Bera, que se utiliza para muestras grandes cuando la distribución inicial es desconocida y sus parámetros deben ser estimados. Se evaluaron los niveles CR de manera independiente en cada lote reclasificado. Se concluye que en su mayoría los lotes siguen una distribución normal, considerando niveles de calidad y rangos unitarios de operación, por lo tanto, se toma esta como base para la estimación de piezas en los repositorios.

A partir de los 5 experimentos se concluye que las distribuciones de lotes de dos tipos son estadísticamente diferentes y los datos se ajustan a una Gaussiana. Este conocimiento permitió obtener la función de la densidad de probabilidad para cada nivel CR y así obtener los resultados del clasificado y estimar estrechamente el comportamiento de los datos con respecto a los repositorios sin considerar las piezas defectuosas.

Se analizó la probabilidad de que una pieza sea depositada en un bin. Se obtuvo la función de densidad de probabilidad y se estimó la cantidad de piezas en cada repositorio.

4. PLANIFICACIÓN DE LA PRODUCCIÓN EN LA OPERACIÓN ESTOCÁSTICA DE CLASIFICADO

Teniendo el conocimiento acerca del ajuste de distribución que revela el modo de estimar el número de piezas por repositorios a partir de un lote, el siguiente paso es el determinar la cantidad mínima de los mismos. La meta principal de la investigación es utilizar el número mínimo de lotes para cumplir con una lista de pedidos. La optimización es condicionada por las diferencias en distribuciones de varios tipos de lotes y los resultados estocásticos de la clasificación. El uso eficiente de material implica que el inventario debe mantenerse en un mejor nivel. Con el inventario controlado, el WIP es más bajo, así el sistema JIT se aplica de un mejor modo dentro de la empresa.

El problema es resuelto en dos etapas: estimación y ajuste (Marti, 2008). En la primera etapa, los pedidos se unen, y entonces se estima la cantidad necesaria de lotes. Así se obtiene una solución inicial, o una pre-planificación. En la segunda etapa, la solución se ajusta a la situación real de una manera iterativa, después del procesamiento de cada lote y utilizando un criterio de re-planificación.

4.1 Estimación de cantidad de lotes

Esta etapa es donde se calcula la cantidad estimada de lotes aprovechando diferencias en distribución esperada para cada tipo de lote. Esta distribución se revela como la cantidad de piezas depositadas en cada repositorio a partir de un lote. A cada tipo de lote corresponde una distribución propia. El resultado de la etapa es la cantidad mínima de lotes de cada tipo que cumple con todos los pedidos, sin disturbios provocados por el carácter estocástico de los resultados en las mediciones.

Se crea una lista de pedidos a completar. Cada pedido tiene su propia fecha límite de terminación, un total de piezas, y un número de parte. La secuencia de los pedidos se basa en la demanda que sea de mayor tamaño; la cantidad de piezas a producir es necesaria para calcular el número de lotes; y el número de parte indica repositorio donde se tomarán las piezas. El total de piezas sólo se obtiene a partir de las piezas depositadas en repositorios marcados en número de

parte del pedido. Esta información se utiliza para estimación y secuenciación de los lotes de producción. A continuación se introduce un modelo que permite estimar las cantidades de lotes de cada tipo.

4.1.1 Notaciones

Índices

j	pedido, $j = 1, \dots, N$;
r	repositorio, $r = 1, \dots, R$;
i	nivel CR, $i = 1, 2, 3$;
g	tipo de lote, $g = 1, \dots, G$;
m	máquina, $m = 1, \dots, M$;
b	bin, $b = 1, \dots, B$;
l	lote, $l = 1, \dots, L$.

Entrada del modelo

$K = [k_{gr}]$	Matriz de distribución de piezas cuyo elemento k_{gr} es la cantidad de piezas depositadas en el repositorio r a partir de un lote de tipo g .
D_{jr}	Tamaño de pedido j que se cumple a partir del repositorio r .

Variables de decision

n_g	Cantidad de lotes de tipo g .
-------	---------------------------------

4.1.2 Modelo de la Programación Lineal

A continuación se describe el modelo que busca la cantidad mínima de lotes para el cumplimiento de los pedidos:

$$L = \sum_{g=1}^G n_g \rightarrow \min \quad (4.1)$$

$$\text{s.a} \quad \sum_{g=1}^G n_g k_{gr} \geq \sum_{j \in \{1..N\}} D_{jr}, \quad r = 1, \dots, R \quad (4.2)$$

$$R \gg G; \quad (4.3)$$

$$\{n_g, k_{gr}, D_j\} \in \Sigma^0, \quad \forall j, g, r. \quad (4.4)$$

$$n_g \geq 0, \quad \forall g. \quad (4.5)$$

Normalmente, el tamaño de un pedido es mucho mayor que el tamaño de un repositorio, por lo tanto, un pedido se asigna a un repositorio por partes (hasta el tope del pedido) y se cumple a partir de varios lotes. Consecutivamente, un lote representa un grupo de piezas asignadas para el cumplimiento de uno o más pedidos, los cuales se procesan simultáneamente dentro de un lote. De otro lado, un pedido es asignado a un solo repositorio, entonces se permite agrupar los pedidos recibidos del mismo repositorio (con el mismo número de parte). La restricción (4.2) de R desigualdades y G incógnitas describe la relación entre la producción y la demanda por repositorios: la demanda no debe ser mayor de la producción. De acuerdo a las condiciones reales de la manufactura, $R \gg G$, es decir, el número de desigualdades es mayor que el número de incógnitas (4.3). Los n_g , k_{gr} , D_{jr} , son los números enteros no negativos (4.4). La función a minimizar es $L = n_1 + \dots + n_g + \dots + n_G$, donde las incógnitas no tienen coeficientes puesto que no hay costos (4.1).

El modelo representa un Problema de Programación Lineal (Kolman & Beck, 1995). La solución del problema (4.1) se utiliza para la colocación de pedidos a los lotes.

Lema 1: El número mínimo de lotes corresponde a una cantidad mínima de piezas en el inventario.

Se comprueba con un simple cálculo.

Lema 2: La cantidad mínima de lotes minimiza el tiempo de procesamiento de los pedidos makespan.

Lema 3: La cantidad de lotes en la solución del problema (4.1) es invariante al orden de procesamiento de pedidos.

Como indica el Lema 3, la resolución del modelo sólo permite encontrar la cantidad de lotes

por tipos, y no indica un orden de procesamiento de los lotes ni los pedidos. Estos órdenes se definen a través de las siguientes reglas.

4.1.3 Orden de procesamiento de pedidos

Los pedidos se procesan de acuerdo a las siguientes suposiciones:

- Los pedidos se arreglan de acuerdo a prioridad en una lista para definir los cuales que se procesan primero. Estas prioridades indican la secuencia para completar los pedidos y se calculan utilizando algunas reglas de prioridad basadas en la demanda mayor.
- El orden de asignación de las piezas sigue el orden de pedidos en la lista, es decir, primero se cumplen los pedidos de la prioridad más alta.
- Una vez empezando, al primer pedido de la lista se le asignan piezas hasta terminar el mismo.
- Un pedido completo se excluye de la lista.
- Además de cumplimiento de primer pedido de la lista, se descartan piezas de subsecuentes pedidos que se cumplen a partir de otros repositorios de la máquina.
- En siguiente turno se procesa el lote con el tipo de material que beneficia más al primer pedido de la lista en el caso de elección entre varios tipos.

4.1.4 Resolución del modelo

En la Sección 4.1.2, se definió un modelo matemático para la estimación de lotes. Resolviendo inecuaciones, se encuentra la cantidad mínima de lotes necesarios para cumplir con los pedidos. El siguiente paso es dar una solución a las inecuaciones. Cada inecuación debe ser cumplida para asegurar que todas las demandas se completarán a tiempo. Ya que es clave la solución de las inecuaciones, la manera de calcular el valor mínimo de lotes cada vez toma mayor fuerza y se vuelve fundamental. Es importante aclarar que el problema no se resuelve en su totalidad con las técnicas de la programación matemática, por lo tanto, sólo es una solución inicial en la etapa de estimación.

Debido al modelo (4.1), es necesario aplicar técnicas de la programación lineal para calcular la cantidad de lotes de cada tipo. La programación lineal ha surgido como un medio de

optimización en varios ámbitos; en el ámbito de la industria, un problema se basa en la optimización de recursos. Estos recursos pueden ser mano de obra, material, insumos, etc. En nuestro estudio lo que se optimiza es el material y, por consecuencia, el tiempo dedicado en la producción de interruptores eléctricos, ya que al minimizar la cantidad de lotes requeridos, los pedidos se cumplen en un tiempo mínimo (Lema 2). Un escenario de esta problemática se presenta a continuación utilizando los términos de la programación lineal. Definiremos cada parte de los elementos de un problema de programación lineal desde el punto de vista del problema.

Dado que la etapa de estimación tiene características de un problema de minimización, usando las fórmulas de programación lineal, este se describe como (Kolman & Beck ,1995):

$$\mathbf{z} = \mathbf{c}^T \mathbf{x} \rightarrow \min \quad (4.6)$$

$$\text{S. a:} \quad \mathbf{Ax} \geq \mathbf{b}; \quad (4.7)$$

$$\mathbf{x} > 0; \quad (4.8)$$

donde

- z** es el valor a minimizar;
- c** es un vector que indica el costo por tipo de lote, en nuestro caso el costo no existe y por lo tanto es igual a 1;
- x** es el vector de tamaño G que contiene la cantidad de lotes a minimizar por tipos.
- A** es una matriz $R \times G$ que determina la cantidad estimada de piezas en un repositorio para cada tipo de vidrio;
- b** es el vector de tamaño R de las cantidades de piezas necesarias para cumplir con los pedidos, por repositorios.

La expresión (4.7) representa un sistema de restricciones del problema. Usando esta definición, se buscan las soluciones factibles para la función objetivo **z**, la cual se requiere minimizar. Las inecuaciones o restricciones se obtienen a partir de las distribuciones esperadas y de tamaños de los pedidos.

Por razones de simplicidad se utiliza el método SIMPLEX, para la búsqueda de soluciones en una región factible. El algoritmo SIMPLEX consiste de dos pasos (Kolman & Beck, 1995):

- (1) Una manera de averiguar si una determinada solución básica factible es una solución optima, y

(2) Una forma de obtener una solución básica factible adyacente con el mismo o mejor valor para la función objetivo.

Se tiene el siguiente ejemplo ilustrativo:

En términos del problema estudiado se define como $z = x_1 + x_2$, donde x_1 y x_2 son la cantidad de lotes de tipos G1 y G2 respectivamente.

Se ingresa un conjunto de demandas.

$D_{1,8}$: 6000, Valor de operación: 10 – 15, Repositorio $r = 8$.

$D_{2,9}$: 12000, Valor de operación: 15 -20, Repositorio $r = 9$.

$D_{3,10}$: 3000, Valor de operación: 20 - 25, Repositorio $r = 10$.

La matriz A obtenida a partir de las funciones de densidad de probabilidad es como sigue a continuación:

$$A = \begin{bmatrix} 451 & 281 \\ 1729 & 1759 \\ 483 & 281 \end{bmatrix} \quad x = \begin{bmatrix} x_1 \\ x_2 \end{bmatrix}$$

Se crean las inecuaciones:

$$541x_1 + 281x_2 \geq 6000$$

$$1729x_1 + 1759x_2 \geq 12000$$

$$483x_1 + 281x_2 \geq 3000$$

$$x_1 \geq 0, x_2 \geq 0.$$

Contamos con nuestro sistema de inecuaciones y las restricciones, ahora se utiliza el método SIMPLEX, para calcular la cantidad mínima de lotes necesarios y cumplir con los pedidos. En la Figura 4.1 se observa la región factible de la solución recibida a partir de las inecuaciones.

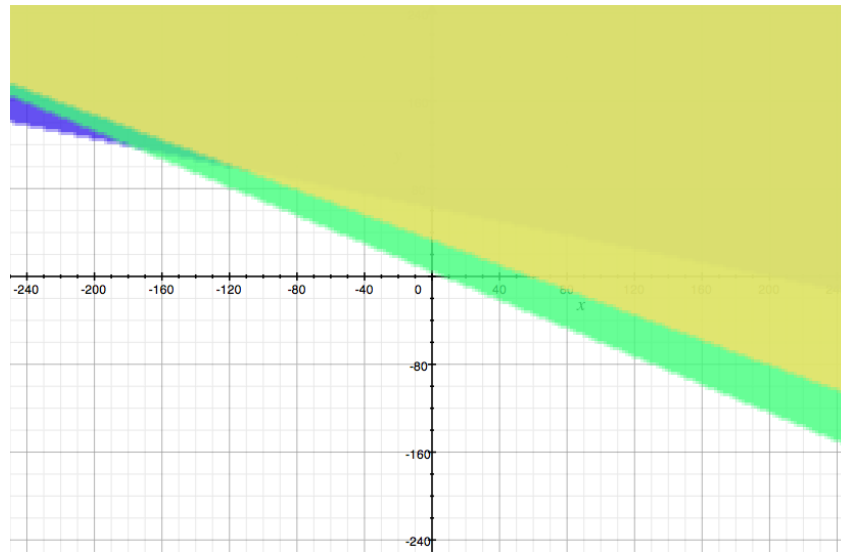


Figura 4.1. Representación gráfica de la región factible de la solución para el ejemplo.

Utilizando módulo de programación lineal de MATLAB, se recibió el resultado: 9 lotes del tipo G1 y 5 lotes del tipo G2.

Obteniendo el resultado, se deberán ir cumpliendo los pedidos de acuerdo al orden de la lista. El siguiente paso es el ajustar los lotes calculados con las distribuciones reales a las esperadas.

4.2 Ajuste de la aproximación al proceso estocástico

4.2.1 Análisis de desviaciones

Debido a la naturaleza del proceso de clasificación, la estimación de lotes no es suficiente para alcanzar el objetivo.

Al momento de clasificar un lote, se sabe que nunca se obtendrá la misma distribución en los repositorios, por lo tanto, de forma constante existirán desviaciones tanto positivas como negativas con respecto a las distribuciones esperadas. En la Figura 4.2, se muestra un ejemplo de la distribución de un lote, comparando los valores esperados con los obtenidos, y las desviaciones positivas y negativas. El diagrama de barras representa una muestra de la distribución de un lote; con la línea roja se traza la distribución esperada. El color verde indica sobreproducción, o la desviación positiva, mientras con el azul claro está pintada la falta de las piezas recibidas, en comparación con las esperadas, o la desviación negativa.

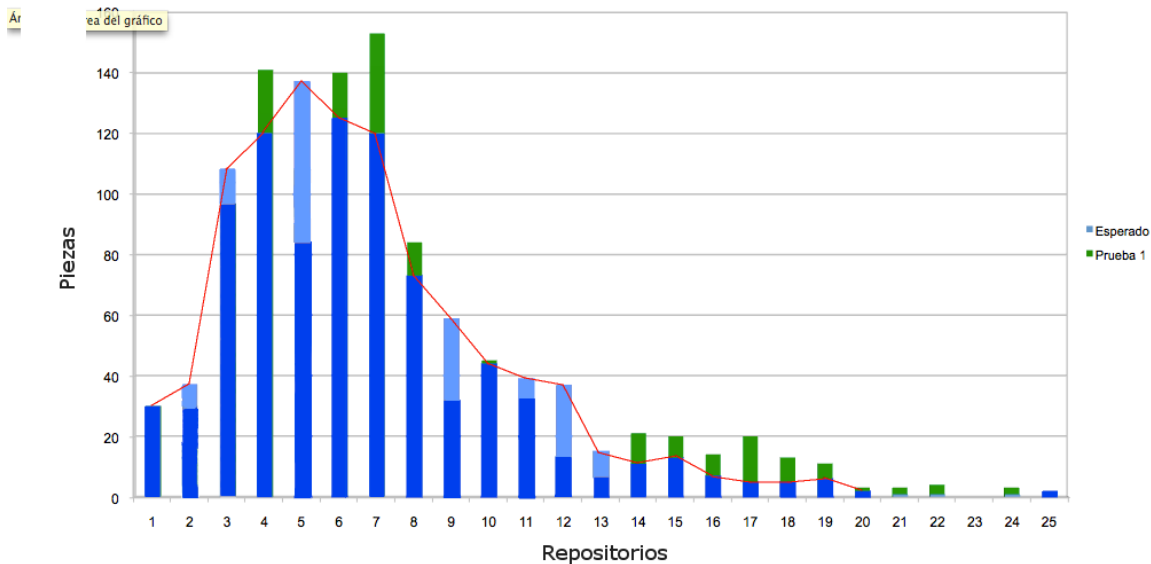


Fig. 4.2. Comparación de la distribución esperada con la distribución de una muestra

Las desviaciones significan una diferencia entre la cantidad de piezas esperadas contra las obtenidas en un repositorio; estas diferencias se vuelven de suma importancia para asegurar el cumplimiento de pedidos. Se considera a la desviación positiva como el WIP, ya que son piezas que no fueron esperadas en el cálculo. El problema de tener demasiada desviación positiva es el acumulamiento de piezas. Aunque no es un problema grave, debido a que estas piezas se utilizan para futuros pedidos, es preferible disminuir esta desviación, debido a la filosofía de la empresa, de acuerdo a la cual la disminución de la desviación positiva, indica un mejor aprovechamiento de material.

Por otro lado, la desviación negativa trae diversos problemas, el más significativo es la falta de piezas para el cumplimiento de uno o más pedidos. Al existir una desviación negativa, el valor obtenido es menor o mucho menor al esperado, por lo tanto en esos casos, se debe hacer correcciones en la planeación de los lotes para poder recuperar las piezas faltantes. Esto en caso recurrente, se vuelve un grave problema, ya que al existir en los repositorios una desviación negativa constante, lo que sucederá es la elaboración de nuevos lotes no contemplados, lo que trae como consecuencias la pérdida de tiempo, atrasos en los otros pedidos, y además en mayor grado acumulación del WIP en repositorios que no son necesarios. Por lo tanto, las desviaciones negativas, se vuelven un punto crucial para el cumplimiento de los pedidos.

4.2.2 Criterio de ajuste

Para la corrección de desviaciones negativas, se propone como un caso extremo volver a planificar los lotes, pero solo en situaciones cuando no es posible completar el pedido con los lotes restantes. En ese caso se deben considerar el siguiente criterio.

Sean:

L' el numero esperado de lotes, recalculado en proceso de ejecución de pedidos;

l' el número de lotes procesados, $l = \{1, \dots, l', \dots, L'\}$;

$(L' - l')$ el numero de lotes faltantes;

D_{jr} el número de piezas en el primer pedido de la lista;

$D_{jr, \text{acumulado}}$ la cantidad de piezas acumuladas en el pedido j a partir de l' lotes;

$D_{jr, \text{resto}}$ la cantidad de piezas faltantes en el pedido j después de procesar l' lotes.

Entonces, según estas definiciones:

$$D_{jr} = D_{jr, \text{acumulado}} + D_{jr, \text{resto}}.$$

Criterio: Se hace recalcu lo si $[(L' - l')k_{gr}]$ piezas son insuficientes para cumplir con $D_{jr, \text{resto}}$:

$$D_{jr, \text{resto}} < (L' - l')k_{gr}.$$

De tal manera, el criterio debe tomar en consideración:

- El número de lotes procesados;
- El número de lotes a procesar;
- El estado de los pedidos, es decir, cuantas piezas han sido registradas.
- El estado del WIP en todos los repositorios.

4.3 Algoritmo LCA

4.3.1 Estructura del algoritmo

La solución propuesta que ayuda a minimizar la desviación entre la cantidad de piezas esperada y la acumulada a través de una mejor elección de tipo de lote en cada iteración es un algoritmo que estima el número de lotes faltantes utilizando las probabilidades por repositorio. El algoritmo LCA (*Lot Calculation Algorithm*) consta de dos etapas: estimación y ajuste (Figura 4.3). La etapa de estimación se enfoca en dar una solución inicial, mientras la segunda, la etapa de ajuste, se enfoca en resolver la parte estocástica del problema, esto quiere decir, que corrige la desviación entre el valor esperado y el valor obtenido de piezas. Esta desviación se encuentra en los repositorios marcados por los pedidos. Así los resultados del algoritmo ayudan principalmente en dos áreas del proceso de producción, la parte de planeación con la estimación de lotes, y la parte del uso eficiente de materiales, con la etapa de ajuste. El algoritmo, además mejora el control de inventario, ayudando en el área de planificación. El algoritmo LCA está planteado y descrito a continuación, como la secuencia de dos etapas: la estimación y el ajuste.

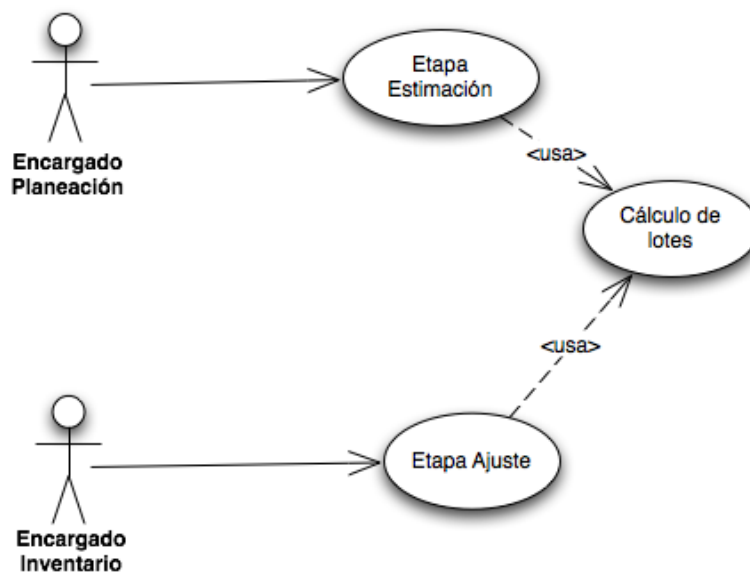


Fig.4.3. Diagrama de casos.

4.3.2 Etapa de estimación

En secciones anteriores, se ha comentado la problemática existente en la planeación de lotes. La etapa de estimación se basa en el modelo matemático y la función de densidad de probabilidad de los repositorios calculada en el Capítulo 3. En la Figura 4.4 se observan las fases que el algoritmo realiza para la estimación de lotes.

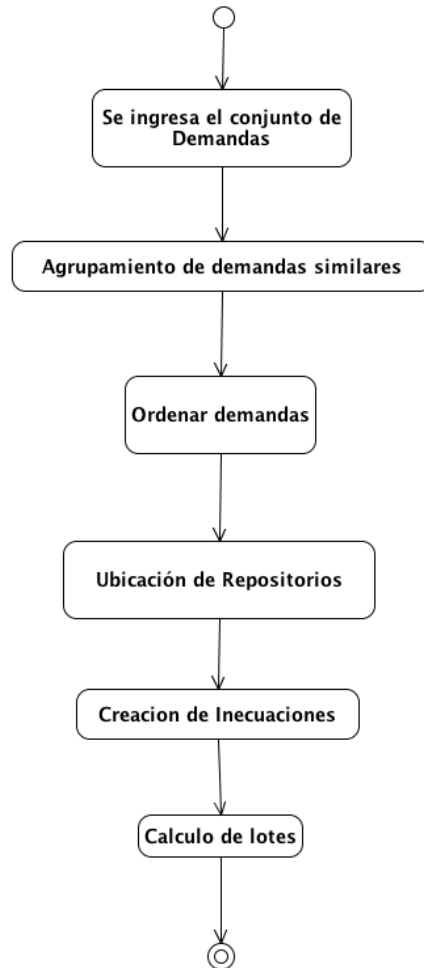


Figura 4.4. Fases de la etapa de estimación del algoritmo

Para entender cada una de las fases, estas se explican a continuación:

- *Se ingresa el conjunto de pedidos.* En esta parte del algoritmo, el encargado del inventario ingresa el total de pedidos solicitadas por los clientes. La información que debe tener una demanda incluye: la fecha de entrega, la cantidad de piezas, y el número

de parte.

- *Agrupamiento de demandas similares.* En esta fase, el algoritmo toma el conjunto de demandas y las agrupa en base al número de parte. Se toma el número de parte ya que es el indicador de los repositorios que se van a utilizar para el cumplimiento de pedidos.
- *Ordenar demandas.* En base a la regla de mayor número de piezas, se ordenan los pedidos para el momento que se deban cumplir. Esto se hace en base al total de piezas de los pedidos.
- *Ubicación de repositorios.* En esta fase, se localizan los repositorios de los cuales se toman piezas para cada pedido.
- *Creación de inecuaciones.* Debido a que el modelo matemático para la estimación de lotes está basado en un conjunto de inecuaciones lineales, en esta fase, se desarrollan estas inecuaciones utilizando como parámetros los valores de las funciones de densidad para los repositorios marcados en el número de parte del pedido. Las incógnitas de las inecuaciones representan los números de los lotes de cada tipo, tomando en cuenta que para este estudio se manejan dos tipos de vidrio.
- *Cálculo de lotes.* En esta fase, se toman las inecuaciones creadas y utilizando un método de programación lineal basado en SIPLEX, se obtiene la cantidad mínima de lotes. Es importante recalcar que esta es una solución inicial, ya que es un problema estocástico, y la cantidad de piezas que cae en cada repositorio varía de un lote a otro, por lo tanto, en su momento, la estimación puede quedarse no suficiente para el cumplimiento de los pedidos.

4.3.3 Etapa de ajuste

En la etapa de ajuste se busca disminuir las desviaciones negativas de piezas en los repositorios de los lotes. Esto se realiza, comparando la distribución del lote obtenida contra la esperada, es decir, por cada repositorio comparar la cantidad de piezas obtenidos contra la cantidad de piezas esperadas. Si en base a las estimaciones, las piezas serán insuficientes es necesario realizar el ajuste en la cantidad de lotes. En caso de ser necesario el ajuste, se modifica la estimación inicial de lotes, volviendo a calcular la cantidad de lotes para el resto de los pedidos.

Para realizar estas acciones correctivas se recaba la información de los lotes que han sido procesados. En base a estos datos, se toman en cuenta varios factores para la conclusión de recalcular el número de los lotes. En la Figura 4.5 se muestran las fases de la etapa de ajuste.

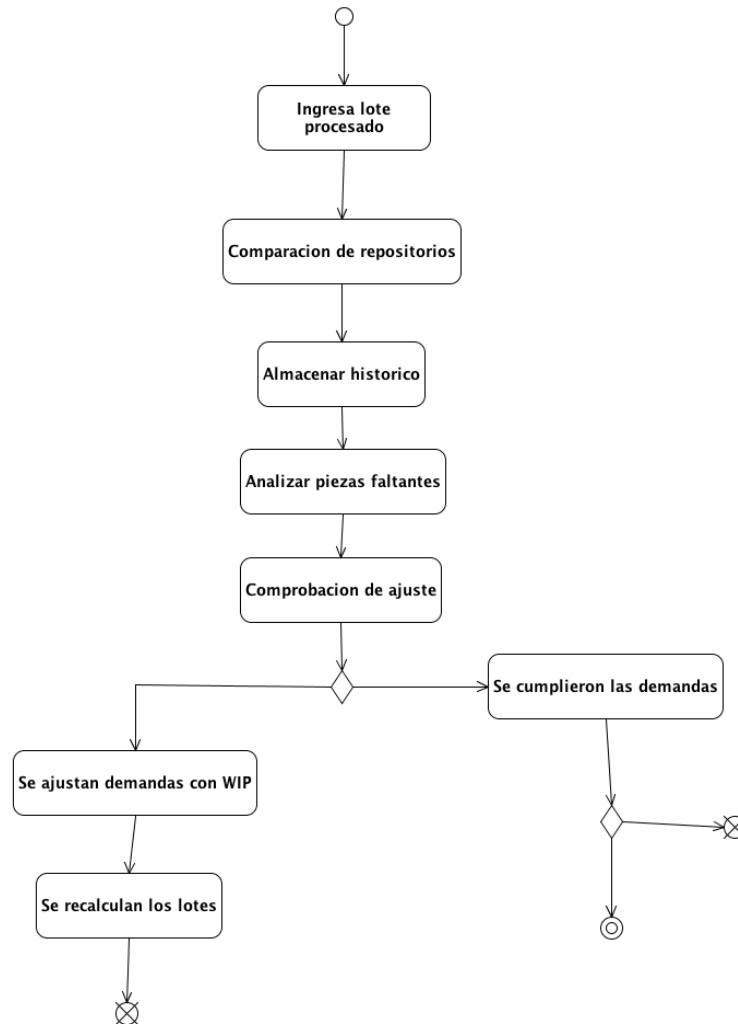


Figura 4.5. Fases de la etapa de ajuste del algoritmo

Las fases de la etapa de ajuste se explican a continuación:

- *Ingreso lote procesado.* En esta fase, el lote recién procesado se ingresa en el algoritmo donde lleva a una revisión general del estado actual de las demandas procesadas. Los datos capturados son todas las piezas clasificadas. La información requerida por pieza es: valor de operación, valor de calidad y repositorio donde fue depositada.
- *Comparación de repositorios.* Se analiza la desviación entre los repositorios necesarios para el pedido. Esta desviación se obtiene entre la cantidad de piezas esperada y obtenida.

- *Se guarda el histórico.* Se almacenan los valores esperados, obtenidos para todos los repositorios. Además, se calculan los lotes restantes y las piezas faltantes para el cumplimiento de los pedidos.
- *Análisis de piezas faltantes.* En esta fase se calculan las piezas faltantes por cada pedido.
- *Comprobación de ajuste,* se realiza la política de ajuste, la cual consiste en verificar si los lotes restantes serán suficientes para completar la cantidad de piezas faltantes.
- *Ajustar demandas con WIP.* Para poder realizar el reajuste de lotes, se debe primero ver el estado del WIP en los demás repositorios y conocer el estado real de piezas faltantes en todos los pedidos, con el estado del WIP. Los pedidos actuales se disminuyen o se complementan. Las piezas faltantes servirán para hacer el recálculo de lotes, lo que significa una nueva estimación inicial.
- *Volver a calcular lotes.* Se regresa a la etapa de estimación.

4.4 Conclusiones del capítulo

Basándose en el análisis previo, se propone realizar la planificación de la producción en dos etapas: la estimación y el ajuste.

La etapa de estimación se considera como una solución inicial que después, de manera iterativa se ajusta a la real.

Se propone un modelo de la programación lineal para calcular la cantidad de lotes de cada tipo necesarios para el cumplimiento de los pedidos y de tal manera, satisfacer a los objetivos de eficiente uso de material y el mínimo de tiempo de la producción makespan.

Se utiliza el método SIMPLEX para la resolución del modelo.

El modelo es invariante con respecto al orden de procesamiento de diferentes tipos de lotes. Para obtenerlo, se utilizan las siguientes reglas de prioridad: la lista de pedidos se ordena para de acuerdo a la disminución de las piezas solicitantes. Entre los tipos de lotes disponibles, primero se elige el que beneficia al máximo este pedido, tomando en cuenta la distribución esperada.

Debido al carácter estocástico de clasificado, siempre existirán desviaciones positivas y negativas de la distribución real de un lote con respecto a la estimada. Las desviaciones negativas provocan el crecimiento del WIP, mientras las negativas afectan a la planificación.

Se propone un criterio de ajuste que se aplica después de procesamiento de cada lote de la lista en caso si la desviación negativa acumulada es crítica para el cumplimiento del primer pedido. Como un caso extremo, se vuelva calcular la solución inicial para el resto de pedidos si las piezas esperadas de los lotes subsiguientes en el repositorio marcado por el pedido son insuficientes para cumplir con todos los pedidos asignados a este repositorio. Un pedido cumplido se elimina de la lista.

Se propone un algoritmo que verifica la etapa de estimación y realiza el ajuste.

5 VERIFICACIÓN DEL ALGORITMO EN EL EXPERIMENTO COMPUTACIONAL

En este capítulo se describe un experimento computacional realizado con el propósito de verificar el funcionamiento del algoritmo descrito anteriormente tomando en consideración las dos partes: la estimación y el ajuste. Para realizar este experimento, se utilizaron preferentemente los datos obtenidos durante la estancia en la empresa. Sin embargo, el número de las muestras reales de la empresa era insuficiente. Con el propósito de abastecer el experimento con una cantidad abundante de las muestras, se diseñó y se creó un banco de pruebas.

5.1 Generación de las muestras

5.1.1 Diseño del banco de pruebas

Las muestras de prueba representan los ejemplos de distribución de un número fijo de objetos idénticos entre 25 repositorios, entonces la combinación de 25 números enteros no negativos, el total de los cuales es igual a un número (tamaño de lote). Para crear las muestras de prueba, se toma como referencia la matriz $K[4, 25]$, desarrollada para los cuatro tipos de vidrio (Romero & Burtseva, 2010). La matriz está creada a partir de datos reales de la empresa, es fija y tiene una forma integral, incluye toda población de piezas con el mismo tipo de vidrio. Un renglón g , $g = 1, \dots, 4$, de la matriz corresponde a la distribución de 1000 piezas con el tipo de vidrio g entre 25 repositorios.

La distribución de una muestra real es estocástica, varía en cada realización con respecto a la distribución de referencia. Una muestra de prueba debe exponer el comportamiento parecido a una muestra real. Una combinación de números se acepta para el banco de pruebas, si tiene las tres propiedades que revelan las muestras reales: (1) la media y (2) la desviación estándar de la muestra de prueba no deben salir fuera del rango real de la media y desviación estándar respectivamente, para su tipo de vidrio (tipo de lote); (3) el tamaño de la muestra debe ser fijo (tamaño de lote). Se seleccionó el tamaño de la muestra de prueba $n = 100$ por siguientes razones:

- Es fácil para escalar;

- Tiempo de creación de una muestra debe ser razonable: una muestra admisible se selecciona entre las posibles combinaciones de las cantidades de piezas por repositorios. El número de combinaciones es elevado ya para $n = 100$, y para valores mayores son prácticamente innumerables.

La complejidad computacional de la selección de una muestra válida se analiza a continuación.

5.1.2 Complejidad computacional

Existen k^n modos de depositar n distinguibles objetos en k distinguibles repositorios (*Occupancy Problem*) cuando algunos repositorios permiten ser vacíos. El número de combinaciones se precisa cuando los objetos son indistinguibles. En el caso considerado, los repositorios son distinguibles y se permite que algunos de estos sean vacíos. Entonces, el número de combinaciones posibles es (Anderson, 2000):

$$C_{n+k-1, n} = C_{n+k-1, k-1} = \frac{(n+k-1)!}{n! (k-1)!}.$$

Precisamente, para $n = 100$ y $k = 25$:

$$C_{124, 100} = C_{124, 24} = \frac{101 \cdot 102 \cdot \dots \cdot 124}{24!} > 10^{24}.$$

Este valor incluye todas las combinaciones de la distribución de 100 objetos entre 25 repositorios sin considerar las condiciones (1) y (2), como, por ejemplo, se crean las combinaciones cuando todos 100 objetos caen en un solo repositorio.

Hay dos caminos principales para construir combinaciones válidas:

1) Iterativamente, enumerando los valores de 0 a n en cada posición de la combinación utilizando variables de control de ciclos.

2) Utilizando algoritmos que forman solamente las combinaciones relevantes.

El algoritmo iterativo enumerativo tiene grandes desventajas: además de ser lento y excesivo, produce miles de muestras de gran semejanza, lo que mostró una prueba computacional con duración de 21 días, realizada en la computadora *SunMicrosystems* sin un avance considerable. Con el propósito de disminuir el número de las combinaciones no

relevantes, para cada tipo de distribución se definen los límites de variación por cada repositorio. Como resultado, se ejecutaron aproximadamente $1E+13$ ciclos, sin conseguir un cambio de la variable de control en los primeros 4 ciclos (repositorios). Se generaron 120,064 archivos de tamaño 500 Kb, se formaron 778,254,848 muestras válidas. En todas muestras coinciden los primeros cuatro números (ver Figura 5.1). El avance en 3 semanas fue realmente poco productivo. Para adquirir solamente combinaciones que tienen un total fijo de elementos, son necesarios métodos directos, como, por ejemplo, el método de ramificación y acotamiento (*branch-and-bound*). Pero estos tampoco resuelven problema en un tiempo razonable. Para un k fijo, el número de combinaciones se evalúa con una función exponencial. Cualquier algoritmo que construye las combinaciones requiere entonces un tiempo no menor que una función exponencial, por tanto, su realización práctica es complicada.

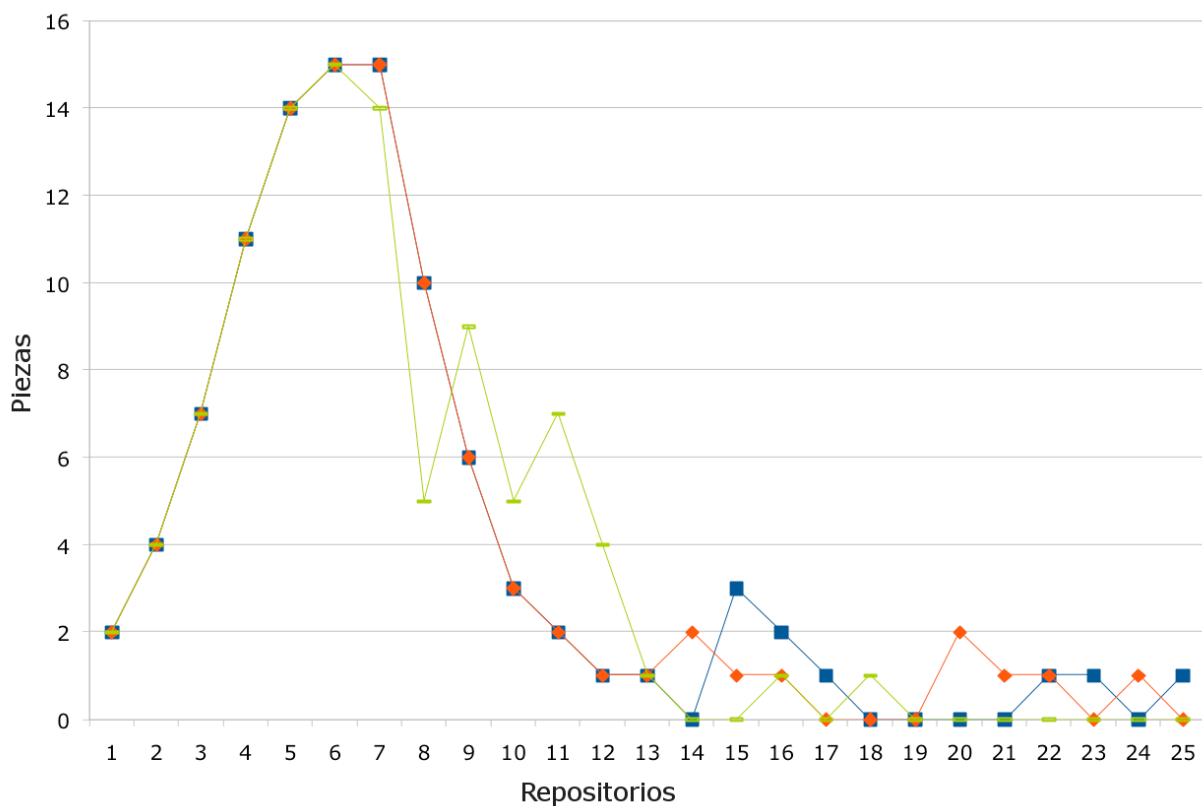


Figura 5.1. Un ejemplo de muestras semejantes adquiridas a partir del algoritmo iterativo enumerativo para el primer tipo de lote.

5.1.3. Algoritmo de generación de muestras

El algoritmo SFSA (ciclos Secuenciales Fijos-Secuenciales Aleatorios) desarrollado para formar combinaciones, dispone ciclos iterativos para un cierto rango de posiciones en la combinación, y además configura otros ciclos a través de valores pseudoaleatorios para evitar las enumeraciones inútiles en registros menores. Así, por ejemplo, para la primera distribución (primer tipo de lote), los ciclos 3-8, que corresponden a los repositorios en los cuales se deposita la mayoría de objetos, se ejecutan iterativamente, mientras las variables de control de los demás ciclos se asignan aleatoriamente. A partir de esta asignación, los ciclos se ejecutan sucesivamente hasta encontrar una combinación válida (que cumple con las 3 condiciones). La muestra se guarda en un *banco de pruebas* y se sube al 1 el valor de variable de control actual en ciclos fijos. El algoritmo utiliza los límites para el valor de variables de control.

A continuación se presenta el pseudocódigo del algoritmo SFSA:

1. Inicialización.
2. $b[25]$ y $B[25]$ son arreglos que contienen límites inferior y superior.
3. Inicialización de las variables de control de inicio de ciclo con la configuración inicial (límite inferior): $gi1 = b[0]$, $gi2 = b[1]$, ..., $gi25 = b[24]$.
4. Inicialización de sumatorias, inferior, superior y actual (sum_b , sum_B , sum_x , respectivamente).
5. Inicialización de bandera *indicaRecalculo* = *false*
6. Inicio de ciclos iterativos, (ciclos del 3 – 8), donde la variable de control varía de manera secuencial.
 - *validación de límites:*
 - ❖ Si la suma de las variables de control entre sum_x y sum_B es menor a 100, entonces continua con este mismo ciclo.
 - ❖ Si no, se valida: si sum_x y sum_b es mayor a 100 entonces se rompe el ciclo.
 - En el ciclo 8 se le asignan a variables de control de inicio de ciclo valores aleatorios entre límites inferior y superior. Para $gi9 = calculoAleatorio(b[8]$,

$B[8]), gi_{10} = \text{calculoAleatorio}(b[9], B[9]), \dots, gi_{25} = \text{calculoAleatorio}(b[24], B[24])$

7. Inicio de los ciclos aleatorios:

- En ciclos del 9 al 12 se valida el límite; si *indicaRecalculo*= *true*, entonces se rompe el ciclo, también se realiza *validación de límites*.
- En ciclos del 13 al 25 se valida el límite; si *indicaRecalculo*= *true*, entonces se rompe el ciclo.

8. En el ciclo 25, se calcula el total de valores de variables de ciclo. Si es igual a 100, entonces se calcula la media. Si la media pertenece al rango deseado, se calcula la desviación estándar, y si la desviación estándar se encuentra entre el rango deseado, se almacena la muestra.

9. Retorno al paso 6.

Este algoritmo es finito; el cambio obligatorio de variables de control acelera la ejecución de ciclos fijos. La asignación de valores aleatorios en ciclos restantes introduce la variación en las muestras. El algoritmo incluye un Generador Lineal Congruente (*LCG*) introducido por D.H. Lehmer en 1949 (Knuth, 1981). Este método es empleado en la mayoría de los lenguajes de programación. Se asignan los siguientes valores:

m el modulo;

a el multiplicador;

c el incremento;

X_0 el valor inicial.

La secuencia de valores aleatorios surge de la formula recurrente:

$$X_{n+1} = (aX_n + c) \bmod m, n \geq 0.$$

El algoritmo se realizó en lenguaje *Java*, y se implementó en el ambiente *Solaris* de *SunMicrosystem*. Desde la obtención de las primeras combinaciones relevantes, se nota una variabilidad considerable entre las muestras (ver Figura 5.2), y la cantidad de muestras generadas es razonable en comparación con el algoritmo iterativo.

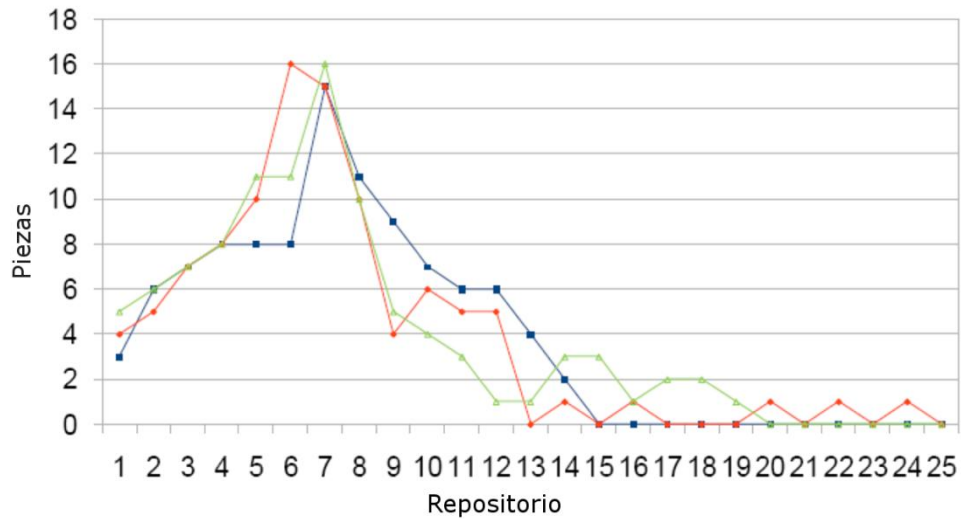


Figura 5.2 Un ejemplo de muestras variables adquiridas a partir del algoritmo SFSA para el primer tipo de lote.

5.2 Diseño de experimento

Las muestras reales recolectadas de la empresa y las muestras de prueba adquiridas a partir del algoritmo SFSA fueron utilizadas en el experimento computacional que simula el proceso real de la empresa apoyándolo con la herramienta, descrita en el Capítulo 4, para la elección de tipos de lotes a procesar. Las características que sigue el experimento se explican a continuación:

Las máquinas

La empresa utiliza máquinas que realizan la clasificación. El experimento toma en cuenta solamente *una máquina clasificadora* que realiza la clasificación de las piezas por lotes.

Los lotes

El experimento utiliza los datos reales de la empresa. Cuando estos son insuficientes, se toman muestras del banco de pruebas. La simulación de la clasificación consiste en tomar de manera aleatoria un lote y utilizarlo como resultado de la etapa de clasificación.

Tamaño de lote

Las muestras obtenidas a partir del banco de pruebas, son normalizadas al tamaño fijo de 100 elementos (piezas). Para utilizar tal muestra en el experimento, primeramente, el número de

piezas se multiplica al coeficiente 55 por cada repositorio para llegar al tamaño de lote en la empresa.

En capítulos anteriores se mencionó que un lote inicialmente contiene 5500 piezas. Al pasar por las diferentes etapas y llegar a la etapa de clasificado ocurre el fenómeno llamado “desecho” (*scrap*), que consiste en rechazo de las piezas por fallas diversas en cada una de las etapas. Otra característica del desecho es que nunca es del mismo valor. Es decir, un lote que llegó a la etapa de clasificado, varía en el tamaño. Se definió el rango del desecho [30%, 35%] que fue el observado en las muestras de la empresa. De este rango se toma de manera aleatoria un porcentaje que representa la parte de desechos que se resta del total del lote, esto para aleatorizar los desechos.

Las demandas

Se consideró una entrada fija de demandas en este experimento. Estas demandas representan pedidos que solicitan los clientes a la empresa. Una demanda incluye la cantidad de piezas y el número de parte. La lista de las demandas se conoce con anterioridad, por lo tanto se considera que la fecha de entrega de las demandas es la misma. El tamaño de demanda se toma de manera aleatoria dentro del rango de [1000, 3000].

El número de parte

Debido a que en la empresa existen diversos números de parte que indican cuales repositorios se utilizarán para el cumplimiento de las demandas, se tomaron aquellos que son los más solicitados. Para poder realizar una estimación más concreta, se definieron las características que cada número de parte debe cumplir para ser utilizado en el experimento. Dado la complejidad del problema, se consideraron las siguientes razones:

- El repositorio que utiliza el número de parte debe ser único.
- A los diferentes números de parte corresponden diferentes repositorios.

En el experimento se consideran tres diferentes números de parte. Debido a la investigación que se realizó en la empresa, se encontró que la mayoría de los lotes fueron ensamblados con dos tipos de vidrio: G1 y G2. Por lo tanto, la investigación se enfocó en encontrar la estimación basada únicamente en estos dos tipos. En conclusión, la estimación de lotes será reflejada para *dos tipos de vidrio únicamente*.

Debido a las características anteriores, se definen las condiciones iniciales del experimento (Tabla 5.1), los cuales se utilizan para calcular la solución inicial, que se obtiene en la etapa de estimación.

Tabla 5.1. Representación de un conjunto de demandas utilizadas en el experimento.

No. Demanda	No. Parte	No. Piezas
D1	NP3	1316
D2	NP3	2942
D3	NP1	2915
D4	NP3	1971
D5	NP2	2601
D6	NP1	1284
D7	NP1	1844
D8	NP2	2832
D9	NP3	2585
D10	NP3	2919

Esta etapa recibe el conjunto de demandas para así obtener como resultado la cantidad de lotes necesarios para su cumplimiento. Con los resultados de la etapa de estimación, se continúa con la simulación de la clasificación y la etapa de ajuste, donde en caso de cumplir con el criterio de ajuste, se realiza el re-cálculo de lotes hasta completar la demanda.

5.3 Resultados

La finalidad de este experimento es el asegurar que los resultados obtenidos del algoritmo, serán útiles en la empresa, ya que los datos obtenidos estarán como guía para la mejora de la planeación de ensamblado de interruptores eléctricos.

Para la realización del experimento se consideran cuatro variables que al modificarlas se obtienen resultados distintos. Las variables son: número de demandas, el tamaño de la demanda, el número de parte utilizado en la demanda y tamaño del lote. Estas cuatro variables asemejan el comportamiento real del proceso de producción. Los valores que toman estas variables están

ligados a su participación en la planeación de los lotes. El número de demandas se considera variable por que se incrementa o disminuye según las necesidades de los clientes. El número de parte también es elegido por los clientes, y varía de una demanda a otra, por lo tanto, se considera aleatorio. El tamaño de lote, como se mencionó anteriormente, es variable por la existencia de desechos. El tamaño de lote se expresa con un porcentaje de pérdida de piezas, este se resta del lote ideal de 5500 piezas, provocando así que el tamaño de lote cambie de una estimación a otra. Y por último, como ya se explicó, la demanda está ligada a los requerimientos de los clientes, por lo tanto, así como varía la cantidad de demandas, a su vez, varía la cantidad de piezas en cada demanda.

La información del algoritmo se divide en dos partes: información de entrada y de salida. La información de entrada concentra las demandas (cantidad de piezas y números de parte) que se obtienen de la planeación. La información de salida contiene número de lotes estimados para cada tipo de vidrio, número de ajustes, número de lotes necesarios para cumplir la demanda, número de lotes ajustados, número de lotes completados en los momentos que se necesitó ajuste, historial de WIP, historial de piezas obtenidas por repositorios, historial de piezas esperadas por repositorios, porcentaje de desecho. Cada punto nos sirve para analizar los resultados del algoritmo, y sacar conclusiones.

El resultado de un experimento se muestra en Tabla 5.2, donde se observa un historial de los lotes procesados, las piezas obtenidas en cada repositorio y el cumplimiento de la demanda; también se muestran los momentos donde se aplicó un ajuste. En esta tabla presentan los datos siguientes:

- ❖ **No. Lote:** número secuencial de lote procesado;
- ❖ **Tipo de vidrio:** el tipo de vidrio con el cual fue ensamblado el lote (G1 y G2);
- ❖ **No. Parte:** indicador del número de parte. Se utilizan tres números de parte de la empresa (NP1, NP2 y NP3). Indica los repositorios de donde se tomaran las piezas;
- ❖ **Esperado:** piezas esperadas para el repositorio considerando el porcentaje de desechos (utilizado en la etapa de estimación);
- ❖ **Obtenido:** piezas obtenidas para el repositorio que se obtienen del lote real;

- ❖ Δ : diferencia entre el obtenido y el esperado en el repositorio solicitado por la demanda a partir del lote; $\Delta > 0$ demuestra sobreproducción, $\Delta < 0$ corresponde a la producción insuficiente, y $\Delta = 0$ cuando el valores esperado coincide con el obtenido;
- ❖ **Demanda**: historial de la demanda;
- ❖ **WIP**: el sobrante de piezas en el repositorio, se genera hasta el momento que se completa la demanda en algún repositorio;
- ❖ **Historial de lotes**: representa la cantidad de los lotes faltantes calculada en la etapa de estimación.

Tabla 5.2. Historial de un experimento

No. Lote	Tipo de Vidrio	No. Parte	Esperado	Obtenido	Δ	Demanda	WIP	Historial Lotes	
								G1	G2
0	-	-	-	-	-	6188	0	1	27
						7774	0		
						4383	0		
1	G1	NP1	600	571	-29	5617	0	0	27
		NP2	1060	841	-219	6933	0		
		NP3	58	35	-23	4348	0		
2	G2	NP1	282	165	-117	5452	0	0	26
		NP2	1650	988	-662	5945	0		
		NP3	167	139	-28	4209	0		
3	G2	NP1	282	409	127	5043	0	0	25
		NP2	1650	1456	-194	4489	0		
		NP3	167	383	216	3826	0		
4	G2	NP1	282	341	59	4702	0	0	24
		NP2	1650	1217	-433	3272	0		
		NP3	167	14	-153	3812	0		
5	G2	NP1	282	307	25	4395	0	0	23
		NP2	1650	2311	661	961	1350		
		NP3	167	175	8	3637	0		
6	G2	NP1	282	1249	967	3146	0	0	22
		NP2	1650	1155	-495	0	2505		
		NP3	167	89	-78	3548	0		
7	G2	NP1	282	266	-16	2880	0	0	21
		NP2	1650	1544	-106	0	4049		
		NP3	167	49	-118	3499	0		
8	G2	NP1	282	165	-117	2715	0	0	20
		NP2	1650	988	-662	0	5037		
		NP3	167	139	-28	3360	0		
9	G2	NP1	282	266	-16	2449	0	0	19
		NP2	1650	1544	-106	0	6581		
		NP3	167	49	-118	3311	0		
10	G2	NP1	282	570	288	1879	0	0	18
		NP2	1650	2686	1036	0	9267		

		NP3	167	371	204	2940	0		
11	G2	NP1	282	627	345	1252	0		
		NP2	1650	808	-842	0	10075	0	17
		NP3	167	11	-156	2929	0		
12	G2	NP1	282	1249	967	3	0		
		NP2	1650	1155	-495	0	11230	0	16
		NP3	167	89	-78	2840	0		
13	G2	NP1	282	803	521	0	803		
		NP2	1650	1280	-370	0	12510	0	15
		NP3	167	14	-153	2826	0		
Ajuste	-						Estimados	1	27
							Ajuste	1	2
14	G1	NP1	650	1239	589	0	2042		
		NP2	1147	806	-341	0	13316	0	17
		NP3	63	22	-41	2804	0		

Los ajustes ocurren cuando en base a los cálculos el criterio de ajuste no se cumple. En otras palabras, el criterio de ajuste es el indicador de la desviación crítica entre la cantidad de piezas esperadas contra las piezas obtenidas. Por lo tanto, si hay una marcada desviación entre la cantidad de piezas esperadas y la cantidad de piezas obtenidas, probablemente se necesitarán ajustes; al contrario, si existe una desviación estable, la necesidad de ajustar la estimación inicial de lotes será mínima o nula. La etapa de ajuste se vuelve fundamental en los primeros casos, tomando en cuenta que el ajuste no existirá en todos los casos.

Al analizar el ejemplo de la Tabla 5.2, se observa que los repositorios que corresponden a NP1 y NP2 fueron completados al llegar al lote número 13. Se observa el lote donde fue necesario un ajuste. En la tabla observamos que las 2826 piezas faltantes no serían completadas por los 14 lotes restantes, recordando que únicamente se toman en cuenta las piezas del repositorio que le corresponde a la demanda (en este caso NP3), por lo tanto, considerando la demanda faltante por completar y realizando nuevamente la estimación, se obtiene la cantidad de los lotes complementarios. Otro detalle a considerar es el WIP generado en los repositorios NP1 y NP2, a partir del lote 5. El WIP empieza a generarse en el repositorio NP1, con esto las piezas almacenadas servirán para futuras planeaciones de pedidos, donde disminuirán la cantidad de lotes necesarios para el cumplimiento de una demanda.

El algoritmo fue probado en diferentes escenarios variando diversas condiciones, para así observar los resultados y sacar conclusiones. La información que se obtuvo del experimento muestra la manera en la que el algoritmo reacciona si las condiciones son variadas, esto es

necesario para validar una futura implementación del algoritmo en un ambiente no controlado. Cada escenario representa la variación de alguna variable anteriormente explicada (número de demandas, la demanda total, el número parte y los desechos).

Escenario 1

El primer escenario se considera como estático ya que la única variación existente es en *la cantidad de demandas*, es decir, manteniendo fijas las demás variables. En este escenario se observa cómo afecta el proceso estocástico a la estimación inicial, y por ende a la planeación de los lotes. El descubrimiento en este escenario fue la afectación del proceso estocástico directamente a la estimación inicial de lotes, debido a que la solución final (cantidad de lotes obtenida) varía de manera considerable de un experimento a otro. Los valores que toman las variables en cada experimento son las siguientes:

- total de demandas 50,
- tamaño total de la demanda: 50000 piezas (1000 piezas por demanda),
- porcentaje de piezas en cada repositorio (NP1-40%, NP2-40%, NP3-20%),
- desechos variables.

Los resultados de este escenario se observan en la Tabla 5.3.

Tabla 5.3 Escenario 1: variables estáticas

Ajustes	% de avance	Estimación	Ajustada	Resultante	ΔL	Factor P/L
9	0,46	65	92	89	24	561,79
11	0,52	67	91	84	17	595,23
4	0,51	65	77	74	9	675,67
4	0,70	66	74	68	2	735,29
3	0,50	68	76	70	2	714,28
4	0,45	64	82	76	12	657,89
6	0,52	64	79	74	10	675,67
4	0,51	68	79	73	5	684,93
0	1,00	68	68	51	-17	980,39
3	0,55	66	74	67	1	746,26

4,80	0,57	66,10	79,20	72,60	6,50	702,75
3,16	0,16	1,60	7,50	10,22	11,01	113,23

En este escenario se obtiene el número mayor de piezas por lote (factor Pieza / Lote), ya que con una mayor cantidad de lotes se obtiene mayor cantidad de piezas. En este escenario se muestra cómo le afecta el algoritmo la aleatoriedad del proceso, ya que la cantidad de lotes estimados varía muy poco; al contrario, en la cantidad resultante de lotes se tiene una variabilidad mayor. La diferencia del estimado contra el resultante a su vez está relacionada con estos dos factores. Lo esperado al realizar este escenario era encontrar aproximadamente la misma cantidad de lotes resultantes, ya que la demanda es igual para todos los experimentos.

Escenario 2

En el escenario 2 la variable que se modifica es *el tamaño de la demanda*, es decir, cada experimento tiene una demanda total diferente; esto provoca también que el porcentaje de piezas requeridas por repositorio varíe un poco. Los rangos que se manejan por demanda se encuentran entre 1000 a 3000 piezas. Los resultados de este escenario con un total de 50 demandas se presentan en la Tabla 5.4.

Tabla 5.4 Escenario 2: tamaño de demanda variable

Ajustes	% de avance	Estimación	Ajustada	Resultante	ΔL	Factor P/L
3	0,56	145	162	145	0	735,41
4	0,55	127	143	130	3	747,07
7	0,53	138	153	140	2	717,31
23	0,30	133	198	197	64	539,76
9	0,44	123	148	144	21	621,75
9	0,46	136	161	155	19	673,60
8	0,57	139	159	149	10	756,13
4	0,54	132	148	137	5	664,71
1	0,53	128	136	126	-2	794,61
10	0,42	143	176	167	24	620,73
7,80	0,49	134,40	158,40	149,00	14,60	687,11
6,12	0,09	7,12	17,90	20,60	19,68	77,63

A diferencia del Escenario 1, el Escenario 2 presenta cambios importantes en la cantidad de lotes, esto se debe a que la cantidad de la demanda es más grande, por lo tanto, la relación entre la cantidad de lotes y la cantidad de piezas es directa; a mayor número de piezas corresponde mayor número de lotes. La cantidad de ajustes aumentó, pero no fue tan significativo como en la cantidad de piezas. La conclusión es: al haber mayor número de lotes, también habrá mayor número de ajustes pero con un crecimiento moderado. Se debe notar que el factor de piezas por lote no se vio afectado de manera drástica, esto nos indica que el tamaño de la demanda solo afecta a la cantidad de lotes necesarios para su cumplimiento.

Escenario 3

En el Escenario 3 se modifica la variable del *porcentaje de piezas por repositorio*. Esto quiere decir que de las 50 demandas utilizadas para cada experimento, la probabilidad de que corresponda a NP1 es similar a que corresponda NP2, por lo tanto, su probabilidad es uniforme ya que los tres repositorios tienen la misma probabilidad de ocurrencia. Donde se encuentra una diferencia significativa, es entre la cantidad de lotes resultantes y lotes estimados. Se anota que hay un aumento en comparación del Escenario 1. El total de piezas de la demanda es de 50000, con tamaño de lote variable. Los resultados se presentan en la [Tabla 5.5](#).

Tabla 5.5. Escenario 3: porcentaje de piezas por repositorio

Ajustes	% avance	Estimación	Ajustada	Resultante	ΔL	Factor P/L
4	0,86	95	111	110	15	454,55
1	0,96	110	118	115	5	434,78
0	1,00	103	103	103	0	485,44
2	0,53	63	70	60	-3	833,33
8	0,23	94	119	117	23	427,35
9	0,29	92	114	106	14	471,70
4	0,29	95	103	98	3	510,20
1	0,98	120	123	123	3	406,50
6	0,36	92	104	102	10	490,20
1	0,95	106	112	112	6	446,43
3,60	0,64	97,00	107,70	104,60	7,60	496,05
3,17	0,33	15,05	14,97	17,41	7,89	122,65

Por los resultados de los experimentos correspondientes se observa que el factor de piezas por lote disminuyó aun más que en el Escenario 2; esto nos indica que la variabilidad en el porcentaje de número de parte afecta más al rendimiento de los lotes que la cantidad de piezas en las demandas. Otra cosa que se observa es que el número de ajustes es menor en relación al Escenario 2, por lo que se concluye que la relación entre el porcentaje de piezas en los repositorios está relacionada con la cantidad de ajustes necesarios para un conjunto de demandas. Por ultimo, la cantidad de lotes en comparación del Escenario 1, aumentó en un 50% aproximadamente, lo que explica la disminución en el factor P/L.

Escenario 4

Por ultimo, encontramos el Escenario 4, en este *todas las variables son aleatorias*, por lo tanto, arroja resultados más acordes a un escenario real, en donde no se tiene control de las variables expuestas. Los resultados de esta simulación se presentan en la Tabla 5.6.

Tabla 5.6 Escenario 4: simulación con variación en todos los parámetros

Ajustes	% avance	Estimación	Ajustada	Resultante	ΔL	Factor P/L
1	0,95	208	219	219	11	447,82
13	0,30	175	210	208	33	482,54
3	0,62	228	292	289	61	348,65
2	0,78	214	270	270	56	359,14
6	0,28	185	204	201	16	495,32
2	0,37	147	158	156	9	645,01
2	0,82	219	265	263	44	370,74
7	0,52	179	193	192	13	544,70
2	0,90	280	296	295	15	324,35
1	0,73	245	277	275	30	410,76
3,90	0,63	208,00	238,40	236,80	28,80	442,90
3,78	0,25	38,43	47,47	47,52	19,26	100,97

Los resultados de los experimentos parecen ser la fusión de los demás escenarios, con pequeñas diferencias, en la cantidad de lotes observamos que la variabilidad de los porcentajes del número de parte indica la cantidad de ajustes, mientras la cantidad de piezas para una

demanda afecta al total de lotes necesarios para su cumplimiento. Por ultimo, el factor P/L es muy parecido al obtenido en el Escenario 3, lo que nos indica que para tener un mejor rendimiento de los lotes se debe observar que porcentaje es el óptimo para disminuir el número de ajustes y aumentar el rendimiento de piezas por lote.

Escenario 5

En este escenario, se consideran todas las variables aleatorias manteniendo el porcentaje de piezas para el repositorio número 3 no más que 5% del total de piezas en un total de demandas para garantizar que el algoritmo disminuya considerablemente el número de ajustes.

Tabla 5.7 Escenario 5: Simulación con 5% de nivel de porcentaje

Ajuste	% avance	Estimación	Ajustada	Resultante	ΔL	Factor P/L
0	1,00	82	82	54	-28	1834,44
0	1,00	70	70	48	-22	1958,31
0	1	78	78	54	-24	1887,6667
0	1	95	95	71	-24	1501,7183
0	1	84	84	59	-25	1757,4068
0	1	75	75	58	-17	1691,7241
0	1	86	86	58	-28	1708,0345
0	1	89	89	62	-27	1581,629
0	1	98	98	68	-30	1550,25
0	1	86	86	55	-31	1867,9818
0,00	1,00	84,30	84,30	58,70	-25,60	1733,92
0,00	0,00	8,60	8,60	6,85	4,14	154,89

Los resultados en la Tabla 5.7 muestran que no hubo necesidad de ajustes, y las demandas fueron perfectamente complementadas. La relación de piezas por lote se aumentó considerablemente, y la diferencia entre la cantidad de lotes esperada y la obtenida está debajo de lo previsto. Con estos resultados, se concluye que al disminuir la cantidad de demandas con este número de parte (NP3), los resultados del algoritmo estarán adecuados a la planeación. Este escenario puede servir como guía para futuras planeaciones.

5.4 Conclusiones del capítulo

Para disponer de una cantidad suficiente de muestras en el experimento funcional, se diseñó un banco de pruebas con características especiales: cada muestra válida representa una combinación de números enteros no negativos con un total de 100; la media y desviación estándar de la distribución que corresponde a tal combinación deben permanecer en ciertos rangos. Es difícil generar de forma manual muestras con tales características, y no existen algoritmos computacionales para generalas.

Se propone el algoritmo SFSA que combina los ciclos sucesivos con una asignación aleatoria a variables de control de ciclos y utiliza sus límites dados. Este algoritmo es finito; el cambio obligatorio de variables de control acelera la ejecución de ciclos fijos. La asignación de valores aleatorios en ciclos restantes introduce la variación en las muestras. El desarrollo del banco de pruebas para la obtención de muestras de distribución de piezas entre repositorios se requiere para la resolución de problemas en la planificación de la producción de interruptores.

Para verificar el funcionamiento del algoritmo de estimación y ajuste que calcula la cantidad de lotes para demandas finitas, fue realizado un experimento computacional. El experimento es un conjunto de experimentos bajo 5 escenarios, los cuales tienen características similares, ya que varían en los datos de entrada.

La variabilidad de los tipos de lotes provoca que la cantidad de lotes estimados no sea siempre la misma, pero se mantiene con una variación muy pequeña. La cantidad de lotes con el mismo tipo de vidrio está relacionada con la cantidad de piezas pertenecientes a un número de parte. La problemática mas grande en esta estimación se encuentra con el repositorio número 3, donde la mayoría de los ajustes ocurren por piezas faltantes, por lo que provoca una descompensación de lotes.

El criterio de ajuste apoyó a la planeación de los pedidos dando la menor cantidad posible de lotes para solucionar las demandas. La cantidad de ajustes no está relacionada con la cantidad de piezas en las demandas, si bien es cierto que al aumentar la cantidad de demandas por ende la cantidad de piezas, aumentó la cantidad de ajustes en los experimentos, el aumento no es significativo. Lo que realmente afecta el aumento de ajustes o en otras palabras lo que provoca una mayor cantidad de desviaciones negativas en los repositorios, es el aumento de piezas que utilicen el número de parte NP3.

Al encontrar que el número de ajustes disminuye, entonces encontramos que la problemática se centra en encontrar la solución de estimación de lotes para este repositorio o buscar un tipo de vidrio de otras dimensiones más cercano a esos repositorios.

Una solución para este tipo de vidrio es utilizar tipos de vidrio con configuraciones que den piezas con números de parte más grandes, y realizar una adaptación del modelo para ese tipo de lote.

Para garantizar que el algoritmo disminuya considerablemente el número de ajustes se requiere que el porcentaje de piezas para el repositorio número 3 no pase del 5% del total de piezas en un total de demandas, esto se debe a que en el repositorio número 3 la cantidad de piezas esperadas es del 5% en promedio tomando en cuenta únicamente los tres repositorios.

Los ajustes de lotes se asocian a la variabilidad de piezas obtenidas de los lotes, aunque el número de ajustes puede ser disminuido si se controla el porcentaje de piezas.

El WIP no es considerado en la etapa de estimación sino que, las demandas que se ingresan al algoritmo, llevan las piezas que no fueron surtidas por el WIP.

6 CONCLUSIONES

A lo largo de los capítulos se demuestra de manera detallada una solución para mejorar la planeación de pedidos en la producción de interruptores eléctricos. El problema surge de un problema que tiene características de un flujo de trabajos. Investigando cada una de las etapas se ratificó en base a estudios anteriores que la etapa de mayor conflicto en el actual proceso de producción es la etapa de clasificado, donde los resultados obtenidos son estocásticos, por lo que es complicado definir una cantidad fija de lotes que sirvan para la planeación de pedidos. Para resolver el problema de la variabilidad de resultados en esta etapa se basó en la política de la empresa, que es utilizar todos los recursos en el momento que se requieran, sin generar grandes cantidades de inventario. Esta política impulsó la consideración del uso de WIP para poder satisfacer los pedidos de los clientes. El WIP entonces se colocó como una opción para la solución del problema de piezas faltantes para el cumplimiento de pedidos. Se propuso un algoritmo de estimación de lotes que fuera capaz de indicar la cantidad de lotes necesarios para satisfacer un conjunto de pedidos, considerando la etapa estocástica de clasificado, donde los resultados siempre son diferentes. Este algoritmo contemplaría además de dar una solución a la cantidad de lotes, el uso de WIP para los pedidos siguientes evitando grandes cantidades de inventario.

Para descifrar la problemática existente en la etapa de clasificado se necesitó encontrar la manera en que los lotes se distribuían en los repositorios, es decir la distribución que seguía el proceso de clasificado. El primer paso fue definir si el uso de cualquier tipo de vidrio influía en la variabilidad de los resultados de la etapa de clasificado. Por lo tanto se demostró que la utilización de un tipo de vidrio en comparación de otro afecta directamente en el cumplimiento de los pedidos. Se demostró que un lote en su forma original no corresponde a ninguna distribución conocida, por lo tanto, se necesitó reclasificar los lotes de una manera que permitiera encontrar la distribución a la que corresponde el proceso de clasificado, sin perder la información básica de un lote: tamaño del lote y el vidrio con que se ensamblaron los interruptores. La reclasificación permitió volver a realizar el experimento de comprobación de la distribución con lotes reclasificados. Se concluyó que después de reclasificación los lotes correspondían a una distribución normal. Al conocer la distribución de los lotes en su forma reclasificada, se necesitó

encontrar un modelo que indicará el comportamiento de las piezas al ser depositadas en los repositorios. Se definió el modelo probabilístico que define la probabilidad de que una pieza corresponda a algún repositorio, considerando que para cada tipo de vidrio el valor de la probabilidad es diferente. Teniendo el modelo probabilístico del proceso de clasificado, se propuso un modelo de estimación de lotes, el cual considera el calculo de piezas en los repositorios para cada tipo de vidrio (modelo probabilístico).

Basándose en el método de estimación de lotes, se propuso el algoritmo para el cálculo de lotes de dos etapas, donde cada etapa se encarga de resolver un problema en específico. Este algoritmo enfrenta el problema del cumplimiento de pedidos con la estimación de lotes y la variabilidad de en la etapa de clasificado con la etapa de ajustes. La etapa de estimación es el punto inicial o de partida para la planeación, y la etapa de ajuste define los cambios necesarios para cumplir con todo el conjunto de demandas. En esta etapa se define la política de ajuste, la cual indica si será necesario hacer un reajuste de los lotes calculados inicialmente en la etapa de estimación, basándose en los lotes obtenidos después de la etapa de clasificación.

Para comprobar el funcionamiento del algoritmo de estimación de lotes fue necesario validar los resultados que genera, en diversos escenarios que reflejan el proceso real de ensamblado de interruptores. Cada escenario es la variación de elementos importantes dentro del proceso de ensamblado por ejemplo, el tamaño del lote. Los elementos que se consideraron fueron: porcentaje de piezas que utilizan algún número de parte, la cantidad de piezas de la demanda, y la cantidad de demandas. Se encontró que al variar estos elementos, los resultados que realiza el algoritmo se ven afectados, por lo tanto, al encontrar patrones en estos elementos se podrán definir los rangos donde el algoritmo genera mejores resultados. Uno de los indicadores de este experimento fue que el número de parte (indicador de repositorios donde se tomaran piezas para surtir la demanda) es el elemento que afecta más en la cantidad de ajustes realizados por la etapa de ajuste del algoritmo. Otra conclusión del experimento fue que la cantidad de lotes necesarios para un conjunto de pedidos pueden ser escalables, si se logra controlar el efecto de los faltantes de material dentro del proceso.

7. FUTURO TRABAJO

Muchos problemas desprenden del presente modelo de investigación. Algunos de estos son:

- Investigación del problemática de la re-planificación o la planificación dinámica de la producción, tomando en cuenta que un lote tiene tamaño variable.
- La asignación de los pedidos a los lotes requiere la resolución de problemas como la partición de trabajos y el batching de sublotes en el contexto considerado.
- Estudio la posibilidad de aplicación de redes bayesianas para enriquecer el análisis probabilístico de distribuciones.
- Desarrollo del procedimiento de cálculo de la cantidad de lotes de producción con diferentes tipos de material para el caso de varias máquinas de clasificado y traslape de repositorios en diferentes números de parte.

Estos problemas tienen un interés tanto teórico, como práctico. Se supone la entrega de los resultados de investigación a la empresa y la publicación de resultados. Con este propósito, es necesario entender diversos procesos de esta naturaleza así como buscar métodos de solución, y comprobar si es posible utilizarlos en el caso de estudio de producción de interruptores eléctricos.

REFERENCIAS

- Allahverdi, A., Ng, C. T., Cheng, T. C. E. & Kovalyov, M. Y. (2008). A survey of scheduling problems with setup times or costs. *European Journal of Operational Research*, 187(3) : 985–1032.
- Azadeh, A., Bidokhti, B., Sakkaki, S.M.R. (2005). Design of practical optimum JIT systems by integration of computer simulation and analysis of variance. *Computers & Industrial Engineering*, 49(4): 504-519.
- Bonney, M.C., Zhang, Z., Head, M.A., Tien, C.C., Barson, R.J. (1999). Are push and pull systems really so different? *International Journal of Production Economics*. 59 (1): 53–64.
- Conway, R., Maxwell, W., McClain, J.O., Thomas L.J. (1988). The role of work-in-process inventory in serial production lines. *Operations Research*, 36: 229–41.
- Crama, Y. (1997). Combinatorial optimization models for production scheduling in automated manufacturing systems. *European Journal of Operational Research*, 99,(1): 136–153.
- Gershwin, S.B. (2000). Design and operation of manufacturing systems: the control- point policy. *IIE Transaction on Design and Manufacturing*. 32: 891–906.
- Glover, F. & Laguna, M. (1997). Tabu search. Kluwer Academic Publishers, Boston.
- Gupta, J.N.D. (1988). Two-stage, hybrid flowshop scheduling problem. *Journal of the Operational Research Society*. 39(4): 359-364.
- Hopp, W.J. & Roof, M.L. (1998). Setting WIP levels with statistical throughput control (STC) in CONWIP production lines. *International Journal of Production Research*, 36(4): 867-882
- Jarque, C. M., & Bera, A. K. (1987). A test for normality of observations and regression residuals. *International Statistical Review*. 55(2): 163–172.
- Jungwattanakit, J., Reodecha, M., Chaovalitwongse, P. & Werner, F. (2009). A comparison of scheduling algorithms for flexible flow shop problems with unrelated parallel machines, setup times, and dual criteria. *Computers & Operations Research*. 36(2): 358–378.
- Kilic, O.A. & Tarim S.A. (2011). An investigation of setup instability in non-stationary stochastic inventory systems. *International Journal of Production Economics*. 133: 286-292.
- Kim, R.S. & Leachman R.C. (1994). Decomposition method application to a large scale linear programming WIP projection model. *European Journal of Operational Research*. 74(1): 152-160.
- Kolman, B. & Beck, R., E. Elementary linear programming with applications. (1995), Academic Press, UK, London, 423 p.
- Kopanos, G. M., Capon-Garcia, E., Espuna, A., Puigjaner, L. (2008). Costs for rescheduling actions: A critical issue for reducing the gap between scheduling theory and practice. *Industrial and Engineering Chemistry Research*, 47(22): 8785-8795.
- Lin, Yu-H., Shie, J.R., Tsai, Ch.H. (2009). Using an artificial neural network prediction model to optimize work-in-process inventory level for wafer fabrication. *Expert Systems with Applications*. 36 (2): 3421-3427.

- Lin, H.-T. & Liao, C.-J. (2003). A case study in a two-stage hybrid flow shop with setup time and dedicated machines. *International Journal of Production Economics*. 86(2): 133–143.
- Marimuthu, S., Ponnambalam, S.G. & Jawahar, N. (2009). Threshold accepting and ant-colony optimization algorithm for scheduling m-machine flow shop with lot streaming, *Journal of Material Processing Technology*. 209(2): 1026-1041.
- Marti, K. (2008). *Stochastic Optimization Methods*. Springer-Verlag, Berlin Heidelberg, 340 p.
- Monthgomery D.C. & Runger, G.C. (2001) *Probabilidad y Estadística aplicadas a la Ingeniería*. John Wiley & Sons, 895 p.
- Naderi, B., Zandieh, M., Khaleghi, A.G.B. & Roshanaei, V. (2009). An improved simulated annealing for hybrid flowshops with sequence-dependent setup and transportation times to minimize total completion time and total tardiness. *Expert Systems with Applications*. 36(6) : 9625–9633.
- Papadopoulos, H.T. & Vidalis, M.I. (2001). Minimizing WIP inventory in reliable production lines. *International Journal of Production Economics*. 70(2): 185-197.
- Pettersen, J.-A. & Segerstedt, A. (2009). Restricted work-in-process: A study of differences between Kanban and CONWIP. *International Journal of Production Economics*. 118 (1): 199-207.
- Pinedo, M.L. (2008). *Planning and Scheduling in Manufacturing and Services*. Springer, 506 p.
- Potts, C.N. & Baker, K.R. (1989). Flow shop scheduling with lot streaming. *Operations Research Letters*. 8(6): 297-303.
- Qiu, R. (2005). Virtual production line based WIP control for semiconductor manufacturing systems. *International Journal of Production Economics*. 95(2): 165-178.
- Romero Parra, R. & Burtseva, L. (2010). Implementation of Bin Packing Model for Reed Switch Production Planning. *IAENG Transactions on Engineering Technologies*. AIP, Melville, NY. 1247: 403-412.
- Rosner, B. *Fundamentals of Biostatistics* (2010). Nelson Education, Ltd., 875 p.
- Ruiz, R. & Vázquez-Rodríguez, J.A. (2010). The hybrid flow shop scheduling problem. *European Journal of Operational Research*. 205(1): 1-18.
- Ruiz, R.; Serifoglu, F.S. & Urlings, T. (2008). Modeling realistic hybrid flexible flowshop scheduling problems. *Computers & Operations Research*. 35 (4): 1151-1175.
- Sivakumar, B. & Arivarignan, G. (2009). A stochastic inventory system with postponed demands. *Performance Evaluation*. 66 (1): 47-58.
- Tsourveloudis, N.C. (2010). On the evolutionary-fuzzy control of WIP in manufacturing systems. *Neurocomputing*. 73 (4-6): 648-654
- Vieira, G.E., Herrmann, J.W., Lin, E. (2003). Rescheduling Manufacturing Systems: A Framework of Strategies, Policies, and Methods. *Journal of Scheduling* (6): 39-62, 2003.

- Vairaktarakis, G. (2004). Flexible Hybrid Flowshop, *In: Handbook of Scheduling: Algorithms, Models, and Performance Analysis*. J. Y.-T. Leung, Editor, Boca Raton, Florida: 5-1 – 5-33.
- Yang, J. & Posner, M.E. (2010). Flow shops with WIP and value added costs. *Journal of Scheduling*. 13(1): 3-16
- Yasin, M.M., Small, M.H, Wafa, M.A. (2003) Organizational modifications to support JIT implementation in manufacturing and service operations. *Omega*. 31(3): 213-226.
- Zandieh, M., Fatemi Ghomi, S.M.T., Moattar Hussein, S.M. (2006). An immune algorithm approach to hybrid flow shops scheduling with sequence-dependent setup times. *Applied Mathematics and Computation*. 180(1): 111–127.