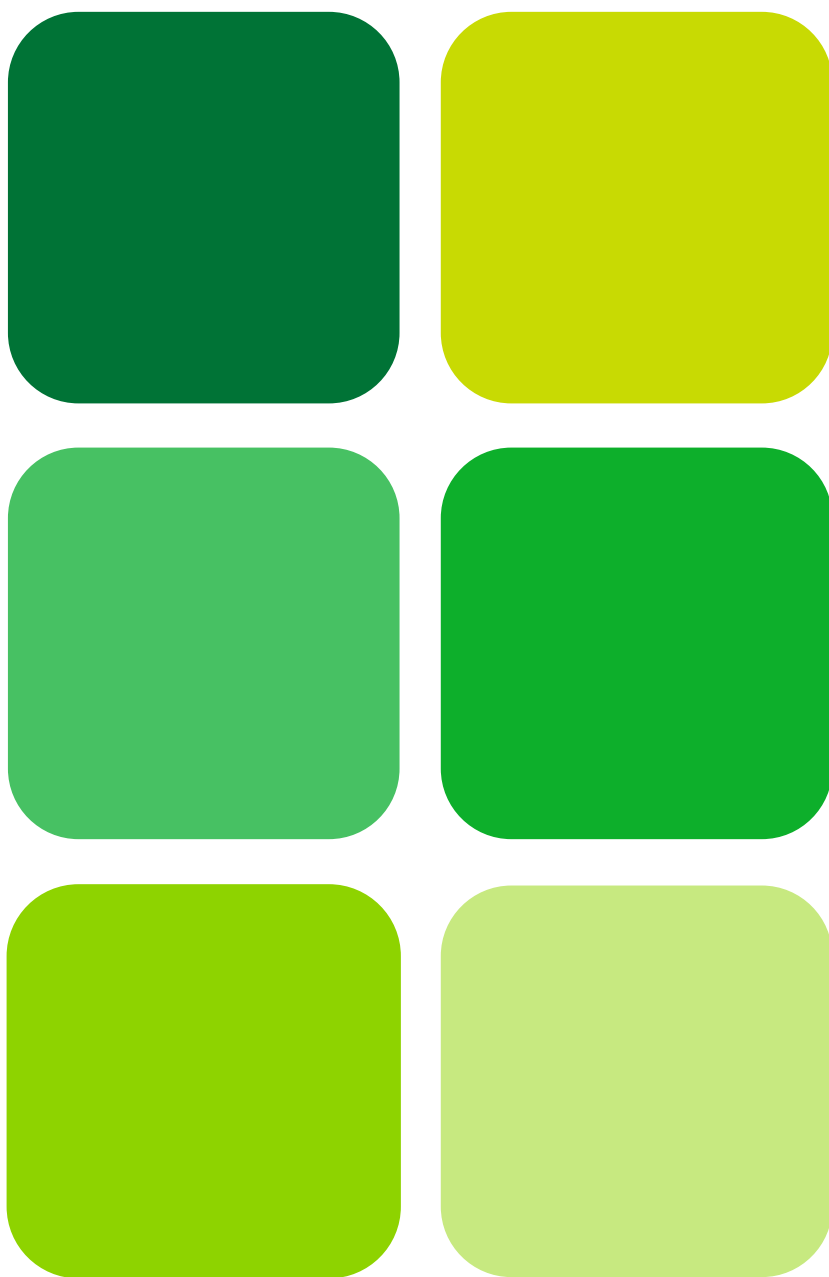




Automatización industrial

Roberto Sanchis Llopis
Julio Ariel Romero Pérez
Carlos Vicente Ariño Latorre

Automatización industrial

Roberto Sanchis Llopis
Julio Ariel Romero Pérez
Carlos Vicente Ariño Latorre



UNIVERSITAT
JAUME I

INGENIERÍA DE SISTEMAS INDUSTRIALES Y DISEÑO

■ Codi d'assignatura 345

Edita: Publicacions de la Universitat Jaume I. Servei de Comunicació i Publicacions
Campus del Riu Sec. Edifici Rectorat i Serveis Centrals. 12071 Castelló de la Plana
<http://www.tenda.uji.es> e-mail: publicacions@uji.es

Col·lecció Sapientia, 31
Primera edició, 2010
www.sapientia.uji.es

ISBN: 978-84-693-0994-0



Aquest text està subjecte a una llicència Reconeixement-NoComercial-Compartir Igual de Creative Commons, que permet copiar, distribuir i comunicar públicament l'obra sempre que especifique l'autor i el nom de la publicació i sense objectius comercials, i també permet crear obres derivades, sempre que siguin distribuïdes amb aquesta mateixa llicència.
<http://creativecommons.org/licenses/by-nc-sa/2.5/es/deed.ca>

Índice general

Índice general	3
1. Introducción a los Automatismos	5
1.1. Introducción	5
1.2. Definición de Automatismo	6
1.3. Clasificación tecnológica	7
1.4. La seguridad en los automatismos	9
1.5. Representación de los automatismos	9
1.6. Automatismos combinacionales y secuenciales	12
2. Tecnología de la Automatización Industrial	14
2.1. Descripción tecnológica del automatismo	14
2.2. Detectores de proximidad	15
2.3. Finales de carrera	15
2.4. Características de los detectores sin contacto. Tipos de salida . .	18
2.5. Conexión de detectores en serie y en paralelo	23
2.6. Detectores ópticos o fotoeléctricos	27
2.7. Detectores de proximidad inductivos	40
2.8. Detectores de proximidad capacitivos	44
2.9. Otros detectores de proximidad	45
2.10. Otros detectores (de nivel, de presión, de temperatura)	46
2.11. Transductores de posición	47
2.12. Otros transductores	54
2.13. Actuadores	55
3. Modelado de sistemas de eventos discretos... GRAFCET	68
3.1. Sistemas de eventos discretos	68
3.2. Modelado mediante diagrama de estados	68
3.3. Las redes de Petri	69
3.4. Diagrama de etapa transición: El Grafcet	72
3.5. Ejemplos	86
3.6. Diseño estructurado de automatismos	92
3.7. Forzado de Grafcets	92
3.8. Jerarquía entre Grafcets parciales	93

4. Guía de estudio de los modos de marcha y paro: GEMMA	99
4.1. Introducción	99
4.2. Descripción de la Guía GEMMA	99
4.3. Utilización de la Guía GEMMA	102
4.4. Ejemplos	107
5. El Autómata Programable Industrial	110
5.1. Definición	110
5.2. Arquitectura del API	110
5.3. Constitución física	113
5.4. Descripción funcional	115
5.5. Módulos de entrada/salida y módulos de comunicación	119
5.6. Unidades de programación	126
5.7. Programación de los autómatas programables industriales	127
5.8. Criterios de selección de un autómata programable industrial	136
5.9. Características de un PLC comercial. CQM1 de OMRON	137
6. Implementación de sistemas de control secuencial...	147
6.1. Introducción	147
6.2. Implementación del algoritmo de control a partir del Grafcet	147
6.3. Implementación del algoritmo de control cuando hay forzados	172
6.4. Implementación del algoritmo de control en otras plataformas	181
7. Sistemas de control distribuido	185
7.1. Introducción	185
7.2. Niveles de un sistema de control distribuido	185
7.3. Redes de comunicación	187
7.4. Enlaces estándares para comunicaciones digitales	195
7.5. Buses de campo	197
7.6. Ethernet como red de campo: Ethernet conmutada	206
7.7. Redes Inalámbricas	208
7.8. Sistemas SCADA	213
A. Ejercicios sobre sensores y actuadores	227
B. Problemas de diseño de automatismos	234
Bibliografía	256

CAPÍTULO 1

INTRODUCCIÓN A LOS AUTOMATISMOS

1.1 INTRODUCCIÓN

En las últimas décadas se ha seguido la tendencia de automatizar de manera progresiva procesos productivos de todo tipo. Esta tendencia ha sido y sigue siendo posible gracias al desarrollo y abaratamiento de la tecnología necesaria. La automatización de los procesos de producción persigue los objetivos:

- Mejorar la calidad y mantener un nivel de calidad uniforme.
- Producir las cantidades necesarias en el momento preciso.
- Mejorar la productividad y reducir costes.
- Hacer más flexible el sistema productivo (facilitar los cambios en la producción).

Estos objetivos se han convertido de hecho en requisitos indispensables para mantener la competitividad, por lo que el aumento del nivel de automatización de los procesos es simplemente una necesidad para sobrevivir en el mercado actual.

Se pueden distinguir varios niveles en la automatización de un proceso productivo:

1. Nivel de máquina. En este nivel se considera la automatización de una máquina que realiza una tarea productiva simple determinada.
2. Nivel de célula (de grupo). En este nivel se considera el control automatizado de un conjunto de máquinas que trabajan conjunta y coordinadamente para realizar un proceso de producción más complejo.
3. Nivel de planta. En este nivel se considera el control automatizado de toda la planta de producción que trabaja de forma coordinada para cumplir unos objetivos de producción global de la fábrica.
4. Nivel de empresa. En este nivel se considera el conjunto de la empresa (gestión, ventas, producción).

Los niveles 3 y 4 requieren de una red informática que permita el flujo de todos los datos de la empresa relacionados con la producción y la gestión. En esencia estos niveles se implementan mediante ordenadores conectados entre sí y con las células de producción del nivel 2.

En el nivel 2 puede haber una red local de comunicación entre los distintos elementos de una célula (si las máquinas están muy separadas). La implementación de los niveles 1 y 2 se realiza mediante sensores, accionadores y equipos de control.

Un automatismo es en esencia una máquina o un proceso automatizado. Los automatismos definen, por tanto, los niveles 1 y 2.

1.2 DEFINICIÓN DE AUTOMATISMO

Se define un sistema (máquina o proceso) automatizado como aquel capaz de reaccionar de forma automática (sin la intervención del operario) ante los cambios que se producen en el mismo, realizando las acciones adecuadas para cumplir la función para la que ha sido diseñado. La figura 1.1 muestra la estructura típica de un sistema automatizado.

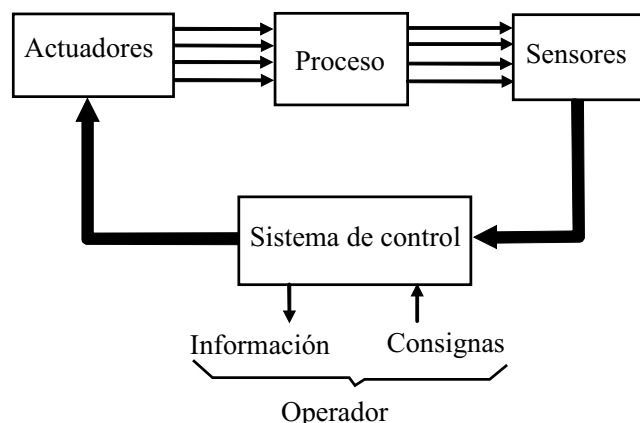


Figura 1.1: Estructura de un sistema automatizado.

Como se observa, se trata de un sistema en bucle cerrado, donde la información sobre los cambios del proceso captada por los sensores es procesada dando lugar a las acciones necesarias, que se implementan físicamente sobre el proceso por medio de los actuadores. Este sistema de control se comunica eventualmente con el operario, recibiendo de éste consignas de funcionamiento, tales como marcha, paro, cambio de características de producción, etc... y comunicándole información sobre el estado del proceso (para la supervisión del correcto funcionamiento).

Se denomina automatismo al sistema completo, aunque con este término suele hacerse referencia fundamentalmente al sistema de control, ya que es éste el que produce de forma automática las acciones sobre el proceso a partir de

la información captada por los sensores. Las señales de entrada y de salida pueden ser de cualquier tipo, sin embargo el concepto tradicional de automatismo se utiliza para sistemas de eventos discretos (también llamados sistemas secuenciales) en los que esas señales son binarias, es decir, solo pueden tomar 2 valores, activa o inactiva (estos valores suelen representarse como un 1 ó un 0). En ese caso el sistema de control implementa el algoritmo de lógica binaria que relaciona los valores que van tomando en cada instante las entradas (1 ó 0) con los valores que deben ir tomando en cada instante las salidas (también 1 ó 0) para que el sistema funcione adecuadamente.

1.3 CLASIFICACIÓN TECNOLÓGICA

En función de la tecnología empleada para la implementación del sistema de control, se puede distinguir entre automatismos cableados y automatismos programados.

Automatismos cableados.

Se implementan por medio de uniones físicas entre los elementos que forman el sistema de control (por ejemplo, contactores y relés unidos entre sí por cables eléctricos). La estructura de conexionado entre los distintos elementos da lugar a la función lógica que determina las señales de salida en función de las señales de entrada. Se pueden distinguir tres tecnologías diferentes:

- Fluídica (neumática o hidráulica).
- Eléctrica (relés o contactores).
- Electrónica estática (puertas lógicas y biestables).

Los inconvenientes fundamentales de los automatismos cableados son:

- Ocupan mucho espacio.
- Son muy poco flexibles. La modificación o ampliación es difícil.
- Solo permiten funciones lógicas simples. No sirven para implementar funciones de control o de comunicación complejas.

Las ventajas de los automatismos cableados son:

- Pueden ser muy robustos.
- Bajo coste para sistemas muy sencillos.
- Es una tecnología muy fácil de entender por cualquier operario.

En general se puede afirmar que los automatismos cableados solo tienen utilidad para resolver problemas muy sencillos (por ejemplo un arranque estrella-triángulo de un motor de inducción).

Automatismos programados.

Se implementan por medio de un programa que se ejecuta en un microprocesador. Las instrucciones de este programa determinan la función lógica que relaciona las entradas y las salidas. Se pueden distinguir 3 formas de implementación:

- **Autómata programable industrial.** Hoy por hoy es el que más se utiliza en la industria. Es un equipo electrónico programable en un lenguaje específico, diseñado para controlar en tiempo real y en ambiente de tipo industrial procesos secuenciales. Se utilizan para el control de máquinas y procesos.
- **Ordenador (PC industrial).** Cada vez se utilizan más. Son ordenadores compatibles con los PC de sobremesa en cuanto a software, pero cuyo hardware está especialmente diseñado para ser robusto en entornos industriales.
- **Microcontrolador.** Son circuitos integrados (“chips”) programables, que incluyen en su interior un microprocesador y la memoria y los periféricos necesarios. Para utilizarlos, normalmente se diseña una tarjeta electrónica específica para la aplicación, que incluye el propio microcontrolador y los circuitos electrónicos de interfaz necesarios para poder conectarse a los sensores y actuadores. Se utilizan sobre todo para sistemas de control de máquinas de las que se van a fabricar muchas unidades, de forma que la reducción de coste por el número de unidades fabricadas justifica la mayor dificultad (y mayor coste) del diseño.

Las ventajas más importantes de los automatismos programados son:

- Permiten una gran flexibilidad para realizar modificaciones o ampliaciones.
- Permiten implementar funciones de control y de comunicación complejas.
- Ocupan poco espacio.

Los inconvenientes respecto de los sistemas cableados son fundamentalmente el mayor coste (solo si el sistema es muy sencillo), la menor robustez y la mayor complejidad de la tecnología. Sin embargo estos inconvenientes cada vez lo son menos, pues el coste se reduce continuamente, cada vez se diseñan equipos más robustos, y los sistemas de programación son cada vez más sencillos.

En resumen, se puede afirmar que la tecnología programada (y en especial los autómatas programables) es superior a la tecnología cableada, salvo en automatismos que sean extremadamente simples.

La naturaleza física de las señales de entrada y salida depende de la tecnología del automatismo. Por ejemplo, un automatismo puramente neumático tiene como entradas señales de presión de aire, y dan como salida señales de presión de aire. Lo más habitual en la industria son los automatismos de naturaleza eléctrica (cableados o programados). En este caso las señales de entrada y de salida son señales eléctricas. Los sensores se encargarán de convertir las magnitudes físicas en señales eléctricas, mientras los actuadores transforman las señales eléctricas en acciones físicas sobre el proceso.

1.4 LA SEGURIDAD EN LOS AUTOMATISMOS

Un aspecto que tiene especial importancia en la implementación de automatismos es la seguridad. Ésta se debe tener en cuenta de dos formas. Por una parte se debe definir la secuencia de operaciones del proceso de forma que se garantice en todo momento la seguridad de los operarios. Por ejemplo, una prensa en la que el operario introduce una chapa para después darle al pulsador de marcha. La secuencia del automatismo debería permitir la puesta en marcha de la prensa solo cuando el operario pulsa de forma simultánea dos pulsadores separados. De esta forma se garantiza que las dos manos quedan fuera de la prensa cuando ésta actúa.

Tener en cuenta la seguridad en la secuencia del automatismo no es, sin embargo suficiente, ya que si por alguna razón falla el sistema de control pueden producirse situaciones de peligro. En función del nivel de riesgo puede ser necesario utilizar en la implementación una tecnología adecuada que garantice la seguridad. Por ejemplo, si el automatismo se implementa mediante un autómatas programable y se quiere garantizar que la apertura de una puerta produzca la parada instantánea de la máquina, no basta con definir la secuencia para que así sea, sino que hay que utilizar un interruptor y un relé de seguridad que corte la alimentación de la máquina independientemente del autómatas programable que la controla.

La tecnología utilizada en la implementación del automatismo deberá tener en cuenta la seguridad, pudiendo ser necesaria la utilización de elementos especiales para implementar alguna de las funciones del automatismo.

1.5 REPRESENTACIÓN DE LOS AUTOMATISMOS

La función lógica implementada por un automatismo se puede representar de diversas formas. Las 2 formas tradicionales son la lógica de contactos y las puertas lógicas, que permiten representar funciones lógicas sencillas. Hay otras formas de representación del automatismo que están a un nivel superior. Sirven para definir funcionalmente un automatismo secuencial completo. Entre ellas se pueden destacar los diagramas de flujo y el GRAFCET.

Lógica de contactos (de relés).

Las variables binarias se representan mediante contactos que están cerrados ó abiertos según esté activa (1) o inactiva (0) la variable en cuestión. La combinación (conexión) de contactos en serie y paralelo permite representar una función lógica. Por ejemplo, la figura 1.2 representa la función lógica $y = a \cdot b + c$.

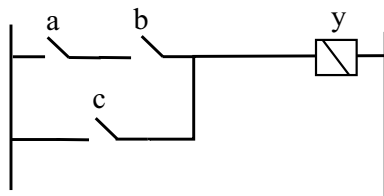


Figura 1.2: Función lógica $y = a \cdot b + c$.

La línea vertical izquierda representa tensión (por ejemplo 24 V) mientras la línea vertical derecha representa la masa (0 V). Si a y b están activas los contactos están cerrados, y la salida (bobina) y queda sometida a 24 V, con lo que se activa. Lo mismo sucede si c está activa.

El origen de esta representación está en la implementación física mediante contactores, que fue la primera forma que se utilizó para implementar los automatismos. De hecho la representación del diagrama de contactos es directamente el cableado de los contactores. Hoy en día la mayoría de automatismos son programados, sin embargo ésta sigue siendo la forma más habitual de representar las ecuaciones lógicas de los automatismos.

La representación de contactos suele utilizar los símbolos de la figura 1.3, en lugar de los interruptores.

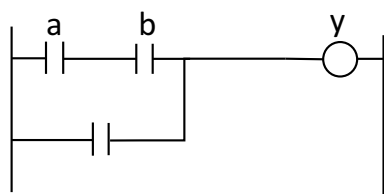


Figura 1.3: Representación en diagrama de contactos de $y = a \cdot b + c$.

Puertas lógicas.

La función lógica se representa mediante compuertas lógicas (puertas AND, NAND, OR y NOR). Por ejemplo, la la figura 1.4 representa la función lógica $y = a \cdot b + c$.

Es la representación utilizada cuando el automatismo se implementa con circuitos electrónicos digitales.

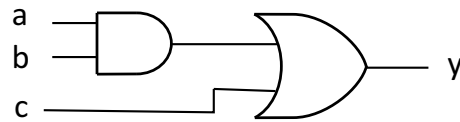


Figura 1.4: Representación mediante puertas lógicas de $y = a \cdot b + c$.

Diagrama de flujo.

Es una forma de representación de más alto nivel que las dos anteriores. Consiste en un diagrama con dos tipos de elementos, los nodos de decisión y los nodos de actuación o tratamiento. Por ejemplo, el diagrama de flujo de un arrancador estrella-triángulo se puede representar mediante diagrama de flujo tal y como muestra la figura 1.5.

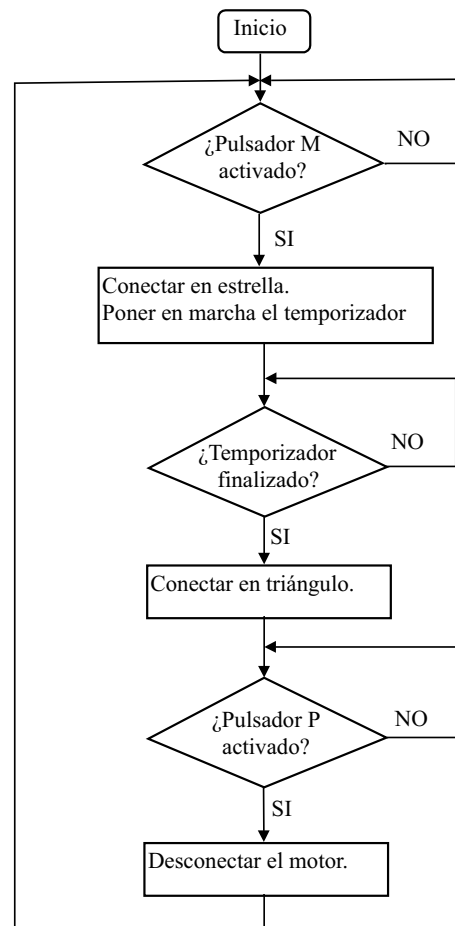


Figura 1.5: Diagrama de flujo del arrancador estrella triángulo.

Esta forma de representación de alto nivel explica el funcionamiento del automatismo. Sin embargo, para la implementación final en el equipo de control (autómata programable por ejemplo) es necesario traducir el diagrama de flujo a una representación de bajo nivel (como el diagrama de contactos, por ejemplo).

GRAFCET.

También se trata de una representación de alto nivel de los automatismos. Se llama también diagrama de etapa-transición. En esta representación se tienen cuadrados que representan las etapas del automatismo, y transiciones entre ellas. El ejemplo del arrancador anterior se puede representar como se muestra en la figura 1.6:

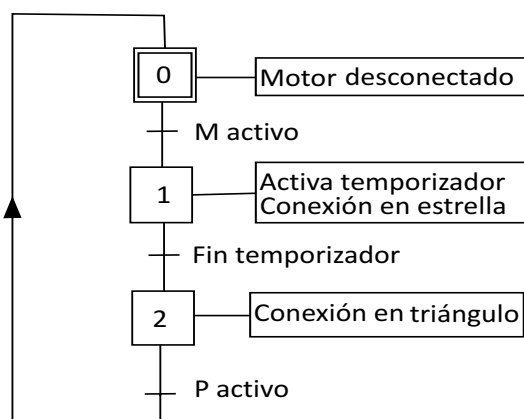


Figura 1.6: Diagrama Grafcet del arrancador estrella triángulo.

Estas formas de representación del automatismo de alto nivel son muy útiles para explicar el funcionamiento del proceso y facilitar el diseño. La implementación final, sin embargo, se hace generalmente con la representación de contactos, por lo que es necesario traducir estos diagramas a esa forma de representación de bajo nivel. Algunos autómatas programables industriales permiten ser programados utilizando diagramas GRAFCET de alto nivel, realizándose de forma automática la traducción al lenguaje de bajo nivel de forma transparente para el usuario.

1.6 AUTOMATISMOS COMBINACIONALES Y SECUENCIALES

Un automatismo combinacional es aquel en el que el valor de las salidas en un instante depende únicamente del estado en ese mismo instante de las entradas (y no de los valores pasados). Los métodos para diseñar estos automatismos son los propios de la electrónica digital combinacional: tablas de verdad y métodos de simplificación (Karnaugh). Estos automatismos son muy limitados y rara vez sirven por sí solos para resolver problemas de automatización reales.

Un automatismo secuencial es aquel en el que el valor de las salidas en un instante depende del valor de las entradas no solo en el instante actual, sino del valor que han ido tomando en instantes anteriores, es decir, las salidas actuales dependen de la historia del proceso. Existen métodos tradicionales en electrónica digital para diseñar circuitos que implementen esos automatismos

secuenciales. En esencia se trata de describir el automatismo mediante un diagrama de estados. A cada estado se le asigna una combinación de variables binarias. El valor de los estados sirve para memorizar la historia pasada del proceso. A continuación se construye la tabla de transiciones, y a partir de ésta se obtienen los circuitos combinacionales que completan el sistema. El esquema final es el representado en la figura 1.7

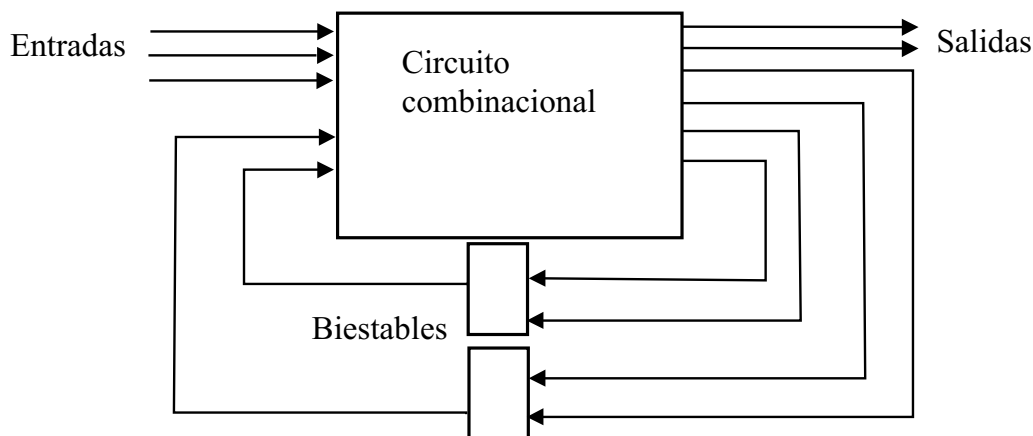


Figura 1.7: Esquema Automatismo combinacional-secuencial.

Los biestables memorizan el estado actual del proceso, y el circuito combinacional determina las salidas hacia el proceso y el cambio de esos estados en función de las entradas y los propios estados. El circuito combinacional tiene, por tanto dos partes. Una de ellas se encarga de ir cambiando los valores de los biestables según va cambiando el proceso (para saber en todo momento el estado del mismo). La otra parte se encarga de activar las salidas hacia los actuadores en función de las entradas y del estado actual. Los métodos de diseño de circuitos secuenciales utilizados en electrónica digital son útiles para resolver problemas sencillos. Si el automatismo es muy complejo el circuito que se obtiene puede llegar a ser complicado, con la consiguiente dificultad en su diseño e implementación, aunque en la práctica el mayor inconveniente es la falta de flexibilidad para realizar modificaciones o ampliaciones posteriores.

Los automatismos secuenciales implementados habitualmente en la industria son programados. Los métodos de diseño de estos automatismos son, en cierta forma, parecidos al anterior. En primer lugar se realiza una descripción de alto nivel del automatismo (mediante el GRAFCET, por ejemplo), y a partir de él se definen las funciones lógicas necesarias para implementarlo. El estado del proceso se memoriza en una serie de variables, siendo esas funciones lógicas las que determinan los cambios de esas variables (para actualizar el estado del proceso conforme éste evoluciona) y de las salidas. Sin embargo, hay diferencias importantes que simplifican en gran medida el diseño y la implementación de automatismos complejos en dispositivos programables. Además, la modificación o ampliación del automatismo resulta muy sencilla.

CAPÍTULO 2

TECNOLOGÍA DE LA AUTOMATIZACIÓN INDUSTRIAL

2.1 DESCRIPCIÓN TECNOLÓGICA DEL AUTOMATISMO

La descripción tecnológica del automatismo es el conjunto de elementos físicos que lo forman. En concreto, estos elementos son los sensores, los actuadores y el sistema de control. Por otra parte está la descripción funcional, que se refiere a las características de funcionamiento del sistema automatizado.

Los sensores son los elementos que permiten obtener información de lo que sucede en el proceso. Se pueden distinguir dos tipos de sensores, según la información que proporcionan:

- **Detectores.** Son los sensores que proporcionan una salida binaria (activa o inactiva). Son los que más se utilizan en los automatismos secuenciales. Los más frecuentes son los detectores de proximidad, que normalmente detectan la presencia de un objeto, aunque también son frecuentes los detectores de nivel, de temperatura o de presión.
- **Captadores.** Son los sensores que proporcionan una salida continua proporcional a una magnitud física. Esta salida puede ser analógica (en tensión o en corriente), o digital (codificada en binario, o en forma de pulsos). Los captadores se utilizan en los sistemas de control continuo (como los PID), en los que se controla una variable continua. En automatismos secuenciales también son frecuentes, utilizándose el valor continuo para obtener un valor binario mediante comparación con un límite determinado (la temperatura es superior o inferior a 70°, por ejemplo).

Los actuadores son los elementos que permiten traducir las señales eléctricas de salida del sistema de control en actuaciones físicas sobre el proceso. Fundamentalmente pueden ser neumáticos, hidráulicos o eléctricos. En los tres casos se distingue entre preactuadores y actuadores. Los preactuadores neumáticos e hidráulicos son las electroválvulas, mientras los actuadores son principalmente los cilindros (aunque también hay actuadores de giro neumáticos o hidráulicos). Los preactuadores eléctricos pueden ser los relés o contactores, o equipos más

complejos como variadores de frecuencia, mientras que los actuadores eléctricos por excelencia son las máquinas eléctricas (principalmente rotativas aunque también las hay de movimiento lineal). También son actuadores eléctricos las resistencias calefactoras.

El sistema de control está formado normalmente por un autómata programable industrial (aunque también puede ser un PC industrial o una tarjeta basada en microcontrolador).

2.2 DETECTORES DE PROXIMIDAD

Definición: los detectores de proximidad son aquellos que se activan o desactivan en función de la presencia o ausencia de un objeto.

Clasificación: los detectores de proximidad se pueden clasificar por el principio físico en que se basan en:

- Final de carrera (mecánico).
- Fotoeléctrico (óptico).
- Inductivo.
- Capacitivo.
- Magnético.
- Ultrasonidos.

Los más simples son los finales de carrera, ya que se basan en la apertura o cierre de un interruptor por el contacto físico del objeto a detectar. El resto de detectores no necesitan el contacto del objeto, pero en cambio requieren de una electrónica adicional que procesa la señal correspondiente para decidir si hay o no un objeto en la proximidad, activando en su caso la salida correspondiente. En el caso de los finales de carrera, la salida del detector es siempre un contacto que abre o cierra, mientras que el resto de detectores puede tener una salida de contacto (o relé), o bien de otros tipos (transistor o triac, por ejemplo).

2.3 FINALES DE CARRERA

Son interruptores que se abren o cierran debido al contacto físico del objeto a detectar. Pueden tener un solo contacto, o varios de ellos. Es habitual que tengan un contacto normalmente cerrado y otro normalmente abierto. Estos contactos suelen tener una tensión nominal de 240V, y una corriente de varios amperios. Se pueden conectar entre la alimentación y la carga, o entre la carga y masa.

La carga puede ser la bobina de un contactor o relé, o cualquier elemento que se active al conectar sus bornes a una diferencia de tensión. Un ejemplo típico de carga es la entrada digital de un autómata programable. La figura 2.1 muestra una entrada digital a 24 V de un autómata programable y la bobina de un contactor o electroválvula representando la carga estándar.

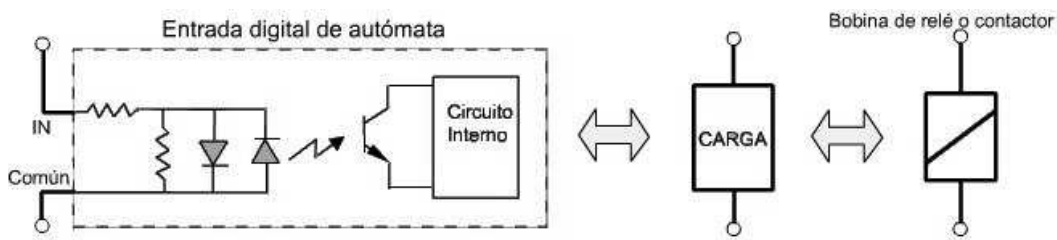


Figura 2.1: Representación de una carga estándar.

La carga se activará cuando quede sometida a su tensión nominal. La entrada del autómata se activa a 24 V, tanto de tensión continua como de alterna. La bobina puede activarse a 24 V de continua, 24 V de alterna o 220 V de alterna, según el tipo de bobina. El circuito de salida del detector será el que, al activarse, pondrá la carga a tensión tal y como muestra la figura 2.2.

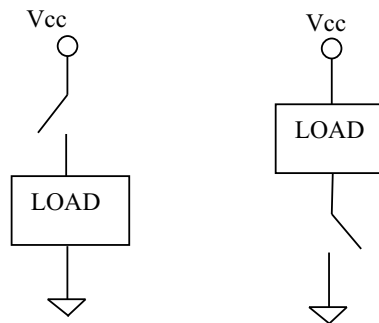


Figura 2.2: Activación de la carga.

La forma de actuación (el elemento de mando con el que el objeto contacta) puede ser:

- Pulsador.
- Pulsador con roldana.
- Palanca.
- Palanca con roldana.
- Palanca flexible.

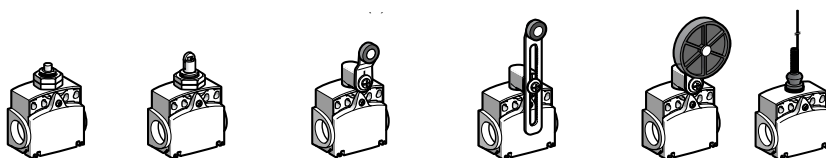


Figura 2.3: Finales de carrera con diferentes elementos de mando, de Schneider Electric.

Los finales de carrera pueden tener un único contacto NO o NC, pero normalmente incorporan al menos dos contactos tal y como se muestra en la figura 2.4

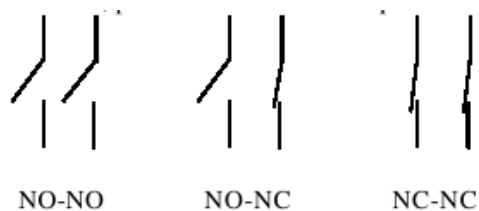


Figura 2.4: Contactos NO y NC.

Según la forma en que se produce la apertura y cierre de los contactos se pueden distinguir dos tipos:

- **De ruptura lenta.** Es el movimiento del objeto detectado el que directamente produce el movimiento de los contactos provocando su apertura o cierre. Funciona bien si el objeto no se mueve a una velocidad demasiado baja. El diagrama de funcionamiento es el indicado en la figura 2.5

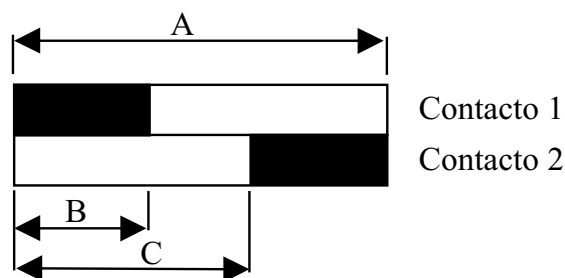


Figura 2.5: Contacto de ruptura lenta: A = carrera total, B = carrera de accionamiento/desaccionamiento del Contacto 1, C = carrera de accionamiento/desaccionamiento del Contacto 2.

A = carrera total. B = carrera de accionamiento/desaccionamiento del contacto 1. C = carrera de accionamiento/desaccionamiento del contacto 2. En la posición del extremo izquierdo (donde el objeto no toca el final de carrera), el Contacto 1 está cerrado y el 2 abierto. Cuando el objeto ha desplazado al elemento de mando una distancia B, el Contacto 1 se abre, y el 2 permanece abierto. Cuando el objeto desplaza al elemento de mando una distancia C, el Contacto 2 se cierra, permaneciendo abierto el Contacto 1. Cuando el objeto se retira, el proceso de activación/desactivación sigue el mismo esquema: al llegar a C se abre el Contacto 2, y al llegar a B se cierra el Contacto 1.

- **De ruptura brusca.** El mecanismo de apertura y cierre dispone de unos muelles, de forma que cuando el objeto sobrepasa una posición, estos muelles saltan y producen el cierre (o apertura) brusca de los contactos. Funcionan bien aunque el objeto a detectar se mueva muy despacio.

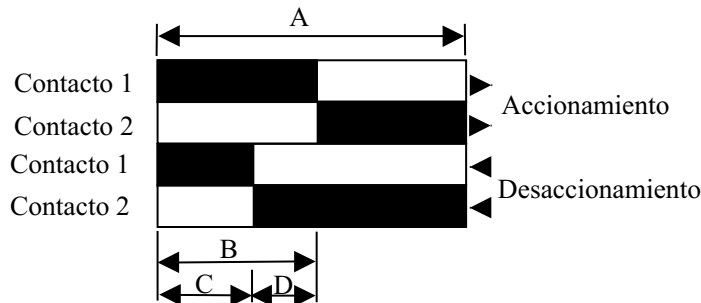


Figura 2.6: Contacto de ruptura brusca: A = carrera total, B = carrera de accionamiento, C = carrera de desaccionamiento, D = histéresis.

En este caso, hay un diagrama de apertura para el accionamiento (cuando el objeto se aproxima al elemento de mando) y otro para el desaccionamiento (cuando el objeto se aleja del elemento de mando), ya que la posición en que saltan los muelles es distinta en los dos casos tal y como se muestra en la figura 2.6. En el momento en que saltan los muelles, un contacto se abre y el otro se cierra de forma instantánea.

Los finales de carrera se utilizan normalmente para detectar el movimiento de mecanismos (el final de carrera de un mecanismo), es decir, para detectar la posición de una parte móvil de una máquina, no para detectar objetos (productos) que se manipulan, aunque en alguna ocasión pueden utilizarse para detectar objetos (por ejemplo, cajas grandes que se manipulan).

2.4 CARACTERÍSTICAS DE LOS DETECTORES SIN CONTACTO. TIPOS DE SALIDA

El final de carrera es el único elemento que necesita que el objeto a detectar contacte físicamente con el sensor. Los demás tipos detectan el objeto sin contacto físico. Para estos detectores sin contacto se definen dos conjuntos fundamentales de características:

- **Características de detección.** Materiales que puede detectar, distancia de detección, histéresis. Estas características dependen mucho del tipo de detector, por lo que se describirán para cada tipo por separado.
- **Características de salida y alimentación.** Tipo de señal de salida que proporcionan (transistor, relé, triac, etc.), frecuencia (o retardo) de

conmutación, tensión de alimentación. Estas características son comunes para todos los tipos (excepto los finales de carrera).

Respecto a las características de salida y alimentación, éstas están relacionadas con la forma de conexión de la carga sobre la que actúa la salida del detector.

En el mercado hay muchos modelos con salidas y alimentación diferentes. Los siguientes tipos son los más utilizados:

- 3 hilos NPN, con alimentación continua (normalmente 24 V).
- 3 hilos PNP, con alimentación continua (normalmente 24 V).
- **NPN / PNP programable** con alimentación continua.
- 2 hilos con alimentación continua o alterna.
- Con salida a relé. Alimentación alterna o continua.
- Salida analógica.
- 2 hilos NAMUR.
- Otros. Salida triac, salida MOSFET, etc.

3 hilos NPN, con alimentación continua

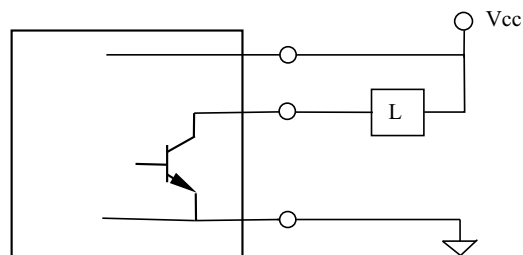


Figura 2.7: Detector con salida tipo NPN.

La carga se conecta tal y como muestra la figura 2.7 entre el colector del transistor y V_{cc} (que suele ser 24 Vdc), de forma que cuando el transistor está cortado, la carga está en circuito abierto, mientras que cuando el transistor satura, la carga queda sometida prácticamente a la tensión de alimentación (la caída de tensión del transistor es menor de 1 V en saturación).

Normalmente suele ser NO. En ese caso el transistor satura cuando se detecta el objeto, estando en corte cuando no se detecta. También puede ser NC, actuando al contrario, o tener dos salidas (en ese caso tiene 4 hilos), una NO y otra NC.

El tiempo de retardo de este tipo de salida (tiempo transcurrido entre que se detecta el objeto y se satura el transistor) es muy bajo, estando en el rango de $30 \mu s$ a $1 ms$.

La corriente máxima que puede conducir el transistor en saturación suele ser del orden de 100 mA. Evidentemente la carga debe consumir una corriente inferior a ésta.

La caída de tensión en saturación es muy baja (menor de 1 V), por lo que a efectos de la carga se puede considerar como un cortocircuito.

3 hilos PNP, con alimentación continua

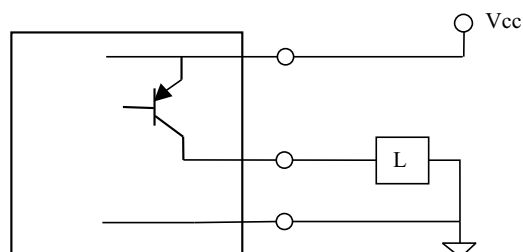


Figura 2.8: Detector con salida tipo PNP.

La carga se conecta tal y como muestra la figura 2.8 entre el colector del transistor y masa, de forma que cuando el transistor está cortado, la carga está en circuito abierto, mientras que cuando el transistor satura, la carga queda sometida prácticamente a la tensión de alimentación, normalmente de 24 Vdc (la caída de tensión del transistor es menor de 1 V en saturación).

Existen modelos NO, NC y ambos (con 4 hilos), exactamente igual que en los NPN.

Las demás características (tiempo de retardo y corriente máxima) son iguales a los NPN.

NPN / PNP programable con alimentación continua

Tiene dos transistores de salida, uno NPN y otro PNP, de forma que se puede seleccionar cuál de ellos se satura cuando se produce la detección. Tiene 4 hilos: los 2 de alimentación, el de salida y el de configuración. En función de que el hilo de configuración se conecte a masa o a tensión, se posibilita la activación del transistor NPN o del PNP, con lo que la salida es NPN o PNP.

2 hilos con alimentación continua o alterna

La carga se conecta tal y como se muestra en la figura 2.9 en serie (en el hilo que va a tensión o en el hilo que va a masa, indistintamente). El detector a 2 hilos trata de asemejarse al máximo a un contacto que se abre o se cierra (como el final de carrera). La diferencia con aquel es que cuando está en estado abierto (no detecta objeto) necesita consumir una corriente residual para que los circuitos electrónicos internos de detección funcionen. Por otra parte,

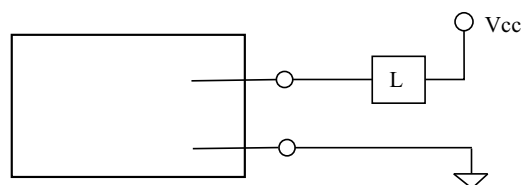


Figura 2.9: Detector con salida de 2 hilos.

cuando está en estado cerrado (se detecta un objeto) necesita tener una caída de tensión mínima para que los circuitos funcionen. De esta forma se define:

- Corriente residual en estado abierto: suele ser del orden de 1 a 2 mA.
- Tensión residual en estado cerrado: suele ser del orden de 4 a 10 V.

Hay que comprobar que tanto la tensión como la corriente residual son compatibles con la carga, es decir, si el detector está abierto, la corriente residual no debe activar la carga, mientras que si está cerrado, la tensión residual debe ser lo bastante baja para que la carga se active con la tensión que le queda (la alimentación menos la tensión residual en estado cerrado). En el caso de alimentación alterna a 240 Vac, eso no suele ser un problema. Cuando la alimentación es de 24 Vdc, podría haber un problema, por lo que es necesario comprobar que lo anterior se cumple.

El elemento interno que produce la conmutación puede ser un transistor, en cuyo caso el tiempo de respuesta es rápido (menor de 1 ms). Si el detector se alimenta en alterna a 240 V y el elemento que conmuta es un triac o un tiristor, el retardo es mayor (del orden de 10 ms) debido a que la conmutación se produce en el paso por cero (retardo máximo de medio ciclo).

Con salida a relé. Alimentación alterna o continua

Suele tener 5 hilos, 2 de ellos son para la alimentación del sensor, y los otros 3 son los terminales del contacto del relé (que está aislado de la alimentación). Cuando se detecta el objeto, el contacto del relé conmuta. La carga se puede conectar como en una salida PNP (figura 2.11) o como en una salida NPN (figura 2.10). Además, se puede elegir conectarlo como NO o NC.

Una característica importante es el aislamiento galvánico que hay entre la salida y la alimentación. La ventaja de este tipo de salida es que puede conmutar corrientes de intensidad elevada (típicamente hasta 5 A a 240 V), por lo que puede atacarse directamente una carga de bastante potencia. El inconveniente fundamental es que el número de operaciones está limitado por el desgaste de los contactos, por lo que se debe sustituir cada cierto tiempo. Otro inconveniente es el elevado tiempo de retardo, que es del orden de 15 ms.

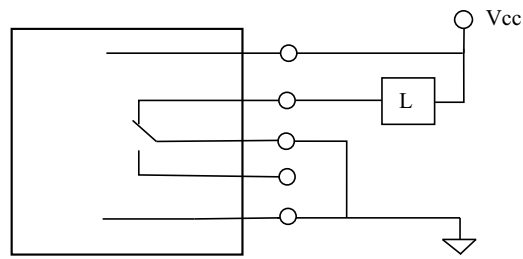


Figura 2.10: Detector con salida a relé conectado como NPN, NC.

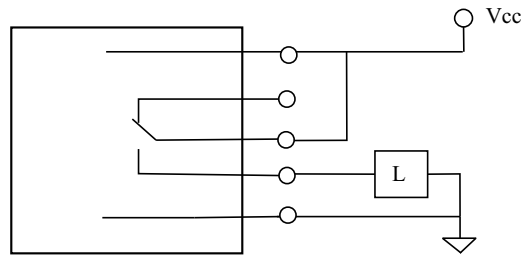


Figura 2.11: Detector con salida a relé conectado como PNP, NO.

Salida analógica

La salida varía de forma continua según cambia la magnitud detectada (que depende del tipo de sensor). Según esto, ya no se trataría de un detector sino de un captador. Sin embargo, se clasifican dentro de los detectores porque su objetivo no es medir la magnitud (no tienen precisión para ello), sino simplemente dar una salida variable que permita distinguir entre varios objetos detectados. Por ejemplo, un detector inductivo con salida analógica da una salida mayor cuanto más cerca esté el objeto a detectar. Esta salida se puede utilizar para distinguir si se trata de un objeto u otro.

La salida puede ser en tensión o en corriente:

- Tensión 0 -10 V. La salida requiere una impedancia de carga mínima. Un valor típico mínimo puede ser de $1\text{ K}\Omega$. El esquema de conexión es el indicado en la figura 2.12.

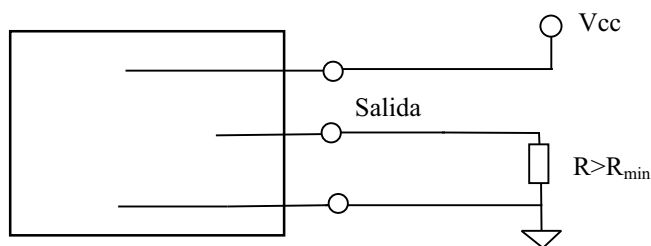


Figura 2.12: Detector con salida analógica de tensión.

- Corriente 4 - 20 mA. La salida en corriente requiere una impedancia de carga menor que un valor máximo admisible. Un valor típico máximo es de 500 Ω . El esquema de conexión es el indicado en la figura 2.13.

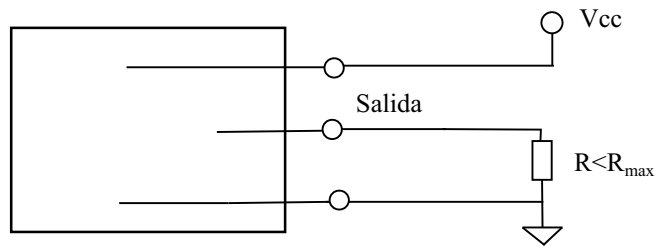


Figura 2.13: Detector con salida analógica de corriente.

Otros. Salida triac, salida MOSFET, etc.

Se utilizan menos. La salida triac requiere alimentación alterna de 240 V, y permite conmutar una carga alterna con bastante corriente. La salida MOSFET suele ir acompañada de alimentación continua, y permite conmutar una corriente mayor que las salidas NPN o PNP.

2 hilos NAMUR

Es un detector a 2 hilos especial, que no conmuta, sino que simplemente varía la corriente que absorbe en función de la cantidad de magnitud física detectada. Se utiliza para aplicaciones especiales, fundamentalmente para ambientes con riesgo de incendio o explosión. Necesita una etapa electrónica de amplificación para dar lugar a una salida de alguno de los tipos descritos.

2.5 CONEXIÓN DE DETECTORES EN SERIE Y EN PARALELO

Para conectar en serie o en paralelo varios detectores hay que tener en cuenta varias consideraciones en función del tipo de salida de esos detectores.

Los detectores con salida de contacto (finales de carrera o detectores con salida a relé) se pueden conectar en serie o en paralelo sin ningún problema.

La conexión en paralelo tiene como resultado la función lógica OR, es decir, la carga se activa si se activa cualquiera de los detectores como se indica en la figura 2.14.

La conexión en serie implementa la función lógica AND, es decir, la carga se activa si se activan todos los detectores como se indica en la figura 2.15.

Si la salida de los detectores no es de relé, hay que comprobar que la conexión eléctrica en serie o paralelo es compatible. Por ejemplo, dos salidas NPN se pueden conectar en paralelo sin problemas, ya que el emisor de las dos se conecta a masa, y los colectores pueden unirse sin problemas. Lo mismo

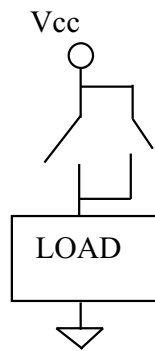


Figura 2.14: Conexión en paralelo de sensores.

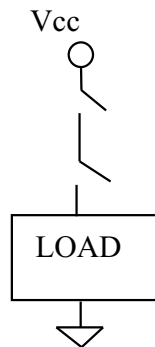


Figura 2.15: Conexión en serie de sensores.

sucede con dos salidas tipo PNP. Sin embargo, no se puede conectar en paralelo un NPN con un PNP, pues uno de ellos va unido a masa y el otro a tensión.

Una cuestión importante a tener en cuenta cuando se conectan varios detectores en serie o paralelo es el **retardo a la disponibilidad**. Es decir, el tiempo que tarda el sensor en estar disponible para actuar desde que se activa la alimentación del mismo. Cuando la alimentación de un sensor depende de que se active otro (un detector no está alimentado hasta que el otro no se activa), el sensor tiene un retardo desde que el primero se activa hasta que el segundo está en disposición de detectar.

Se pueden describir algunas conexiones posibles, aunque en una aplicación particular hay que estudiar el problema eléctrico con detalle.

Conexiones posibles:

Conexión en paralelo de varios NPN, o de varios PNP

No presenta ningún problema si los detectores son iguales (todos NPN o todos PNP), como muestra la figura 2.16. Sin embargo, no sería posible si se tiene un NPN y un PNP, pues la conexión de la carga es diferente en los dos.

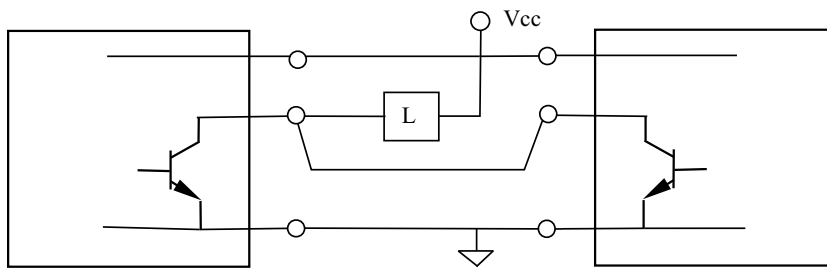


Figura 2.16: Conexión en paralelo de dos detectores tipo NPN.

Conexión en serie de varios detectores de 2 hilos

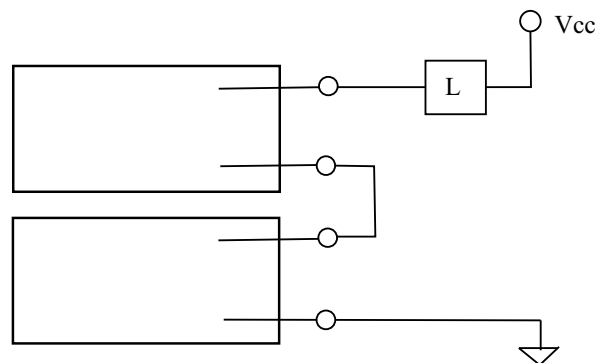


Figura 2.17: Conexión en serie de dos detectores de 2 hilos.

Es posible siempre que las tensiones y corrientes residuales sean compatibles, y que la tensión de alimentación admisible tenga un rango amplio. Cuando están en estado abierto la tensión de alimentación se reparte entre los detectores en serie en función de su corriente residual, por lo que es conveniente que éstas sean parecidas. Por otra parte, si uno está cerrado y el otro abierto, el abierto queda sometido a casi toda la tensión de alimentación. Debido a esto, es necesario que la tensión de alimentación admisible tenga un rango amplio que cubra los dos casos extremos. Cuando se encuentran en estado cerrado, las tensiones residuales se suman. Este efecto se ha de tener en cuenta comprobando que la carga puede activarse a pesar de esa caída de tensión.

Conexión en serie de un detector de 2 hilos con un contacto mecánico

Es posible, aunque cuando el contacto mecánico está abierto, el detector de 2 hilos no funciona debido a que no tiene la corriente residual mínima. Por este motivo, hay que tener en cuenta el retardo a la disponibilidad.

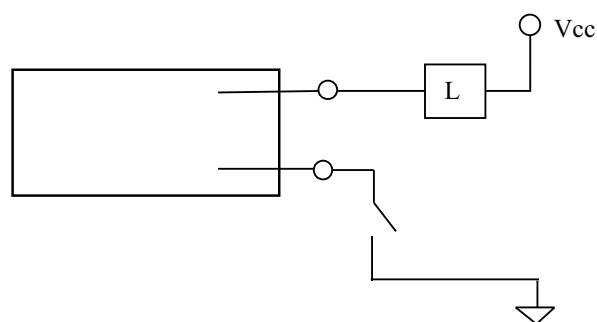


Figura 2.18: Conexión en serie de un detector de 2 hilos con un contacto mecánico.

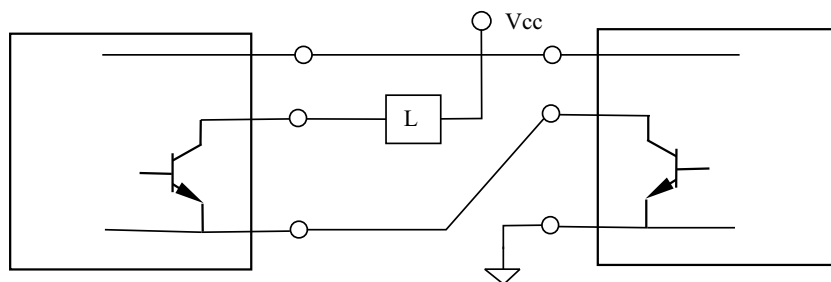


Figura 2.19: Conexión en serie de varios detectores NPN.

Conexión en serie de varios detectores NPN, o de varios PNP

Como se puede observar en la figura 2.19, el primer detector no está alimentado (la conexión de masa está al aire) hasta que el segundo detector se activa. Esto implica que hay que tener en cuenta el retardo a la disponibilidad.

Además, la corriente que absorbe el segundo detector es la suma de la de la carga más la de alimentación del primer detector. Esta corriente total debe ser menor que la máxima admitida por el detector. Por otra parte la caída de tensión cuando están todos activos es la suma de las caídas de tensión, por lo que debe ser lo bastante baja para que se active la carga.

Conexión en paralelo de detectores de 2 hilos

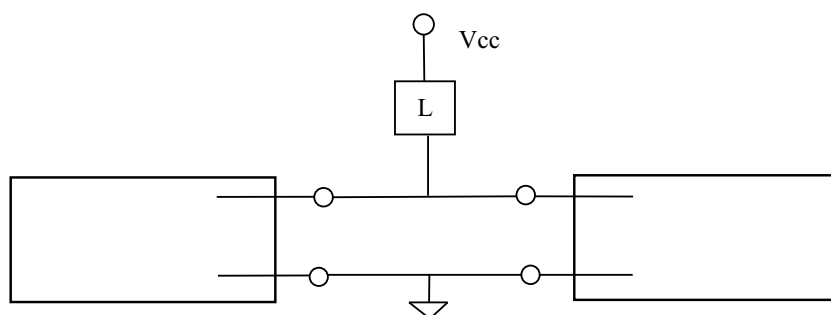


Figura 2.20: Conexión en paralelo de detectores de 2 hilos.

Es posible, pero si uno de ellos está cerrado, los demás, quedan sometidos únicamente a la tensión residual, que no es suficiente para que funcionen, por lo que hay que tener en cuenta el retardo a la disponibilidad. Por otra parte, cuando están abiertos, las corrientes residuales se suman, por lo que hay que tener muy en cuenta la corriente mínima de activación de la carga.

Conexión en paralelo de un detector de 2 hilos con un contacto mecánico

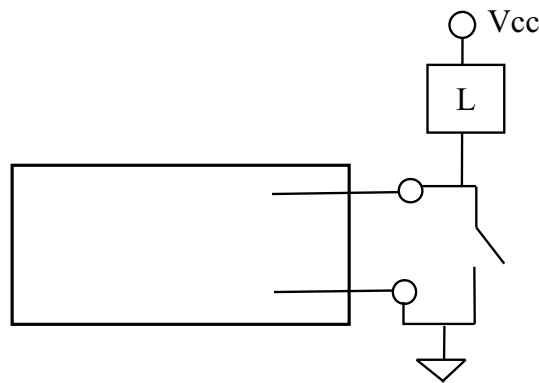


Figura 2.21: Conexión en paralelo de un detector de 2 hilos con un contacto mecánico.

Es posible también, pero si el contacto está cerrado el detector a dos hilos no está alimentado (tiene tensión cero) por lo que hay que tener en cuenta el retardo a la disponibilidad.

2.6 DETECTORES ÓPTICOS O FOTOELÉCTRICOS

También llamados fotocélulas. Constan de un emisor de luz, y un receptor que detecta la luz emitida. El emisor es un diodo LED, mientras el receptor es un fotodiodo o fototransistor, con la electrónica necesaria de amplificación. Normalmente, disponen de lentes adecuadas para enfocar la luz emitida y recibida y mejorar así el alcance de detección. Se utiliza para la detección de todo tipo de objetos. La detección se produce por el cambio en la cantidad de luz recibida por el receptor debido a la presencia del objeto.

El tipo de luz suele ser infrarroja, aunque en determinadas aplicaciones se usa también luz roja (cuando se desea que el haz de luz sea visible). El LED se suele atacar con una señal pulsante de frecuencia elevada (5 kHz) y un ciclo de trabajo bajo (5 %). De esta forma, se optimiza la intensidad de luz máxima emitida, y por tanto el alcance. El receptor suele sincronizarse con esta señal pulsante del emisor, lo que mejora la robustez frente a la detección de luz ambiente o de luz de otros detectores.

Se puede considerar el detector fotoeléctrico formado por 3 partes: parte emisora y receptora de luz, parte de amplificación y tratamiento electrónico,

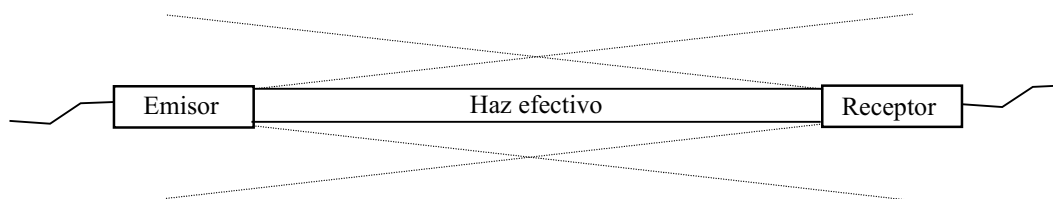


Figura 2.22: Funcionamiento de un detector de barrera y definición del haz efectivo.

y parte de salida. La parte de salida es de alguno de los tipos descritos en la sección anterior. La parte de amplificación es la encargada de decidir si el cambio en la cantidad de luz detectada es suficiente para provocar un cambio del estado lógico de la salida.

Se define el **margen de operación** (también llamado ganancia) como el cociente entre la cantidad de luz recibida y la cantidad de luz mínima necesaria para producir la activación de la salida. Cuanto más alto sea el margen, más robusto será el detector a la suciedad.

En muchos casos los detectores disponen de un potenciómetro de ajuste que permite modificar la sensibilidad, es decir, la cantidad mínima de luz necesaria para que se active la salida.

Se pueden distinguir tres tipos de detectores fotoeléctricos en función de cómo interviene el objeto en la modificación de la cantidad de luz recibida: de barrera, reflex y difuso (o de proximidad).

De barrera

Se tiene por una parte el emisor y por otra el receptor. En funcionamiento normal el haz emitido llega al receptor. Si se interpone un objeto entre ellos interrumpiendo el haz de luz, el receptor deja de recibir, y se detecta el objeto.

El campo de visión del emisor y del receptor tiene forma cónica (con un ángulo pequeño, normalmente). La anchura efectiva del haz coincide con el diámetro de la lente del emisor y del receptor. La detección se produce cuando el objeto es opaco e interrumpe una parte suficiente del haz efectivo (del orden del 50 %). La anchura del haz efectivo varía de unos modelos a otros, y depende de las lentes utilizadas.

Ventajas:

- Muy largo alcance (los hay de hasta 200m).
- Margen de operación muy grande, por lo tanto, muy robusto a la suciedad. La figura 2.23 muestra el margen de operación en función de la distancia para una fotocélula de barrera comercial.

Inconvenientes:

- Más caro, por necesitar 2 elementos.

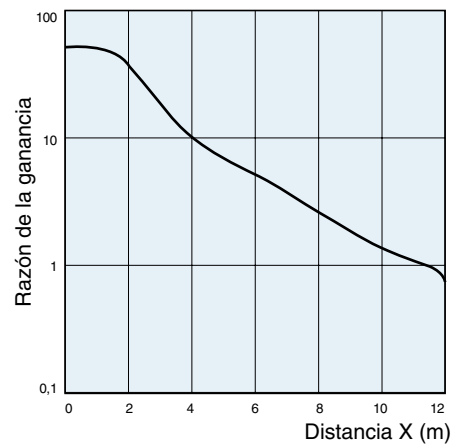


Figura 2.23: Margen de operación en función de la distancia para una fotocélula de barrera de Omron Electronics.

- Necesita más cableado.
- Necesita poner un elemento a cada lado del objeto a detectar.
- No es adecuado para objetos translúcidos.

Reflex

El emisor y el receptor se encuentran integrados en el mismo dispositivo. El haz de luz emitido se refleja en una placa especial reflectora, siendo detectado ese reflejo por el receptor, en condiciones normales. Cuando un objeto se interpone entre el detector y el reflector, el haz se interrumpe y el objeto es detectado. El reflector no necesita estar perfectamente perpendicular al haz emitido.

El haz efectivo tiene el tamaño del reflector, siempre que esté dentro del campo de visión del emisor y del receptor:

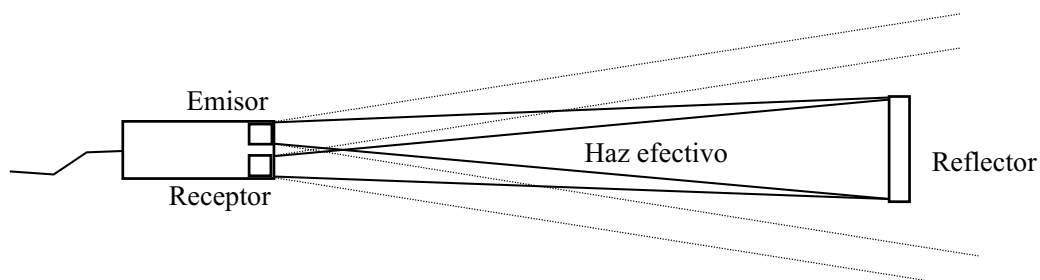


Figura 2.24: Funcionamiento de un detector reflex y definición del haz efectivo.

El objeto a detectar debe interrumpir una parte suficiente de este haz efectivo para que se produzca la detección. La detección queda garantizada si el objeto es mayor que el reflector.

Ventajas:

- Es más barato por tener solo un elemento.
- El cableado es más sencillo que el de barrera.

Inconvenientes:

- Alcance menor que el de barrera (aunque puede ser de varios metros).
- Menor margen de operación que el de barrera (recibe menos cantidad de luz).
- Sigue necesitando el montaje del reflector al otro lado del objeto a detectar.
- Puede tener errores de detección si el objeto es brillante.

Para evitar el inconveniente de la reflexión en objetos brillantes se utiliza una variante de los detectores reflex en la que el emisor tiene un filtro que polariza la luz. El reflector despolariza esa luz, de forma que el reflejo que vuelve al receptor está despolarizado. Si la luz se refleja en el objeto no se despolariza, por lo que al receptor le llega luz polarizada. Para distinguir ambos reflejos el receptor solo detecta la luz no polarizada, y por lo tanto no detecta la reflejada en el objeto. Este detector se llama **reflex polarizado**. La ventaja respecto del reflex normal es que no da errores aunque el objeto sea brillante o translúcido. El inconveniente es que la distancia de detección y el margen de operación suelen ser menores.

Aun así puede haber problemas si el objeto tiene una superficie capaz de despolarizar la luz. Esto sucede, por ejemplo, si se trata de una superficie brillante recubierta de un material transparente como el plástico.

Difuso (o de reflexión sobre objeto, o de proximidad)

En este tipo de detector, el emisor y el receptor están también en el mismo dispositivo. En este caso, no hay ningún reflector, por lo que si no hay objeto, la luz emitida no es detectada por el receptor. Cuando se sitúa un objeto suficientemente cerca, la luz reflejada en la superficie de este objeto es captada por el receptor, detectándolo. Se denomina difuso porque la reflexión en la superficie del objeto es difusa, en general.

Ventajas:

- No es necesario situar ningún elemento al otro lado del objeto a detectar.
- El cableado es sencillo.

Inconvenientes:

- La distancia de detección es pequeña (menor de un metro, normalmente).

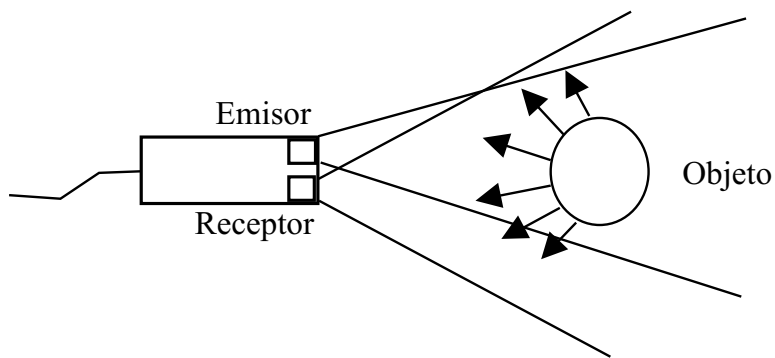


Figura 2.25: Funcionamiento de un detector difuso o de reflexión sobre objeto.

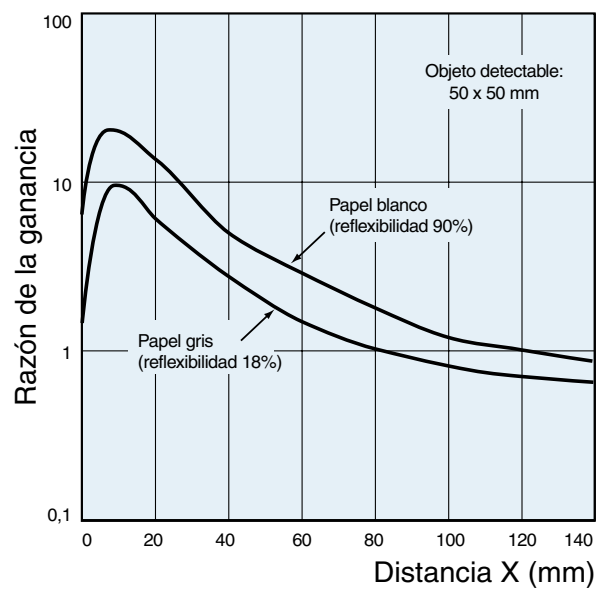


Figura 2.26: Margen de operación en función de la distancia y el color para una fotocélula difusa de Omron electronics.

- El margen de operación es pequeño, y depende mucho de las propiedades de reflexión del objeto que se detecta. La figura 2.26 muestra el margen de operación de sensor difuso comercial en función de la distancia. Se observa que es bastante inferior para papel gris que para papel blanco.

Existen diversos tipos de detectores difusos que presentan características especiales. Entre ellos podemos destacar:

- **Difuso con supresión de fondo.** Permite ignorar el reflejo de los objetos situados a una distancia superior a una umbral tal como se observa en la figura 2.27. Normalmente esa distancia umbral es ajustable (mediante un potenciómetro). Se utiliza cuando detrás del objeto a detectar hay un fondo que refleja luz. Ajustando correctamente el potenciómetro se puede detectar el objeto y no el fondo. Funciona aunque el fondo sea más brillante que el objeto a detectar, pues no se basa en la intensidad reflejada, sino en la distancia donde se produce la reflexión.

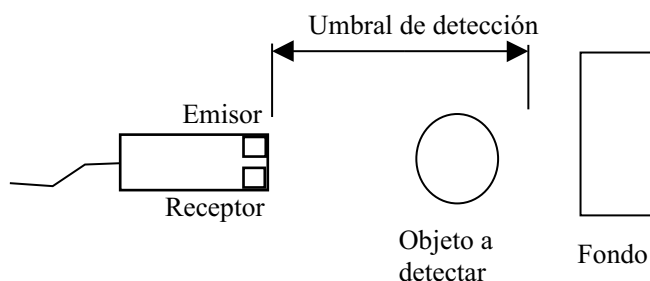


Figura 2.27: Detector difuso con supresión de fondo.

- **Difuso de foco fijo.** También se llama **de haz convergente**. Tal y como se muestra en la figura 2.28, el emisor y el receptor están orientados de forma que el reflejo solo llega al receptor si el objeto está situado a una distancia determinada. En realidad, tiene un pequeño margen de distancias en las que se detecta el objeto. En distancias superiores o inferiores, el objeto no es detectado. Normalmente, el punto de detección está muy próximo al detector (centímetros).

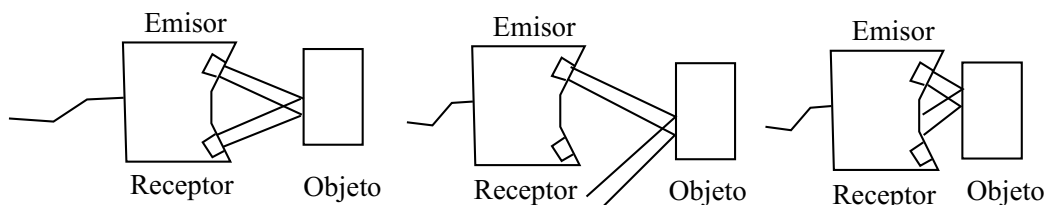


Figura 2.28: Detector difuso de foco fijo.

- **Difuso de gran angular.** Es un detector difuso en el que el campo de visión del emisor y del receptor es muy ancho (tiene un ángulo grande). Puede ser de utilidad para detectar objetos pequeños cuya posición respecto del eje de visión pueda estar desplazada. También puede utilizarse para ignorar agujeros en objetos grandes (un haz muy fino puede no detectar el objeto si coincide justo con un agujero del mismo).

Detectores con fibra óptica

Los tres tipos de detectores tienen versiones con fibra óptica. Éstos se utilizan cuando hay un problema de espacio, o cuando las condiciones de funcionamiento hacen inviable el montaje del sensor completo (por ejemplo, por la vibración o por la temperatura). En los detectores de fibra óptica hay un módulo que contiene toda la electrónica. A ese módulo se conectan las fibras ópticas (una para la luz emitida y otra para la luz recibida). Las puntas de las fibras se montan en el lugar donde se debe realizar la detección, ocupando muy poco espacio.

Hay fibras ópticas de vidrio y de plástico. Las de plástico son más flexibles y más baratas. Las de vidrio soportan temperaturas más elevadas, pero no son tan flexibles, y son más caras.

Especificaciones de los detectores fotoeléctricos (criterios de selección)

Hay muchos parámetros a tener en cuenta cuando se elige un detector fotoeléctrico. Las especificaciones más importantes son:

- **Operación por luz u operación por sombra (L ON o D ON).** Un detector que opera por luz tiene su salida activa si el receptor recibe una parte suficiente de la luz emitida, y se desactiva si esa luz no es suficiente. El que opera por sombra funciona al revés. Por ejemplo, un detector de barrera o reflex que opera por luz está activo si no hay objeto, y se desactiva si se interpone un objeto. En cambio, un detector difuso que opera por luz está desactivado si no hay objeto, y se activa cuando se aproxima un objeto que refleje la luz suficiente.
- **Distancia máxima de detección.** Es la distancia máxima a la que se produce la detección (con margen de operación igual a 1). Para detectores reflex o de barrera, esta distancia se define como la que hay entre el reflector y el sensor, o entre el emisor y el receptor, por lo que no depende del objeto detectado. Para detectores difusos, esta distancia se define como la distancia entre el sensor y el objeto detectado, por lo que depende del tamaño, el color y la textura del objeto. Debido a esto, en los detectores difusos se suele definir la distancia máxima para un objeto estándar (superficie plana blanca de un tamaño determinado, por

ejemplo de 20x20 cm). Normalmente, no se pueden utilizar los sensores con esta distancia máxima, ya que para que sean robustos a la suciedad deben trabajar con un margen de operación bastante mayor que 1.

- **Distancia nominal de detección** (S_n). Suele ser algo menor que la distancia máxima, y es un indicativo de la distancia máxima a la que se puede utilizar el detector con un margen suficientemente mayor que 1 (suele ser de 2 a 4). En una aplicación práctica se podrá utilizar a esta distancia o no en función de la suciedad y del tipo de objeto a detectar (si es un detector difuso).
- **Distancia mínima de detección.** Los detectores reflex y los difusos tienen una pequeña área de ceguera en la cual la luz reflejada no llega al receptor, como se observa en la figura 2.29. Se debe a que el emisor y el receptor están algo separados. Debido a esto existe una distancia mínima por debajo de la cual no se detecta el objeto o el reflector.

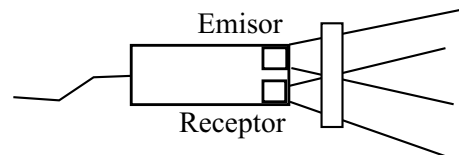


Figura 2.29: Distancia mínima de detección.

- **Curva de respuesta.** Es una curva que representa el margen de operación respecto de la distancia. Para detectores difusos se utiliza un objeto blanco plano estándar. La distancia máxima de detección es la distancia en la que la curva vale 1.
- **Tiempo de respuesta.** Es el tiempo que tarda en activarse (o desactivarse) la salida desde que se detecta (o deja de detectarse) el objeto. El tiempo de activación y el de desactivación pueden ser diferentes. Depende fundamentalmente del tipo de salida. A veces se define la frecuencia máxima de conmutación en lugar del tiempo de respuesta. Es la frecuencia máxima a la que puede activarse y desactivarse, y por lo tanto es algo menor que la inversa de la suma de los tiempos de activación y de desactivación.
- **Curva de detección (contorno del haz).** Es una curva que representa la zona del espacio en la que se produce la detección. Se obtiene aproximando el objeto con un movimiento perpendicular al eje de detección hasta que es detectado. La curva define la posición del borde interior del objeto. Para los reflex lo que se aproxima es el reflector, para los de barrera el receptor, y para los difusos un objeto estándar blanco.

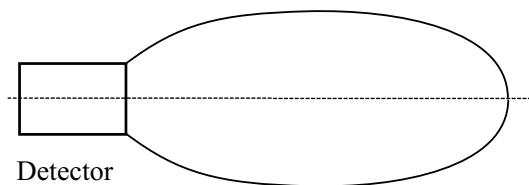


Figura 2.30: Curva de detección.

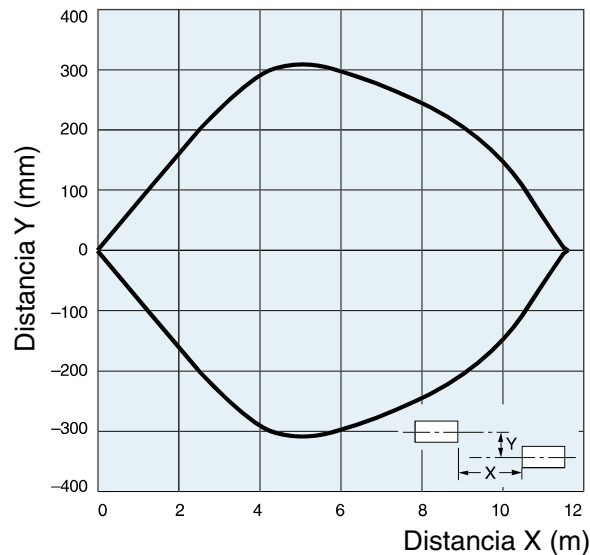


Figura 2.31: Curva de detección de una fotocélula de barrera de Omron Electronics.

- **Tipo de salida.** Es el tipo de salida que da el detector. Puede ser de cualquiera de las ya descritas.
- **Forma y tamaño.** La forma y el tamaño de los detectores determina las posibilidades de montaje, por lo que puede ser un parámetro importante a la hora de seleccionar un detector. Las formas más habituales son:
 - Cilíndricos. Los más sencillos suelen ser cilíndricos, pudiendo ser de barrera, reflex o difusos. El diámetro más habitual es de 18 mm.
 - De petaca (paralelepípedica). Suelen tener esta forma los que tienen funciones más complejas (como supresión de fondo, salida analógica, etc.).
 - Amplificador más cabeza óptica por separado. Pueden ser de cualquier tipo. La ventaja que tienen es que la cabeza óptica ocupa menos espacio que el detector completo. En la cabeza solo está el LED emisor y el fototransistor receptor (con las lentes necesarias), quedando toda la electrónica en el amplificador, que puede montarse en otro sitio. Un mismo amplificador puede admitir cabezas de

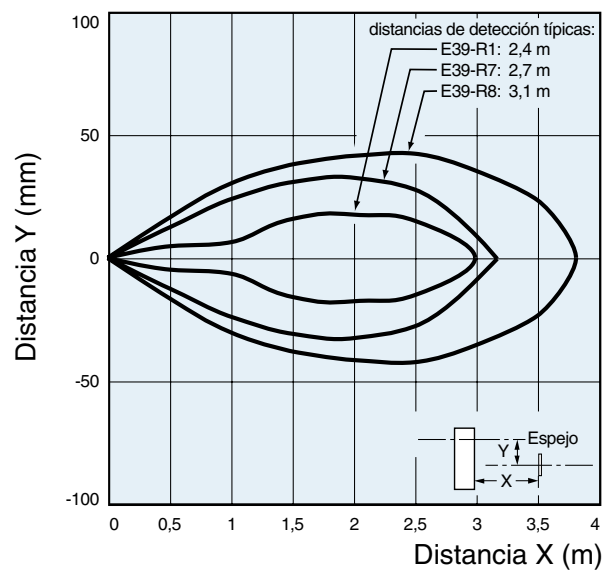


Figura 2.32: Curva de detección de una fotocélula reflex de Omron Electronics.

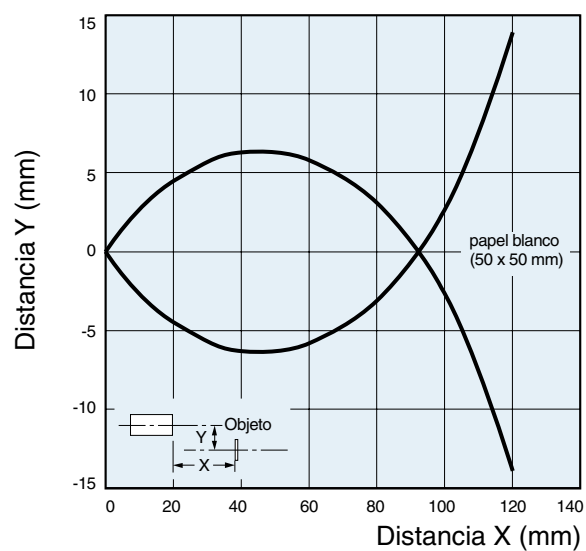


Figura 2.33: Curva de detección de una fotocélula de reflexión sobre objeto de Omron Electronics.

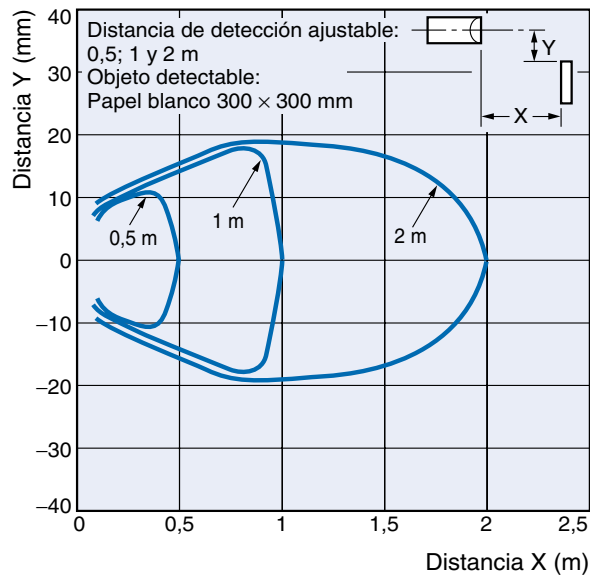


Figura 2.34: Curva de detección de una fotocélula de supresión de fondo ajustable de Omron Electronics.

distintas formas y tamaños.

- De fibra óptica. En el punto de detección solo se coloca la punta de la fibra óptica, por lo que es la solución que menos espacio ocupa. En el amplificador está toda la electrónica (incluido el LED emisor y el fototransistor receptor).



Figura 2.35: Fotocélula cilíndrica de Omron Electronics.

- **Temporización y lógica adicional.** Algunos detectores disponen de funciones adicionales, como el retardo a la conexión o a la desconexión, o la generación de un único pulso de anchura determinada cada vez que se detecta un objeto.
- **Ajuste de sensibilidad.** Algunos detectores disponen de un potenciómetro que permite ajustar la sensibilidad del receptor de luz. Esto



Figura 2.36: Fotocélula de petaca de Omron Electronics.



Figura 2.37: Fotocélula miniatura de forma plana de Omron Electronics.



Figura 2.38: Fotocélula de horquilla de Omron Electronics.



Figura 2.39: Amplificador de fotocélula de fibra óptica de Omron Electronics.



Figura 2.40: Cabezas de fibras ópticas para fotocélulas de Omron Electronics.

permite ajustar el detector a las características de reflexión del objeto a detectar (si es difuso).

- **Histéresis.** Representa la diferencia entre la distancia de detección cuando el objeto se acerca y la distancia de detección cuando el objeto se aleja. Cuando se acerca, el objeto es detectado a una distancia. Si después se aleja, deja de detectarse a una distancia mayor. De esta forma, cuando es detectado el objeto, la señal de salida del sensor es estable. Si no hubiera histéresis, en el momento en que el objeto está justo en el límite de detección, la salida pasaría de activa a inactiva varias veces de forma aleatoria. El valor de este parámetro no suele venir especificado.
- **Campo de visión.** Es la región cónica que forma la luz emitida, o la región cónica que “ve” el receptor. Normalmente no es un parámetro de selección, puesto que la detección no depende solo de que el objeto (o el reflector) esté en el campo de visión, sino de que la cantidad de luz recibida sea suficiente.

2.7 DETECTORES DE PROXIMIDAD INDUCTIVOS

Sirven para detectar la presencia de objetos metálicos, tanto ferromagnéticos (como el hierro o acero), como no ferromagnéticos (como el cobre o el aluminio). Se basan en el cambio en la inductancia producido por la presencia del metal en las proximidades del detector. Si el objeto es ferromagnético, el cambio de inductancia se debe al aumento de la permeabilidad magnética del medio por el que se cierran las líneas de campo producidas por la bobina del detector. Esto produce un aumento de la inductancia.

Si el metal no es ferromagnético, el cambio en la inductancia se debe a las corrientes de Foucault inducidas en el metal por el campo generado por la bobina del detector. Estas corrientes producen un campo contrario al del detector, dando como resultado en una disminución de la inductancia. El efecto de los metales ferromagnéticos es mayor que el de los no ferromagnéticos, por lo que se detectan a distancias mayores.

Para detectar el cambio en la inductancia, el detector tiene un oscilador electrónico, cuya frecuencia de oscilación depende de la inductancia. En ausencia de objetos metálicos la frecuencia de oscilación es una determinada. Cuando se aproxima un metal cambia la inductancia, y por tanto la frecuencia de oscilación. Un circuito electrónico detecta el cambio de frecuencia y produce la activación de la salida si ese cambio es suficientemente grande.

Especificaciones de los detectores inductivos (criterios de selección)

Hay muchos parámetros a tener en cuenta cuando se elige un detector inductivo. Las especificaciones más importantes son:

- **Distancia de detección.** Es la distancia entre la superficie sensora y la parte más cercana del objeto a detectar. Se define la **distancia nominal**, S_n , como la distancia a la que se detecta un cuadrado de acero templado de 1 mm de espesor y de lado igual al diámetro del sensor, en condiciones nominales de humedad y temperatura. Esta distancia depende del diámetro del detector, y puede variar entre 1 mm y 15 mm. Si el objeto es de un metal diferente (especialmente si no es ferromagnético) la distancia real de detección puede ser significativamente menor (menos de la mitad para el aluminio, por ejemplo). En la figura 2.41 se observa cómo cambia la distancia de detección con el material para un sensor comercial. Se define la distancia mínima de operación, S_a , a la que se garantiza para el cuadrado de acero en cualquier condición de humedad y temperatura (de las que soporta el sensor).

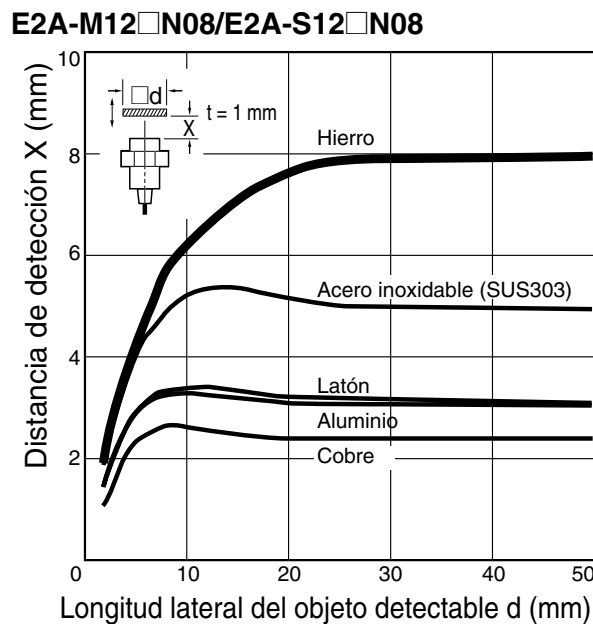


Figura 2.41: Influencia del material en la distancia de detección del sensor inductivo E2A-M12 de Omron Electronics.

- **Histéresis.** Igual que en los detectores fotoeléctricos. Cuando se acerca el objeto, la salida se activa a una distancia determinada. Si después se aleja, la distancia a la que se desactiva la salida es mayor.
- **Curva de detección.** Es una curva que se obtiene acercando lateralmente el objeto a detectar (el cuadrado de acero) a distintas distancias del sensor. Esta curva define la región en la que el objeto es detectado, figura 2.42. Suele haber dos curvas, una de acercamiento y otra de alejamiento (debido a la histéresis), aunque algunos fabricantes solo incluyen una, como en la figura 2.43.

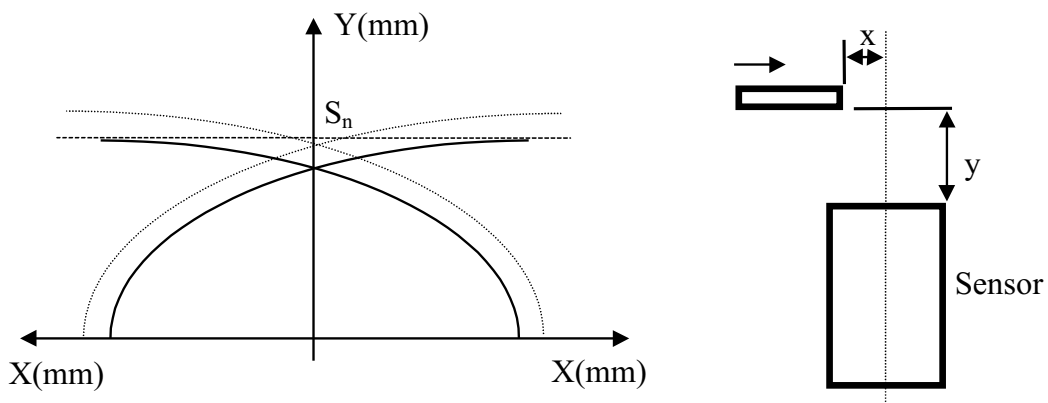


Figura 2.42: Curva de detección lateral de un sensor inductivo.

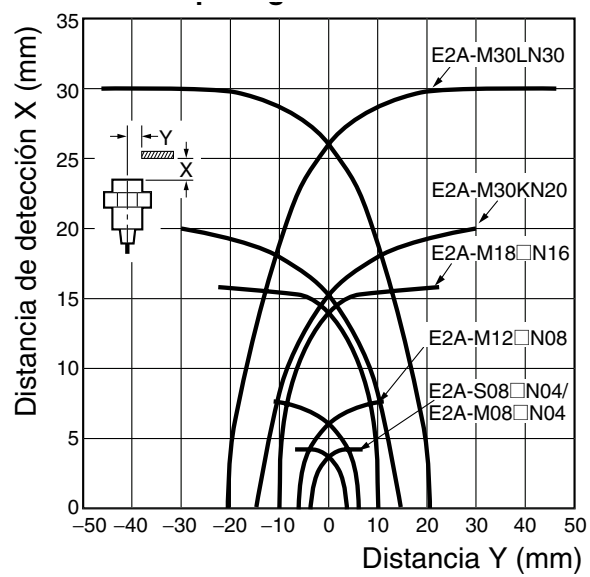


Figura 2.43: Curva de detección de sensores inductivos E2A de Omron Electronics.

- **Tipo de salida.** Al igual que los detectores fotoeléctricos, los inductivos pueden tener cualquiera de los tipos de salida que se han descrito (a transistor NPN o PNP, a relé, a 2 hilos, etc.). También los hay con salida analógica (en tensión o en corriente).
- **Frecuencia de conmutación (tiempo de respuesta).** Se define igual que en los detectores ópticos. Las frecuencias máximas son del orden de los 3 kHz.
- **Forma y tamaño.** La forma y el tamaño de los detectores determina las posibilidades de montaje, por lo que puede ser un parámetro importante a la hora de seleccionar un detector inductivo. La forma más habitual, con diferencia, es la cilíndrica. Una de las bases del cilindro es la superficie detectora, por la que salen al exterior las líneas de campo magnético. Hay diferentes diámetros, siendo los más habituales los de 8, 12, 18 y 30 mm. A mayor diámetro, mayor distancia de detección. También hay detectores inductivos de otras formas (extraplanos, paralelepípedos, etc.), aunque no son tan habituales. Estas formas especiales pueden tener utilidad cuando los cilíndricos dan problemas de espacio.



Figura 2.44: Sensor inductivo cilíndrico de Omron Electronics.



Figura 2.45: Sensor inductivo de forma plana de Omron Electronics.

- **Restricciones de montaje.** Respecto al montaje, la propiedad más importante es la de poder colocarse enrasado en una superficie metálica o no. Se distinguen **detectores enrasables y no enrasables** en metal cuyo montaje se indica en la figura 2.46. Los enrasables tienen la ventaja de montarse más fácilmente, pero su distancia de detección es menor. Los no enrasables requieren dejar un hueco circular alrededor del sensor, pero su distancia de detección es mayor.

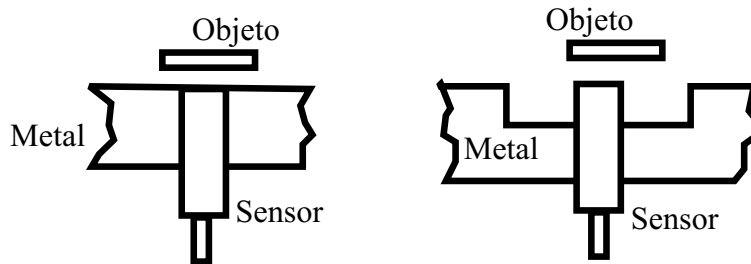


Figura 2.46: Montaje enrasable y no enrasable de un sensor inductivo.

Tanto en los enrasables como en los no enrasables las masas metálicas cercanas pueden afectar a la detección del objeto, por lo que existen unas distancias mínimas que se deben respetar para el funcionamiento correcto. Lo mismo sucede con la colocación próxima de varios sensores, que deben respetar unas distancias mínimas para no interferirse.

2.8 DETECTORES DE PROXIMIDAD CAPACITIVOS

Se basan en la modificación de la capacidad vista por el sensor debido a la presencia del objeto. Para que el objeto modifique la capacidad es necesario que su constante dieléctrica sea diferente a la del aire. Sirven, por tanto, para detectar la presencia de todo tipo de objetos, metálicos y no metálicos, siempre que su constante dieléctrica sea mayor a la del aire. Por ejemplo, la constante dieléctrica respecto a la del aire es de 80 para el agua, 2 a 7 para madera seca, 10 a 30 para madera verde, 2 a 5 para el papel, de 4 a 7 para la porcelana.

Cuanto mayor sea esa constante dieléctrica, mayor el tamaño del objeto y menor la distancia al detector, mayor será el aumento de la capacidad, y por tanto más fácil será la detección.

La forma de detectar el cambio de capacidad es idéntica a la de los sensores inductivos. Se trata de un oscilador electrónico cuya frecuencia de oscilación depende de la capacidad. Al variar la capacidad varía la frecuencia, y ésta es detectada por un circuito que activa la salida.

Especificaciones de los detectores capacitivos (criterios de selección)

Son muy similares a las de los detectores inductivos. Las diferencias más importantes son:

- **Distancia de detección.** Se define de la misma forma. La distancia nominal se define para el mismo objeto que en los inductivos (cuadrado de acero templado de 1 mm de espesor y de lado igual al diámetro del sensor). Esta distancia puede ser de hasta 20 mm, según el tamaño del sensor. Si el objeto es de un material con constante dieléctrica menor, o su tamaño es menor, la distancia de detección se reduce.
- **Frecuencia de conmutación (tiempo de respuesta).** La frecuencia de conmutación suele ser menor que en los inductivos. Las frecuencias máximas están del orden de los 200 Hz.
- **Forma y tamaño.** La forma más habitual también es la cilíndrica. Una de las bases del cilindro es la superficie detectora, por la que salen al exterior las líneas de campo eléctrico. Hay diferentes diámetros, aunque no hay tan delgados como los inductivos, siendo los más habituales los de 18 y 30 mm. A mayor diámetro, mayor distancia de detección.

2.9 OTROS DETECTORES DE PROXIMIDAD

Aunque los detectores de proximidad más habituales en la industria son los finales de carrera, los ópticos, los inductivos y los capacitivos, también existen otros cuyo uso es relativamente frecuente. Entre ellos se pueden destacar:

Detectores de proximidad magnéticos

Sirven para detectar objetos de material magnético (imanes). Hay dos tipos fundamentales: magnetorresistivos y contactos reed. Los magnetorresistivos se basan en el aumento de la resistencia de un conductor cuando se somete a un campo magnético. Al aproximarse el imán, el campo magnético produce un aumento de la resistencia, que es detectado por un circuito, que activa la salida. El detector de contactos reed tiene un contacto que se cierra al acercarse un imán, debido a la fuerza de atracción magnética del imán sobre una placa. La aplicación fundamental de este tipo de detectores es el posicionado de cilindros neumáticos (o hidráulicos). El pistón tiene un pequeño imán, que es detectado por el sensor colocado en el exterior del cilindro cuando el pistón está en una posición determinada.

Detectores de efecto Hall

Se basan en el efecto Hall, que consiste en la aparición de una diferencia de tensión en un conductor (o semiconductor) por el que circula una corriente eléctrica cuando se somete a un campo magnético. En principio, sirven por tanto para detectar objetos magnéticos, aunque con una pequeña variación permiten detectar objetos ferromagnéticos, que producen un cambio del campo magnético que actúa, y en consecuencia de la diferencia de tensión. Se utilizan por ejemplo para detectar la posición del rotor en motores *brushless* (sin escobillas).

Detectores de proximidad por ultrasonidos

Sirven para detectar todo tipo de objetos. Se basan en la emisión de un ultrasonido, y la recepción del rebote del mismo contra los objetos próximos. Si no hay objeto próximo, el rebote llega muy atenuado, y la salida no se activa. Si hay un objeto de tamaño suficientemente grande situado suficientemente cerca del detector, el rebote del ultrasonido llega al sensor con una amplitud grande, y se activa la salida.

Suelen tener distancias de detección mayores que los otros detectores (incluidos los fotoeléctricos difusos), siendo del orden de un metro.

Los más sofisticados miden el tiempo que tarda el ultrasonido desde que es emitido hasta que es detectado el rebote, calculando así la distancia a la que se encuentra el objeto. De esta forma, se puede realizar la supresión de fondo, o bien dar una salida analógica proporcional a la distancia. El inconveniente es que la velocidad del sonido depende de muchos factores, especialmente de la temperatura.

2.10 OTROS DETECTORES (DE NIVEL, DE PRESIÓN, DE TEMPERATURA)

Además de los detectores de proximidad utilizados para detectar objetos, existen otros que también se utilizan mucho en la industria. Los más comunes son los de nivel de líquidos, los de presión o los de temperatura. Las salidas de los mismos suelen ser de tipo relé, aunque pueden ser de cualquiera de los tipos descritos en el capítulo anterior.

Detectores de nivel

Activan o desactivan su salida en función de que el nivel del líquido sea superior o inferior a la posición en que están colocados. Hay de varios tipos:

- **Boya con contacto.** Consiste en una boya flotante y un contacto que se abre o cierra en función de que el líquido haga flotar la boya o no.

- **Conductivo.** Sirve para líquidos conductores, como el agua no destilada. En esencia, consta de dos conductores situados a distinta altura. Cuando el líquido cubre los dos, se establece una corriente entre ellos, que conmuta una salida (generalmente un relé).
- **Capacitivo.** Consta de un detector capacitivo como los descritos en la sección anterior, colocado en la pared del depósito. Si el líquido cubre el sensor, éste lo detecta.
- **Ultrasónico.** Consiste en utilizar un detector de ultrasonidos como el descrito en la sección anterior. Éste se coloca encima del depósito, midiendo la distancia a la superficie. Normalmente, se suele utilizar uno con salida analógica, con lo que se obtiene el nivel continuo.

Detectores de presión

Sirven para conocer si la presión de una instalación de gas o líquido es mayor o menor que una determinada presión. Se llaman también presostatos. Los hay mecánicos, de forma que cuando la presión supera el valor umbral, la fuerza ejercida por el líquido o gas vence un resorte y acciona un contacto. Los hay también electrónicos, que incorporan un captador continuo de presión y que activan o no la salida en función de que la medida del sensor sea mayor o menor que un valor de referencia. Éstos suelen ser ajustables.

Detectores de temperatura

Sirven para saber si la temperatura es o no superior a una determinada temperatura. Se llaman también termostatos. Los hay mecánicos, basados en la utilización de un elemento bimetal (formado por dos metales con diferente coeficiente de dilatación) que se curva al calentarse, de tal forma que al alcanzar una temperatura se cierra un contacto. También los hay electrónicos, que incorporan un captador continuo de temperatura, y que activan o no la salida en función de que la medida del sensor sea mayor o menor que un valor de referencia. Éstos suelen ser ajustables.

2.11 TRANSDUCTORES DE POSICIÓN

Se denomina transductores de posición a aquellos sensores (captadores) que permiten medir la posición relativa de un objeto respecto a una referencia fija.

Los hay para medir la posición angular (ángulo girado por un eje), y para medir la posición lineal (distancia lineal hasta la referencia fija).

Tanto los angulares como los lineales se pueden clasificar en dos grupos, según su salida sea analógica o digital. Los sensores de posición que dan directamente una salida digital se llaman codificadores (*encoders*).

Potenciómetros

Son sensores analógicos. Los hay para medir posición lineal y angular. En esencia, son una resistencia variable (potenciómetro) en los que el terminal móvil se une al objeto cuya posición se quiere medir. Conforme el objeto se mueve, cambia la resistencia entre los tres terminales tal y como se muestra en las figuras 2.47 y 2.48. Para medir la posición basta con alimentar el potenciómetro con una tensión constante conocida. De esta forma, en el terminal móvil se tiene una tensión proporcional a la posición.

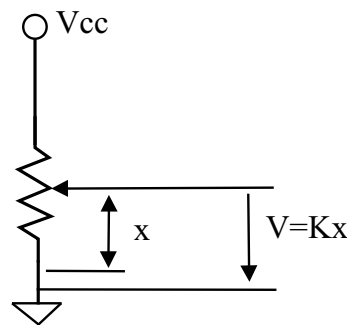


Figura 2.47: Transductor de posición lineal mediante potenciómetro.

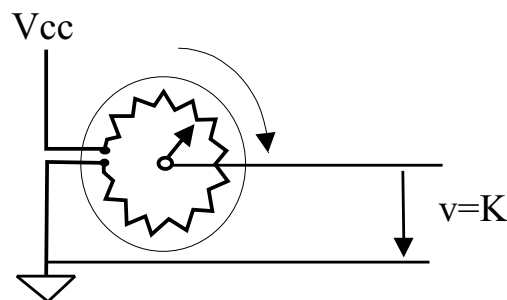


Figura 2.48: Transductor de posición angular mediante potenciómetro.

Los potenciómetros que más se usan como sensores de posición suelen ser de plástico conductivo. El plástico conductivo permite que el contacto con el elemento móvil sea suave (tenga poco rozamiento), y que el desgaste sea mínimo, además de tener una buena resolución.

La ventaja de los potenciómetros es que son relativamente baratos. Además, dan una señal con una gran resolución (en teoría infinita).

El inconveniente es que se desgastan con el uso, por lo que tienen un número de ciclos limitados. Además, si el movimiento es rápido, el contacto del elemento rozante no es perfecto, por lo que la señal puede tener un ruido importante. Por otra parte, necesitan una tensión de alimentación muy estable, pues cualquier error en esa tensión aparece a la salida.

En cuanto a los potenciómetros angulares, los hay de una sola vuelta y de varias vueltas. En los de una sola vuelta hay una pequeña zona muerta en la

que no se mide nada. La ventaja que tienen es que el potenciómetro puede dar vueltas indefinidamente (sólo que en cada vuelta mide lo mismo). Los de varias vueltas no tienen una zona muerta, pudiendo medir ángulos mayores de 360°. En cambio, su rango de movimiento está limitado por el número de vueltas del potenciómetro.

LVDT

El transformador diferencial de variación lineal es un sensor de posición que da también una salida analógica. Los hay lineales y angulares, aunque los más habituales son lineales.

Constan de un transformador con un bobinado primario y dos secundarios. El primario se alimenta con una tensión alterna de amplitud y frecuencia fija. Por el interior de los bobinados del transformador se mueve un núcleo ferromagnético que va unido al objeto cuya posición se quiere medir. La posición del núcleo modifica la inductancia mutua entre el primario y los dos secundarios, de forma que existe una posición en la que esta inductancia mutua es igual para los dos. En esa posición, la tensión en los dos secundarios es la misma, y la diferencia es cero. Cuando se desplaza el núcleo, aumenta la tensión en un secundario y disminuye en el otro, de forma que la amplitud de la diferencia de tensión entre los dos es proporcional al desplazamiento del núcleo. La salida es, por tanto, una señal alterna de frecuencia fija y amplitud proporcional a la posición. Evidentemente, se necesita una electrónica adicional que trate esa señal para dar una salida estándar (en tensión o en corriente).

Normalmente, estos sensores incorporan toda la electrónica necesaria para generar la señal alterna y para tratar la salida, de forma que externamente se alimentan en continua (por ejemplo a 24 V), y se obtiene una salida en tensión (0 a 10 V) o en corriente (4 a 20 mA). Se utilizan para medir posiciones lineales desde rangos de menos de 1 cm hasta varias decenas de centímetros. Una aplicación típica es la medición de la posición de un cilindro neumático.

La gran ventaja respecto de los potenciómetros es que no hay contacto entre el núcleo y las bobinas, por lo que no hay desgaste, y la señal es mucho más limpia (con menos ruido). Tienen resolución teóricamente infinita, limitada solo por la electrónica de tratamiento de la señal.

El inconveniente es que son mucho más caros que los potenciómetros. Otro inconveniente es que el tratamiento de la señal de salida implica una limitación del ancho de banda del sensor. Esta limitación es tanto menor cuanto mayor es la frecuencia del primario.

Resolvers

Se utilizan para medir la posición angular de un eje. Constan de un transformador con un devanado primario situado en el rotor del resolver, y dos devanados secundarios situados en el estátor. El devanado primario se alimenta con una tensión alterna de frecuencia y amplitud constantes. Al girar el

rótor (que va solidario al eje cuya posición se quiere medir) cambia la relación de transformación entre el primario y cada uno de los secundarios, por lo que cambia la amplitud de la tensión en cada secundario. Éstos están contruidos y situados en el estátor de tal forma que la amplitud de la tensión es proporcional al seno del ángulo en un devanado, y al coseno del ángulo en el otro. Si la tensión en el primario es:

$$v_1(t) = V_1 \text{sen}(\omega t)$$

la tensión en los secundarios tiene la forma:

$$v_{2a}(t) = K \text{sen}(\omega t) \text{sen}(\theta)$$

$$v_{2b}(t) = K \text{sen}(\omega t) \text{cos}(\theta)$$

donde θ es la posición angular del eje (solidaria al rótor del resolver). A partir de estas dos señales eléctricas se puede obtener la posición θ del eje. Esta operación la realiza un circuito electrónico especial, denominado convertidor resolver-digital, que recibe como entradas las dos señales del resolver, y da como salida un número binario que representa la posición angular con una resolución determinada. Resoluciones típicas van desde $2^{12} = 4096$ puntos por vuelta hasta $2^{16} = 65536$ puntos por vuelta.

El resolver permite medir la posición absoluta del eje si el rango de movimiento del mismo es una sola vuelta. Si el eje da varias vueltas (el motor ataca un engranaje o una banda transportadora), no se puede distinguir en principio entre una vuelta y otra. Para hacerlo es necesario tener un contador que almacene el número de vueltas que se han dado. Por otra parte, esto plantea un problema en la inicialización, puesto que cuando se conecta el equipo no se sabe en principio en qué número de vuelta está el eje. Para poder inicializarlo se utiliza un detector externo, que se activa cuando el sistema está en una posición conocida. La primera tarea del equipo cuando se conecta será mover el eje hasta que se active el detector, inicializando en ese momento el número de vueltas.

Existen sensores de posición lineal que se basan en la misma propiedad que el resolver. Un ejemplo es el inductosyn. En éste se tiene una regla fija con un devanado que se alimenta con una tensión alterna de frecuencia y amplitud constantes. El elemento móvil (que se mueve solidario al objeto que se quiere medir) contiene dos devanados en los que se inducen dos tensiones cuya amplitud depende del coseno y del seno del desplazamiento lineal respectivamente. Las señales se repiten cada pocos milímetros, por lo que es necesario almacenar en un contador el número de pasos. Las dos señales permiten subdividir cada paso hasta en mil posiciones, con lo que se obtienen resoluciones de micras. Se utiliza en el posicionado de máquinas de control numérico.

Codificadores incrementales

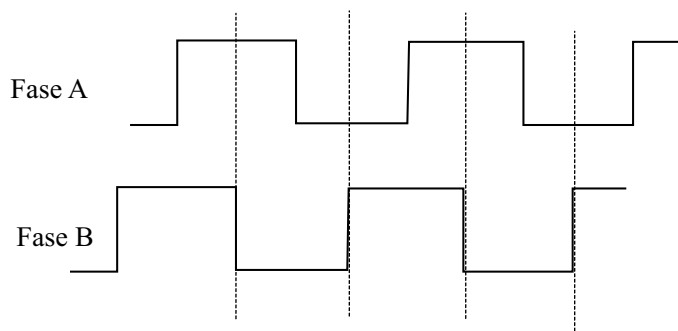
Sirven también para medir la posición angular de un eje, aunque a diferencia del resolver el principio en que se basan no es analógico, sino directamente digital. Constan de un disco ranurado, que se coloca solidario al eje que gira, y

un fotoemisor y fotoreceptor colocados a ambos lados del eje que gira. La luz emitida por el LED emisor llega o no al fotoreceptor en función de la posición de las ranuras del disco. Conforme gira el disco se tiene, por tanto, una señal que cambia de cero a uno. Contando el número de pulsos se obtiene el número de ranuras que han pasado, y por tanto el ángulo girado. La resolución depende, por tanto, del número de ranuras del disco. También existen codificadores incrementales magnéticos, donde los pulsos se deben al paso de los dientes de un disco dentado de acero delante de un detector inductivo, aunque los más utilizados son los ópticos.



Figura 2.49: Codificador incremental de Omron Electronics.

El inconveniente de tener un solo emisor y un solo receptor es que no se distingue si el eje gira en un sentido o en el otro. Para poder distinguir esto se utilizan dos parejas emisor-receptor, situadas de tal forma respecto de las ranuras que las señales producidas están desfasadas 90° . Si el eje gira en sentido horario a velocidad constante las señales de salida serían las indicadas en la figura 2.50.



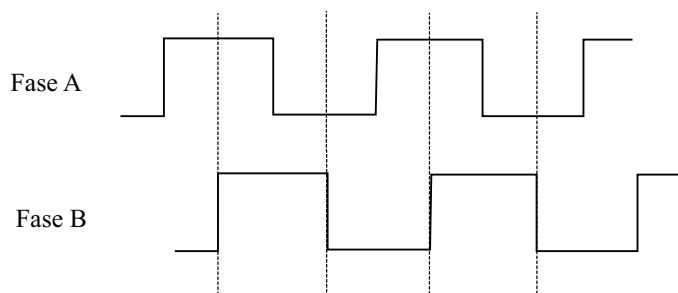


Figura 2.51: Evolución de las salidas de un encoder girando en sentido antihorario.

girando en un sentido o en otro). La utilización de dos señales en cuadratura hace que la resolución final (número de puntos por vuelta) sea 4 veces el número de ranuras, ya que el circuito electrónico incrementa o decrementa el contador cada vez que se produce un flanco de subida o de bajada en alguna de las dos fases. Así pues, un codificador incremental de 1024 ppr permite distinguir entre 4096 posiciones diferentes en una vuelta.

El inconveniente que tienen los codificadores incrementales es que no dan la posición absoluta del eje, ya que los pulsos son idénticos en cualquier ángulo. Eso significa que para tener la posición absoluta hay que inicializar el contador en algún valor determinado. Para que este reset del contador se produzca siempre en la misma posición cuando se conecta el equipo, los codificadores suelen incorporar una tercera señal representada en la figura 2.52 (denominada fase Z), que únicamente da un pulso por vuelta, de forma que cuando se produce ese pulso la posición absoluta del eje es conocida. Bastará, por tanto, con poner el contador a cero cuando se detecta ese pulso de la fase Z.

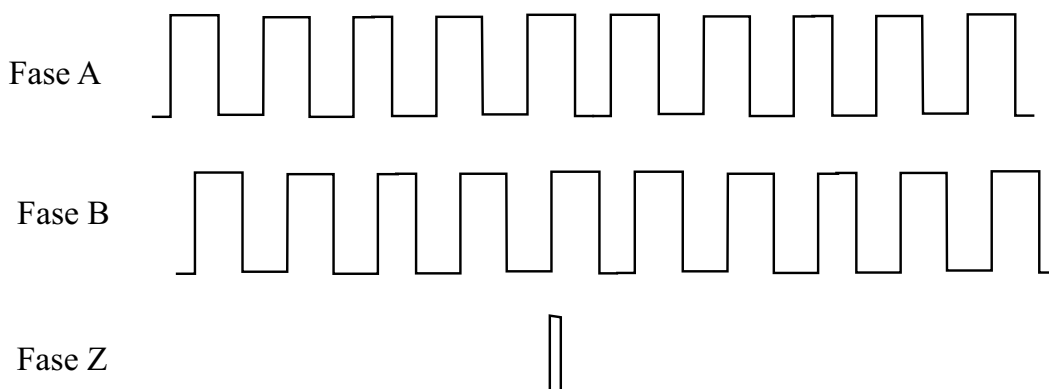


Figura 2.52: Evolución de las salidas de un encoder de tres fases; A, B y Z.

Cuando el rango de movimiento del eje es de una sola vuelta el pulso de cero es suficiente para inicializar el contador. En ese caso, cuando se conecta el equipo, se realiza un movimiento del eje hasta encontrar el pulso, y se inicializa el contador, teniendo a partir de ahí la posición absoluta.

Cuando el rango de movimiento del eje es de varias vueltas (por ejemplo, si se ataca un engranaje, o una banda transportadora), el pulso no es suficiente para conocer la posición absoluta inicial. Para resolver en estos casos el problema, se sitúa externamente al codificador un detector (fotoeléctrico o inductivo, por ejemplo), de forma que cuando se activa ese detector, el sistema se encuentra en una posición conocida. Al conectarse el equipo, se movería el eje hasta que se activase ese detector, momento en el cual se inicializa el contador.

Las resoluciones más habituales de los codificadores incrementales angulares van desde 500 ppr hasta $2^{14} = 16384$ ppr, siendo muy comunes los de 1024 ppr y los de 4096 ppr.

Existe un tipo especial de codificador incremental óptico en el que las señales que se obtienen en los fotoreceptores son senoidales (en lugar de cuadradas). Estos codificadores senoidales son más caros, pero con el tratamiento electrónico adecuado permiten subdividir cada pulso en varios, aumentando la resolución hasta en 20 veces.

Existen también codificadores incrementales lineales. Constan de una regla que tiene las ranuras o marcas, y de los fotoemisores y fotoreceptores que se mueven solidarios al objeto cuya posición se quiere medir. Las ranuras pueden estar muy próximas entre sí, llegando a resoluciones muy buenas (del orden de micras con codificadores senoidales). Se utilizan, por ejemplo, en las máquinas de control numérico para medir con precisión la posición. Tienen los mismos inconvenientes que los angulares respecto a la medida absoluta de la posición, por lo que algunos incorporan además un sistema de medida absoluta (de mucha menos precisión), que permite inicializar el contador en muchos puntos del recorrido de la regla.

Codificadores absolutos

También sirven para medir la posición angular de un eje. Su principio de funcionamiento físico es igual al de los codificadores incrementales ópticos. La diferencia es que el disco ranurado tiene varias franjas de ranuras, cada una de las cuales es detectada por un par emisor-receptor. Cada franja tiene un espaciado de ranuras diferente, de forma que el conjunto de las señales de todos los detectores codifica en binario la posición absoluta del eje. La franja más externa tiene el ranurado más fino que define la resolución (es el bit menos significativo).

La salida del codificador absoluto es, por tanto, un número binario que representa el ángulo absoluto girado respecto de una referencia fija.

Las franjas suelen construirse de forma que de una ranura a la siguiente sólo cambie 1 bit. De esta forma se evitan errores transitorios (si cambian varios bits a la vez es prácticamente imposible que cambien exactamente al mismo tiempo, por lo que durante un instante muy breve se podría leer una posición equivocada). Este tipo de codificación binaria se llama código Gray. La salida final del codificador puede ser directamente este código Gray, o un código

binario normal (tras la correspondiente conversión realizada por un circuito electrónico).

Los codificadores absolutos pueden ser de una sola vuelta, o multivoltas. Además pueden tener resoluciones desde 8 bits (256 posiciones) hasta 19 bits (524288 posiciones).

El inconveniente de estos codificadores es su alto precio.

2.12 OTROS TRANSDUCTORES

Existen transductores para medir todo tipo de magnitudes físicas. En esta sección se describen algunos de los más usuales.

Velocidad

Los sensores de velocidad más habituales lo son de velocidad de giro de ejes. El sensor más utilizado antiguamente era la **dinamo tacométrica** (o tacómetro). En esencia es un generador de imanes permanentes, en el que la tensión generada es proporcional a la velocidad de giro.

Hoy en día es mucho más común utilizar **codificadores incrementales** para medir la velocidad. El codificador incremental da una señal de pulsos cuya frecuencia es proporcional a la velocidad de giro. La conversión de los pulsos a una tensión continua se puede realizar mediante un convertidor frecuencia-tensión; que es un circuito electrónico especial que realiza esta función. Otra posibilidad es contar los pulsos que se producen en un tiempo determinado, calculando después la velocidad. Esta posibilidad es la que se utiliza cuando el sistema que trata la señal está basado en microprocesador.

Aceleración

Los sensores encargados de medir la aceleración se denominan acelerómetros. Se basan en la utilización de una masa conocida fijada a un elemento elástico. La aceleración produce sobre la masa una fuerza igual al producto de la masa por la aceleración. El sensor mide esa fuerza, con lo que se tiene una medida de la aceleración. Hay muchas formas de medir esa fuerza, pero las más utilizadas son el efecto piezoeléctrico y las galgas extensiométricas.

Fuerza

Se basan en medir la deformación que produce la fuerza aplicada sobre un elemento elástico. Los más comunes son: el sensor piezoeléctrico y el de galgas extensiométricas. El sensor piezoeléctrico consta de un cristal de material piezoeléctrico, que al deformarse por la aplicación de la fuerza produce una diferencia de tensión proporcional a dicha fuerza. El sensor de galgas tiene un elemento elástico (que se deforma muy poco al aplicarle la fuerza), y unas galgas extensiométricas que permiten medir la deformación de ese elemento

(del orden de micras). Las galgas extensiométricas son unas resistencias que varían al alargarse o encogerse, y que permiten medir estas deformaciones tan pequeñas.

Otros

- **Temperatura.** Existen multitud de sensores para medir la temperatura. Los más utilizados son: termopares, resistencia de platino (PT100), pirómetros, sensores de semiconductor y termistores.
- **Nivel.** Existen diversos sensores para medir el nivel de líquidos en depósitos. Los más comunes son: capacitivos, de presión diferencial y de ultrasonidos.
- **Caudal.** Los más comunes son: de turbina, de presión diferencial, electromagnético, de hilo caliente y ultrasónico.
- **Presión.** Se basan en medir la fuerza producida por el fluido sobre un elemento elástico. Si son diferenciales, miden la diferencia de fuerza entre dos puntos del fluido. Si son absolutos, miden la fuerza del fluido contra un vacío de presión nula. La medida de la fuerza puede ser piezoeléctrica o de otro tipo.

2.13 ACTUADORES

Los actuadores son los elementos que convierten las señales eléctricas del equipo de control en acciones físicas sobre el proceso. Se puede distinguir entre preactuadores y actuadores. El preactuador es el elemento que actúa de interfaz, recibiendo como entrada la señal eléctrica y actuando sobre el actuador. Los más utilizados en la industria son aquellos que permiten regular el paso de fluidos por un conducto (válvulas, bombas y ventiladores), y aquellos que permiten realizar un movimiento en objetos que se manipulan o en partes de una máquina (cilindros neumáticos e hidráulicos, y motores eléctricos).

Actuadores y preactuadores neumáticos

Los actuadores neumáticos por excelencia son los cilindros, cuyo esquema básico está representado en la figura 2.53. Éstos se mueven por acción del aire comprimido que actúa a uno de los lados del pistón. Los hay de doble efecto (los más habituales) y de simple efecto. Los de doble efecto pueden hacer fuerza en los dos sentidos, en función de que se conecte a presión una u otra cara del émbolo. Los de simple efecto solo pueden hacer fuerza en un sentido, incorporando un resorte para llevar el pistón a su posición de reposo cuando deja de actuar la fuerza del aire.

La figura 2.54 muestra la estructura interna de un cilindro de doble efecto.

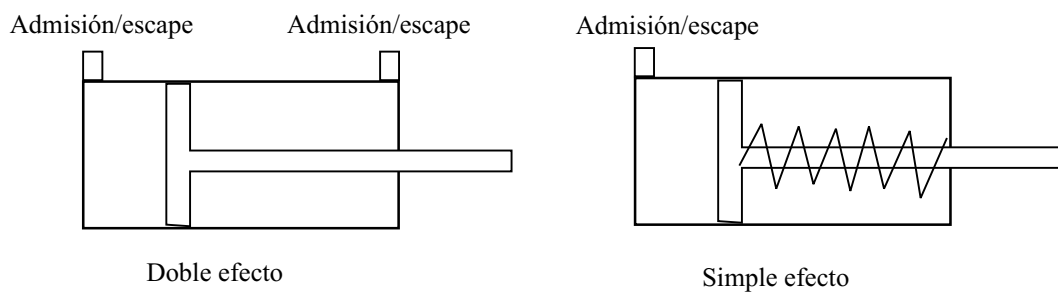


Figura 2.53: Representación esquemática de cilindros neumáticos de simple y de doble efecto.

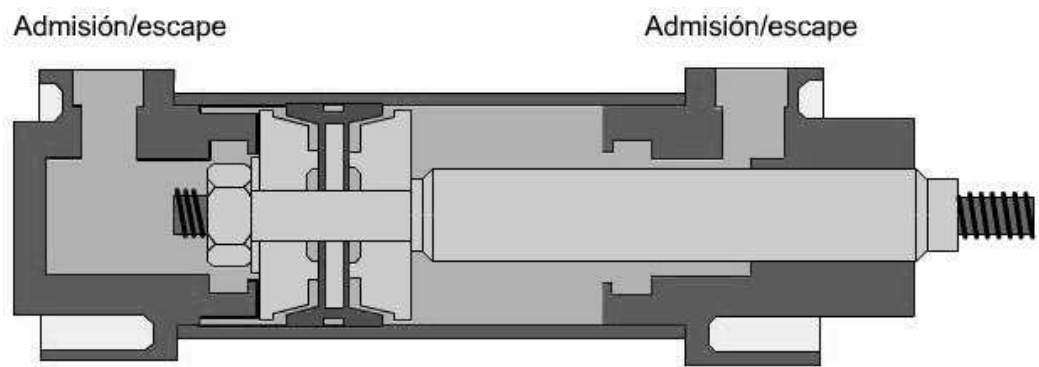


Figura 2.54: Estructura interna de un cilindro de doble efecto.

El vástago del cilindro es el que actúa sobre el sistema, empujando o estirando para producir el movimiento. Pueden hacer fuerzas importantes, ya que la presión del aire suele estar entre 6 y 10 kg/cm², y la sección del cilindro puede ser grande (100 cm² o más).

Existen cilindros de muchos tipos. Algunos llevan un mecanismo antigiro para evitar que el émbolo gire. Otros no llevan vástago, y se utilizan cuando las carreras son largas porque de esta forma la longitud total es la de la carrera, y no el doble.

Para que los cilindros actúen es necesario poner a presión una parte del cilindro y a escape (al aire) la otra. Esto se consigue generalmente mediante electroválvulas, que son los preactuadores neumáticos por excelencia.

Las electroválvulas se caracterizan por el número de vías y el número de posiciones. El número de vías es el número de conductos que se pueden conectar a la electroválvula. El número de posiciones representa las distintas posiciones que puede tener la corredera (elemento móvil que produce la apertura o cierre) de la válvula.

Por ejemplo, una electroválvula de 3 vías y 2 posiciones se representa como 3/2. Las diferentes interconexiones que se pueden obtener entre las vías de una electroválvula se logran mediante las distintas posiciones. Así pues, una válvula 3/2 solo puede tener dos estados de conexión, en función de que esté en una posición o en otra. El número de posiciones de una electroválvula depende del número de bobinas que actúan sobre la corredera. Si solo hay una bobina, únicamente puede haber dos posiciones, pues la bobina solo estará activa o inactiva.

La representación gráfica de la válvula se hace dibujando un cuadrado para cada posición, poniendo dentro de ese cuadrado el estado de conexión correspondiente a esa posición. Por ejemplo, una válvula 2/2 tiene el esquema representado en la figura 2.55.

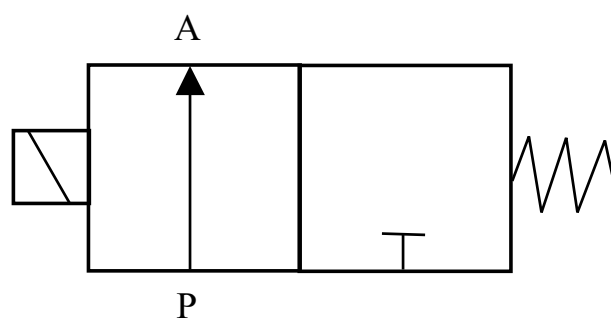


Figura 2.55: Esquema de una electroválvula 2/2 (2 vías y 2 posiciones).

Solo hay dos vías, es decir, un conector de entrada y uno de salida. La P indica que ahí se conecta la presión. La A indica que es la salida que actúa sobre el sistema. La bobina de la izquierda indica que es una electroválvula,

de forma que cuando se activa la bobina se tiene la posición de la izquierda. El muelle de la derecha indica que cuando se desactiva la bobina, la corredera recupera su posición de reposo (la de la derecha) gracias a un muelle. Esta válvula se denomina monoestable, pues si se desactiva la bobina solo tiene una posición de reposo posible. La válvula 2/2 no se utiliza para mover un cilindro, pues solo puede dejar pasar la presión o no, y esto no es suficiente.

Los cilindros de simple efecto (los menos habituales) se pueden accionar mediante una electroválvula 3/2 representados en la figura 2.56.

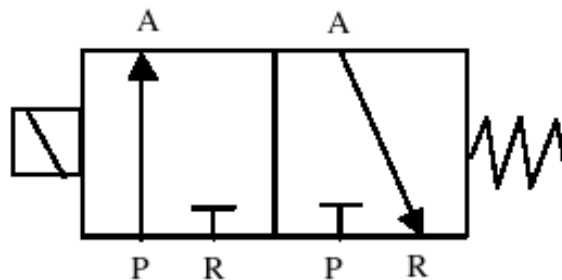


Figura 2.56: Esquema de una electroválvula 3/2 (3 vías y 2 posiciones).

Esta válvula es normalmente cerrada, pues en la posición de reposo, la salida (A) se encuentra desconectada de la entrada de presión (P) y conectada al escape (R). La conexión A se conectaría a la parte del cilindro de simple efecto opuesta al resorte. De esta forma, cuando se activa la bobina, se pone esa parte del cilindro a presión, y el cilindro se mueve venciendo el resorte. Cuando se desactiva la bobina, se cierra la presión, y se comunica esa parte del cilindro con el escape (con el exterior), permitiendo que se vacíe de aire y que el resorte del cilindro mueva a éste hasta su posición de reposo.

Las más habituales son las electroválvulas 5/2 representadas en la figura 2.57, que sirven para accionar cilindros de doble efecto.

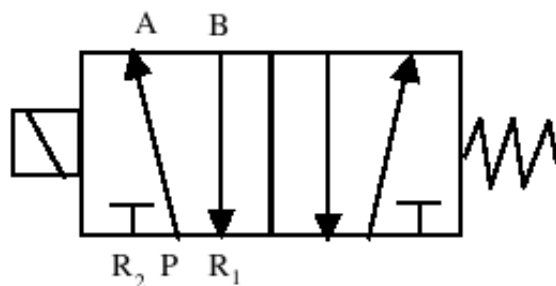


Figura 2.57: Esquema de una electroválvula 5/2 (5 vías y 2 posiciones).

Ésta es una válvula monoestable (si la bobina está desactivada, la posición es la de la derecha, mientras que si está activada, la posición es la de la izquierda). Las salidas A y B se conectan a las dos partes del cilindro de doble efecto.

Si está en una posición, A está a presión y B a escape, por lo que el cilindro se mueve en un sentido. Si está en la otra posición, las conexiones se invierten y el cilindro se mueve en sentido contrario. Con esta válvula el cilindro solo puede estar en una de las dos posiciones extremas.

También hay electroválvulas 5/2 que son biestables, es decir, que tienen dos bobinas sin ningún resorte. Si se activa una bobina, la válvula pasa a una posición, quedándose en ella aunque se desactive la bobina. Si se activa la otra bobina, se pasa a la otra posición, quedándose en ella aunque se desactive. En este caso, con la válvula desactivada hay dos posiciones posibles (biestable). El símbolo es parecido, pero con una bobina a la derecha y otra a la izquierda.

Otra electroválvula muy usada es la 5/3. En este caso se tiene 3 posiciones posibles, lo cual indica que habrá necesariamente 2 bobinas. Pueden ser de centro cerrado, de centro abierto, o con centro a presión, como se observa en la figura 2.58.

Se utilizan también para accionar cilindros de doble efecto, con la ventaja adicional de que si se desactivan las dos bobinas, la válvula se queda en la posición central, en la que según la configuración puede conseguirse que el cilindro se quede parado en una posición intermedia (con la válvula de centro cerrado).

Las electroválvulas se activan conectando la bobina a la tensión correspondiente. Ésta suele ser de 24 V de continua, aunque las hay de 220 V de alterna. Las bobinas de las electroválvulas actuales consumen muy poca corriente (del orden de 100 mA), por lo que pueden ser activadas directamente por la salida de un detector. Sin embargo, es mucho más habitual que se activen a partir de la salida de un autómatas programable, tal y como se explicará en el tema 3.

Además de los cilindros existen los motores neumáticos (actuadores de giro). Son mucho menos frecuentes, y producen el giro de un eje. Pueden ser unidireccionales, en cuyo caso giran si tienen presión y dejan de girar si no la tienen. Éstos pueden accionarse mediante una electroválvula 2/2 o 3/2. También los hay bidireccionales que giran en un sentido o en otro según se conecte a presión una entrada u otra. Éstos se pueden accionar con una electroválvula 5/2.

Los actuadores neumáticos tienen la gran ventaja de la comodidad y limpieza del fluido que utilizan. Al utilizar aire, éste se vierte directamente al exterior, no necesitando un circuito de recirculación del mismo. Además, las pequeñas pérdidas (en cilindros o válvulas) no originan problemas, pues el escape de aire es limpio.

Actuadores y preactuadores hidráulicos

Los actuadores hidráulicos se utilizan para realizar fuerzas muy grandes, cuando los neumáticos no son capaces. Se basan en la acción del aceite a presión. Los actuadores básicos son también cilindros, cuyos émbolos empujados por el aceite a presión ejercen la fuerza sobre el sistema a mover. Las presiones

son del orden de 50 a 300 kg/cm², por lo que pueden hacer fuerzas mucho mayores que los neumáticos.

En cambio, tienen el inconveniente de que el fluido que utilizan no se puede verter al exterior, sino que hay que recircularlo hasta un tanque para reutilizarlo. Además, las pequeñas fugas son problemáticas, debido a que el aceite ensucia.

En cuanto a los cilindros son siempre de doble efecto.

También son frecuentes los motores hidráulicos (actuadores de giro). Pueden ser de paletas, de engranajes o de pistones. En cualquier caso, tienen una entrada y una salida de aceite, de forma que cuando se conecta la entrada a presión, el eje del motor gira. Pueden ser unidireccionales o bidireccionales.

Igual que en los actuadores neumáticos, para accionar los cilindros y motores hidráulicos se utiliza electroválvulas. Las más habituales son la 3/2, que se utiliza para accionar un motor hidráulico unidireccional, y la 4/3 cuyo esquema se representa en la figura 2.59, que se utiliza para accionar un cilindro de doble efecto o un motor bidireccional, y que permite parar el cilindro en una posición intermedia cualquiera. El significado de los símbolos es el mismo, salvo que R no significa ahora escape al exterior, sino retorno del aceite.

Actuadores eléctricos

Los actuadores eléctricos por excelencia son los motores eléctricos. Los motores eléctricos más utilizados en automatización son:

- **Motores de inducción.** Son los más económicos. Funcionan con tensión alterna trifásica. Se utilizan, fundamentalmente, para aplicaciones donde se controla la velocidad (no la posición). Aplicaciones típicas son las cintas transportadoras, ventiladores, bombas, compresores, etc.
- **Motores brushless de imanes permanentes.** Son los que se han desarrollado más recientemente, gracias a los potentes imanes de tierras raras descubiertos no hace muchos años. Para su conexión necesitan un equipo electrónico específico, que puede tener alimentación monofásica o trifásica, aunque el motor en sí es trifásico (el equipo de control es el encargado de dar tensión a las tres fases). Tienen unas características de par/velocidad muy buenas, por lo que se utilizan sobre todo en aplicaciones de control de posición (para control de velocidad también funcionan bien, pero los de inducción son más baratos). Prácticamente han sustituido a los motores de continua en las aplicaciones de posicionado, debido a que su mantenimiento es mucho más simple (no hay escobillas rozantes), y a que los equipos de control tienen precios cada vez más bajos.
- **Motores paso a paso.** Son motores que avanzan a pasos, es decir, a incrementos angulares determinados. Se necesita un equipo de control

específico que es el encargado de alimentar los devanados del motor, conmutando las fases de la forma adecuada para que el motor avance o retroceda un paso. Tienen la gran ventaja sobre los demás motores de que si no se pierde el paso, la posición del motor es perfectamente conocida sin necesidad de sensor (encoder o resolver). La posibilidad de pérdida del paso depende del par resistente que tienen que vencer. Se utilizan sobre todo en aplicaciones donde el par es pequeño y no se requieren grandes aceleraciones; fundamentalmente, para el posicionado de ejes.

- **Motores de continua.** Son los que más se han utilizado históricamente para controlar la posición y velocidad de ejes debido a su excelente característica par/velocidad. Pueden funcionar conectándolos directamente a una tensión continua, pudiendo variar la velocidad variando esa tensión. Normalmente, se atacan por medio de un equipo de control específico para motores de continua. Hoy en día se utilizan cada vez menos, ya que están siendo sustituidos por motores de inducción con variadores de frecuencia (para control de velocidad), y por motores brushless (para control de posición).
- **Resistencias calefactoras.** Se utilizan, sobre todo, en sistemas térmicos para calentar (como hornos, por ejemplo).

Los preactuadores eléctricos por excelencia han sido siempre los contactores y los relés. Cuando se activa la bobina de mando de estos elementos se cierran los contactos permitiendo la conexión o desconexión de motores de diversos tipos. El desarrollo de equipos electrónicos de potencia ha aumentado el número de lo que se considera preactuadores eléctricos, así como las funciones que llevan integradas.

- **Relé.** Con un contacto de un relé se puede conectar o desconectar un motor de continua o una resistencia calefactora.
- **Contactor.** Con un contactor trifásico se puede conectar o desconectar directamente un motor de inducción o tres resistencias calefactoras. Cuando no hay problemas de corrientes de arranque este método puede ser suficiente para conectar o desconectar un motor de inducción.
- **Regulador de alterna.** Se trata de un equipo que permite variar la tensión alterna (sin variar la frecuencia). No es adecuado, por lo tanto, para regular la velocidad de motores trifásicos de inducción. Su utilidad

fundamental es regular la potencia suministrada por una resistencia calefactora.

- **Arrancador estático.** En muchas situaciones prácticas es conveniente limitar la corriente de arranque del motor de inducción, para lo cual se puede utilizar un arrancador estático, que va aumentando progresivamente la tensión que aplica al motor conforme éste acelera.
- **Variador de frecuencia.** Es un equipo electrónico de potencia que permite atacar un motor de inducción con una señal trifásica de frecuencia programable. De esta forma, se logra cambiar a voluntad la velocidad de giro del motor. Hay muchos modelos en el mercado, con prestaciones diferentes, desde los más simples, que únicamente permiten modificar la frecuencia de salida, hasta los que funcionan en bucle cerrado (midiendo la velocidad mediante un encoder, por ejemplo), implementando un regulador PID, y que incorporan funciones de comunicación (para comunicar con un autómatas, por ejemplo). De hecho, la forma de actuar sobre el variador desde un autómatas es por medio de un puerto de comunicaciones (a través del cual se le pasa la velocidad deseada, o la orden de marcha y paro), o por medio de salidas digitales (para la orden de marcha o paro) o salidas analógicas (para la velocidad deseada).
- **Equipo de control de motor brushless.** Es similar al variador de frecuencia pero sirve para atacar motores brushless. La parte de potencia es idéntica a la del variador de frecuencia. De hecho, algunos equipos sirven para los dos tipos de motores. Suelen incorporar funciones de posicionado (bucle de control PID con realimentación por encoder o resolver) y de comunicaciones.
- **Equipos de control de motores paso a paso.** Sirven para actuar sobre los motores paso a paso. Normalmente, incorporan funciones de avance (o retroceso) de medio paso o de un paso, o funciones de control de velocidad o de posición (en bucle abierto o en bucle cerrado con realimentación por encoder), así como funciones de comunicación.
- **Equipos de control de motores de continua.** Permiten atacar al motor de continua con una tensión variable, lo que permite modificar su velocidad. Pueden tener funciones de control de velocidad o posición, o funciones de comunicaciones.

Actuadores para control de fluidos

Los más importantes son las válvulas (para líquidos y para gases), las bombas (para líquidos) y los ventiladores y compresores (para gases).

Las válvulas permiten o impiden el paso del fluido por el conducto gracias al movimiento de un elemento que produce la obturación. En función de las posibles posiciones de ese elemento las válvulas se clasifican en:

- **Válvulas todo/nada.** Únicamente pueden estar totalmente abiertas o totalmente cerradas. Requieren por lo tanto una señal binaria (activo o inactivo).
- **Válvulas proporcionales.** Pueden adoptar cualquier posición intermedia entre la apertura completa y el cierre completo. Requieren por lo tanto una señal continua (en tensión o en corriente, por ejemplo) que le indique el grado de apertura.

En función de la forma en que se mueve el elemento que produce la obturación en la válvula, éstas se clasifican principalmente en:

- **Electroválvulas de mando directo.** El elemento móvil se mueve por la acción directa de la fuerza del campo magnético producido por una bobina. En las todo/nada, la apertura o cierre se consigue poniendo en tensión o desconectando la bobina. En las proporcionales, la apertura intermedia se consigue regulando la corriente por la bobina. La electroválvula suele venir acompañada de un circuito electrónico que recibe una entrada analógica estándar (0 - 10 V o 4 - 20 mA) y que ataca a la bobina con una corriente proporcional. Estas válvulas admiten caudales pequeños.
- **Electroválvulas autopilotadas.** Constan de una pequeña electroválvula de mando directo en el interior, que abre un pequeño conducto por el que la presión del propio fluido produce el movimiento y apertura de la válvula principal. Necesita, por tanto, una presión mínima en el fluido para funcionar. Pueden ser también todo/nada o proporcionales. Admiten caudales más elevados que las de mando directo.
- **Válvulas neumáticas.** Son aquellas en las que el elemento móvil es accionado por un elemento neumático (un pequeño cilindro de membrana accionado mediante aire comprimido). Normalmente, se utilizan junto con un preactuador que recibe una señal eléctrica y la transforma en señal neumática. Si la válvula es todo/nada, el preactuador es una electroválvula, que deja pasar el aire a presión hacia la válvula o no. Si la válvula es proporcional, el actuador suele ser un convertidor corriente-presión (electroválvula proporcional de presión de aire), que recibe una señal analógica (4 - 20 mA) y da como salida una presión de aire que

actúa sobre la válvula abriéndola más o menos. Estas válvulas admiten caudales mayores a las autopilotadas.

- **Válvulas motorizadas.** Son aquellas en las que el elemento móvil es accionado por un motor eléctrico. Suelen incorporar el equipo electrónico que ataca al motor eléctrico. Pueden ser todo/nada o proporcionales. Las proporcionales suelen tener un sensor de posición que mide la posición del elemento móvil (y por tanto el grado de apertura). En ese caso la señal de mando suele ser una señal analógica estándar (0 - 10 V o 4 -20 mA).

El flujo de un fluido por un conducto se puede modificar también por medio de una bomba (para líquidos) o de un ventilador o compresor (para gases). Estos dispositivos se accionan mediante un motor eléctrico (normalmente de inducción), por lo que se utilizan los preactuadores ya descritos para estos motores. Una configuración muy utilizada (cuando se desea un control continuo) es la utilización de un variador de frecuencia junto con una bomba (o un ventilador, o un compresor), lo que permite variar de forma continua el flujo.

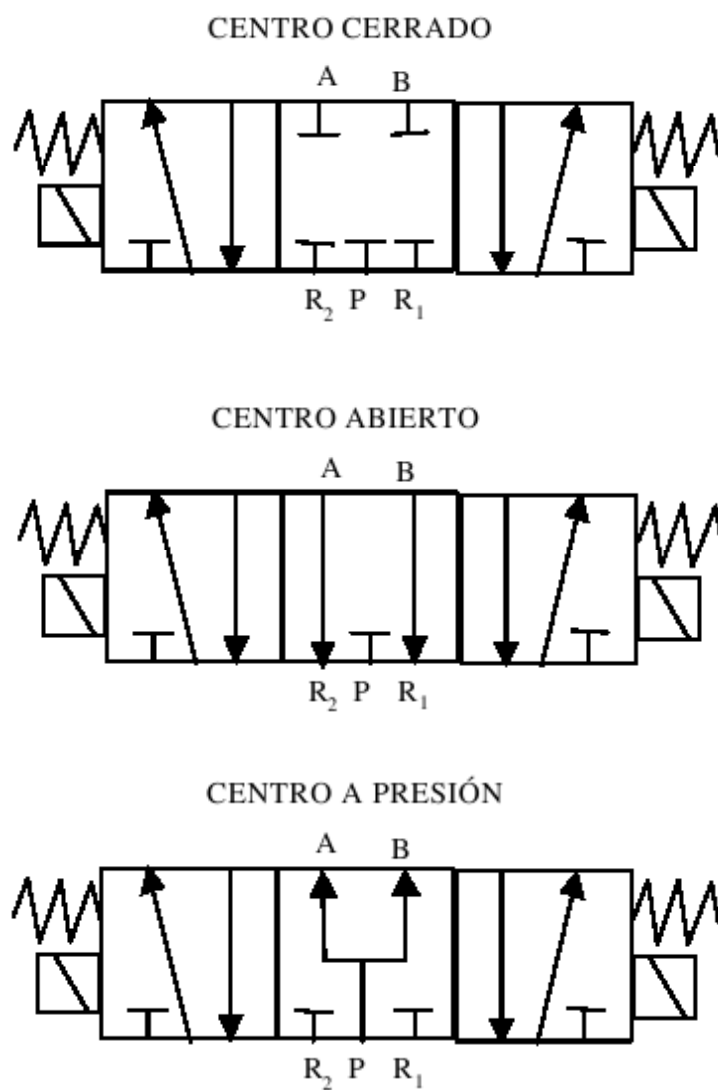


Figura 2.58: Esquema de una electroválvula 5/3 (5 vías y 3 posiciones).

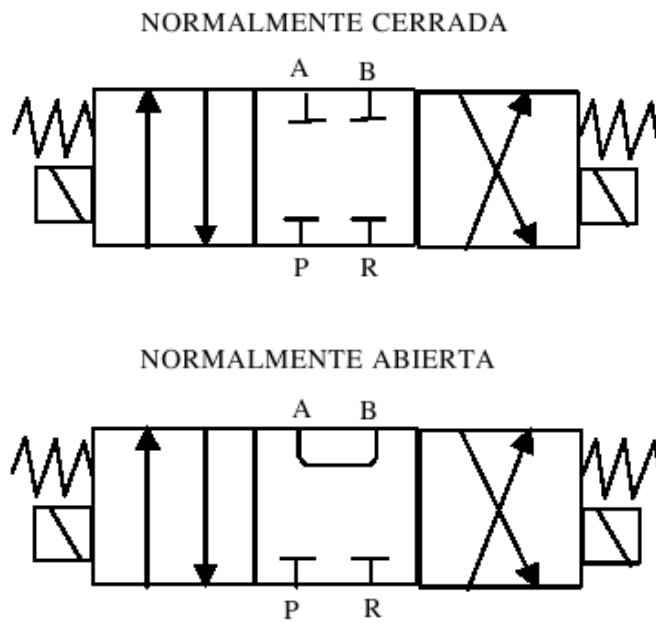


Figura 2.59: Esquema de una electroválvula 4/3 (4 vías y 3 posiciones).



Figura 2.60: Arrancador estático Altistart 01 de Schneider Electric.



Figura 2.61: Variador de frecuencia Altivar 61 de Schneider Electric.

CAPÍTULO 3

MODELADO DE SISTEMAS DE EVENTOS DISCRETOS MEDIANTE EL DIAGRAMA ETAPA TRANSICIÓN GRAFCET

3.1 SISTEMAS DE EVENTOS DISCRETOS

Se define un sistema dinámico de eventos discretos como aquél que evoluciona entre un número finito de estados discretos y cuyo paso de un estado a otro depende del valor de determinadas variables binarias.

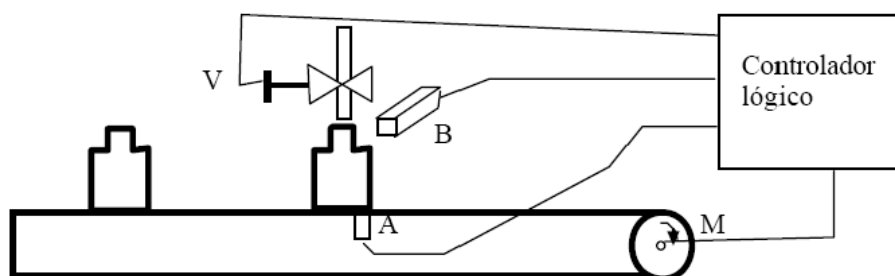


Figura 3.1: Proceso de llenado de botellas.

Como ejemplo considérese la embotelladora que se muestra en la figura 3.1. Existen 2 señales dadas por 2 sensores. La señal A es 1 si tiene una botella encima y 0 en caso contrario. La señal B es 1 si el líquido que llena la botella ha llegado al nivel máximo y 0 en caso contrario. Asimismo existen otras 2 señales. Cuando M se pone a 1, el motor se pone en marcha, y cuando se pone a 0, se para. Cuando V se pone a 1, la electroválvula se abre empezando a llenar la botella, cerrándose cuando $V = 0$.

3.2 MODELADO MEDIANTE DIAGRAMA DE ESTADOS

En el sistema de eventos discretos anterior se distinguen 3 estados: botella acercándose, botella llenándose y botella llena alejándose. El paso de un estado a otro se produce según los valores de las variables binarias A y B. El diagrama de estados está descrito en la figura 3.2.

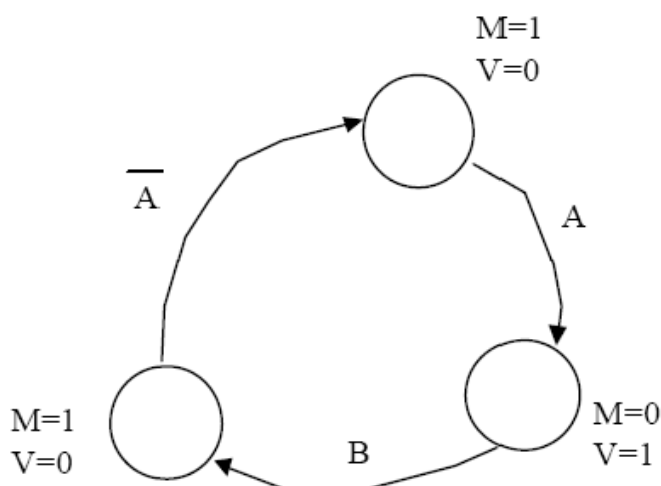


Figura 3.2: Diagrama de estados de la embotelladora.

En realidad, tanto el proceso a controlar como el controlador lógico son sistemas de eventos discretos que interaccionan entre sí, dando lugar a un sistema global que cumple unos objetivos. La interacción se produce a través de las señales de entrada y salida. En este caso, A y B son señales de entrada para el controlador lógico, y de salida para el proceso. M y V son señales de salida para el controlador lógico, y señales de entrada para el proceso.

El controlador lógico tiene una serie de variables de entrada. Estas variables son las que determinan el paso de un estado al otro del controlador. Los cambios en esas variables son los eventos discretos que hacen evolucionar al controlador. Se denomina condiciones de transición a las condiciones lógicas que producen el paso de un estado al siguiente.

Por otra parte, el controlador tiene una serie de variables de salida. Estas variables de salida son las que determinan la actuación sobre el sistema. Normalmente, su activación se asocia a un estado determinado (cuando el controlador está en un estado, activa la salida o salidas asociadas a ese estado). La activación de las salidas puede estar condicionada además a alguna variable. Por ejemplo, una salida puede activarse en un estado sólo si además una entrada está activa.

3.3 LAS REDES DE PETRI

Introducción

Los diagramas de estados son una herramienta válida para modelar sistemas de eventos discretos. Sin embargo, tienen algunos inconvenientes importantes. Considérese el siguiente ejemplo.

En la figura 3.3 los carros están inicialmente en A y C. Cuando se pulsa P, se ponen en marcha hacia la derecha hasta llegar a B y D. Cada carro, al llegar

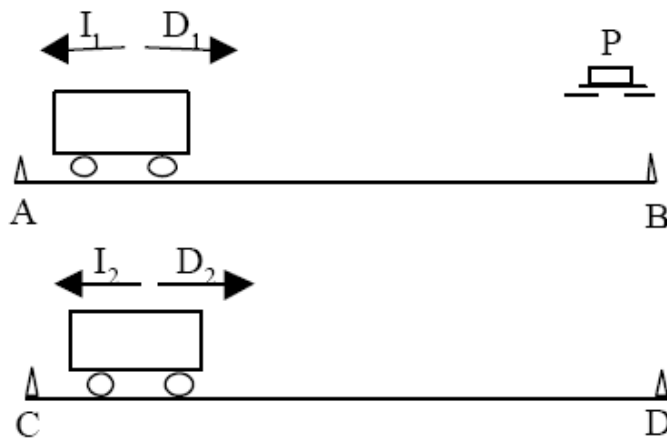


Figura 3.3: Problema de las vagonetas.

al extremo derecho, cambia de sentido de forma independiente poniéndose en marcha hacia la izquierda hasta llegar a C o A. Sólo cuando los dos carros han llegado al extremo izquierdo y se pulsa P, se vuelven a poner en marcha hacia la derecha. Las velocidades de cada carro no se saben en principio, por lo que hay que tener en cuenta todas las posibilidades. Si se intenta dibujar el diagrama de estados se obtiene un esquema muy complicado. Este sencillo ejemplo muestra que el diagrama de estados no es una buena herramienta cuando hay acciones paralelas o concurrentes. Las redes de Petri permiten modelar estos casos de forma mucho más simple.

Definición

Una red de Petri es un grafo orientado formado por 2 tipos de nodos, los **lugares** (simbolizados por un círculo) y las **transiciones** (simbolizadas por una línea recta) unidos alternativamente por **flechas**. Las flechas unen lugares con transiciones y viceversa, pero nunca elementos iguales. Tanto los lugares como las transiciones pueden tener varias flechas de entrada o de salida. En un momento determinado cada lugar puede tener una o varias marcas (representadas por puntos). El conjunto de marcas de la red define el estado del sistema.

La evolución de la red de Petri es como sigue: cuando todos los lugares anteriores a una transición tienen marcas y la condición asociada a la transición está activa, la transición se dispara y el sistema evoluciona reduciendo en una las marcas de los lugares anteriores y aumentando en una las marcas de los lugares posteriores.

Una transición está validada (o está habilitada) cuando todos los lugares anteriores están marcados. Se denomina receptividad asociada a una transición a la condición asociada a su disparo.

Las redes de Petri utilizadas en automatización suelen tener solo una marca

en cada lugar (o ninguna). Si tiene una marca se dice que el lugar está activo y si tiene cero, que está inactivo. La figura 3.4 representa el ejemplo de las vagonetas resuelto mediante redes de Petri.

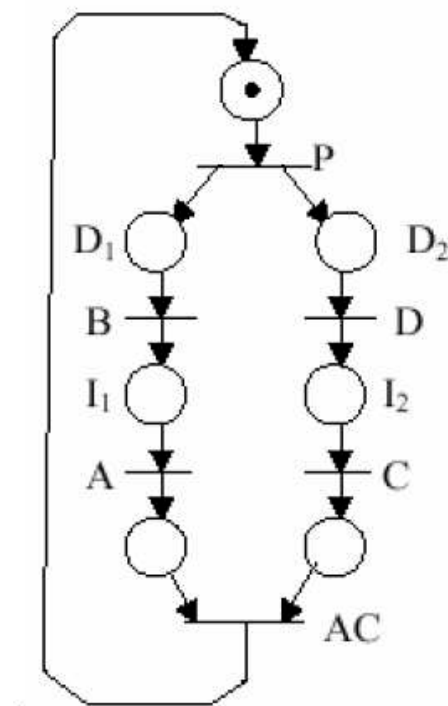


Figura 3.4: Red de Petri que modela el problema de las vagonetas.

La diferencia fundamental con los diagramas de estados son las transiciones de bifurcación. En este caso, cuando se pulsa P, la transición de distribución hace que se marquen simultáneamente dos lugares (uno para cada carro) de forma que hay dos secuencias funcionando en paralelo. Cuando los dos carros han llegado al punto inicial, la transición de conjunción (o sincronización) impone la condición de que los dos estén en A y C antes de volver a empezar un ciclo.

La diferencia fundamental con el diagrama de estados es que en aquél sólo hay un estado activo (que representa el estado del sistema). En la red de Petri puede haber varios lugares activos al mismo tiempo. El estado del sistema queda definido por los lugares marcados (activos), por lo que con pocos lugares se puede representar muchos estados.

Las redes de Petri se utilizan para modelar el comportamiento de los controladores lógicos. Éstos reciben una serie de entradas del proceso y dan como resultado una serie de salidas que actúan sobre el proceso.

Al igual que en el diagrama de estados, las condiciones asociadas a las transiciones suelen depender de las entradas. Lo más normal es que sean una función lógica de las entradas. A veces, puede ser una condición de flanco (de cambio) de una combinación lógica de entradas.

Las salidas suelen definirse asociadas a la activación de los lugares, es decir, una salida se activa cuando el lugar asociado está activo. Algunas salidas pueden, además, estar condicionadas a otras variables (entradas o temporizadores), de forma que solo se activan si el lugar está activo y la otra condición se cumple. Cuando el lugar está activo se dice que la salida es **sensible** a esa condición adicional. También puede haber salidas impulsionales (que solo se activan durante un pequeño pulso). Estas salidas pueden asociarse a la transición entre un lugar y otro.

Algunas características de las redes de Petri que las hacen especialmente indicadas para modelar sistemas de eventos discretos son:

- La red de Petri permite modelizar sistemas de eventos discretos con evoluciones simultáneas, permitiendo una descripción sencilla de los mismos.
- La estructura de la red de Petri contiene información inmediata de las propiedades funcionales de sistema, por lo que se puede validar fácilmente la corrección del modelo.
- La red de Petri es una herramienta de modelado independiente de la tecnología que se utilice para implementar el controlador lógico.
- Además permite describir el sistema por refinamientos sucesivos.

3.4 DIAGRAMA DE ETAPA TRANSICIÓN: EL GRAFCET

El Grafcet es una herramienta de modelado de sistemas de eventos discretos derivada de las redes de Petri. En realidad, es como una red de Petri en la que los lugares solo pueden tener una marca. Es una herramienta adecuada para representar sistemas con evoluciones simultáneas.

Los lugares se llaman **etapas**. Una etapa está activa si tiene una marca, y está inactiva si no tiene marca. Las etapas tienen acciones asociadas (normalmente salidas), que se activan cuando la etapa está activa. Un diagrama de etapa-transición (GRAFCET) es un grafo orientado formado por dos tipos de nodos, las **etapas** (simbolizadas por un cuadrado) y las **transiciones** (simbolizadas por una línea recta) unidos alternativamente por otras líneas rectas perpendiculares a las transiciones. Se puede unir etapas con transiciones y viceversa, pero nunca elementos iguales. La orientación del grafo es siempre de arriba abajo (se pasa de una etapa a la transición de abajo y de ésta a la etapa de abajo), salvo excepciones en las que la línea que une la etapa y la transición debe tener una flecha indicando el sentido de evolución.

En un momento determinado, cada etapa puede tener una marca (representada por un punto) como se puede observar en la figura 3.5 o no tenerla, indicando que la etapa está activa o no. El conjunto de marcas del Grafcet (de etapas activas) define el estado del sistema.

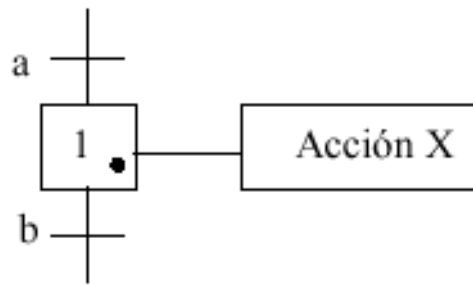


Figura 3.5: Etapa marcada como activa.

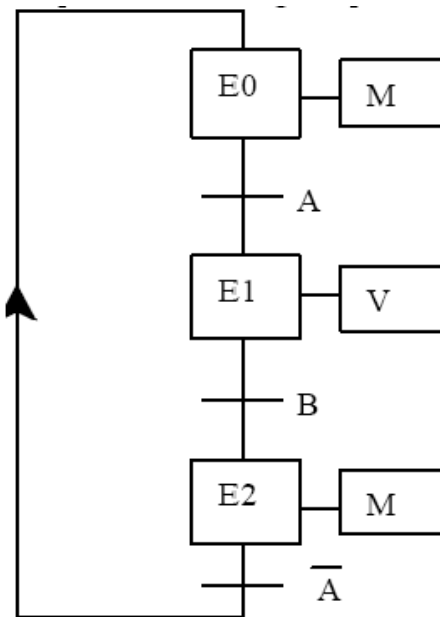


Figura 3.6: Grafcet que modela el ejemplo de la embotelladora.

El ejemplo de la embotelladora quedaría prácticamente igual que en el diagrama de estados como se puede ver en la figura 3.6.

Los rectángulos asociados a cada etapa representan las acciones a realizar cuando la etapa está activa. Se indican únicamente las señales de salida que se tiene que poner a 1 (las que no están incluidas en el rectángulo están a 0 durante esa etapa).

La gran ventaja del Grafcet respecto del diagrama de estados se aprecia, sin embargo, en el ejemplo de las vagonetas, en el que hay actividades que transcurren en paralelo. El proceso quedaría modelado por el Grafcet de la figura 3.7.

La evolución del Grafcet es como sigue: cuando todas las etapas anteriores (por tanto, dibujadas encima) a una transición están activas y la condición asociada a la transición está también activa, la transición se dispara y el sistema evoluciona desactivando las etapas anteriores y activando las etapas posterior-

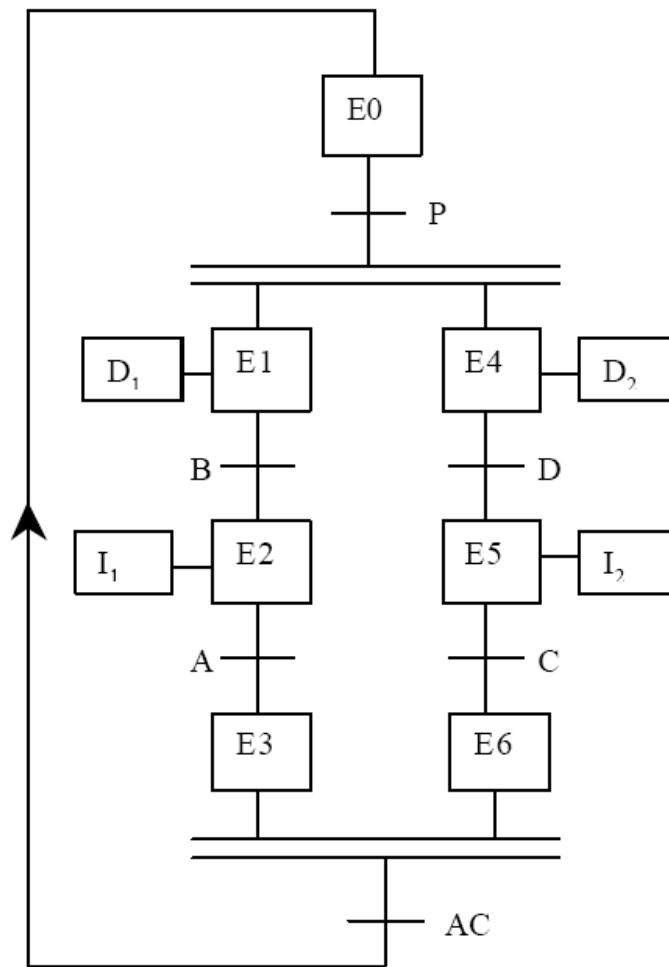


Figura 3.7: Grafcet que modela el ejemplo de las vagonetas.

res.

Se dice que una transición está validada (o está habilitada) cuando todas las etapas anteriores están activas. Se denomina receptividad asociada a una transición, a la condición lógica asociada a su disparo.

La diferencia fundamental con los diagramas de estados son las bifurcaciones. En este caso, cuando se pulsa P, la transición de distribución hace que se activen simultáneamente dos etapas (una para cada carro) de forma que hay dos secuencias funcionando en paralelo. Cuando los dos carros han llegado al punto inicial, la transición de conjunción (o sincronización) impone la condición de que los dos estén en A y C antes de volver a empezar un ciclo.

La diferencia fundamental con el diagrama de estados es que en aquél sólo hay un estado activo (que representa el estado del sistema). En el Grafcet puede haber varias etapas activas al mismo tiempo. El estado del sistema queda definido por las etapas activas, por lo que con pocas etapas se puede representar muchos estados.

Niveles de descripción

Un diagrama Grafcet puede realizarse en 3 niveles:

- Nivel 1. En él se describen las operaciones a realizar, sin hacer mención a la tecnología de sensores o accionadores.

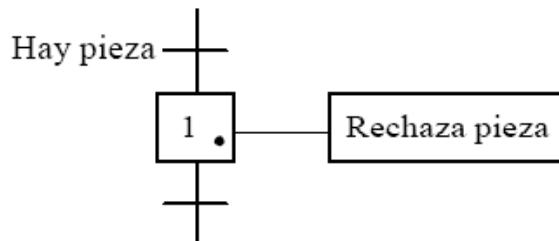


Figura 3.8: Modelado de Grafcet de nivel 1.

- Nivel 2. En él se describen las operaciones a realizar detallando las variables que se activan y que se leen del proceso (sensores y actuadores).

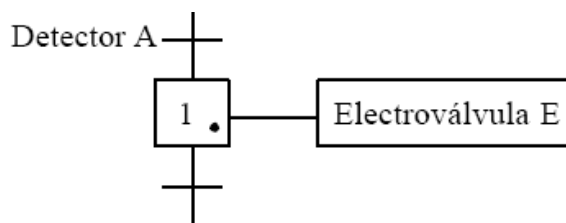


Figura 3.9: Modelado de Grafcet de nivel 2.

- Nivel 3. En este nivel se describen las variables del autómata programable que activan o leen las variables externas.

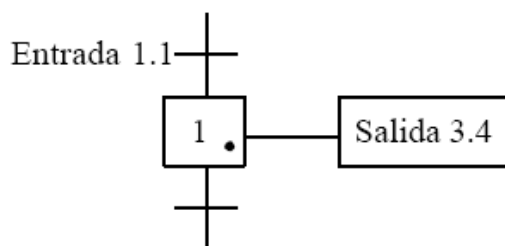


Figura 3.10: Modelado de Grafcet de nivel 3.

En general, bastará con los niveles 1 y 2, pues el software de programación del autómata permite nombrar las variables del autómata con nombres que describan los sensores y actuadores relacionados.

Elementos del Grafcet

Etapa

Se representa por un cuadrado tal y como se indica en la figura 3.11. Puede estar activa o inactiva (con marca o sin ella). Si es una etapa inicial tiene doble cuadrado. En ese caso, se activa cuando se inicializa por primera vez el Grafcet.

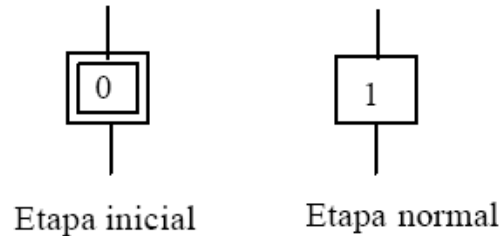


Figura 3.11: Representación de las etapas.

Acción asociada a una etapa

Se representa por un rectángulo unido a la etapa. La acción consiste, normalmente, en la activación de una salida. La salida determinada se activa mientras la etapa esté activa. Puede haber varias acciones a la vez. Una acción puede estar condicionada a una variable. En ese caso se representa con una línea perpendicular donde se incluye la variable adicional que debe estar activa para que se active la salida.

Normalmente, la acción se realiza (la salida permanece activa) mientras la etapa está activa. Sin embargo, puede haber acciones impulsionales, que se ejecutan una sola vez cuando se activa la etapa. La forma de indicar una acción impulsional es incluir una flecha hacia arriba a la izquierda del rectángulo. Entre las acciones impulsionales se pueden encontrar:

- Poner a 1 una variable (Set). Dicha variable quedará a 1 hasta que otra acción impulsional la ponga a 0 (Reset). No se debe utilizar este tipo de acción con una variable de salida que esté definida en otra etapa como acción a nivel.
- Poner a 0 una variable (Reset). Dicha variable quedará a 0 hasta que otra acción impulsional la ponga a 1 (Set). No se debe utilizar este tipo de acción con una variable de salida que esté definida en otra etapa como acción a nivel.
- Incrementar un contador.
- Decrementar un contador.
- Cualquier evento que suponga ejecutar una rutina de programa, como enviar un mensaje, hacer un cálculo, darle un valor a una variable, etc.

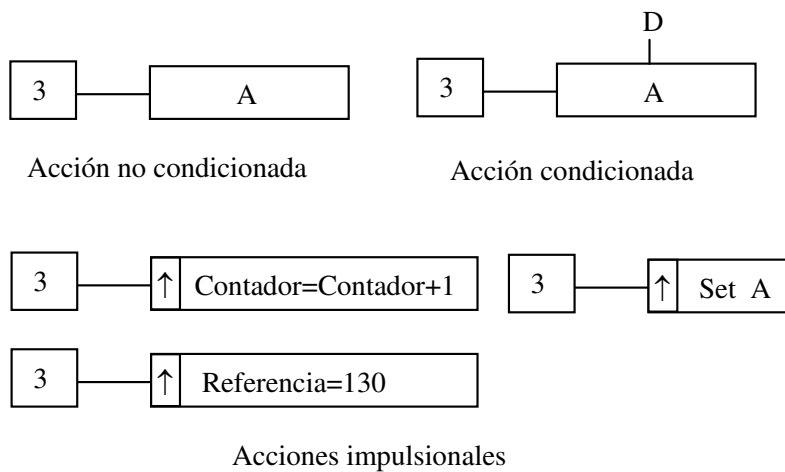


Figura 3.12: Representación de las acciones asociadas a una etapa.

Transición y receptividad

Una transición se representa por un trazo horizontal tal y como se observa en la figura 3.13. Se sitúa siempre entre dos etapas. Tiene asociada una **receptividad**, que puede ser una condición lógica de nivel o de flanco. Cuando es de flanco, se representa con una flecha al lado de la condición (↑ si es de subida, ↓ si es de bajada).

Arco

Es una línea que une una transición con una etapa y viceversa. El sentido es siempre de arriba abajo. Cuando el sentido es el inverso (de abajo arriba), se indica con una flecha, como se puede ver en la figura 3.13.

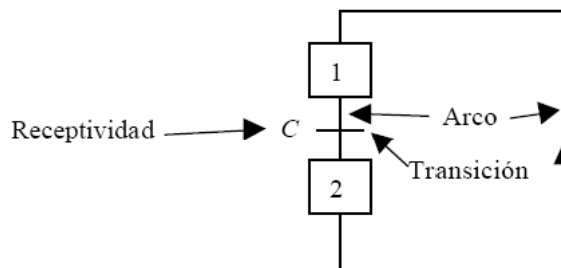


Figura 3.13: Representación de las transiciones y arcos.

Trazos paralelos

Se utilizan para indicar que una transición se une a varias etapas. Sirven para definir secuencias paralelas o simultáneas y para sincronizar los finales de esas secuencias.

--

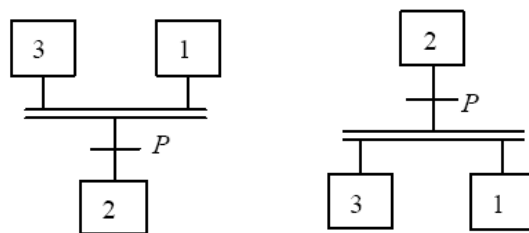


Figura 3.14: Representación de las transiciones entre varias secuencias simultáneas.

Estructuras de programación en Grafcet

Las estructuras que se utilizan con más frecuencia en el diseño de automatismos con Grafcet son:

Secuencia

Es una sucesión alternada de etapas y transiciones. Representa una serie de operaciones secuenciales. Se dice que una secuencia está activa si lo está alguna de sus etapas.

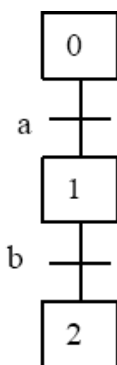


Figura 3.15: Representación de una secuencia.

Selección entre secuencias

Representa una bifurcación en la que según la condición se activa una secuencia u otra. Para evitar problemas suelen ser secuencias excluyentes (nunca se activan simultáneamente).

Salto

Es una selección de secuencias. También suelen ser excluyentes para evitar problemas. El salto puede ser hacia adelante o hacia atrás.

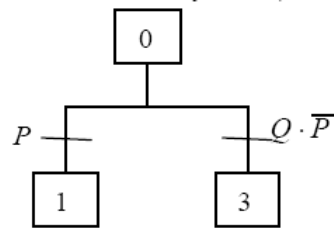


Figura 3.16: Representación de bifurcaciones.

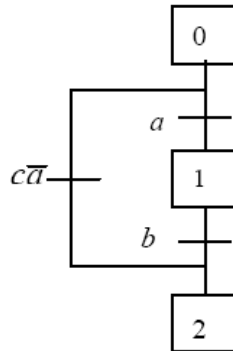


Figura 3.17: Representación de un salto de etapas.

Paralelismo de secuencias

Se utiliza cuando una transición hace que dos o más secuencias se activen al mismo tiempo y, por tanto, evolucionen de forma simultánea. Se representa mediante dos trazos paralelos después de la transición.

Al finalizar las secuencias simultáneas se suele poner una transición de sincronización, formada también por dos trazos paralelos. Indica que sigue la secuencia original sólo cuando han terminado todas las secuencias simultáneas.

Las secuencias simultáneas también podrían producirse con una selección de secuencias que no sea excluyente, pero esto podría producir problemas al finalizar esas secuencias.

Normas especiales de representación Grafcet

Etapas y transiciones fuente y sumidero

- Una etapa fuente no tiene ninguna transición ni etapa anterior. Solo se puede activar al inicializar (o mediante forzado).
- Una etapa sumidero no tiene ninguna transición ni etapa posterior. Una vez activada no se puede desactivar (salvo por forzado).
- Una transición fuente no tiene ninguna etapa anterior. Esta transición está siempre validada, por lo que solo tiene sentido con receptividad activa por flanco.

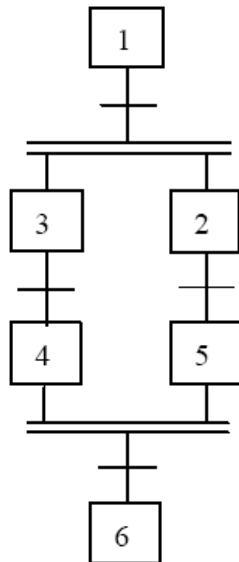


Figura 3.18: Representación de secuencias paralelas.

- Una transición sumidero no tiene ninguna etapa posterior.

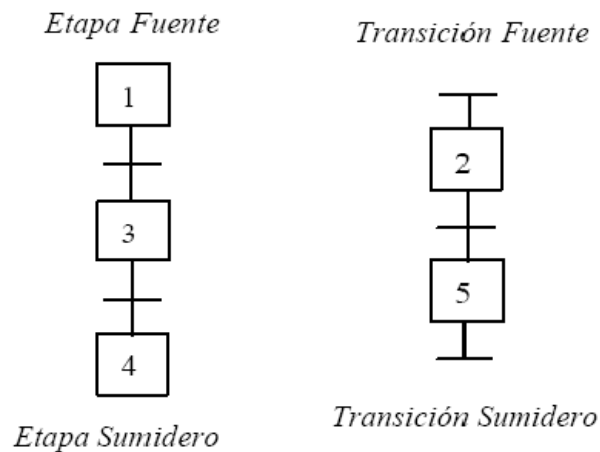


Figura 3.19: Etapas y transiciones fuente y sumidero.

Acciones y receptividades temporizadas

Existen varias formas estándar de representar la medición de tiempos en Grafcet, aunque hay dos formas que son las más utilizadas.

Una de las posibles notaciones tiene la forma:

$$t / N^{\circ}\text{etapa} / \text{tiempo (s)}$$

Por ejemplo: $t/4/2s$ es una señal de temporizador que se activa dos segundos después de que se active la etapa 4. **Esta señal permanece activa hasta**

que se desactive la etapa 4, momento en que se desactivará la señal de temporización.

La otra posible notación es la utilizada por la norma IEC848:

$$t_1 / \text{Variable} / t_2$$

La variable puede ser cualquiera (también una etapa). Por ejemplo: 5s/X4/3s es una señal que se activa 5 segundos después de hacerlo la variable X4, y que se desactiva 3 segundos después de hacerlo X4.

El autómata CQM1 de Omron implementa temporizadores (función TIM) de la forma $t_1/X/0$, es decir, se desactivan al hacerlo la variable X.

Las condiciones de temporización se suelen utilizar en las receptividades, para lograr que una etapa esté activa durante un tiempo determinado. En ese caso, las dos notaciones tienen un efecto similar:

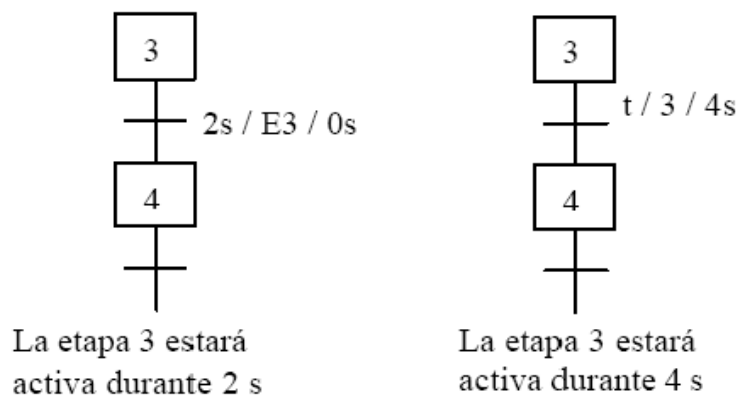


Figura 3.20: Representación de temporizaciones.

También se pueden utilizar las condiciones de temporización para condicionar una salida. Por ejemplo:

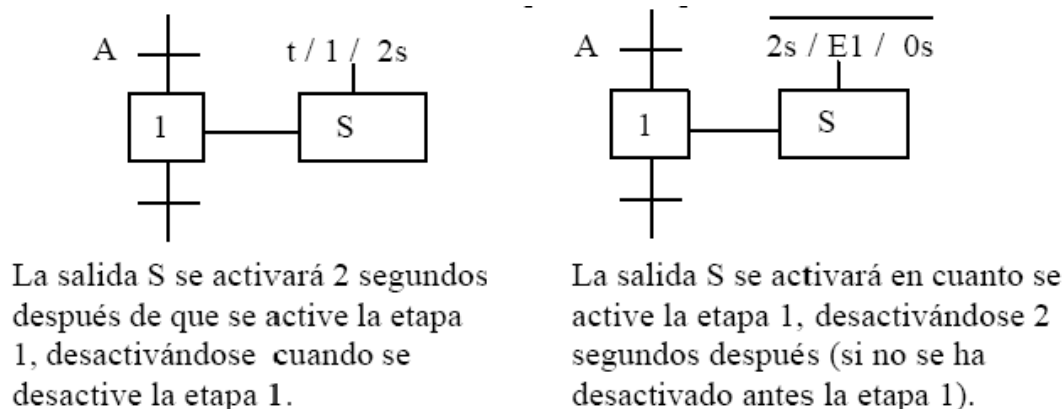


Figura 3.21: Uso de temporizaciones en las salidas condicionadas.

Representación de acciones según IEC848

La acción normal asociada a una etapa representa una variable que estará activa mientras esté activa la etapa. La norma IEC848 contempla otros cinco tipos de acciones, además de la acción normal. Se representan con una letra a la izquierda del rectángulo de la acción. Éstas son:

- **C:** Acción condicionada. Se puede representar también con una línea perpendicular y la condición.
- **D:** Acción retardada (Delayed). Se activa un tiempo después de la activación de la etapa. Ésta se puede representar también con una acción normal y un temporizador.
- **L:** Acción limitada en tiempo. Se activa en cuanto lo hace la etapa, pero solo está activa un tiempo limitado. Se puede representar también con temporizadores.
- **P:** Acción pulso. Se activa cuando se activa la etapa, y dura un pulso muy corto. En la práctica, el pulso debería ser suficiente para producir el efecto deseado. El inconveniente de esta notación es que la duración del pulso queda indefinida. Para representar ese tipo de acción se recomienda utilizar una acción a nivel y un temporizador que defina el tiempo que debe estar activa.
- **S:** Acción memorizada (Set). Se activa cuando se activa la etapa, y no se desactiva aunque se desactive ésta. Para desactivarse se necesita otra acción memorizada (de Reset) en una etapa posterior. Este tipo de acción se puede representar como una acción impulsional en la que se pone a 1 una variable (si la acción es de Set), o a 0 (si la acción es de reset).

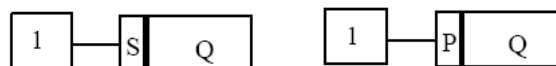


Figura 3.22: Representación de acciones según IEC848.

La notación básica de acciones a nivel y acciones impulsionales, junto con el uso de temporizadores, permite representar todos los tipos de acciones de la norma IEC848, por lo que no es necesario utilizar la notación de esta norma. No obstante, es importante conocer dicha notación para saber interpretar un Grafcet desarrollado según esa norma.

Combinación de paralelismo y selección de secuencia

Cuando se quiere una selección de secuencia inmediata a un paralelismo es necesario introducir una etapa auxiliar intermedia que no tiene ninguna acción.

La situación es que estando en una etapa, al activarse A se quiere empezar dos secuencias, una fija y otra seleccionada entre dos en función del valor de la variable B. En este caso hay que hacer primero una transición paralela con A, y después, en una de las secuencias, poner una etapa intermedia y hacer la selección de secuencias en función de B.

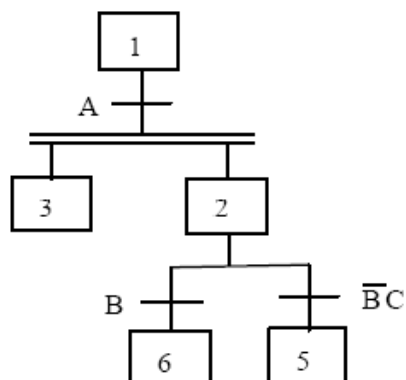


Figura 3.23: Ejemplo de combinación de paralelismo y selección de secuencias.

Macroetapas

Cuando una secuencia de operaciones se utiliza varias veces, se puede definir una macroetapa, que representa toda la secuencia. Se representa gráficamente como un cuadrado con líneas horizontales dobles. Tienen siempre una sola etapa de entrada y una sola etapa de salida.

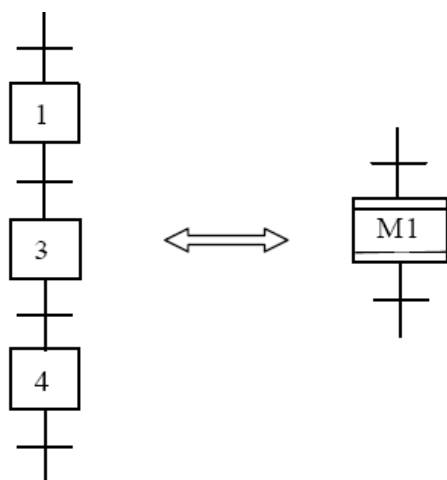


Figura 3.24: Representación de macroetapas.

Reglas de evolución del Grafcet

Las siguientes reglas de evolución permiten interpretar cómo evolucionará el sistema modelado por un Grafcet conforme se vayan activando las diversas

señales que intervienen.

Inicialización

En la inicialización del sistema se tienen que activar todas las etapas iniciales (marcadas con doble cuadrado) y solo las iniciales.

Evolución de las transiciones

Una transición está validada si todas las etapas inmediatamente anteriores están activas. Si además de estar validada, la receptividad asociada es cierta, se dice que la transición es franqueable. Una transición franqueable es disparada (franqueada) inmediata y obligatoriamente.

Evolución de las etapas

Cuando se dispara (o franquea) una transición, todas las etapas inmediatamente anteriores se desactivan y simultáneamente se activan todas las posteriores.

Simultaneidad en el disparo

Si dos transiciones se disparan al mismo tiempo, las activaciones y desactivaciones de etapas se producen de forma simultánea.

Prioridad de la activación

Si en el disparo de una transición, una etapa debe activarse y desactivarse a la vez, quedará activa. En el ejemplo de la figura 3.25, si la receptividad *b* es cierta, hay que volver a la etapa 2. Cuando eso ocurre, se debería desactivar y activar la etapa 2 simultáneamente. Si no se cumpliera la regla de “prioridad de la activación”, el GRAFCET se quedaría sin ninguna etapa activa.

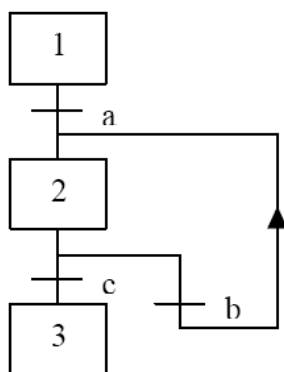


Figura 3.25: Ejemplo de prioridad de la activación.

Estados estables e inestables

Un estado de un Grafcet es estable cuando no cambia ninguna etapa mientras no haya cambios en las entradas. Un estado es inestable si se produce alguna evolución de etapas sin que cambie ninguna entrada.

Los estados inestables duran muy poco tiempo, pues rápidamente pasan al estado siguiente.

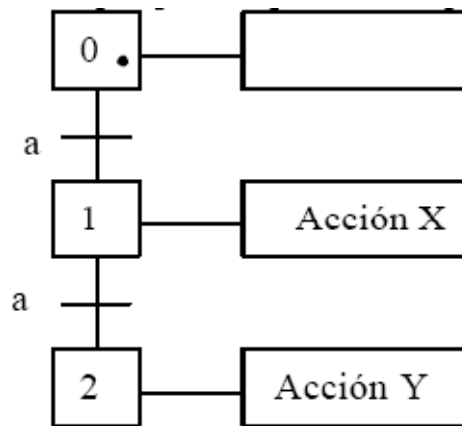


Figura 3.26: Ejemplo de estados inestables.

Por ejemplo, en el Grafcet representado en la figura 3.26, el estado en que solo la etapa 0 está activa es estable, pues mientras $a = 0$ no hay evolución. Cuando $a = 1$, se produce la evolución desactivándose la etapa 0 y activándose la etapa 1. Ese estado es ahora inestable, pues aunque no haya cambios en las entradas se debe producir la evolución de forma que se desactive la etapa 1 y se active la etapa 2. El Grafcet, por tanto, debe evolucionar siempre desde un estado estable hasta otro estado estable.

En teoría, el sistema está en la etapa 1 durante un pulso de duración muy pequeña. El dilema aparece con las acciones asociadas a esa etapa. **Si son acciones de nivel no deberían activarse.** En cambio, **si son acciones impulsionales sí deben activarse**, pues la etapa 1 se activa durante un pulso que es suficiente para producir el efecto de esas acciones. Por ejemplo, si la acción X consiste en hacer el SET de una variable, ese SET debe producirse. Si la acción X es de nivel (poner en marcha un motor mientras la etapa 1 esté activa) no tiene sentido que se active durante un pulso muy pequeño (idealmente casi nulo), por lo que no habría que activarla.

Transiciones activas por flanco

Por otra parte, hay que tener en cuenta las situaciones especiales que pueden suceder cuando existen transiciones activas por flanco.

Por ejemplo, en el Grafcet representado en la figura 3.27, cuando se produce un flanco de subida en a , el Grafcet debe evolucionar, activándose la etapa 1 y desactivándose la etapa 0. La siguiente transición no está activa, pues el flanco

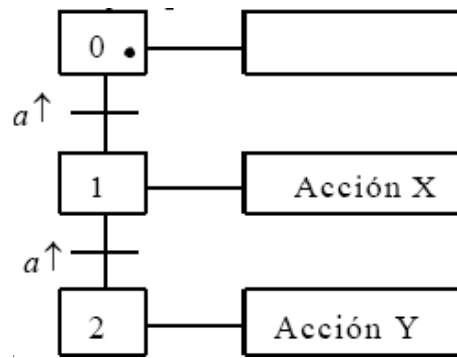


Figura 3.27: Ejemplo de transiciones activas por flanco.

ya se ha producido. Es decir, un flanco sólo puede hacer avanzar una etapa. La etapa 1 es en este caso, estable. Para evolucionar a la etapa 2, la variable a debe bajar a 0 y volver a subir a 1. Si hay problemas en la programación de las transiciones por flanco, se pueden convertir a transiciones por nivel. El ejemplo anterior podría transformarse al Grafcet de la figura 3.28.

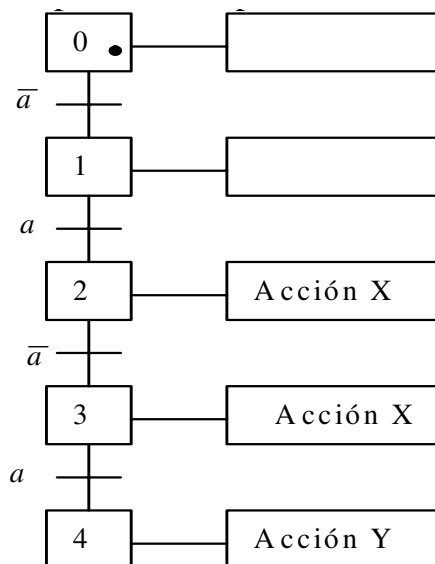


Figura 3.28: Grafcet equivalente al de la figura 3.27.

3.5 EJEMPLOS

EJEMPLO 1

Sea el cilindro de la figura 3.29.

Inicialmente se supone que el cilindro está en la posición más a la izquierda. Al pulsar P se debe mover hacia la derecha hasta que se active el detector C, para volver después a la izquierda. Si se pulsa Q debe hacer lo mismo, pero

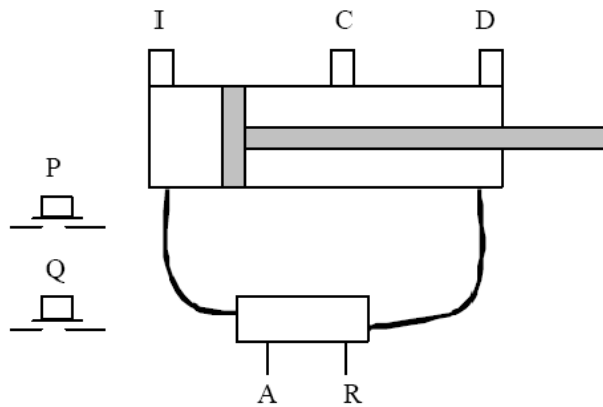


Figura 3.29: Sistema del ejemplo 1.

esperando 1 segundo antes de empezar a moverse, y llegando a D en lugar de C. Para mover el cilindro se actúa sobre las electroválvulas A y R.

Una posible solución es la representada en la figura 3.30.

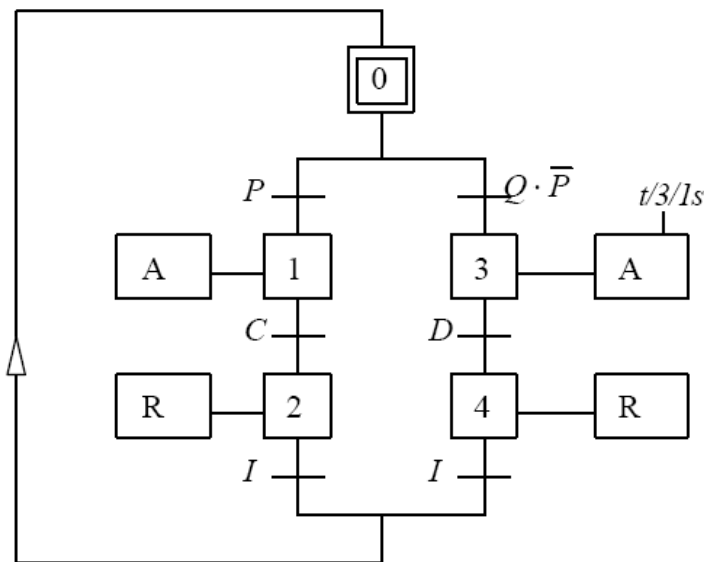


Figura 3.30: Grafcet que resuelve el ejemplo 1.

EJEMPLO 2

Sea el disco de la figura 3.31.

El disco tiene una placa metálica que es detectada por los detectores inductivos C y D. Inicialmente se supone el disco en el punto C. Cuando se pulsa P el disco empieza a girar. Si el interruptor I está desactivado, el disco debe girar hasta que llegue a C. Si el interruptor I está activado cuando la placa pasa por D, el disco debe parar hasta que se desactive I. En cualquier caso,

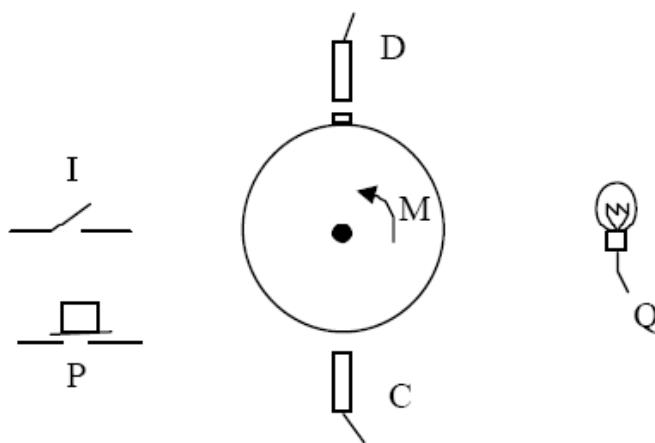


Figura 3.31: Sistema del ejemplo 2.

además, cada vez que se activa D, se debe incrementar un contador de vueltas y activar una salida Q (luz) durante 200 ms.

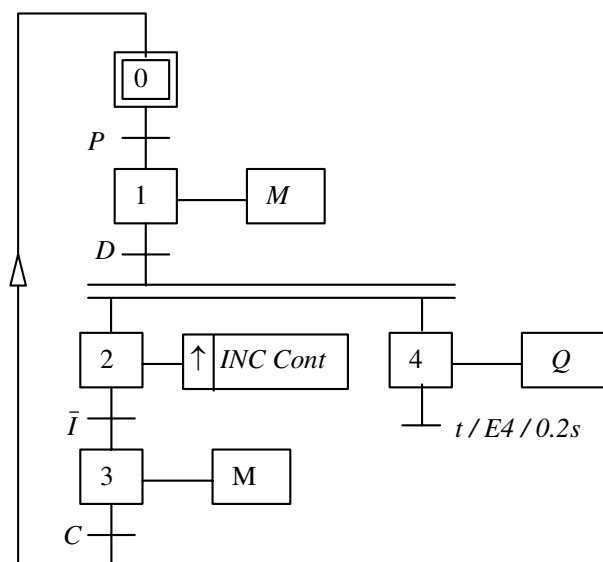


Figura 3.32: Grafcet que resuelve el ejemplo 2.

Una posible solución es la representada en la figura 3.32. En ella, se puede observar que si I no está activado, la etapa 2 es inestable, por lo que se pasa de la 1 a la 3. Eso significa que la variable M (de nivel) debe permanecer activa (no se debe desactivar ni siquiera en un pequeño pulso).

Por otra parte, la acción impulsional de incrementar el contador debe activarse, pues se pasa por la etapa 2.

EJEMPLO 3

Considérese el problema de poner en marcha y apagar un motor con el mismo pulsador. Cuando se pulsa P (en el flanco de subida) el motor se pone en marcha. Cuando se vuelve a pulsar P (el siguiente flanco de subida) el motor debe pararse.

Una posible solución (la más simple) sería el Grafcet de la figura 3.33.

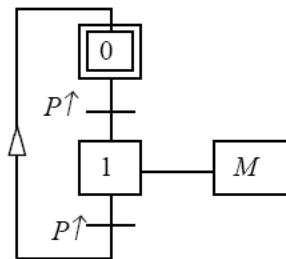


Figura 3.33: Grafcet que resuelve el ejemplo 3.

EJEMPLO 4

Considérese el sistema de la figura 3.34.

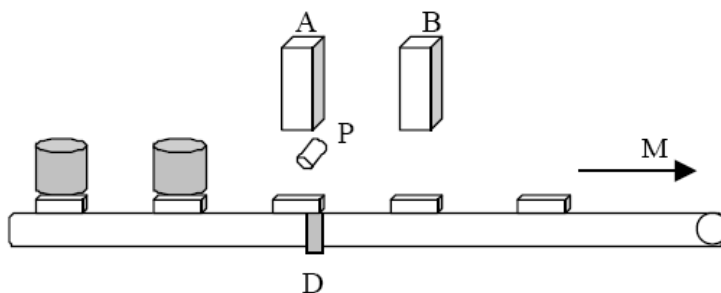


Figura 3.34: Sistema del ejemplo 4.

Por una cinta accionada por el motor M circulan palets espaciados regularmente, de forma que cuando hay un palet en A, hay otro en B. Sobre estos palets puede o no haber un producto sobre el que hay que realizar dos operaciones A y B (podrían ser de llenado, taponado, etiquetado, control, etc). El detector inductivo D detecta el paso de los palets, mientras que la fotocélula P detecta la presencia de producto (si hay producto se activa antes que D). Supondremos, en este caso, que las operaciones A y B están temporizadas (duran 1 y 2 segundos respectivamente).

Una posible solución al problema queda expuesta en la figura 3.35.

En este caso, si vienen dos productos consecutivos se activan dos etapas de la secuencia, realizándose las dos operaciones. Si vinieran dos botes consecutivos y después ninguno, la evolución de las etapas se representa en la figura 3.36.

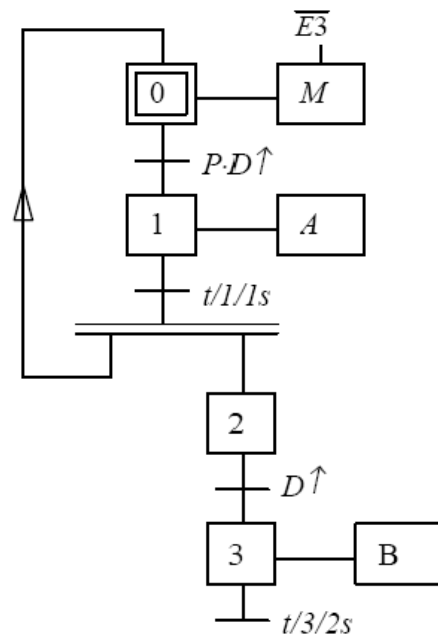


Figura 3.35: Grafcet que resuelve el ejemplo 4. Solución A.

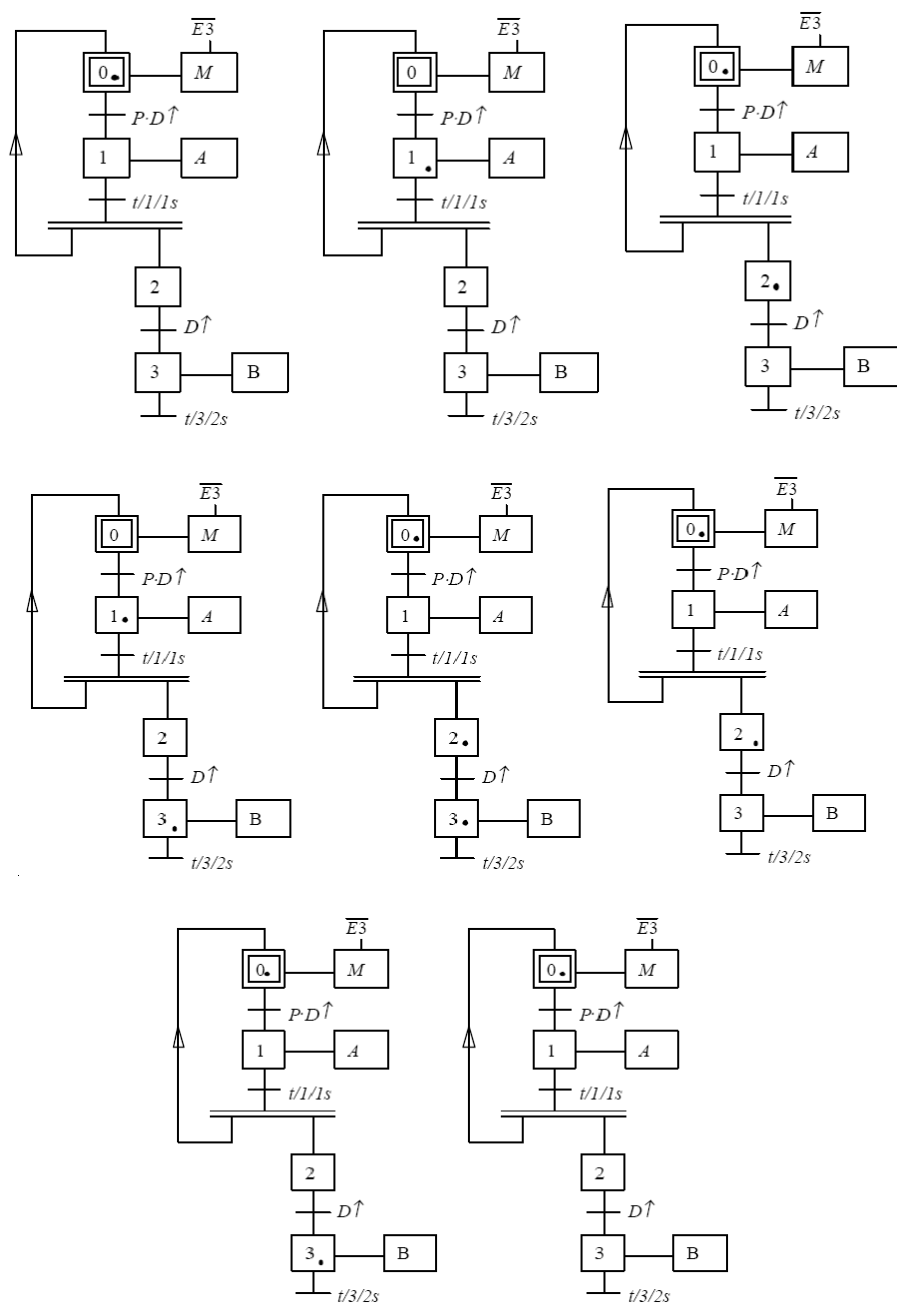


Figura 3.36: Evolución de las etapas en el Grafset de la solución A del ejemplo 4.

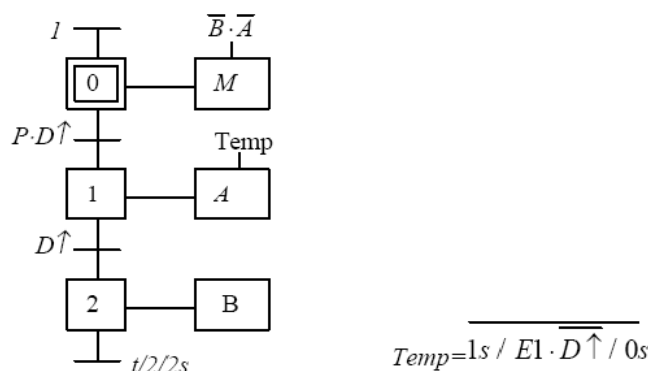


Figura 3.37: Grafcet que resuelve el ejemplo 4. Solución B.

Otra solución podría ser la representada en la figura 3.37. En este caso, la etapa 0 está siempre activa. El temporizador se ha expresado según la norma IEC848. Aquí, no basta con poner la etapa 1, pues cuando pasan dos botes seguidos ésta no se desactiva, y por lo tanto el temporizador no se resetearía. La evolución del Grafcet si llegan dos botes consecutivos y después ninguno está representada en la figura 3.38.

3.6 DISEÑO ESTRUCTURADO DE AUTOMATISMOS

Los ejemplos vistos en secciones anteriores se han resuelto por medio de un solo diagrama Grafcet. Cuando se trata de problemas más complejos es normal utilizar varios diagramas (Grafkets parciales), que pueden funcionar de forma independiente o de forma coordinada.

El ejemplo 4 se puede resolver con los dos Grafkets representados en la figura 3.39. En este caso, la coordinación entre los dos Grafkets se produce por la primera transición de G2, que es función de E1 (cuando ha habido un bote en el primer proceso E1 se ha activado, y como consecuencia también E2, por lo que el segundo Grafcet hace que en B se procese el bote).

3.7 FORZADO DE GRAFCETS

En el diseño de un automatismo pueden utilizarse varios Grafkets parciales. Desde uno de ellos se puede forzar a activar o desactivar una etapa de otro. Esta acción especial se representa dentro de un rectángulo de líneas discontinuas de la forma: F/G3:4,7 (se fuerzan las etapas 4 y 7 del Grafcet G3). El Grafcet G3 permanecerá con esas etapas activas (y el resto inactivas) mientras el Grafcet principal tenga activa la etapa que produce el forzado. Es decir, el Grafcet 3 no podrá evolucionar mientras esté activa la etapa 2.

Hay dos variantes de esta orden: F/G2: (se desactivan todas las etapas de G2) y F/G5:* (el Grafcet G5 permanece congelado en las etapas que tuviera activas).

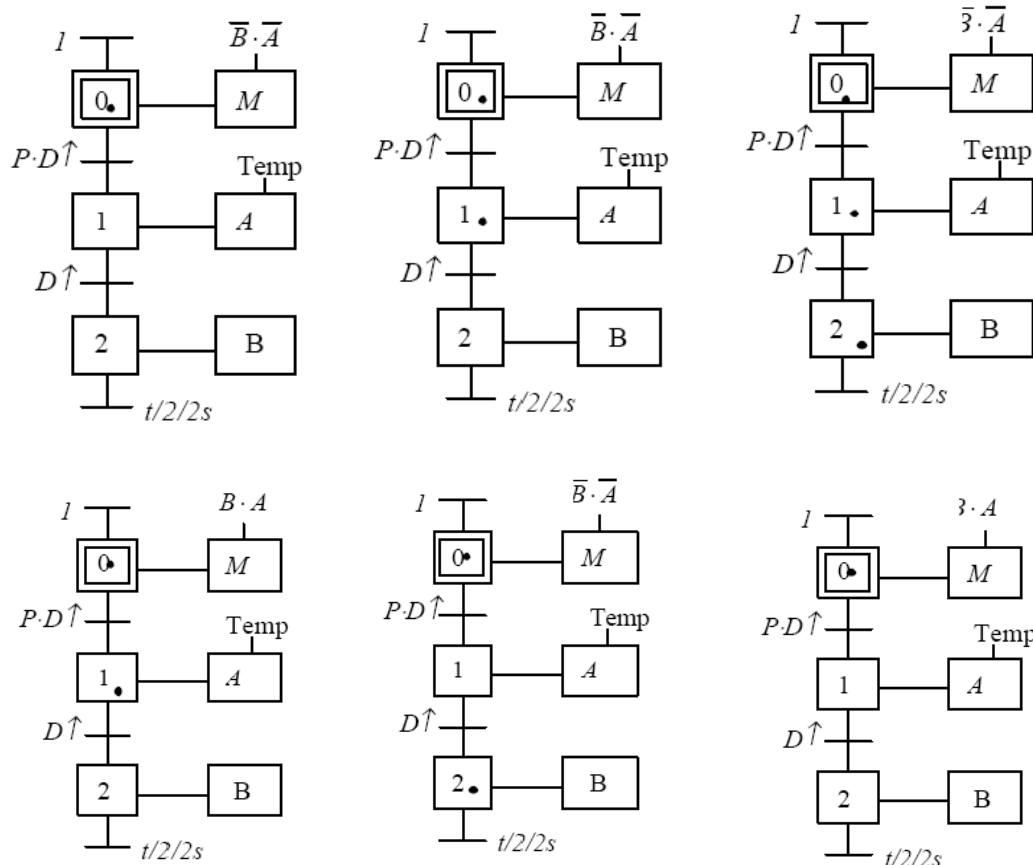


Figura 3.38: Evolución de las etapas en el Grafcet de la solución B del ejemplo 4.

La acción especial de forzado de otro Grafcet también puede ser impulsional, indicándose con una flecha tal y como se puede ver en la figura 3.41.

En este caso, cuando se activa la etapa 2, se activan las etapas 4 y 7 del Grafcet 3, pudiendo después evolucionar libremente este Grafcet aunque permanezca activa la etapa 2.

3.8 JERARQUÍA ENTRE GRAFCETS PARCIALES

Cuando se utilizan varios Grafcets, puede definirse una jerarquía entre ellos. Esta jerarquía puede establecerse por medio de las órdenes de forzado: el Grafcet jerárquicamente superior fuerza al Grafcet inferior.

El ejemplo anterior se podría resolver tal y como se indica en el figura 3.42. En este caso, G1 es el Grafcet principal (o maestro), que fuerza la activación de la etapa 2 del Grafcet subordinado cuando se tiene que procesar un bote en el proceso B. En este ejemplo, la utilización de Grafcets jerárquicos complica el problema, por lo que no tendría mucho sentido hacerlo así.

La acción de forzado puede ser activa por nivel o por flanco (impulsional).

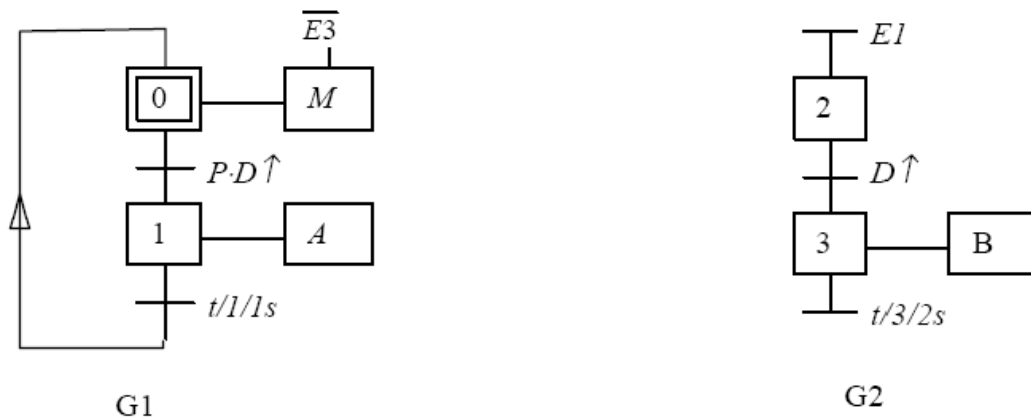


Figura 3.39: Solución C del ejemplo 4 utilizando varios Grafcets.

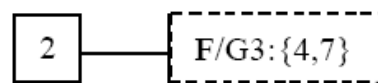


Figura 3.40: Representación del forzado de Grafcets.

El forzado impulsional se representará igual, pero con una flecha después de la F. Por ejemplo: $F\uparrow/G2:\{2\}$ significa que cuando se activa la etapa correspondiente del Grafcet G1, se produce el forzado de la etapa 2 del Grafcet G2, pudiendo después evolucionar de forma libre este Grafcet, aunque el primero no cambie de etapa. El forzado por nivel, sin embargo, mantiene forzado el Grafcet subordinado mientras no se desactive la etapa. El forzado por flanco se puede siempre sustituir por un forzado por nivel y una etapa auxiliar con transición de salida siempre activa, tal y como se muestra en el ejemplo 5.

El forzado de Grafcets cumple las siguientes reglas:

1. El forzado es prioritario respecto de cualquier otra evolución (sea de activación o de desactivación).
2. Las reglas de evolución del Grafcet no se aplican a un Grafcet mientras éste permanezca forzado. En el ejemplo anterior, mientras el Grafcet G1 permanezca en la etapa 4, el Grafcet G2 permanecerá con la etapa 2 activa y la 3 inactiva, independientemente de las entradas.

Cuando se utilizan varios Grafcets jerarquizados con acciones de forzado

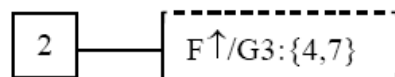


Figura 3.41: Representación del forzado impulsional.

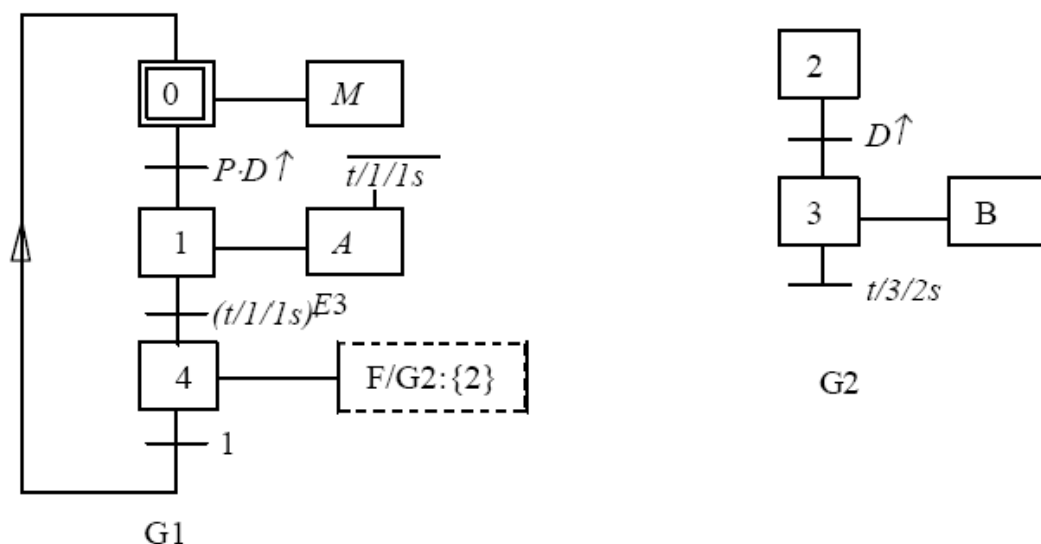


Figura 3.42: Solución D del ejemplo 4.

entre ellos, hay que aplicar las siguientes reglas de jerarquía:

1. Si un Grafcet puede forzar a otro, éste no puede forzar al primero.
2. Un Grafcet no puede ser forzado en ningún instante por más de un Grafcet a la vez.
3. No debe haber bucles en el forzado (G1 fuerza a G2, G2 fuerza a G3 y G3 fuerza a G1), pues podría dar lugar a un funcionamiento inestable.

Se puede utilizar la jerarquía entre Grafcets para tener en cuenta los diferentes modos de funcionamiento del proceso (manual, automático, ciclo a ciclo, test, procedimiento de arranque, procedimiento de parada, etc.) y las emergencias. Se puede utilizar, por ejemplo, un Grafcet maestro que controla el modo en el que se encuentra el proceso, y varios Grafcets subordinados que implementan el control del proceso en esos modos.

También es bastante habitual utilizar un Grafcet subordinado para cada subproceso cuando el sistema tiene subprocesos bastante delimitados. La relación entre estos Grafcets parciales se suele implementar por la utilización de variables de un Grafcet en las transiciones del otro.

La utilización de macroetapas también contribuye a estructurar los automatismos y hacerlos más fáciles de entender y verificar. Una macroetapa no es un Grafcet diferente, sino una forma condensada de representar una secuencia dentro de un Grafcet.

EJEMPLO 5

Considérese el ejemplo 4. Se desea añadir las siguientes funciones:

- Habrá un pulsador de MARCHA y otro de PARADA. Al iniciarse el sistema todo estará desactivado hasta que se pulse el pulsador de marcha. Cuando se pulse el pulsador de parada el sistema acabará de procesar los elementos que lleguen hasta que se active D sin que haya pieza ($P = 0$) y sin que quede producto en B que procesar. En ese momento se pasará a la situación de inicio.
- Además, habrá un pulsador de EMERGENCIA (pulsador con enclavamiento). Este pulsador produce un corte de la alimentación del motor de la cinta y de los procesos A y B, independientemente del automatismo. Cuando se accione el pulsador de emergencia hay que desactivar los procesos A y B y la cinta (aunque ya estén de hecho apagados por falta de alimentación) y poner en marcha una sirena. Esa sirena se desactivará cuando se accione un pulsador de parada de sirena. Para rearmar el sistema (después de resolver manualmente el problema que hubiera) se deberá rearmar el pulsador de emergencia (ponerlo a off) y después presionar el pulsador de marcha.
- Además, habrá un pulsador de parada de cinta temporal, y un pulsador de reinicio de cinta, que se utilizarán para parar la cinta en cualquier momento.

Una solución es la representada en las figuras 3.43.

El Grafcet maestro es el que controla los modos de marcha y parada y la emergencia. Otra posibilidad es evitar que el Grafcet maestro utilice forzados por flanco (lo sustituye por un forzado por nivel y una etapa auxiliar) esta solución está representada en la figura 3.44.

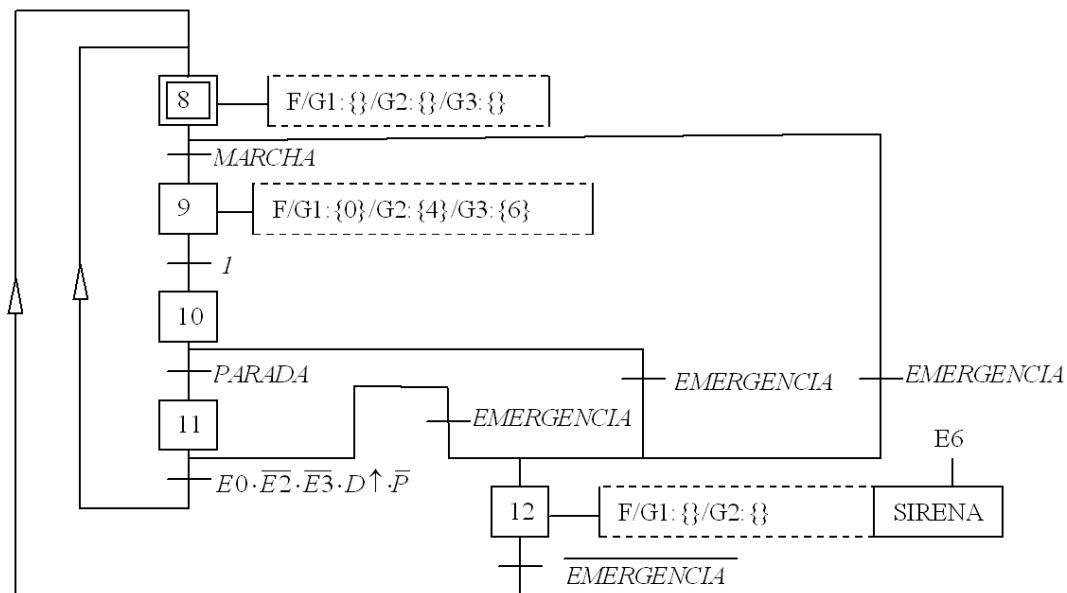


Figura 3.44: Solución del ejemplo 5 con una modificación del Grafset maestro.

CAPÍTULO 4

GUÍA PARA EL ESTUDIO DE LOS MODOS DE MARCHA Y PARO: GEMMA

4.1 INTRODUCCIÓN

En un proceso productivo una máquina no está siempre funcionando en modo automático, pueden surgir problemas que, por ejemplo, conlleven a una parada inmediata.

En la automatización de una máquina es necesario prever todos los modos posibles: funcionamiento manual o semiautomático, paradas de emergencia, puesta en marcha. Además, el propio automatismo debe ser capaz de detectar defectos en la parte operativa y colaborar con el operario o técnico de mantenimiento para su puesta en marcha y reparación.

La guía GEMMA se trata de una representación organizada de todos los modos de Marcha y Parada en que se puede encontrar una máquina o proceso de producción automatizado, y orienta sobre los saltos o transiciones que pueden darse de un modo de funcionamiento a otro.

Uno de los objetivos principales de la guía GEMMA es la utilización de una metodología sistemática y estructurada que ofrezca información precisa del sistema en clave de estados posibles; de ahí que habitualmente se presente en el formato de la descripción completa de todos los estados posibles que puede llegar a tener el sistema.

4.2 DESCRIPCIÓN DE LA GUÍA GEMMA

Un automatismo consta de dos partes fundamentales: el sistema de producción y el control de este sistema (ordenador, autómatas programables, etc.). El control puede estar alimentado o sin alimentar; desde el punto de vista de la automatización, el estado sin alimentar no es interesante pero sí el paso de este estado al de control alimentado.

Cuando el control está alimentado, una máquina o proceso automatizado puede encontrarse en tres situaciones (en las cuales puede estar o no produciendo):

- En funcionamiento normal (por lo tanto en producción).

- Parado, o en proceso de parada.
- En defecto, situación en la que o bien no está produciendo, o bien el producto no es aprovechable o sólo lo es si se manipula adecuadamente a posteriori.

La guía GEMMA muestra estas tres situaciones (funcionamiento, parada y defecto) con rectángulos. Además, un quinto rectángulo, que indica que el sistema productivo está en producción, figura 4.1.

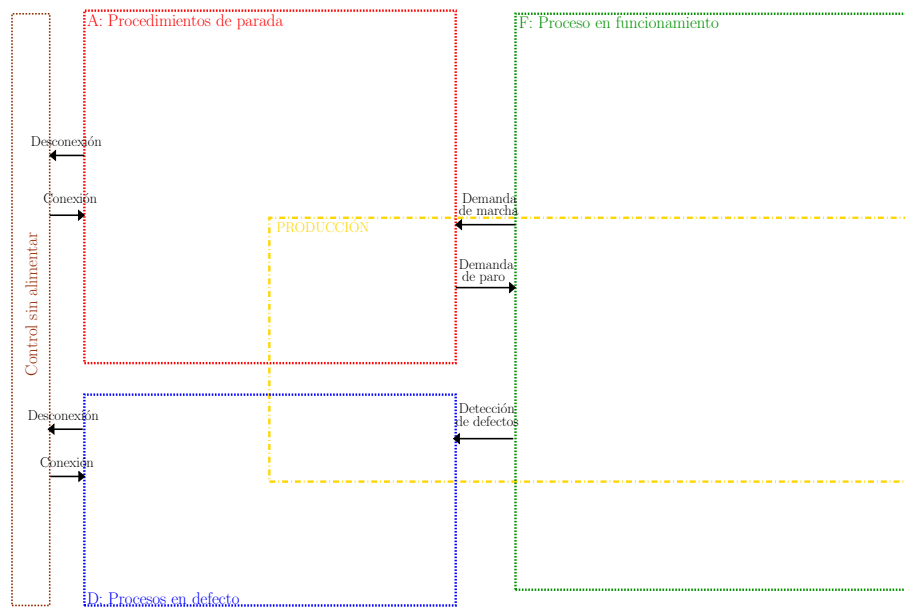


Figura 4.1: Representación de los modos de funcionamiento en la Guía GEMMA.

Cada una de estas situaciones se subdivide de forma que al final la guía GEMMA presenta 16 estados de funcionamiento posible, los cuales se describen a continuación.

Grupo F. Procedimientos de funcionamiento

- F1** - Producción normal. Estado en que la máquina produce normalmente. Es el estado más importante y en él se deben realizar las tareas por las cuales la máquina ha sido construida.
- F2** - Marcha de preparación. Son las acciones necesarias para que la máquina entre en producción (precalentamiento, preparación de componentes...).
- F3** - Marcha de cierre. Corresponde a la fase de vaciado y/o limpieza que en muchas máquinas debe llevarse a cabo antes de la parada o del cambio de algunas de las características del producto.

- F4** - Marchas de verificación sin orden. En este caso, la máquina, normalmente por orden del operario, puede realizar cualquier movimiento o unos determinados movimientos preestablecidos. Es el denominado control manual y se utiliza para funciones de mantenimiento y verificación.
- F5** - Marchas de verificación con orden. En este caso, la máquina realiza el ciclo completo de funcionamiento en orden pero al ritmo fijado por el operador. Se utiliza también para tareas de mantenimiento y verificación. En este estado la máquina puede estar en producción. En general, se asocia al control semiautomático.
- F6** - Marchas de test. Sirve para realizar operaciones de ajuste y mantenimiento preventivo, por ejemplo: comprobar si la activación de los sensores se realiza en un tiempo máximo, curvas de comportamiento de algunos actuadores, etc.

Grupo A. Procedimiento de paradas y puestas en marcha

- A1** - Paradas en el estado inicial. Se corresponde con el estado de reposo de la máquina. La máquina normalmente se representa en este estado en los planos de construcción y en los esquemas eléctricos.
- A2** - Parada solicitada al final del ciclo. Es un estado transitorio en el que la máquina, que hasta el momento estaba produciendo normalmente, debe producir solo hasta acabar el ciclo y pasar a estar parada en el estado inicial.
- A3** - Parada solicitada en un estado determinado. Es un estado en el que la máquina se detiene en un estado determinado que no coincide con el final de ciclo. Es un estado transitorio de evolución hacia A4.
- A4** - Parada obtenida. Es un estado de reposo de la máquina distinto al estado inicial.
- A5** - Preparación para la puesta en marcha después de un defecto. Es en este estado donde se procede a todas las operaciones, de vaciado, limpieza, reposición de un determinado producto, etc., necesarias para la puesta de nuevo en funcionamiento de la máquina después de un defecto.
- A6** - Puesta del sistema en el estado inicial. En este estado se realiza el retorno del sistema al estado inicial (reinicio). El retorno puede ser manual (coincidiendo con F4) o automático.
- A7** - Puesta del sistema en un estado determinado. Se retorna el sistema a una posición distinta de la inicial para su puesta en marcha, puede ser también manual o automático.

Grupo D. Procedimientos de defecto

- D1** - Parada de emergencia. Es el estado que se consigue después de una parada de emergencia, en donde debe tenerse en cuenta tanto las paradas como los procedimientos y precauciones necesarias para evitar o limitar las consecuencias debidas a defectos.
- D2** - Diagnóstico y/o tratamiento de fallos. Es en este estado cuando la máquina puede ser examinada después de un defecto y, con o sin ayuda del operador, indicar los motivos del fallo para su rearme.
- D3** - Producción a pesar de los defectos. Corresponde a aquellos casos en que se deba continuar produciendo a pesar de los defectos. Se incluye en estas condiciones casos en que, por ejemplo, sea necesario finalizar un reactivo no almacenable, en que se pueda substituir transitoriamente el trabajo de la máquina por la de un operario hasta la reparación de la avería.

La figura 4.2 muestra la Guía GEMMA incluyendo los estados en cada uno de los modos de funcionamiento.

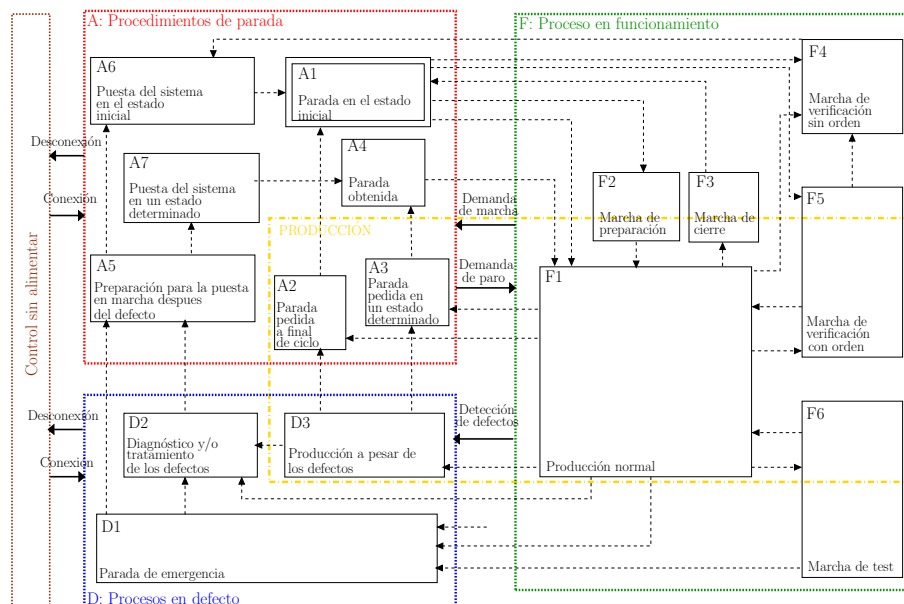


Figura 4.2: Representación de la guía GEMMA, incluyendo los 16 estados.

4.3 UTILIZACIÓN DE LA GUÍA GEMMA

Como ya se ha indicado, la guía GEMMA es un gráfico de soporte al diseñador de automatismos. El procedimiento a seguir en su utilización es el siguiente:

1. Estudiar los estados necesarios de la máquina a automatizar, anotando en cada uno de los rectángulos la descripción correspondiente y posibles variantes, si las hay. Aquellos estados que no serán utilizados se marcan con una cruz, indicando así que no se han considerado.
2. Estudiar entre cuáles de los estados será posible la evolución. La guía permite mostrar de forma gráfica todos los caminos que se quiere permitir, marcando éstos con una línea continua.
3. Finalmente, de forma parecida a como se indican las transiciones en GRAFCET, se marcan las condiciones necesarias para poder seguir un determinado camino. En algunas ocasiones, un determinado camino no tiene una condición específica o determinada, en ese caso puede no ponerse indicación o bien utilizar la condición de que la acción anterior se haya completado.

Una vez finalizado el paso 3, la guía Gemma resultante se traduce a un diagrama de Grafcet de un nivel jerárquico superior, el cual estará formado por diversas etapas, de forma que una única etapa debe estar activa cada vez. Cada etapa actuará mediante órdenes de forzado en diagramas de Grafcet localizados en un nivel jerárquico inferior que implementan los detalles de los distintos modos de Funcionamiento, Parada y Defecto.

A continuación veremos, simplificadaamente, algunos de los casos más corrientes.

Marcha por ciclos y parada a fin de ciclo

El sistema está parado en el estado inicial (A1). Cuando las condiciones de puesta en marcha se verifican (modo de marcha, pulsador de arranque, etc.), se pasa a funcionar en modo normal (F1). Cuando el operador pulsa el pulsador de parada a fin de ciclo, la máquina pasará al estado de parada a fin de ciclo (A2) y, cuando acabe el ciclo, pasará al estado inicial (A1).

Fijémonos en que el paso de A2 a A1 es directo al acabarse el ciclo, pero hemos querido indicarlo (condición "Fin ciclo") para una mayor claridad. Si se selecciona el modo de funcionamiento ciclo a ciclo, el paso de F1 a A2 es directo inmediatamente después de comenzar el ciclo y no necesita la actuación sobre ningún pulsador. El modo ciclo a ciclo puede ser con antirrepetición, en cuyo caso el paso de A2 a A1 sólo se puede hacer en el caso de que el pulsador de arranque no esté pulsado; de esta forma, se garantiza que el operador presiona el pulsador cada vez que ha de comenzar un ciclo y que, por tanto, el ciclo no puede recomenzar en caso de que el pulsador esté encallado.

Marcha de verificación con orden

En este caso, la máquina puede pasar a funcionar en este modo (F5) cuando está parada (A1) o cuando está en producción normal (F1) si se selecciona el

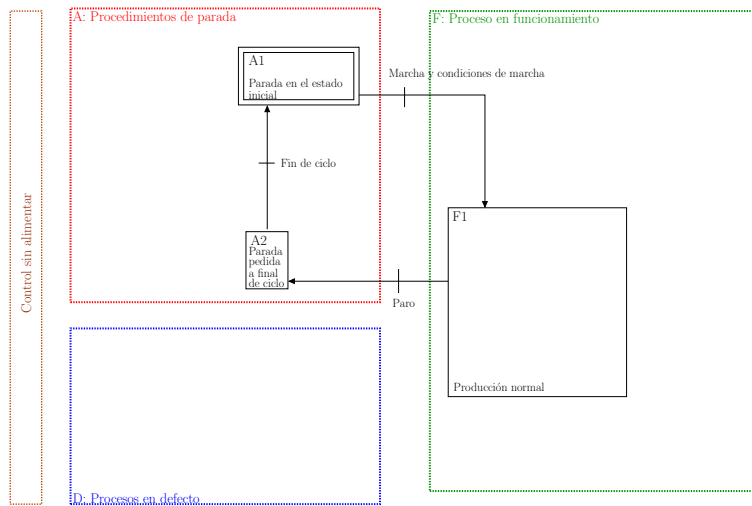


Figura 4.3: Representación en la guía GEMMA de la marcha por ciclos

modo etapa a etapa.

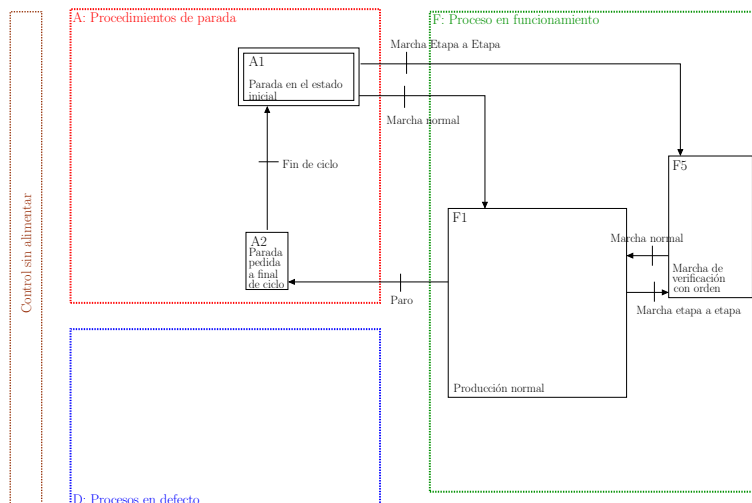


Figura 4.4: Representación en la guía GEMMA de la marcha de verificación con orden.

Mientras la máquina funcione etapa a etapa, será necesario accionar un pulsador para pasar de una etapa a la siguiente. Seleccionando el modo normal, la máquina pasará al estado de producción normal (F1). Si se selecciona el modo normal cuando la máquina está en la última etapa y se presiona el pulsador de parada, la máquina se parará (A2 seguido de A1).

Marcha de verificación sin orden

Se puede pasar al modo de verificación sin orden (conocido habitualmente como funcionamiento manual) tanto desde el estado inicial (A1) como desde el funcionamiento normal (F1).

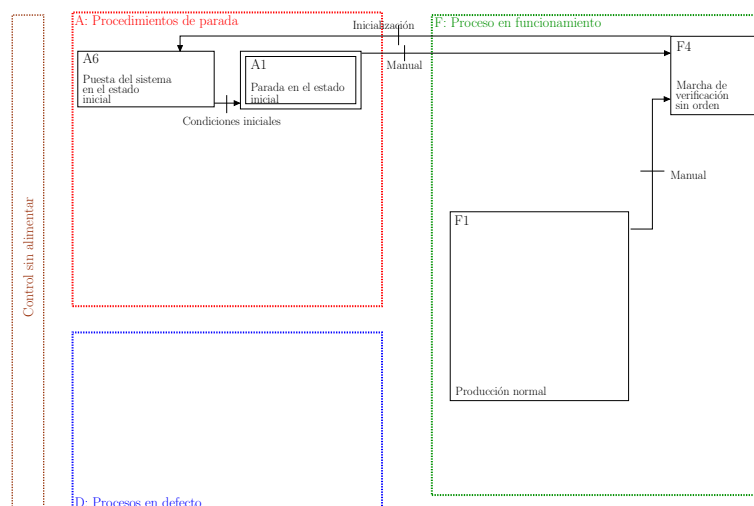


Figura 4.5: Representación en la guía GEMMA de la marcha de verificación sin orden.

Allí el operador puede realizar todos los movimientos por separado y en un orden cualquiera (en algunas instalaciones sólo son posibles algunos movimientos en modo manual). En algunos casos, el operador tiene mandos adecuados en el panel para ordenar los movimientos deseados mientras que en otros, hay que actuar directamente en los mandos locales de los preaccionadores. Apresando el pulsador de inicialización se pasa a poner el sistema en el estado inicial (A6) y, una vez alcanzado, se pasa al estado inicial (A1).

Paradas de emergencia

El sistema está funcionando normalmente (F1) y se acciona el pulsador de parada de emergencia. Esto, en los sistemas habituales, implica normalmente dejar sin alimentación (físicamente, sin intervención del sistema de control) todo el sistema de producción que, por diseño, quedará en posición segura al quedarse sin alimentación. El mismo pulsador de parada de emergencia informa al control de que pasará al estado de parada de emergencia (D1). Al desenclavar el pulsador de emergencia se pasa a preparar la puesta en marcha (A5). En este caso, hay dos posibilidades de uso habitual según el tipo de sistema que se está controlando. En el primer caso, se lleva al sistema hasta el estado inicial (A6), lo que a menudo requiere la intervención del operador y, una vez alcanzado (A1), el sistema espera una nueva puesta en marcha presionando el pulsador de marcha que hará recomenzar el proceso de producción (F1).

La segunda posibilidad consiste en llevar al sistema hasta a un estado determinado (A7), lo que a menudo requiere la intervención del operador y, una vez alcanzado (A4), el sistema espera la nueva puesta en funcionamiento cuando el operador accione el pulsador de marcha que hará continuar el proceso (F1) a partir de la etapa alcanzada.

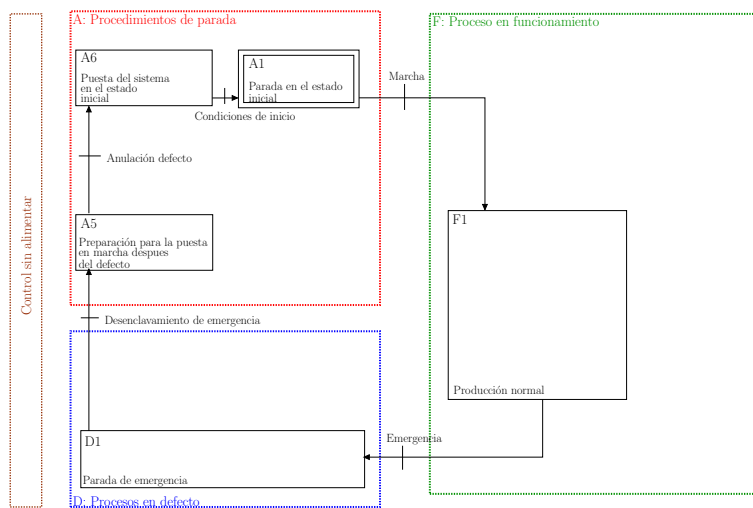


Figura 4.6: Representación en la guía GEMMA de paradas de emergencia.

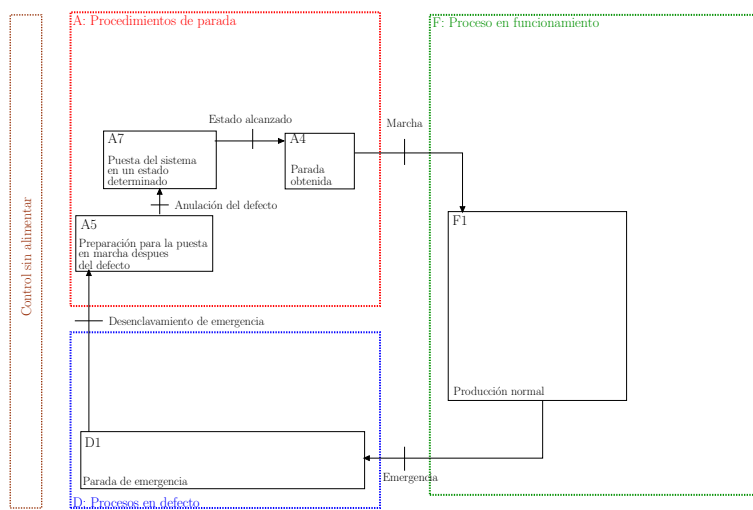


Figura 4.7: Representación en la guía GEMMA de paradas de emergencia.

Parada en un punto

El sistema está funcionando en producción normal (F1) y el operador acciona el pulsador de parada; entonces se pasa a la situación de parada pedida (A3) y, una vez alcanzado el punto deseado, el sistema se para (A4).

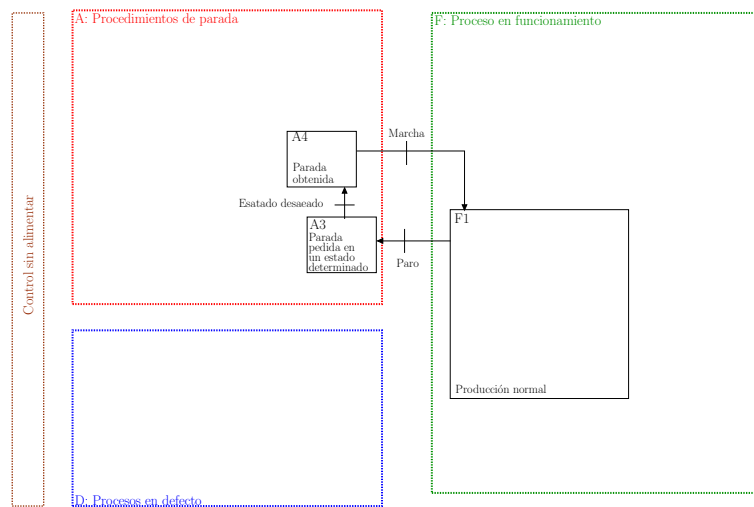


Figura 4.8: Representación en la guía GEMMA de la parada en un punto.

Se debe presionar el pulsador de arranque para que el sistema siga funcionando (F1) a partir del punto de parada.

4.4 EJEMPLOS

Ejemplo 1

La guía Gemma para el EJEMPLO 5 del capítulo anterior se muestra en la figura 4.9.

Ejemplo 2

En el ejemplo 5 del tema anterior considere que fallo_A y fallo_B son señales que indican que ha ocurrido un fallo en los procesos A y B respectivamente. Cuando se activa una de estas dos señales, se debe desactivar el proceso correspondiente y el producto sobre el cual ocurrió el fallo debe avanzar hasta la posición siguiente al proceso B para que pueda ser retirado. Además, se deberá activar una señal de alarma que indique al operador el fallo ocurrido. El operador accionará el pulsador de Emergencia en primer lugar, y después retirará el producto y reparará el fallo. Después, el sistema se pondrá nuevamente en funcionamiento al desactivar el pulsador de Emergencia.

La guía Gemma correspondiente se muestra en la figura 4.10. El conjunto de diagramas de grafcet que resuelven el automatismo anterior se muestra en la figura 4.11.

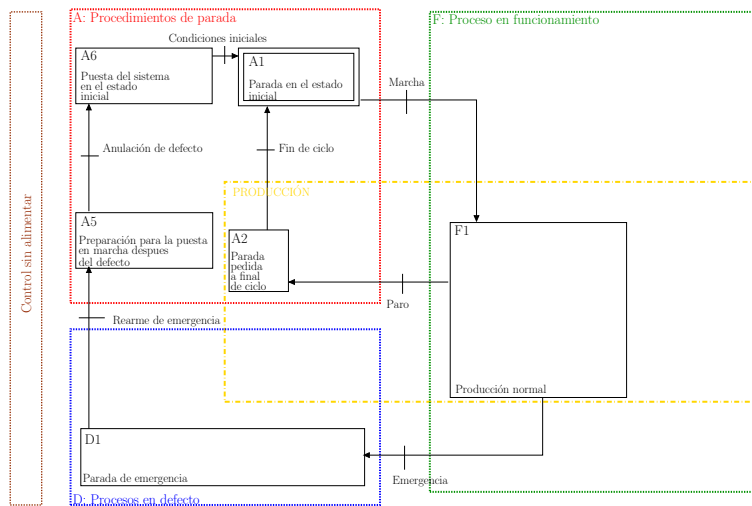


Figura 4.9: Representación del EJEMPLO 5 en la guía GEMMA: un automatismo con marcha por ciclo y parada por emergencia.

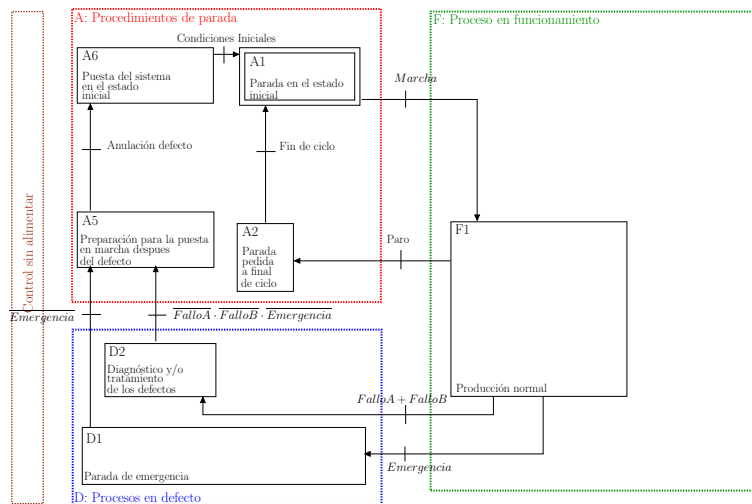


Figura 4.10: Representación del Ejemplo 2 en la guía GEMMA: un automatismo con marcha por ciclo y parada por emergencia, con y sin diagnóstico de fallos.

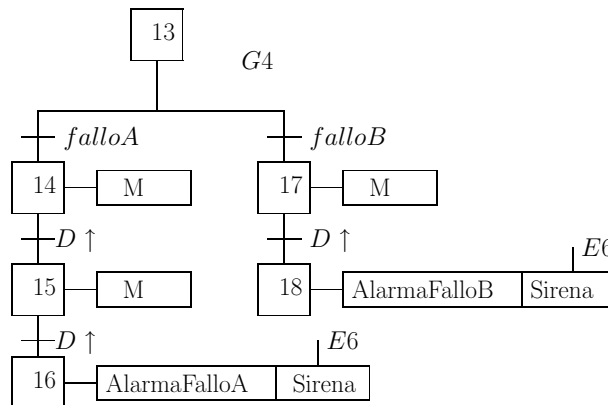
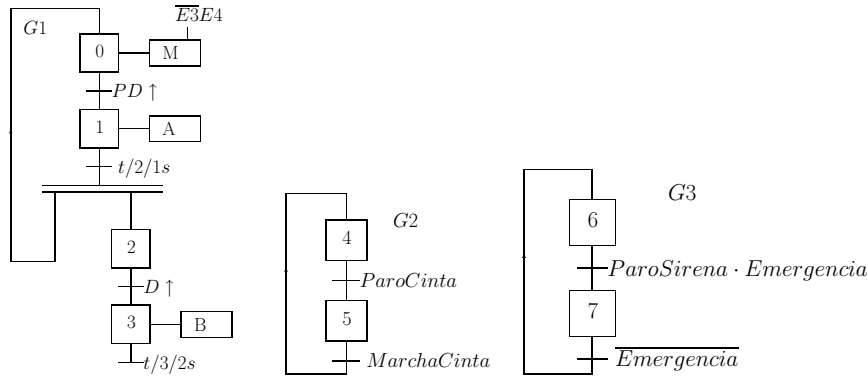
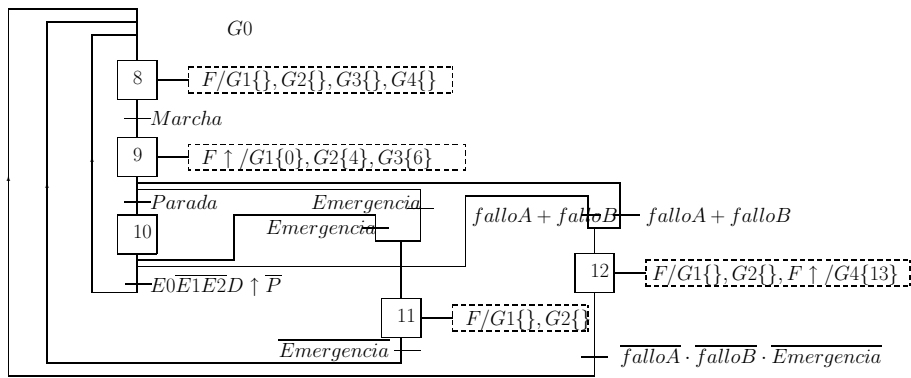


Figura 4.11: Diagrama de Grafset para el Ejemplo 2: un automatismo con marcha por ciclo y parada por emergencia, con y sin diagnóstico de fallos.

CAPÍTULO 5

EL AUTÓMATA PROGRAMABLE INDUSTRIAL

5.1 DEFINICIÓN

Un autómata programable industrial (API) es un equipo electrónico, programable en lenguaje no informático, diseñado para controlar en tiempo real y en ambiente de tipo industrial procesos secuenciales.

El autómata programable también se conoce como PLC, que es la sigla de *Programmable Logic Controller*. Tal y como se resume en la definición, se trata de un computador especial, tanto en el software como en el hardware. En el software, porque se programa en un lenguaje especial diseñado específicamente para generar de forma sencilla el programa que implementa el algoritmo de control de procesos secuenciales (de sistemas de eventos discretos), y porque el algoritmo de control programado es ejecutado de forma periódica en un ciclo temporal que es lo bastante breve como para poder controlar los procesos en tiempo real. En el hardware, porque utiliza componentes robustos que soportan condiciones de trabajo adversas, como las que se dan en ambientes industriales (polvo, temperatura, vibraciones, etc.), y porque su constitución física incluye los circuitos de interfaz necesarios para conectarlo de forma directa a los sensores y actuadores del proceso.

5.2 ARQUITECTURA DEL API

Arquitectura

La figura 5.1 muestra de forma esquemática la arquitectura interna de un autómata programable industrial típico. Como computador que es, tiene un procesador que es el que ejecuta el programa almacenado en la memoria de programa. La memoria de programa y la de datos están físicamente separadas, constituyendo una arquitectura tipo Harvard. Además, la memoria de datos está separada en dos tipos, que en la figura se denominan memoria de datos y memoria interna. Ésta última se utiliza para almacenar los valores de las señales de entrada y salida, por lo que están conectadas con los módulos de entradas y salidas, que son los elementos de interfaz donde se conectan los sensores y

actuadores del proceso. También dispone de periféricos para comunicar con otros dispositivos, como pantallas táctiles, ordenadores u otros autómatas.

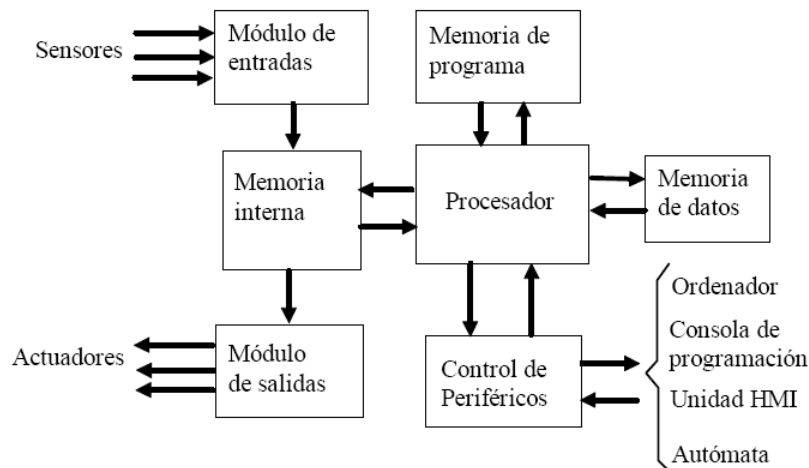


Figura 5.1: Arquitectura de un Autómata Programable Industrial.

Tipos de memoria. Clasificación física

Antes de describir cómo es cada uno de los bloques de memoria que se pueden encontrar en un autómata programable, presentamos un pequeño resumen de los tipos de memoria existentes desde el punto de vista de su constitución física:

- Memoria RAM (*Random Acces Memory*). Es una memoria volátil, es decir, que se borra si se quita la alimentación. Se puede distinguir entre RAM estática (que permanece mientras no se quite la alimentación) y RAM dinámica, que se va borrando aunque se mantenga la alimentación, por lo que requiere de un sistema que refresque (vuelva a grabar) los valores que almacena cada cierto tiempo. La memoria RAM de los PC es de este tipo. La memoria RAM utilizada en los autómatas programables suele ser estática.
- Memoria RAM con batería. La batería mantiene los datos aunque se apague la alimentación. Es bastante usual en los autómatas programables para mantener el programa y algunos datos críticos aunque se quite la alimentación. El inconveniente que tienen es que la batería se tiene que cambiar cada cierto tiempo (varios años).
- Memoria ROM (*Read Only Memory*). Grabable sólo una vez en fábrica (cuando se fabrica el chip). No se puede borrar.

- Memoria PROM (*Programmable Read Only Memory*). Grabable sólo una vez, pero se graba después de fabricar el chip (la puede grabar el propio usuario). Se llama también OTP (*One Time Programmable*) PROM. Una vez programada no se puede borrar.
- Memoria EPROM (*Electrically Programmable Read Only Memory*). Es una memoria no volátil (no se borra aunque se quite la alimentación), pero que, además, se puede borrar con luz ultravioleta y volver a grabar eléctricamente.
- Memoria EEPROM (*Electrically Erasable and Programmable Read Only Memory*). Es una memoria no volátil, es decir, que no se borra si se quita la alimentación, pero que es grabable y borrrable eléctricamente. La memoria FLASH EEPROM es una memoria de este tipo. Cada vez se utiliza más, sustituyendo a la RAM con batería, ya que no necesita batería para mantener el programa, y éste se puede modificar.

Bloques de memoria en un API. Clasificación funcional

A continuación se describen los bloques de memoria que suelen tener los autómatas programables industriales, incluyendo el tipo de memoria física y su función:

- **Memoria de programa.** Contiene el programa (instrucciones) que se ejecutan en el procesador. Se puede dividir en dos partes:
 - Una parte ROM que contiene el programa monitor (para comunicar el autómata con los módulos de programación). Este programa monitor es fijo.
 - Una parte RAM con batería (puede ser también FLASH EEPROM) en la que se almacena el programa del usuario, que implementa el algoritmo de control del proceso. Evidentemente, este programa se mantiene aunque se desconecte el autómata.
- **Memoria interna.** Almacena los valores de entradas y salidas, además de otras variables internas del autómata. En ella se almacenan variables de 1 solo bit, es decir, variables que, aunque estén organizadas en bytes (grupos de 8 bits), se puede acceder a cada uno de los bits de forma independiente para leer o escribir. En esta zona de memoria se leen los valores de las entradas (donde están conectados los sensores), y se escriben los valores de las salidas (donde están conectados los actuadores).
- **Memoria de datos.** Contiene datos de configuración o parámetros de funcionamiento del autómata y del proceso, o datos de propósito general. En ella se almacenan variables tipo byte (8 bits) o word (16 bits).

Tanto la memoria interna como la de datos suelen tener una parte de RAM normal (volátil) y una parte de RAM con batería o EEPROM para almacenar datos que no se deben perder cuando se desconecta el autómatas.

5.3 CONSTITUCIÓN FÍSICA

La mayoría de autómatas programables del mercado son modulares, es decir, están formados por módulos que pueden conectarse entre sí, permitiendo una gran flexibilidad en la configuración. Las figuras 5.2 y 5.3 muestran dos autómatas modulares comerciales.

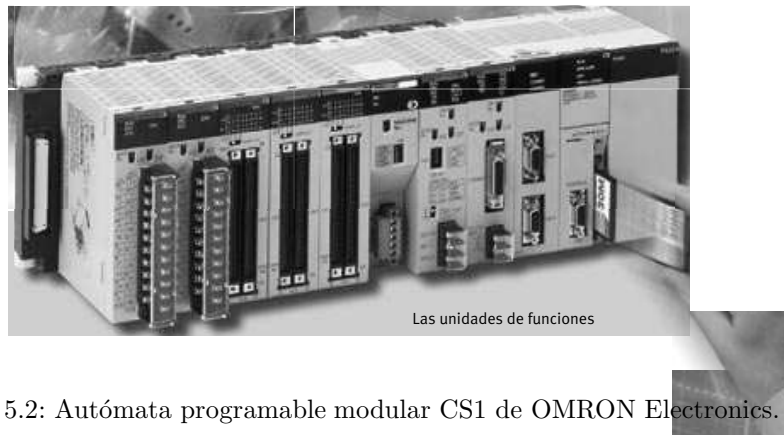


Figura 5.2: Autómatas programables modulares CS1 de OMRON Electronics.



Figura 5.3: Autómatas programables modulares Modicon m340 de Schneider Electric.

El módulo principal es el de CPU, que contiene el procesador, la memoria y algunos controladores de periféricos (puerto serie, por ejemplo). En algunos modelos, el módulo de CPU contiene además algunas entradas y/o salidas digitales.

El módulo de fuente de alimentación da la tensión (normalmente 24 V) para alimentar a todos los módulos del equipo. Junto con el módulo de CPU

forman la estructura mínima necesaria. A veces hay módulos que requieren una fuente de alimentación especial, además de la general del equipo.

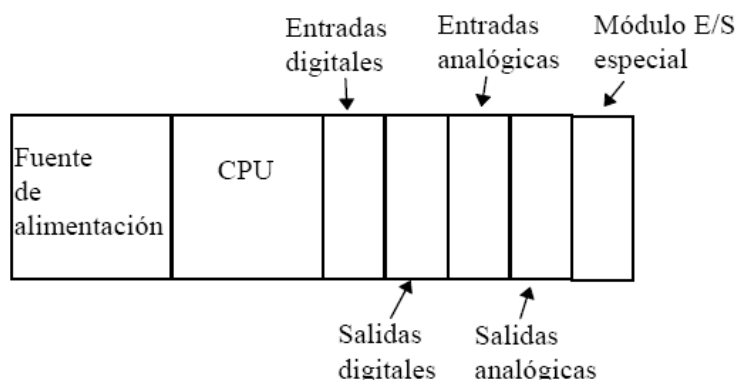


Figura 5.4: Elementos constitutivos de un autómata modular.

Los módulos adicionales se conectan al módulo de CPU por medio del bus. Respecto de los módulos entre los que se puede elegir para definir la configuración necesaria para controlar un proceso, se pueden distinguir:

- Módulos de entradas digitales.
- Módulos de salidas digitales.
- Módulos de entradas analógicas.
- Módulos de salidas analógicas.
- Módulos de comunicación (RS232, RS485, Devicenet, Profibus, Ethernet, etc.).
- Módulos de control PID.
- Módulos de control de posición de ejes.

También hay autómatas que se denominan compactos, que incluyen en un único elemento, la fuente de alimentación, la CPU y algunas entradas y salidas digitales y/o analógicas, además de algún puerto de comunicación. Éstos no son tan flexibles, aunque en algunos casos también admiten la ampliación de entradas o salidas digitales, pero en cambio son más baratos y pueden ser la mejor solución para aplicaciones sencillas. La figura 5.5 muestra un autómata compacto comercial.



Figura 5.5: Autómata programable compacto CP1H de OMRON Electronics.

5.4 DESCRIPCIÓN FUNCIONAL

El API es un computador, y por lo tanto, su funcionamiento consiste en la ejecución de un determinado programa almacenado en la memoria de programa. Se puede distinguir dos modos de funcionamiento: modo de programación y modo de ejecución.

- **Modo de programación.** En este modo de funcionamiento el programa monitor que reside en la parte ROM de la memoria de programa comunica el API con el elemento de programación (normalmente un PC), para que desde éste se le trasvase el programa que implementa el algoritmo de control que se desea ejecutar. En este modo, el API no está controlando el proceso.
- **Modo de ejecución (modo RUN).** En este modo, se ejecuta el programa del usuario que implementa al algoritmo de control, con lo que el autómata controla el proceso. En este modo, cuando se inicializa el autómata, el programa monitor salta a la dirección donde está el programa de control, y se empieza a ejecutar éste. La ejecución del programa de control se lleva a cabo de forma cíclica, ejecutando **ciclos de scan** de forma indefinida

Ciclo de trabajo o de scan

Se llama **ciclo de trabajo (o de scan)** al conjunto de tareas que el autómata lleva a cabo (instrucciones que se ejecutan) de forma cíclica cuando está controlando un proceso. El ciclo de trabajo más habitual tiene la estructura de la figura 5.6.

- La inicialización se lleva a cabo una sola vez siempre que se pone en marcha el autómata.

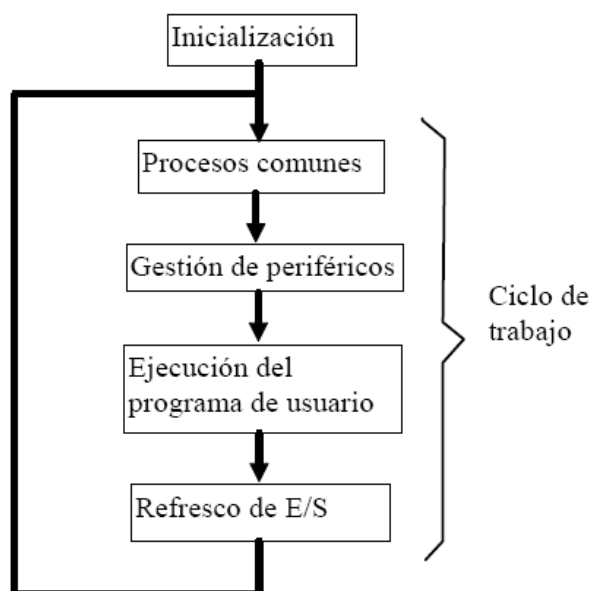


Figura 5.6: Ciclo de trabajo de un API.

- Los procesos comunes son funciones de autodiagnóstico y de programación del Watch Dog (perro guardián). No dependen del proceso que se esté controlando. El perro guardián es un temporizador de seguridad que se utiliza para detectar bucles infinitos en el programa de control. Cada ciclo de scan, ese temporizador se borra, por lo que vuelve a empezar a contar desde cero. Si pasa demasiado tiempo y el temporizador no es borrado, éste llega al final de cuenta, y produce un error de bucle infinito. Tiempos típicos del periodo máximo del perro guardián están del orden de los 100 ms.
- La gestión de periféricos gestiona los puertos de comunicaciones (comunicación con un ordenador, con un terminal HMI, con otro autómatas, etc.)
- El programa de usuario se ejecuta instrucción a instrucción, de forma secuencial. Las instrucciones (líneas) del programa de usuario son las que implementan la lógica de control del proceso. Todas estas instrucciones se ejecutan utilizando los valores almacenados en la memoria interna de entradas, y dan como resultado un cambio en la memoria interna de salidas.
- El refresco de E/S significa la actualización de la memoria interna de entradas con los valores actuales de las entradas (de los sensores), y la actualización de las salidas (actuadores) con los valores almacenados en la memoria interna de salidas. Esto quiere decir que si en medio de un ciclo de scan cambia el estado de un sensor, ese cambio no tendrá ningún efecto

en el programa hasta que se ejecute este bloque de refresco de entradas y salidas. De la misma forma, aunque el programa de usuario modifique el valor de una salida, ese cambio no llegará al actuador correspondiente hasta que se ejecute el bloque de refresco de E/S.

El **tiempo de scan** es el tiempo que tarda el autómatas en realizar un ciclo completo. Éste depende de la velocidad de la CPU, que se puede medir mediante el **tiempo de ejecución de una instrucción** de la CPU (que depende del tipo de instrucción) o por el tiempo que tarda en ejecutar 1024 instrucciones, y de la longitud del programa de usuario. Valores típicos del tiempo de scan están en el orden de varios milisegundos.

El tiempo total de un ciclo de trabajo para un autómatas determinado **es variable**, en función del tamaño (número de instrucciones) del programa de usuario. Incluso para un mismo programa de usuario, el tiempo de scan puede variar en función de las instrucciones que se ejecuten, pues puede haber instrucciones de salto o instrucciones condicionadas que no siempre se ejecutan.

Se define el **tiempo de respuesta** como el tiempo que tarda en cambiar una salida en relación al cambio de una entrada. Este tiempo es importante, pues determina la rapidez con que el autómatas reacciona ante cambios del proceso. El tiempo de respuesta es variable y depende del tiempo de scan, del tiempo de retraso de activación/desactivación de los dispositivos físicos de entradas y salidas, y del momento en que se produce el cambio de la entrada. El peor caso es aquél en el que la entrada cambia justo después de que se haya refrescado la memoria interna de entradas. En ese caso, se ejecuta un ciclo de programa completo en el que el nuevo valor de esa entrada no es tenido en cuenta. Después se actualiza la memoria de entradas, con lo que se lee el valor de la entrada. A continuación, se ejecuta otro ciclo de programa en el que el nuevo valor de la entrada sí se tiene en cuenta, para después actualizar las salidas. El tiempo de respuesta máximo es, pues:

$$T_{respmax} = 2T_{ciclo} + T_{out}$$

donde T_{out} es el tiempo que tarda el dispositivo físico de salida en activar realmente la salida, y que depende del tipo de salida (relé, transistor o triac).

Entradas de interrupción

El tiempo de respuesta suele variar desde unos pocos milisegundos (para salida transistor) hasta unos 15 milisegundos (para salida relé). Normalmente, esto es suficientemente rápido para la mayoría de los procesos. En el caso de que ese tiempo de respuesta sea demasiado elevado (si el proceso es muy rápido), se pueden utilizar unas **entradas especiales de interrupción**, que funcionan de forma independiente del ciclo de trabajo. Cuando se activa una entrada de interrupción, el procesador interrumpe el programa que está ejecutando y en su lugar ejecuta una rutina (la rutina de interrupción) en la que se ejecutan

las instrucciones que se deseen, para después volver al programa principal en el punto en que se había quedado. En esa rutina de interrupción se puede programar la acción que sea necesaria para que el sistema reaccione ante la entrada (por ejemplo modificar una salida). La entrada de interrupción, en realidad, no es atendida de forma inmediata, pues el procesador debe terminar la instrucción que tiene a medias en ese momento además de guardar el estado actual del programa, antes de poder ejecutar la rutina de interrupción. Ese retardo se llama tiempo de latencia de la interrupción, y suele ser muy pequeño. Cuando se utiliza una entrada de interrupción para actuar rápidamente sobre el proceso, es necesario utilizar una salida tipo transistor, pues el retardo de activación/desactivación es pequeño. Si se utilizara una salida tipo relé (que tienen un retardo de activación/desactivación de 10 a 20 ms), el tiempo de reacción total sería muy alto, por lo que no tendría demasiado sentido utilizar una interrupción.

Interrupciones temporizadas

También se pueden programar tareas especiales que se ejecuten de forma síncrona cada cierto tiempo, independientemente del ciclo de trabajo normal. Para eso se tienen las interrupciones temporizadas; un temporizador que está continuamente decrementando su valor al ritmo que le marca el reloj del sistema. Cuando llega a cero vuelve a cargarse al valor máximo que se le ha programado. Se puede programar para que cada vez que el temporizador llegue a cero se ejecute una rutina de interrupción para realizar una tarea determinada.

Las interrupciones temporizadas se utilizan cuando el periodo de la tarea que se quiere realizar cada cierto tiempo es muy bajo (del orden de 100 ms o menos). Un ejemplo podría ser la lectura del valor de una variable del proceso que cambia rápidamente y cuya evolución con el tiempo se quiera almacenar. Otro ejemplo podría ser la ejecución de un bucle de control PID.

Si el periodo de la tarea a realizar periódicamente es elevado (por ejemplo varios segundos o minutos), se utilizan los temporizadores software del autó-mata, que son comprobados cada ciclo de scan. El máximo error que se comete con estos temporizadores es el de un ciclo de scan (unos pocos milisegundos), que en general es un error despreciable si el periodo es grande.

Otros tipos de ciclo de funcionamiento

El ciclo de funcionamiento descrito antes es el más habitual, pero no es la única opción posible. Otras posibilidades son:

- Refresco de las salidas conforme éstas cambian al ejecutar cada instrucción del programa de usuario.
- Refresco de las entradas y salidas conforme éstas van siendo utilizadas en cada instrucción del programa de usuario.

- También cabe la posibilidad de un ciclo de scan síncrono (que tarde siempre el mismo tiempo), que ejecute un ciclo completo cada T milisegundos (fijos). Evidentemente, este tiempo fijo debe ser mayor del máximo tiempo posible de ejecución de un ciclo completo, por lo que lo normal será que, en cada ciclo, el autómatas esté un tiempo de espera sin hacer nada.

5.5 MÓDULOS DE ENTRADA/SALIDA Y MÓDULOS DE COMUNICACIÓN

A continuación se describen los diferentes tipos de módulos que se pueden encontrar en autómatas comerciales, y que pueden combinarse para formar la configuración que sea necesaria para un proceso concreto.

Módulos de entradas digitales

Son módulos que permiten leer señales binarias (digitales) obtenidas a partir de detectores (cuya salida es activa o inactiva). Normalmente, van aisladas por medio de optoacopladores. Suelen funcionar con niveles de tensión de 24 V (de continua o de alterna) o niveles de 220 V (de alterna). La figura representa el circuito de entrada típico. La utilización de 2 diodos emisores en antiparalelo hace que la entrada pueda activarse con cualquier polaridad. Se observa que uno de los bornes es común (es compartido por todas las entradas digitales del módulo). Este borne común se puede conectar o bien a tensión (a 24 Vdc) o bien a masa, en función de que el detector que se conecta sea NPN o PNP. Si se conecta un detector con salida NPN, el borne común se tiene que conectar a tensión, mientras que si el sensor es PNP, el borne común se tiene que conectar a masa. Debido a esto, no se pueden mezclar sensores NPN y PNP en el mismo módulo (todos los sensores que van al mismo módulo tienen que ser del mismo tipo, iguales). Los sensores con salida a relé se pueden conectar como si fueran NPN o como si fueran PNP, por lo que se podrían combinar sin problemas con los NPN o con los PNP.

CONEXIÓN DE DETECTOR CON SALIDA DE CONTACTO

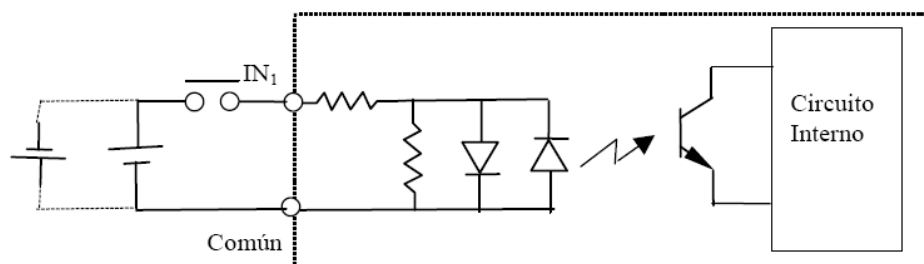


Figura 5.7: Conexión de un sensor con salida de relé a la entrada digital de un autómatas.

CONEXIÓN DE DETECTOR CON SALIDA PNP

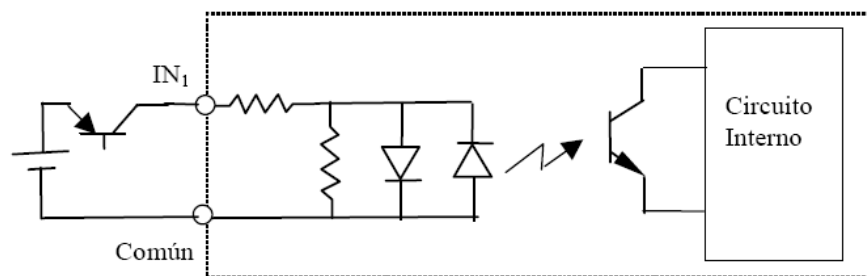


Figura 5.8: Conexión de un sensor con salida PNP a la entrada digital de un autómata.

CONEXIÓN DE DETECTOR CON SALIDA NPN

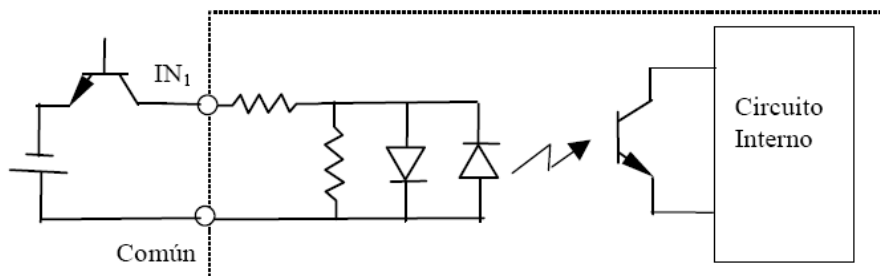


Figura 5.9: Conexión de un sensor con salida NPN a la entrada digital de un autómata.

Algunos módulos de entradas digitales tienen los dos bornes de cada entrada independientes (aislados unos de otros). En esos módulos cada entrada puede conectarse a un detector NPN o PNP de forma independiente. Sin embargo, este tipo de módulo no es el más habitual.

La figura 5.10 muestra el circuito de entrada digital del módulo ID111 del autómata CQM1 de Omron, que admite señales de 24 V. Este autómata también tiene otros módulos de entradas de 240 V de alterna, como el que se muestra en la figura 5.11.

CQM1-ID111/212

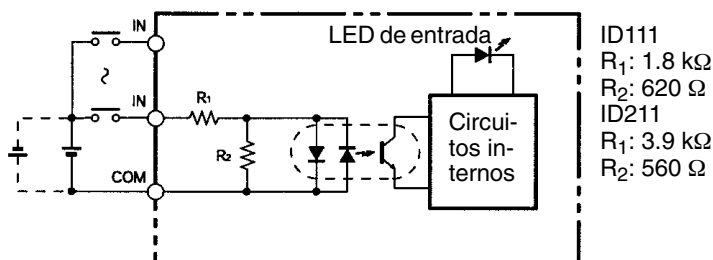


Figura 5.10: Circuito de entrada digital del módulo ID111/212 del autómata CQM1 de Omron Electronics.

CQM1-IA121

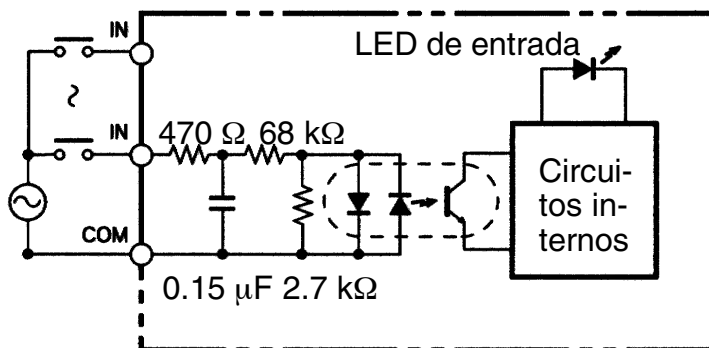


Figura 5.11: Circuito de entrada digital del módulo IA121 del autómata CQM1 de Omron Electronics.

Módulos de salidas digitales

Estos módulos permiten transmitir una señal binaria (activada/desactivada) a un dispositivo externo. Hay tres tipos de salidas digitales que son las más

utilizadas: salida a relé (contacto), salida a transistor (que puede ser NPN o PNP) y salida a TRIAC. En todos los casos, el circuito interno está aislado de los componentes externos (en la salida a relé por el propio relé, en la salida a transistor por medio de un optoacoplador y en la salida a TRIAC por un optotriac). Los circuitos típicos simplificados para los distintos tipos de salidas se muestran en las figuras 5.12 a 5.15. Normalmente, en un módulo hay varias salidas, que comparten uno de los bornes (el borne común). En el caso de la salida a relé, el borne común se puede conectar a tensión o a masa. En el caso de salida a transistor NPN, el borne común se conecta necesariamente a masa, y la carga (bobina de electroválvula, por ejemplo) entre el otro borne de la salida y tensión. En el caso de la salida a transistor PNP, el borne común se conecta necesariamente a tensión, y la carga entre el otro borne de salida y masa.

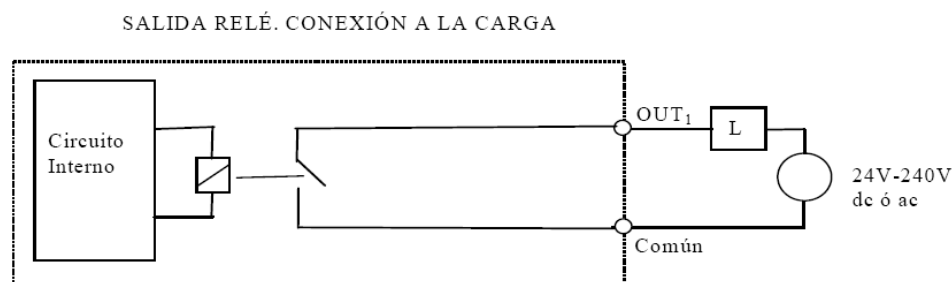


Figura 5.12: Circuito de salida digital a relé y su conexión a la carga.

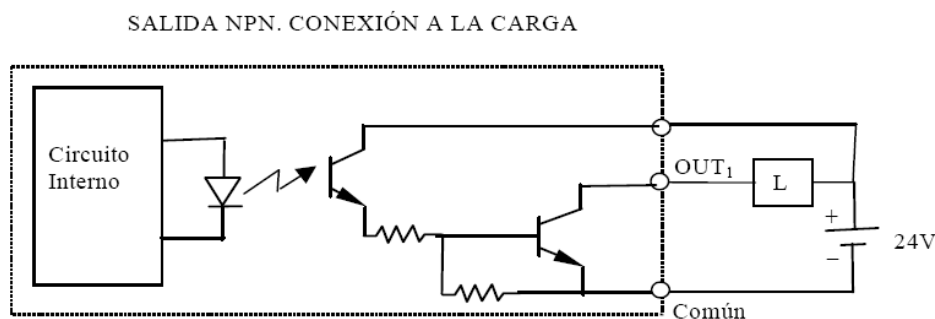


Figura 5.13: Circuito de salida digital a transistor NPN y su conexión a la carga.

La figura 5.16 muestra el circuito correspondiente a un módulo de salida digital a relé del autómata CQM1 de Omron. Este autómata también tiene módulos de salida a transistor NPN (como el que se muestra en la figura 5.17), módulos de salidas a transistor PNP y a triac.

SALIDA PNP. CONEXIÓN A LA CARGA

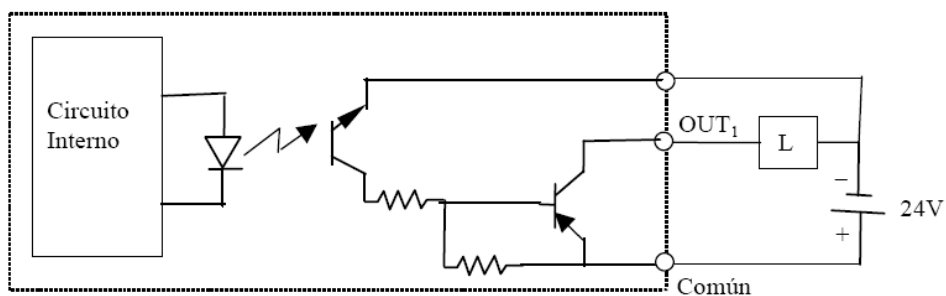


Figura 5.14: Circuito de salida digital a transistor PNP y su conexión a la carga.

SALIDA TRIAC. CONEXIÓN A LA CARGA

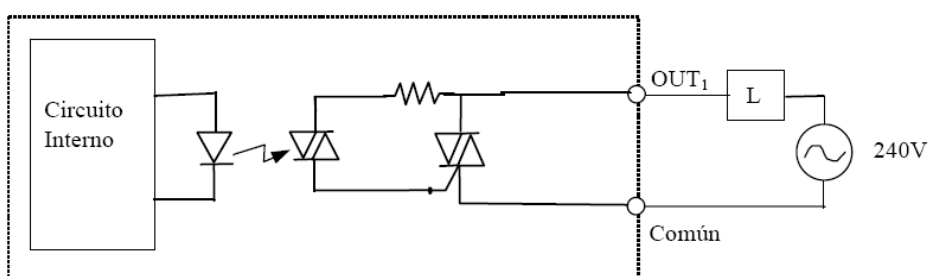


Figura 5.15: Circuito de salida digital a triac y su conexión a la carga.

CQM1-OC222

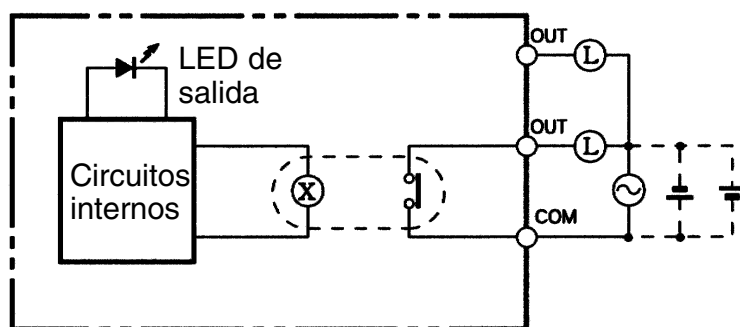


Figura 5.16: Circuito de salida digital del módulo OC222 del autómata CQM1 de Omron Electronics.

CQM1-OD211

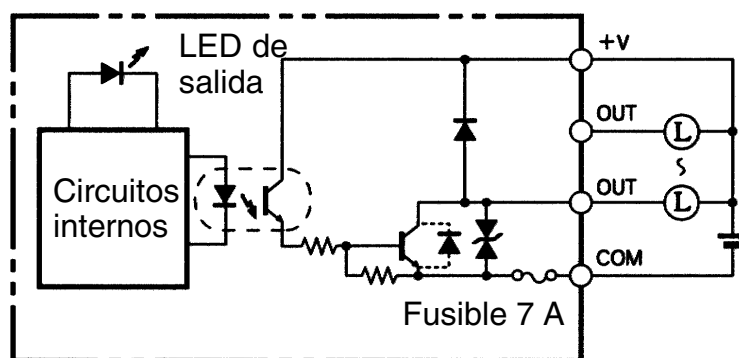


Figura 5.17: Circuito de salida digital del módulo OD211 del autómata CQM1 de Omron Electronics.

Módulos de entrada/salida analógica

Son módulos que permiten medir señales analógicas y dar como salida señales analógicas. De esta forma, permiten controlar procesos que requieran señales continuas.

Las entradas analógicas pueden ser de tensión (normalmente en un rango de -10 a 10 V ó de 0 a 10 V) o de corriente (en un rango de 4 a 20 mA). En realidad, estos módulos tienen un convertidor analógico-digital, que da un número proporcional a la señal. La resolución (número de bits) varía de unos autómatas a otros. Es bastante común una resolución de 10 bits (1024 valores) o 12 bits (4096 valores). Los más precisos tienen resoluciones de 16 bits (distingue entre 65536 valores) o incluso más. Las entradas analógicas suelen incorporar un filtro paso bajo (que puede ser simplemente de primer orden) para reducir el ruido de la señal que se mide.

Las salidas analógicas pueden ser también de tensión o de corriente (en rangos similares a las entradas). Disponen de un convertidor digital-analógico dando como salida una tensión (o corriente) proporcional al número que se escribe en ellos. La resolución también puede oscilar entre 10 bits y 16 bits.

Las entradas y salidas analógicas en corriente (normalmente de 4 a 20 mA) se utilizan, especialmente, cuando los elementos entre los que se transmite la señal están situados a mucha distancia. La señal de corriente es mucho más inmune que la de tensión, a los ruidos electromagnéticos. El ruido electromagnético captado es mayor cuanto más largo es el cable, por lo que cuando éste es largo, la señal de tensión puede verse afectada de forma importante por el ruido.

Con los módulos de entrada y salida analógica se puede implementar un controlador PID programando adecuadamente el autómata. Sin embargo, lo

normal es que se utilice un módulo especial de control PID, que implementa el controlador, la entrada analógica y la salida analógica.

Modulos de utilidad específica

Además de los módulos de entradas y salidas anteriores, los API disponen de una serie de elementos que permiten realizar funciones específicas:

Módulo PID. Realiza el control PID de un proceso. Dispone de entrada analógica (o entrada de temperatura) y salida analógica, y desde el autómata se le pueden programar las constantes (BP, Td y Ti) así como el valor de la referencia.

Módulo de posicionado. Es un módulo que realiza el control de posición de un eje (de un motor paso a paso, de continua o brushless). Recibe como entrada la señal del codificador incremental del motor, y da como salida las señales que actúan sobre el driver de potencia que ataca al motor. El módulo se encarga de tareas complejas como calcular trayectorias. Desde el autómata se pueden programar los distintos parámetros de posicionado (constantes PID, posición de referencia, valor de rampa, etc.).

Módulo de entrada de temperatura. Se utiliza específicamente para medir temperaturas. Puede aceptar entradas de termopar o de PT100.

Módulo de entrada de pulsos de alta velocidad. Permite contar los pulsos de una señal digital que cambie a una frecuencia muy elevada (por ejemplo un codificador).

Módulo de salida de pulsos de alta frecuencia. Permite dar como salida señales PWM de ciclo de trabajo programable y de frecuencia elevada. Suelen utilizarse para atacar equipos que controlan motores, para indicarles el número de pasos que deben avanzar (uno por pulso enviado), o la velocidad a la que deben llevar al motor (en función de la frecuencia o del ciclo de trabajo de la señal).

Módulo de entradas/salidas remotas. Permite conectar entradas y salidas digitales situadas a grandes distancias del autómata. Normalmente, consta de un módulo conectado al rac del autómata, y de uno o varios módulos de entrada/salida situados lejos, conectados mediante un bus de campo (un cable) de 2 o 4 hilos.

Módulo de comunicación con bus de campo. Es un módulo que permite al autómata comunicarse con otros elementos iguales o inferiores (otros autómatas, sensores, variadores de frecuencia, módulos de entradas/salidas remotos, etc) por medio de un bus de campo (que suele ser un cable de 2 hilos).

Módulos de comunicaciones de nivel superior. Permiten al autómata comunicarse con elementos iguales o superiores (otros autómatas, ordenadores, etc.), por ejemplo, a través de una red Ethernet.

5.6 UNIDADES DE PROGRAMACIÓN

Las unidades de programación son los periféricos más comunes que se utilizan al trabajar con los API. Estas unidades son las que permiten introducir en la memoria del autómatas el programa de usuario que servirá para controlar el proceso. Las más utilizadas son: la consola de programación y el PC.

Consola de programación

La consola de programación es un dispositivo específico para programar el autómatas. Consta de un incómodo teclado y una pequeña pantalla. Se comunica con el autómatas por medio del puerto serie o de un puerto especial de periféricos (situados en el módulo de CPU). Introducir programas mediante la consola resulta engorroso y lento. Normalmente tienen utilidad para realizar pequeñas modificaciones en los programas de autómatas que ya están instalados en planta, debido a que son de pequeño tamaño (cómodos de transportar) y se alimentan directamente del autómatas.

Ordenador personal

El elemento más utilizado para programar los API es el ordenador personal. Se conecta al autómatas mediante el puerto serie RS232 o un puerto USB (normalmente). Cada fabricante de autómatas dispone de un software que permite la edición de programas, la depuración, la compilación y la carga del mismo en el autómatas. También permite ejecutar el programa en el autómatas y tener acceso al mismo tiempo a las distintas variables, lo que permite la depuración y detección de errores.

Unidades HMI

Además de las unidades de programación, podemos destacar otros periféricos muy utilizados que son las unidades HMI (*Human Machine Interface*). Son unidades que disponen de un teclado y una pantalla (que puede ser bastante grande), o bien únicamente de una pantalla táctil. Se utilizan para mostrar al operador el estado del proceso y permitir la introducción de consignas (órdenes de marcha, de paro, cambio de parámetros de funcionamiento, etc.). Estas unidades contienen un procesador y una memoria no volátil en la que se programa la aplicación de interfaz con el usuario. Se programan por medio de un ordenador por el puerto serie o un puerto USB con el software adecuado, y una vez programadas se comunican con el autómatas por el puerto serie o por el puerto de periféricos (aunque también se pueden comunicar por medio de una red de mayor nivel, como Ethernet). En la figura 5.18 se muestra un terminal táctil comercial.



Figura 5.18: Terminal táctil programable NS5 de Omron Electronics.

5.7 PROGRAMACIÓN DE LOS AUTÓMATAS PROGRAMABLES INDUSTRIALES

La forma más común de programar un autómatas es mediante un ordenador personal con el software adecuado. En primer lugar se debe configurar el equipo (seleccionar el tipo de autómatas y los módulos de entradas/salidas que se tiene). Después se introduce el programa de control (en el lenguaje adecuado), y por último, se carga este programa en el autómatas para proceder a su verificación (se ejecuta el programa comprobando la evolución de las distintas variables). Antes de introducir el programa es conveniente hacer una lista de las variables de entrada y salida del proceso, y su asignación a variables internas en la memoria del autómatas. Algunos software de programación permiten simular la ejecución del programa sin necesidad de cargarlo en un autómatas, por lo que se puede hacer la depuración de forma más sencilla.

Lenguajes de programación

El lenguaje de programación más utilizado en los autómatas programables es el **diagrama de contactos** (*Ladder Diagram* o **diagrama de escalera**). Está basado en los automatismos cableados por medio de contactores, que fueron los primeros en implementarse. Gráficamente se representan dos líneas verticales largas separadas, de forma que la de la izquierda representa tensión y la de la derecha, masa. Entre esas líneas verticales, se representan las ecuaciones lógicas por medio de contactos. Hay 3 tipos de elementos fundamentales.

- Contacto normalmente abierto (NA): Representa un contacto que está abierto si la variable asociada vale 0, y que se cierra si la variable asociada vale 1. Se representa con dos líneas verticales paralelas.

- Contacto normalmente cerrado (NC): Representa un contacto que está cerrado si la variable asociada vale 0, y que se abre si la variable asociada vale 1. Se representa con dos líneas verticales paralelas cruzadas por una línea oblicua.
- Bobina. Representa el valor de una variable. Cada salida tiene asociada una bobina. Si esta bobina tiene corriente (el diagrama de contactos la activa), la salida está a 1, y si no tiene corriente, la salida está a cero. Además puede haber bobinas asociadas a variables internas (por ejemplo, las que representan cada etapa del proceso). Se representa mediante un círculo.

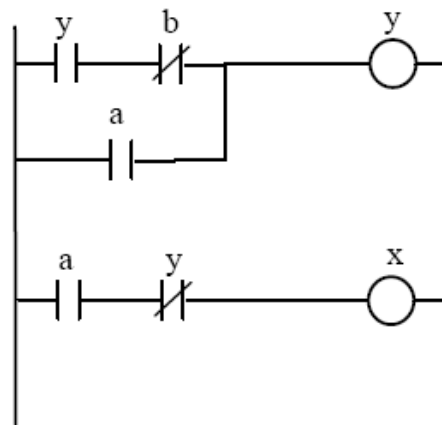


Figura 5.19: Ejemplo de programa en diagrama de contactos (*ladder diagram*).

Cada línea o red (que normalmente contiene una sola bobina, aunque puede contener varias) representa una línea del programa. El autómata ejecuta las líneas de arriba abajo. En el caso del ejemplo de la figura 5.19, calcularía $y \cdot \bar{b} + a$ y el resultado (1 ó 0) lo asignaría a la variable y . Después calcularía $a \cdot \bar{y}$, y el resultado lo asignaría a la variable x . Esta sería la fase de ejecución del programa de usuario, que se repetiría de forma cíclica según se ha visto.

La función lógica implementada por los contactos de una red se suele denominar **condición de activación**, ya que cuando esta condición es 1, la variable de salida (bobina) asociada se activa.

Hay una serie de reglas que se debe seguir para dibujar correctamente un diagrama de contactos. Algunas de ellas son:

- A la derecha de la bobina no puede haber contactos.
- En una red puede haber varias bobinas en paralelo.
- La conducción a través de los contactos solo se produce de izquierda a derecha.

Además de la bobina simple puede haber otros dos tipos de bobinas:

- Bobina normalmente cerrada. Se representa mediante un círculo con una línea transversal. Su valor es el inverso del que resulta de la ecuación lógica de los contactos.
- Bobina de enclavamiento. En algunos autómatas se representa mediante un círculo que en el interior tiene una letra L, de *Latch* (o S de *Set*) o una letra U, de *Unlatch* (o R, de *Reset*). Si es una letra L (*Latch*) o S (*Set*) indica que la variable correspondiente se activa, quedando activa aunque después se haga cero la condición de activación. Para desactivar la variable tiene que utilizarse una bobina de desenclavamiento con la letra U (*Unlatch*) o R (*Reset*).

Esos elementos básicos solo permiten programar relaciones lógicas simples. Evidentemente esto no es suficiente para controlar un proceso complejo. Además de estos elementos hay otros que permiten implementar funciones más elaboradas. Se colocan en lugar de las bobinas simples, tal como se muestra en la figura 5.20.

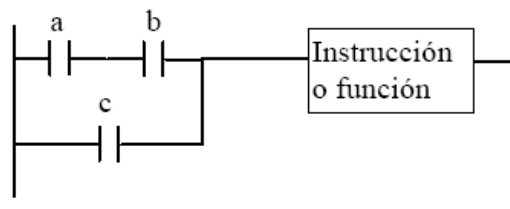


Figura 5.20: Uso de funciones complejas en diagramas de contactos.

El significado es que cuando el autómata llega a la línea correspondiente, ejecuta la instrucción o función compleja en función de que los contactos den un valor 1 ó 0, es decir, en función de que la condición de activación sea 1 ó 0. En algunos casos, si dan un valor 0 la función no se ejecuta, en otros, la ejecución es diferente en función de que se tenga un 1 o un 0. Entre estas funciones se puede destacar:

- Instrucción END. Indica el final del programa.
- Temporizadores.
- Contadores.
- Instrucciones de salto.
- Operaciones binarias (mover bytes, sumar, multiplicar, comparar, etc.).
- Refresco de E/S. Permite refrescar las E/S cuando se desee, independientemente del refresco automático que se produce cada ciclo de scan.
- Funciones matemáticas complejas (senos, cosenos, exponenciales).
- Control de interrupciones.

- Funciones especiales de comunicaciones (transmisión o recepción de datos).

Otro lenguaje muy utilizado es la **lista de instrucciones**. Éste es un lenguaje de aspecto similar al ensamblador, aunque no se codifican instrucciones de código máquina del procesador. El programa se expresa en una lista de instrucciones muy básicas.

Cada instrucción contiene un nemónico (el nombre de la instrucción) y uno o varios datos con los que opera. Estos datos pueden ser un bit o una palabra (registro de 8 ó 16 bits). El conjunto de instrucciones disponible y la nomenclatura utilizada depende del fabricante del autómata.

Hay una serie de instrucciones básicas que permiten implementar las ecuaciones lógicas equivalentes al diagrama de relés. Estas instrucciones básicas son (en autómatas Omron):

LD xxxx: abre una red de contactos con el contacto xxxx abierto. La condición de ejecución tiene en ese momento el valor del bit xxxx.

LD NOT xxxx: abre una red de contactos con el contacto xxxx cerrado. La condición de ejecución tiene en ese momento el valor del bit xxxx negado.

AND xxxx: hace el Y lógico entre la condición de ejecución actual en la red y el operando xxxx. El resultado es la condición de ejecución.

AND NOT xxxx: hace el Y lógico entre la condición de ejecución actual en la red y el operando xxxx negado. El resultado es la condición de ejecución).

OR xxxx: hace el O lógico entre la condición de ejecución actual en la red y el operando xxxx. El resultado es la condición de ejecución.

OR NOT xxxx: hace el O lógico entre la condición de ejecución actual en la red y el operando xxxx negado. El resultado es la condición de ejecución.

Con estas instrucciones se puede escribir la ecuación lógica de una red. Por ejemplo, tomando como partida el diagrama ladder de la figura 5.21. La lista de instrucciones de Omron equivalente se muestra en la figura 5.22.

En este ejemplo, lo que hace el procesador es: carga en un registro el valor de la variable situada en la dirección 1000, después realiza un AND entre ese registro y la variable situada en 0002 (previamente negada), dejando el resultado en el mismo registro. Después hace un OR entre el registro y la variable situada en 0001 dejando el resultado en el mismo registro. Ese registro contiene la condición de ejecución de la rama. La instrucción (o instrucciones) siguiente (en este caso OUT 1000) se ejecuta condicionada a la condición de ejecución. En el caso de la instrucción OUT, se activa la variable si la condición es 1 y se desactiva si es 0. Otras instrucciones, simplemente, se ejecutan o no se ejecutan en función de que la condición sea 1 ó 0.

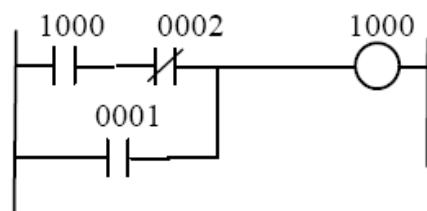


Figura 5.21: Representación de instrucciones en Diagrama de contactos (*ladder*).

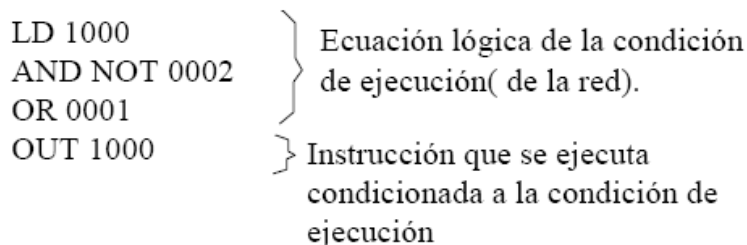


Figura 5.22: Representación de instrucciones en lenguaje de lista de instrucciones.

Hay diagramas cuyas ecuaciones no se pueden expresar con las instrucciones anteriores. Por ejemplo, el representado en la figura 5.23.

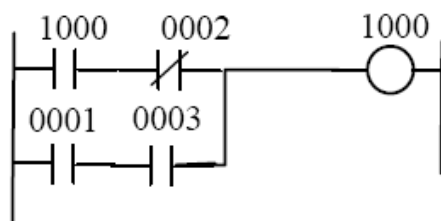


Figura 5.23: Ejemplo de diagrama de contactos con dos ramas.

En este caso hay que abrir dos ramas independientes, para después hacer un OR entre ellas.

Para expresar estas relaciones existen otras dos instrucciones:

AND LD: realiza la función Y lógica entre las dos últimas ramas abiertas

OR LD: realiza la función O lógica entre las dos últimas ramas abiertas

El ejemplo anterior quedaría:

```
LD 1000
AND NOT 0002
LD 0001
AND 0003
```

OR LD

OUT 1000

En este caso el procesador sigue el siguiente procedimiento: carga en el registro 0 el valor de la variable situada en la dirección 1000, después realiza un AND entre ese registro y la variable situada en 0002 (previamente negada), dejando el resultado en el mismo registro 0. Cuando se ejecuta un segundo LD se mueve el valor del registro 0 al registro 1. Después se carga en el registro 0 el valor de la variable situada en 0001 y realiza un AND con la variable situada en 0003 dejando el resultado en el registro 0. La instrucción OR LD realiza un OR entre el registro 0 y 1, dejando el resultado en el registro 0. La condición de ejecución es el valor que queda en ese registro 0. Se pueden encadenar estas instrucciones hasta un nivel de 8. Cada vez que se utiliza una instrucción LD o LD NOT, se desplaza a la derecha todo el byte (8 registros) de forma que el registro 0 se mueve al 1, el 1 al 2, etc., perdiéndose el 7. La instrucción LD además carga un valor en el registro 0. Cada vez que se utiliza la instrucción OR LD o AND LD, se hace un OR (o AND) entre el registro 0 y 1, y se desplazan hacia la izquierda todos los registros.

Es evidente que la lista de instrucciones resulta mucho menos autoexplicativa de las relaciones lógicas entre las variables que el diagrama de contactos. Sin embargo, a diferencia del diagrama de contactos, se aproxima más a las instrucciones que realmente ejecuta el procesador del autómata.

Además de las instrucciones que permiten realizar las ecuaciones lógicas, existen multitud de instrucciones y funciones que son las que se ejecutan condicionadas a la condición de ejecución de cada rama. La más simple es la instrucción OUT. Ésta activa una variable (que puede ser una salida) si la condición de ejecución es 1, y la desactiva si la condición de ejecución es 0. También está la versión negada OUT NOT, que hace lo contrario.

Se llaman instrucciones de control de bit a aquellas que permiten modificar el valor de variables binarias. En los autómatas de Omron éstas son:

OUT xxxx: representa una bobina de salida normal, es decir, a la variable xxxx se le asigna el valor de la condición de ejecución.

OUT NOT xxxx: representa una bobina de salida negada, es decir, a la variable xxxx se le asigna el valor negado de la condición de ejecución.

SET xxxx: representa una bobina de enclavamiento de *Set* o “Latch”, es decir, si la condición de ejecución es 1, se asigna 1 a la variable xxxx, pero si es cero, no se hace nada.

RESET xxxx: representa una bobina de desenclavamiento de *Reset* o “Un-latch”, es decir, si la condición de ejecución es 1, se asigna 0 a la variable xxxx, pero si es cero, no se hace nada.

DIFU xxxx: es un detector de flancos de subida. Asigna un valor 1 a la variable xxxx durante un ciclo de scan cuando la condición de ejecución

cambia de 0 a 1, asignándole después un valor 0 en el siguiente ciclo de scan.

DIFD xxxx: es un detector de flancos de bajada. Asigna un valor 1 a la variable xxxx durante un ciclo de scan cuando la condición de ejecución cambia de 1 a 0, asignándole después un valor 0 en el siguiente ciclo de scan.

Temporizadores y contadores

Unas de esas funciones de importancia especial son los temporizadores y contadores. Un contador es un número que se incrementa (o decrementa) cada vez que hay un pulso en una variable (que puede ser una entrada o una variable interna). Un temporizador es un contador cuya entrada de cuenta es el reloj del sistema (una señal cuadrada de frecuencia constante). Normalmente, se utiliza escribiendo un valor inicial en él. El valor del temporizador va disminuyendo conforme cuenta los pulsos del reloj hasta que llega a cero, momento en que activa una variable de fin de temporización. De esta forma se pueden medir tiempos (se puede activar una salida durante un tiempo determinado, o esperar un tiempo determinado hasta activar una salida, o medir el tiempo transcurrido entre las activaciones de dos entradas, etc.).

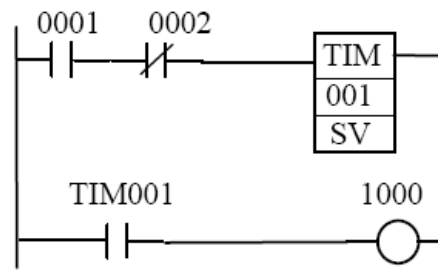


Figura 5.24: Uso de temporizadores en los autómatas Omron.

La figura 5.24 muestra el uso de temporizadores en los autómatas Omron. Es un temporizador de retardo en la conexión. Esta instrucción TIM se activa cuando la condición de ejecución es 1, empezando a contar decrementando su valor de cuenta interno a partir del valor inicial SV. Cuando llega a cero se activa la variable binaria asociada al temporizador (TIM001), que en este ejemplo se utiliza para activar la salida 1000. Cuando se desactiva la condición de ejecución del temporizador, éste se resetea y su salida (variable TIM001) se pone a cero. Existen diversos tipos de temporizadores y contadores, y cada fabricante de autómatas tiene definidos los suyos.

Estándar de programación IEC 61131-3

Cada fabricante de autómatas programables tiene su propia versión del lenguaje de diagrama de contactos y del de lista de instrucciones. En general, un programa escrito para un autómata determinado no sirve para un autómata de otra marca. Las diferencias más importantes están en las instrucciones y funciones especiales, que cada fabricante implementa a su manera (especialmente temporizadores y contadores), aunque también hay diferencias en elementos más simples, como bobinas de enclavamiento.

La norma IEC 61131-3 surgió con el objetivo de establecer una sintaxis de programación estándar que permitiera programar un automatismo sin necesidad de recurrir a la sintaxis específica de ningún fabricante. El objetivo último es, evidentemente, que el programa escrito según la norma IEC 1131-3 se pueda utilizar directamente en cualquier autómata.

Cada vez son más los fabricantes cuyos sistemas de programación se acogen a este estándar. Además, recientemente han aparecido paquetes de software de programación de autómatas basados en este estándar, cuyo objetivo es poder realizar programas independientemente de la marca del autómata que se vaya a utilizar. CoDeSys es una de esas aplicaciones universales de programación de autómatas. Evidentemente, para que la aplicación funcione, es necesario que el autómata sea capaz de ejecutar los programas realizados según el estándar. Para ello se modifica el *firmware*, es decir, el sistema operativo y el programa monitor del autómata para hacerlo compatible. En la actualidad son ya muchos los fabricantes que disponen de autómatas compatibles con este tipo de herramientas de programación basadas en el estándar IEC 61131-3.

El estándar IEC 61131-3 contempla seis lenguajes diferentes, tres basados en gráfico y tres basados en texto. Los lenguajes basados en gráfico son:

- Diagrama de contactos o de escalera. Es muy similar al de cualquier autómata, aunque hay diferencias. Por ejemplo el lenguaje de Omron utiliza solo contactos NA o NC, mientras la norma IEC 61131-3 admite además contactos de flanco de subida o de bajada que se activan solo durante un ciclo de scan cuando la variable asociada cambia de 0 a 1 o de 1 a 0.
- Diagrama de bloque de funciones. Dispone de cuatro tipos de elementos gráficos que se combinan para la creación de un programa. Estos elementos son: conexiones, elementos de control de ejecución, elementos para llamada a funciones y los conectores. No se incluyen otros elementos gráficos como contactos y bobinas que sí aparecen en los diagramas de contactos.
- Diagrama secuencial de funciones (modo gráfico). Permite descomponer programas complejos en unidades más simples y definir la forma en que se realizan las transiciones entre estas unidades. Su apariencia es similar

a la de los diagramas de Grafcet, pues cuenta con dos tipos de elementos: las etapas (o pasos), representadas por cuadrados, y las transiciones, representadas por líneas horizontales, a las que se le asocia una condición de activación. Dentro de cada paso o etapa se incluye el código de programa (usando cualquiera de los lenguajes definidos en la norma) que se ejecutará cuando se active dicha etapa.

Los lenguajes basados en texto son:

- Lista de instrucciones. Es similar a un lenguaje ensamblador. En general, los programas creados con otros lenguajes, gráficos o de texto, son traducidos a listas de instrucciones en el proceso de compilación.
- Texto estructurado. Es un lenguaje de programación de alto nivel que ofrece un amplio rango de instrucciones que simplifican la realización de operaciones complejas.
- Diagrama secuencial de funciones (modo texto).

Una de las diferencias fundamentales del estándar IEC 61131-3 respecto de los sistemas de programación tradicionales de API es que la declaración de variables se hace de la misma manera que en los lenguajes de programación de alto nivel (C, Pascal, Java, etc.), o sea que el programador no tiene que especificar la dirección de cada variable. La asignación de las direcciones de memoria del PLC donde se almacenarán cada una de las variables declaradas se hace de forma automática durante la compilación del programa. En los sistemas de programación tradicionales, las referencias a las direcciones de memoria se hacen de forma directa o a través de símbolos que se asocian a cada localización en la tabla de definición de símbolos. La forma en que se realiza la declaración de variables en el estándar IEC 61131-3 es sin duda es un elemento novedoso que facilita la tarea de programación.

Procedimiento a seguir para programar un autómatas

Una vez se tiene el programa en papel (obtenido a partir del GRAFCET del proceso, o de cualquier otra forma), se debe proceder según los siguientes pasos para introducir el programa en el PLC:

1. Configurar el equipo (tipo de PLC, módulos de entrada y salida, comunicaciones, etc).
2. Identificar el programa (nombre, empresa, fecha, versión, etc.).
3. Crear la estructura básica del programa (bloques y grupos dentro de los bloques).
4. Hacer la lista de etiquetas (nombres) de las variables a utilizar (entradas, salidas, variables internas, datos, temporizadores, contadores), asignándoles su dirección en memoria.

5. Introducir el programa (diagrama de contactos).
6. Comprobar el programa. Si el software lo permite, se puede simular la ejecución del programa en el PC y verificar el funcionamiento correcto.
7. Transferir el programa al autómat.
8. Ejecutar el programa en el PLC en modo monitor y comprobar el funcionamiento correcto.

5.8 CRITERIOS DE SELECCIÓN DE UN AUTÓMATA PROGRAMABLE INDUSTRIAL

Las características fundamentales que definen un PLC, y que se deben tener en cuenta a la hora de elegir el más adecuado para una aplicación son:

- Número máximo de entradas y salidas (digitales).
- Memoria de programa.
- Memoria de datos.
- Velocidad de proceso (μ s por instrucción).
- Posibilidad de entradas/salidas especiales (analógicas, PID, contador de alta velocidad, salidas de pulsos PWM).
- Lenguaje de programación (diagrama de contactos, lista de instrucciones, Grafcet, Basic) y dispositivos de programación (ordenador, consola).
- Comunicaciones. Posibilidad de conexión en redes de comunicación (E/S remotas, bus de campo, red de área local, red de nivel superior).
- Capacidad para implementar funciones matemáticas complejas.

Los criterios fundamentales a tener en cuenta para la selección del autómat más adecuado según las características anteriores son:

- Número y tipo de E/S a controlar. El PLC debe admitir el número de entradas y salidas que se necesitan y algunas más para posibles imprevistos o ampliaciones o modificaciones futuras.
- Complejidad (tamaño) del programa de control que se quiere implementar. El programa de control debe caber en la memoria, y debe sobrar algo de espacio para imprevistos o posibles ampliaciones y modificaciones futuras.
- Potencia de las instrucciones necesarias (funciones matemáticas, rapidez de respuesta). Hay que tener en cuenta si se necesita hacer cálculos complejos (senos, raíces cuadradas, etc.) o si el algoritmo de control necesita ejecutarse especialmente rápido por alguna característica del proceso.

- Necesidad de utilización de periféricos (como terminales táctiles, lectores de código de barras, etc.), módulos especiales de entrada o salida y módulos de comunicación. El autómata elegido debe admitir los módulos especiales y de comunicación que se necesiten.
- Precio. Evidentemente, a mayor precio, mayores prestaciones (en tamaño de memoria, en capacidad de cálculo, en disponibilidad de módulos, etc.). Se debe llegar a un compromiso entre las prestaciones y el precio final. Este compromiso dependerá del número de unidades que se vayan a fabricar. Si el número de unidades es alto, será conveniente ajustar el precio al máximo, pero si se trata de muy pocas unidades, o de un proyecto singular, es conveniente no ajustar tanto el precio y dejar margen para posibles cambios imprevistos o ampliaciones.
- Posibilidad de ampliaciones futuras. Como ya se ha comentado, incluso en el caso de que no estén previstas ampliaciones futuras determinadas, es conveniente dejar algo de margen en el tamaño de memoria y número de entradas/salidas para imprevistos. No obstante, en muchas ocasiones se pueden prever ampliaciones futuras del proceso en las que se puede evaluar las necesidades en cuanto al autómata de control. En ese caso, se debe tener en cuenta en la elección, que esas necesidades futuras se puedan cubrir con el PLC elegido.

5.9 CARACTERÍSTICAS DE UN PLC COMERCIAL. CQM1 DE OMRON

El CQM1 es un autómata de la gama media-baja. El módulo de CPU tiene además del procesador y la memoria, un puerto de periféricos, un puerto serie, y 16 entradas digitales. Dispone de módulos de expansión de todo tipo: E/S digitales, E/S analógicas, E/S remotas, control de temperatura (PID), módulos de comunicación, etc.

Las características básicas del CQM1 (con la CPU21) son:

- Memoria de programa de 3.2 KW (palabras de 16 bits).
- Memoria de datos no volátil de 1 KW.
- 16 entradas digitales integradas en el módulo de CPU.
- Número máximo de entradas/salidas, 128 (número máximo contando entradas y salidas).
- Tiempo de ejecución de instrucciones: instrucciones básicas de 0.5 μ s a 1.5 μ s.
- 4 interrupciones externas. Proceso de interrupción inmediato.

- 512 temporizadores/contadores.
- Temporizadores de intervalo (para ejecutar acciones periódicas).
- Posibilidad de ciclo de scan con salida directa.
- Programación por diagrama de relés (o lista de instrucciones).

Una configuración básica de CQM1 para tener entradas y salidas digitales y analógicas, es la mostrada en la figura 5.25, que dispone de los siguientes módulos: módulo de 16 salidas digitales a relé (OC222), módulo de 4 entradas analógicas (AD041), módulo de 2 salidas analógicas (DA021) y módulo de fuente de alimentación adicional (para los módulos analógicos IPS02). No es necesario un módulo de entradas digitales puesto que la CPU incorpora 16.

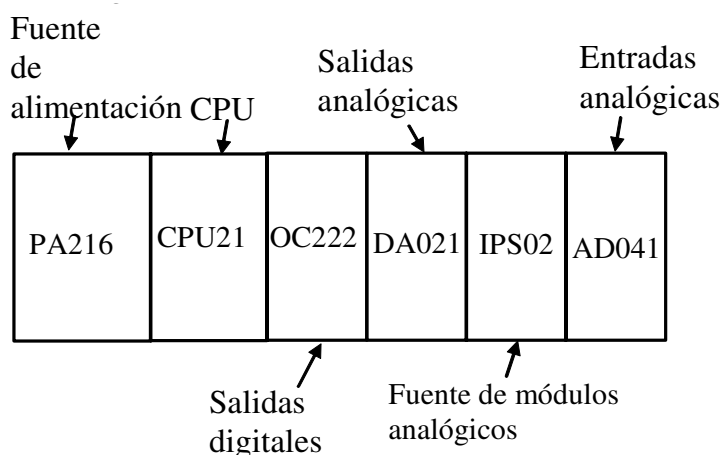


Figura 5.25: Una configuración básica de ejemplo del autómata CQM1.

Área de memoria de datos del CQM1

La memoria de datos es de 16 bits, es decir, está formada por palabras, registros o canales de 16 bits. El direccionamiento (identificación de un canal concreto) se hace por medio del número de dirección. Cuando se programan los contactos y las funciones hay que utilizar la sintaxis correcta de direccionamiento para referirse a una variable de memoria concreta.

El área de memoria de datos se divide en varias categorías. Cada área tiene una función específica y se direcciona de una manera diferente:

IR. Es el área donde están la memoria de entradas y salidas, así como las variables internas utilizadas en forma de bits auxiliares. Se direccionan con cinco números. Los tres primeros indican la palabra o canal (un número del 000 al 229). Los dos últimos indican el bit dentro de esa

palabra (un número del 0 al 15). Por ejemplo, el bit 013.08 es el bit nº 8 del canal 013.

Las entradas están a partir de la dirección 000.00 en el orden en que se colocan los módulos de entradas. En la configuración básica descrita como ejemplo en la sección anterior, **las entradas digitales irían de la dirección 000.00 a la 000.15**. Las entradas analógicas (de 12 bits) irían del canal 001 al canal 004 (en cada uno de esos canales estarían los 12 bits del resultado de la conversión de cada una de las entradas analógicas). Las salidas digitales están situadas desde la dirección 100.00 en adelante. En el ejemplo de configuración propuesto se tendrían **las salidas digitales desde la dirección 100.00 a la 100.15**. Las dos salidas analógicas (de 12 bits) se escribirían en los canales 101 y 102. Los demás bits del área de memoria IR se pueden utilizar como bits internos auxiliares en el programa de usuario.

SR. En realidad es una continuación del área IR, que va del canal 244 al canal 255. Los bits de esta zona tienen funciones específicas, muy útiles para la programación. Hay bits que siempre están a 1, o que siempre están a 0. Hay bits que cambian a frecuencia constante (cada 0.1 seg, o cada 1 seg, etc.), que se pueden utilizar para temporizar (utilizando un contador). Hay un bit que está a 1 sólo durante el primer ciclo de scan, pasando después a 0, por lo que se puede utilizar para ejecutar una rutina de inicialización (que solo se ejecuta una vez cuando se pone en marcha el autómatas). **Se direccionan igual que los IR.**

AR. Es un área de registros auxiliares. Se direccionan como **AR xxyy**, donde yy es el número de canal (entre 00 y 27), y xx es el número de bit dentro del canal. Los bits de esta área tienen funciones especiales.

LR. Es un área de datos utilizada para comunicación 1 a 1 entre dos autómatas por el puerto serie. Si no se utilizan para esto se pueden utilizar como bits internos. Se direccionan como **LR xxyy**, donde xx es el número de canal (entre 00 y 63) e yy es el número de bit dentro del canal (entre 00 y 15).

TIM/CNT. En este área se almacenan los valores de los temporizadores y contadores. Se direccionan como TC 000 a TC 511. Una vez definido en un programa como temporizador o contador, se direcciona como **TIM xxx** o **CNT xxx**. Si se utiliza para designar un contacto, éste representa el bit de salida del temporizador o contador (que se pone a 1 cuando la cuenta llega a 0). Si se utiliza como operando de canal significa el valor actual de ese temporizador (16 bits). No se puede utilizar el mismo número (la misma dirección) para un contador y un temporizador.

HR. Es un área de datos de bits de retención, que no se pierden si se quita la alimentación. Se direccionan como **HR xxyy** donde xx es el canal (entre

00 y 99) e yy es el bit (entre 00 y 15). Se puede utilizar igual que los bits del área IR, pero con la diferencia de que no se borran al desconectarse el autómata.

DM. Es un área de datos no volátil (no se borra al apagar el PLC) a la que solo se puede acceder por canales (no por bits individuales). Se direcciona como **DM xxxx** (donde xxxx es un número del 0000 al 1023). Se utiliza para almacenar datos, teniendo la particularidad de que no se pierden aunque se apague el autómata.

Instrucciones de programación

Las figuras 5.27, 5.28 y 5.29 recogen el conjunto completo de instrucciones (funciones) disponibles para programar el CQM1. Cada instrucción tiene uno o varios operandos, que son las variables (posiciones de memoria) sobre las que actúa. Estos operandos pueden ser bits individuales o canales de 16 bits. Algunos operandos pueden no ser una posición de memoria, sino un valor numérico concreto. En ese caso se escribe el signo # delante del número. Por ejemplo, la figura 5.26 muestra el uso de la función MOV que actúa sobre canales de 16 bits.

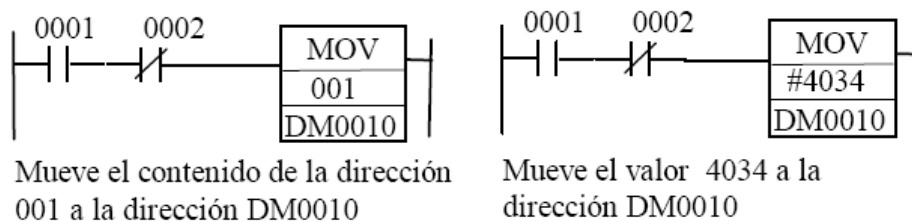


Figura 5.26: Uso de la función MOV.

Algunas funciones pueden ser activas por nivel o por flanco de subida. Cuando son activas por flanco de subida tienen una @ delante. Esas funciones solo se ejecutan una vez; cuando la condición de ejecución cambia de 0 a 1.

Las instrucciones más importantes del CQM1 son:

Instrucciones básicas de relés

Las instrucciones básicas de relés son: LD, LD NOT, AND, AND NOT, OR, OR NOT, OR LD, AND LD; ya descritas en secciones anteriores.

Instrucciones básicas de control de bit

OUT. Asigna a una variable de un bit el valor de la condición de activación.

OUT NOT. Asigna a una variable de un bit el valor negado de la condición de activación.

■ Instrucciones de secuencia

Instrucciones de diagramas de relés

Nombre	Nemónico	Código
LOAD	LD	Nota 1.
LOAD NOT	LD NOT	
AND	AND	
AND NOT	AND NOT	
OR	OR	
OR NOT	OR NOT	
AND LOAD	AND LD	
OR LOAD	OR LD	

Instrucciones de control de bit

Nombre	Nemónico	Código
OUTPUT	OUT	Nota 1.
OUTPUT NOT	OUT NOT	Nota 1.
RELÉ DE ENCLAVAMIENTO	KEEP	11
FLANCO ASCENDENTE	DIFU	13
FLANCO DESCENDENTE	DIFD	14
SET	SET	Nota 1.
RESET	RSET	Nota 1.

Instrucciones de control de secuencia

Nombre	Nemónico	Código
END	END	01
NO OPERACIÓN	NOP	00
ENCLAVAMIENTO	IL	02
BORRAR ENCLAVAMIENTO	ILC	03
SALTO	JMP	04
FIN DE SALTO	JME	05

■ Instrucciones de temporizador/contador

Nombre	Nemónico	Código
TEMPORIZADOR	TIM	Nota 1.
TEMPORIZ. ALTA VELOCIDAD	TIMH	15
TEMPORIZADOR TOTALIZADOR	TTIM	Nota 2/3.
CONTADOR	CNT	Nota 1.
CONTADOR REVERSIBLE	CNTR	12

■ Instrucciones de comparación

Nombre	Nemónico	Código
COMPARAR	CMP	20
COMPARAR DOS A DOS	CMPL	(60: Nota 2)
COMPARACIÓN BINARIA CON SIGNO	CPS	Nota 2.
DOBLE COMPARACIÓN BINARIA CON SIGNO	CPSL	Nota 2.
COMPARACIÓN DE BLOQUE	(@)MCMP	(19: Nota 2)
COMPARAR TABLA	(@)TCMP	85
COMPARAR BLOQUE	(@)BCMP	(68: Nota 2)
COMPARAR RANGO DE ÁREA	ZCP	Nota 2
COMPARACIÓN DE RANGO DOS A DOS	ZCPL	Nota 2

■ Instrucciones de transferencia de datos

Nombre	Nemónico	Código
MOVER	(@)MOV	21
MOVER NEGADO	(@)MVN	22
MOVER BIT	(@)MOVB	82
MOVER DÍGITO	(@)MOVD	83
TRANSFERENCIA DE BITS	(@)XFRB	Nota 2.
TRANSFERENCIA DE BLOQUE	(@)XFER	70
RELLENAR BLOQUE	(@)BSET	71
INTERCAMBIO DE DATOS	(@)XCHG	73
DISTRIBUCIÓN DE DATOS	(@)DIST	80
RECOGIDA DE DATOS	(@)COLL	81

Figura 5.27: Instrucciones del autómata CQM1 de Omron Electronics. Las que incluyen la nota 3 solo están disponibles en el CQM1H.

■ Instrucciones de desplazamiento datos

Nombre	Nemónico	Código
REGISTRO DE DESPLAZAMIENTO	SFT	10
REGISTRO DESPLAZ. REVERSIBLE	(@)SFTR	84
DESPLAZ. UN DÍGITO A IZQDA	(@)SLD	74
DESPLAZ. UN DÍGITO A DRCHA	(@)SRD	75
REGISTRO DESPLAZ. ASÍNCRONO	(@)ASFT	(17: Nota 2)
DESPLAZAMIENTO DE CANAL	(@)WSFT	16
DESPLAZ. BINARIO A IZQUIERDA	(@)ASL	25
DESPLAZ. BINARIO A DERECHA	(@)ASR	26
ROTAR A IZQUIERDA	(@)ROL	27
ROTAR A DERECHA	(@)ROR	28

■ Instrucciones de Incremento/Decremento

Nombre	Nemónico	Código
INCREMENTAR EN BCD	(@)INC	38
DECREMENTAR EN BCD	(@)DEC	39

■ Instrucciones de cálculo

Nombre	Nemónico	Código
SUMA BCD	(@)ADD	30
RESTA BCD	(@)SUB	31
SUMA BCD DOBLE	(@)ADDL	54
RESTA BCD DOBLE	(@)SUBL	55
SUMA BINARIA	(@)ADB	50
RESTA BINARIA	(@)SBB	51
SUMA BINARIA DOBLE	(@)ADBL	Nota 2.
RESTA BINARIA DOBLE	(@)SBSL	Nota 2.
MULTIPLICACIÓN BCD	(@)MUL	32
MULTIPLICACIÓN BCD DOBLE	(@)MULL	56
MULTIPLICACIÓN BINARIA	(@)MLB	52
MULTIPLICACIÓN BINARIA CON SIGNO	(@)MBS	Nota 2.
MULTIPLICACIÓN BINARIA DOBLE CON SIGNO	(@)MBSL	Nota 2.
DIVISIÓN BCD	(@)DIV	33
DIVISIÓN BCD DOBLE	(@)DIVL	57
DIVISIÓN BINARIA	(@)DVB	53
DIVISIÓN BINARIA CON SIGNO	(@)DBS	Nota 2.
DIVISIÓN BINARIA DOBLE CON SIGNO	(@)DBSL	Nota 2.

■ Instrucciones de cálculo especiales

Nombre	Nemónico	Código
OPERACIONES ARITMÉTICAS	(@)APR	Nota 2
CONTADOR DE BIT	(@)BCNT	(67: Nota 2)
RAÍZ CUADRADA	(@)ROOT	72

■ Instrucciones de conversión

Nombre	Nemónico	Código
BCD A BINARIO	(@)BIN	23
BCD A BINARIO DE 2 CANALES	(@)BINL	58
BINARIO A BCD	(@)BCD	24
BINARIO A BCD DE 2 CANALES	(@)BCDL	59
COMPLEMENTO A 2	(@)NEG	Nota 2
DOBLE COMPLEMENTO A 2	(@)NEGL	Nota 2
DECODIFICADOR 4 A 16	(@)MLPX	76
CODIFICADOR 16 A 4	(@)DMPX	77
COVERTIR A ASCII	(@)ASC	86
ASCII A HEXADECIMAL	(@)HEX	Nota 2
COLUMNA A LÍNEA	(@)LINE	Nota 2
LÍNEA A COLUMNA	(@)COLM	Nota 2

■ Instrucciones lógicas

Nombre	Nemónico	Código
PRODUCTO LÓGICO	(@)ANDW	34
SUMA LÓGICA	(@)ORW	35
SUMA LÓGICA EXCLUSIVA	(@)XORW	36
SUMA LÓGICA EXCLUSIVA NEG.	(@)XNRW	37
COMPLEMENTO	(@)COM	29

■ Instrucciones de conversión y matemáticas de coma flotante

Nombre	Nemónico	Código
COMA FLOTANTE A 16-BIT	(@)FIX	Nota 2/3
COMA FLOTANTE A 32-BIT	(@)FIXL	
16-BIT A COMA FLOTANTE	(@)FLT	
32-BIT A COMA FLOTANTE	(@)FLTL	
SUMA EN COMA FLOTANTE	(@)+F	
RESTA EN COMA FLOTANTE	(@)-F	
MULTIPLICACIÓN EN COMA FLOTANTE	(@)*F	
DIVISIÓN EN COMA FLOTANTE	(@)/F	
GRADOS A RADIANES	(@)RAD	
RADIANES A GRADOS	(@)DEG	
SENO	(@)SIN	
COSENO	(@)COS	
TANGENTE	(@)TAN	
ARCO SEÑO	(@)ASIN	
ARCO COSENO	(@)ACOS	
ARCO TANGENTE	(@)ATAN	
RAÍZ CUADRADA	(@)SQRT	
EXPONENCIAL	(@)EXP	
LOGARITMO	(@)LOG	

Figura 5.28: Instrucciones del autómata CQM1 de Omron Electronics. Las que incluyen la nota 3 solo están disponibles en el CQM1H.

■ Instrucciones sobre tablas

Nombre	Nemónico	Código
BÚSQUEDA DE DATOS	(@)SRCH	Nota 2
ENCONTRAR MÁXIMO	(@)MAX	
ENCONTRAR MÍNIMO	(@)MIN	
SUMA	(@)SUM	
CHECKSUM DE BLOQUE	(@)FCS	

■ Instrucciones de control de datos

Nombre	Nemónico	Código
CONTROL PID	PID	Nota 2
ESCALA	(@)SCL	(66: Nota 2)
ESCALA 2	(@)SCL2	Nota 2
ESCALA 3	(@)SCL3	Nota 2
VALOR MEDIO	AVG	Nota 2

■ Instrucciones de subrutina

Nombre	Nemónico	Código
LLAMADA A SUBROUTINA	(@)SBS	91
PRINCIPIO DE SUBROUTINA	SBN	92
FINAL DE SUBROUTINA	RET	93
MACRO	(@)MCRO	99

■ Instrucciones de interrupción

Nombre	Nemónico	Código
CONTROL DE INTERRUPTIÓN	(@)INT	(89: Nota 2)
TEMPORIZADOR DE INTERVALO	(@)STIM	(69: Nota 2)

■ Instrucciones de contador de alta velocidad y control de salida de pulsos

Nombre	Nemónico	Código
CONTROL DE MODO	(@)INI	(61: Nota 2)
LECTURA DE PV	(@)PRV	(62: Nota 2)
REGISTRAR TABLA COMPARACIÓN	(@)CTBL	(63: Nota 2)
NÚMERO DE PULSOS	(@)PULS	(65: Nota 2)
FRECUENCIA DE PULSOS	(@)SPED	(64: Nota 2)
CONTROL DE ACELERACIÓN	(@)ACC	Nota 2
SALIDA DE PULSOS	(@)PLS2	Nota 2
PULSOS RELACIÓN ON/OFF VARIABLE	(@)PWM	Nota 2

■ Instrucciones de paso

Nombre	Nemónico	Código
DEFINIR PASO	STEP	08
INICIAR PASO	SNXT	09

■ Instrucciones de unidad de E/S

Nombre	Nemónico	Código
REFRESCO DE E/S	(@)IORF	97
DECODIFICADOR 7-SEGMENTOS	(@)SDEC	78
SALIDA DE DISPLAY DE 7-SEGMENTOS	7SEG	(88: Nota 2)
ENTRADA DE DÉCADAS DE SELECCIÓN	DSW	(87: Nota 2)
ENTRADA DE TECLADO DECIMAL	(@)TKY	(18: Nota 2)
ENTRADA DE TECLADO HEXADECIMAL	HKY	Nota 2

■ Instrucciones de comunicaciones serie

Nombre	Nemónico	Código
MACRO DE PROTOCOLO	(@)PMCR	Nota 2/3
TRANSMITIR	(@)TXD	(48: Nota 2)
RECIBIR	(@)RXD	(47: Nota 2)
CAMBIAR SETUP DE PUERTO SERIE	(@)STUP	Nota 2/3

■ Instrucciones de comunicaciones de red

Nombre	Nemónico	Código
ENVIAR A RED	(@)SEND	90 (Nota 3)
RECIBIR DE RED	(@)RECV	98 (Nota 3)
ENTREGAR COMANDO	(@)CMND	Nota 2/3

■ Instrucciones de mensaje

Nombre	Nemónico	Código
MENSAJE	(@)MSG	46

■ Instrucciones de reloj

Nombre	Nemónico	Código
HORAS A SEGUNDOS	(@)SEC	Nota 2
SEGUNDOS A HORAS	(@)HMS	Nota 2

■ Instrucciones de seguimiento

Nombre	Nemónico	Código
SEGUIMIENTO DE DATOS	TRSM	45

■ Instrucciones de diagnóstico

Nombre	Nemónico	Código
ALARMA DE FALLO	(@)FAL	06
ALARMA DE FALLO GRAVE	FALS	07
DETECCIÓN DE FALLOS	FPD	Nota 2

■ Instrucciones de acarreo

Nombre	Nemónico	Código
ACARREO ON	(@)STC	40
ACARREO OFF	(@)CLC	41

Figura 5.29: Instrucciones del autómata CQM1 de Omron Electronics. Las que incluyen la nota 3 solo están disponibles en el CQM1H.

SET. Activa una variable de un bit.

RSET. Desactiva una variable de un bit.

DIFU. Activa una variable de un bit durante un solo ciclo de scan cuando la condición de activación pasa de 0 a 1 (flanco de subida).

DIFD. Activa una variable de un bit durante un solo ciclo de scan cuando la condición de activación pasa de 1 a 0 (flanco de bajada).

KEEP. Tiene dos condiciones de activación (una es la entrada Set y la otra la entrada Reset). La salida se activa o desactiva en función del valor de las dos entradas (R y S) actuando como un biestable tipo RS.

Instrucciones básicas de temporizadores/contadores

TIM. Temporizador de retardo a la conexión. La salida se activa al cabo de un tiempo desde que se ha puesto a 1 la condición de activación. El tiempo se cuenta en periodos de 100 ms.

TIMH. Igual que TIM, pero permite contar intervalos de tiempo más cortos (temporizador de alta velocidad). El tiempo se cuenta en periodos de 10 ms.

CNT. Contador descendente. Cada vez que la condición de ejecución pasa de 0 a 1, el contador se decrementa.

CNTR. Contador reversible. Igual que el anterior, pero con dos entradas de conteo: una para incrementar y la otra para decrementar.

Instrucciones básicas de control de programa

END. Fin del programa.

JMP. Salto. Cuando la condición de ejecución es 0 se produce el salto hasta instrucción JME con el mismo número. Si la condición es 1, el programa continúa (no hay salto).

JME. Punto de llegada de un salto.

IL. Si la condición de ejecución es 0, todas las redes entre ésta y la siguiente instrucción ILC se fuerzan a 0, independientemente del valor que tengan. Si la condición es 1, no se hace nada.

ILC. Indica el final de la zona de actuación de la instrucción IL anterior.

Instrucciones de comparación de datos

CMP. Compara dos canales (variables de 16 bits). El resultado de la comparación afecta al bit EQ (se pone a 1 si los valores son iguales), al bit LE (se pone a 1 si el primer valor es menor) y al bit GR (se pone a 1 si el primer valor es mayor). Cualquiera de estos bits se puede utilizar en la instrucción siguiente para hacer algo o no en función de la comparación.

Instrucciones de transferencia de datos

MOV. Transfiere un canal o un valor constante a otro canal.

MOVB. Transfiere un bit de un canal a otro bit de otro canal.

Instrucciones matemáticas (en binario y en BCD)

ADB. Suma binaria de dos canales, o de un canal y una constante.

ADD. Suma en BCD de dos canales o un canal y una constante.

SBB. Resta binaria de dos canales, o de un canal y una constante.

SUB. Resta en BCD de dos canales o un canal y una constante.

MLB. Multiplicación binaria de dos canales, o de un canal y una constante.

MUL. Multiplicación en BCD de dos canales o un canal y una constante.

DVB. División binaria de dos canales, o de un canal y una constante.

DIV. División en BCD de dos canales o un canal y una constante.

INC. Incremento en BCD de un canal.

DEC. Decremento en BCD de un canal.

Instrucciones de conversión de datos

BIN. Convierte de BCD a binario (hexadecimal).

BCD. Convierte de binario (hexadecimal) a BCD.

Instrucciones lógicas (de canales)

ANDW. Realiza la función lógica AND entre dos canales, bit a bit.

ORW. Realiza la función lógica OR entre dos canales, bit a bit.

XORW. Realiza la función lógica XOR (o exclusiva) entre dos canales, bit a bit.

Instrucciones de subrutina e interrupción

INT. Sirve para configurar interrupciones.

STIM. Temporizador para contar el tiempo entre interrupciones periódicas.

SBS. Llama a una subrutina desde el programa principal.

SBN. Indica el comienzo de una subrutina.

RET. Indica el final de una subrutina.

Instrucciones especiales

IORF. Refresca las entradas y salidas del autómata.

APR. Realiza cálculos de seno o coseno.

PID. Implementa un algoritmo de control PID.

PWM. Produce una salida PWM del ciclo de trabajo especificado.

CAPÍTULO 6

IMPLEMENTACIÓN DE SISTEMAS DE CONTROL SECUENCIAL MEDIANTE DISPOSITIVOS PROGRAMABLES

6.1 INTRODUCCIÓN

El Grafcet es un modelo del controlador lógico que resuelve un problema de automatización. Para implementar ese controlador es necesario programar un algoritmo en el equipo de control, de forma que al ejecutarse dicho algoritmo, se cumplan todas las reglas de evolución de ese Grafcet modelo. En particular, hay que prestar mayor atención en dos situaciones especiales, que son: la posible existencia de etapas inestables y la posible existencia de transiciones activas por flanco.

En el programa se asocia una variable interna (1 bit) a cada etapa (el bit estará a 1 si la etapa está activa y a 0 en caso contrario). Las ecuaciones lógicas que forman el programa tratarán de ir modificando esos bits según cambian las entradas de acuerdo a las reglas de evolución del Grafcet.

El algoritmo se ejecutará en un bucle infinito, de forma que estará continuamente comprobando las entradas para modificar las etapas y las salidas. Si el algoritmo se implementa mediante un autómata programable hay que tener en cuenta que el bucle infinito lo realiza internamente el autómata, y únicamente se deben escribir las instrucciones del algoritmo (sin el bucle).

6.2 IMPLEMENTACIÓN DEL ALGORITMO DE CONTROL A PARTIR DEL GRAFCET

Se trata de desarrollar un programa que reproduzca la forma de evolucionar de las distintas variables del Grafcet, siguiendo todas sus reglas de evolución con exactitud. Para ello, hay que poner especial atención en dos situaciones especiales, que son la existencia de etapas inestables y la existencia de transiciones activas por flanco.

Existen varios métodos más o menos sistemáticos para desarrollar el programa cumpliendo lo anterior.

En todos los casos se asocia una variable interna (1 bit) a cada etapa (el

bit estará a 1 si la etapa está activa y a 0 en caso contrario). Las ecuaciones lógicas que forman el programa se encargarán de ir modificando los valores de esos bits, que definen las etapas activas, en función del valor que van tomando las entradas y del valor que tienen esas mismas etapas, cumpliendo en todo momento las reglas de evolución del Grafcet. Las ecuaciones lógicas que forman el programa también se encargarán de definir el valor que deben tomar las salidas en función del valor de las etapas y de las entradas.

El algoritmo se ejecuta en un bucle infinito, de forma que está continuamente comprobando las entradas para modificar las etapas y las salidas.

Se describirán dos métodos. Uno simplificado que únicamente es válido cuando no se dan las dos situaciones especiales comentadas anteriormente (la existencia de etapas inestables y la existencia de transiciones activas por flanco), pero que resulta en un algoritmo muy sencillo. El segundo método, más complejo, sirve para cualquier situación (etapas inestables, flancos, etc.).

En los dos casos se describirá en primer lugar el algoritmo resultante suponiendo que la implementación se realiza en un autómata programable industrial. Posteriormente se comentarán las diferencias para implementarlo en otra plataforma.

MÉTODO 1

Este método es el más simple. A continuación se describen los bloques que contiene el programa en el caso de implementarse en un PLC. En la ejecución del programa, al terminar el bloque 5 se vuelve al bloque 2 y se repite el algoritmo de forma indefinida. Ese bucle infinito no se incluye aquí porque el PLC lo produce de forma automática. Si se utilizara otra plataforma, habría que incluir de forma explícita el salto desde el bloque 5 al bloque 2 para implementar el bucle infinito.

1. **Inicialización** de las etapas iniciales mediante el bit de inicio. Este bloque sólo debe ejecutarse una vez, en el primer ciclo de scan. En este bloque se ponen a 1 las etapas iniciales de todos los Grafcet (las marcadas con doble cuadrado), poniéndose a 0 las demás.
2. **Detección de flancos**, tanto de las entradas que dan lugar a transiciones por flanco como de las etapas que tienen acciones impulsionales asociadas. En este bloque, de cada variable cuyo flanco se tenga que calcular se obtiene una variable de bit que valdrá 1 únicamente durante el siguiente ciclo de scan si la variable ha pasado de 0 a 1 o de 1 a 0 (según sea un flanco ascendente o descendente).
3. **Desactivación/activación de las etapas** anteriores/posteriores a las transiciones franqueables. En este bloque es donde se modifican las etapas, produciéndose la evolución del Grafcet.

4. **Activación de las salidas.** Se ponen a 1 las salidas (acciones a nivel) asociadas a las etapas que están activas. También se incluyen aquí las acciones impulsionales, como por ejemplo, incrementar un contador. Esas acciones se deben ejecutar condicionadas al flanco de subida de la etapa correspondiente.
5. **Definición de temporizadores y contadores.** Aquí se incluyen las funciones que definen los contadores y temporizadores utilizados.

Funciona bien siempre que:

- **No existan estados inestables.** Si los hay, en un solo ciclo de scan pueden avanzarse varias etapas, por lo que las acciones impulsionales no se producirían.
- **No haya transiciones por flanco consecutivas.** Con este método una transición por flanco está activa en el autómata durante un ciclo completo de scan. Esto hace que un solo flanco pueda hacer evolucionar varias etapas consecutivas.
- **No existan transiciones que puedan desactivar y activar a la vez una etapa.**
- **No haya bifurcaciones de selección no excluyentes.** Si hay bifurcaciones no excluyentes, la primera que se evalúe será la única que se active, ya que la etapa anterior quedará desactivada.
- **No haya transiciones que dependan de etapas.**

En esencia, el método simplificado 1 sirve para problemas sencillos que se resuelven normalmente con 1 solo Grafcet. El ejemplo 1 sirve de muestra de este tipo de implementación.

En el bloque 3 **debe haber una ecuación (red) para cada transición** del Grafcet. La ecuación de activación/desactivación de una transición debe tener en cuenta todas las etapas anteriores y todas las posteriores, tal y como se muestra en la figura 6.1.

En el bloque 4 se introduce **una red para cada salida** que esté definida como acción a nivel. En ese caso, la red debe incluir todas las etapas en las que la salida está activa, incluyendo además las condiciones lógicas correspondientes si la salida está condicionada a alguna variable. Por ejemplo, si la salida Y debe estar activa en la etapa 3 y en la etapa 2 condicionada a B, la ecuación de dicha salida sería la expuesta en la figura 6.2.

Si se trata de una acción impulsional (por ejemplo incrementar un contador), la condición de activación debe ser el flanco de subida de la etapa, que se genera en el bloque 2 (con la instrucción DIFU o DIFD en el autómata OMRON, o con un contacto de pulso en otros autómatas). Por ejemplo, si la

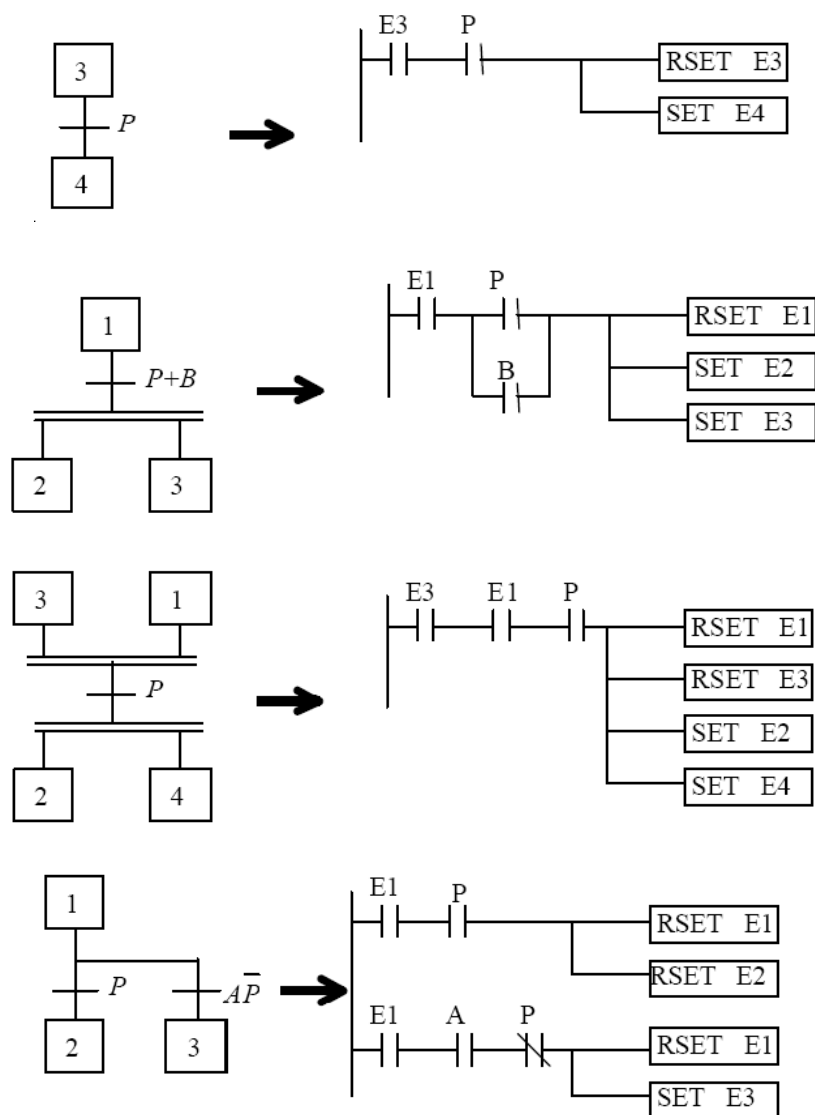


Figura 6.1: Implementación del Grafcet por el método 1. Representación de las transiciones.

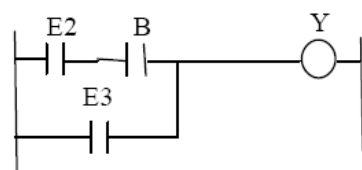


Figura 6.2: Implementación del Grafcet por el método 1. Representación de salidas.

etapa E3 tiene asociada la acción impulsional de incrementar la variable X, en el bloque 2 habría que incluir la detección del flanco de E3, tal y como se muestra en la figura 6.3, mientras que en el bloque 4 se pondría la acción impulsional condicionada al flanco, como se puede observar en la figura 6.4.

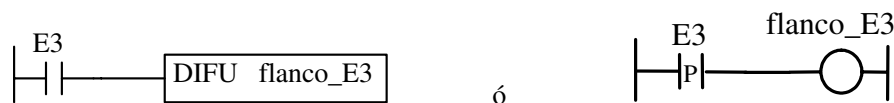


Figura 6.3: Implementación de la detección de un flanco en una etapa.



Figura 6.4: Implementación de una salida impulsional.

Ejemplo 1

Sea el cilindro de la figura 6.5. Inicialmente se supone que el cilindro está

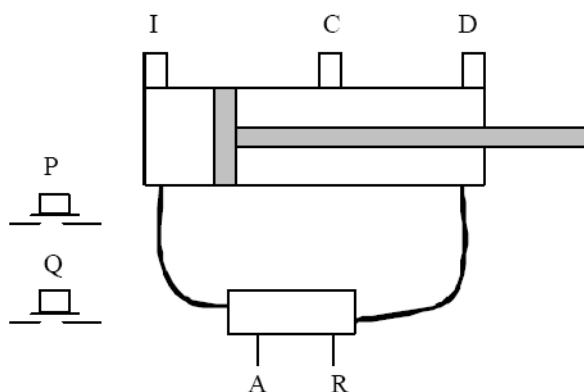


Figura 6.5: Sistema del ejemplo 1. Control de un cilindro neumático.

en la posición más a la izquierda. Al pulsar P se debe mover hacia la derecha hasta que se active el detector C, para volver después a la izquierda. Si se pulsa Q debe hacer lo mismo, pero esperando 1 segundo antes de empezar a moverse, y llegando a D en lugar de C, para mover el cilindro se actúa sobre las electroválvulas A y R.

Una posible solución es la representada en la figura 6.6. En las figuras 6.7 y 6.8, se puede ver el programa que implementa el Grafcet anterior siguiendo

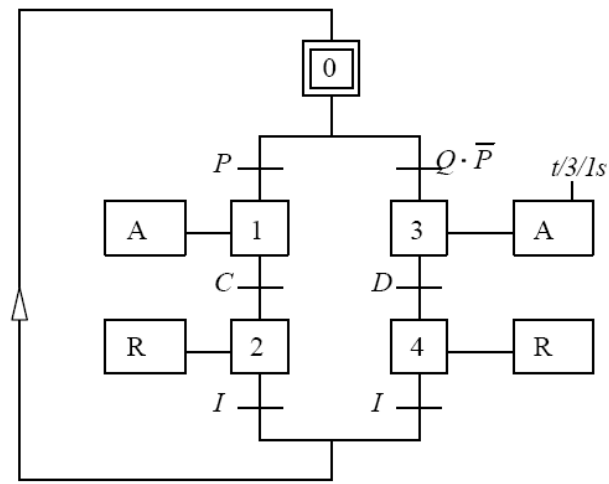


Figura 6.6: Grafcet que resuelve el ejemplo 1.

el método 1, realizado para un autómata CQM1 de Omron. Se observa que la implementación es muy simple. En este caso, no hay problemas por no darse ninguno de los casos conflictivos descritos anteriormente.

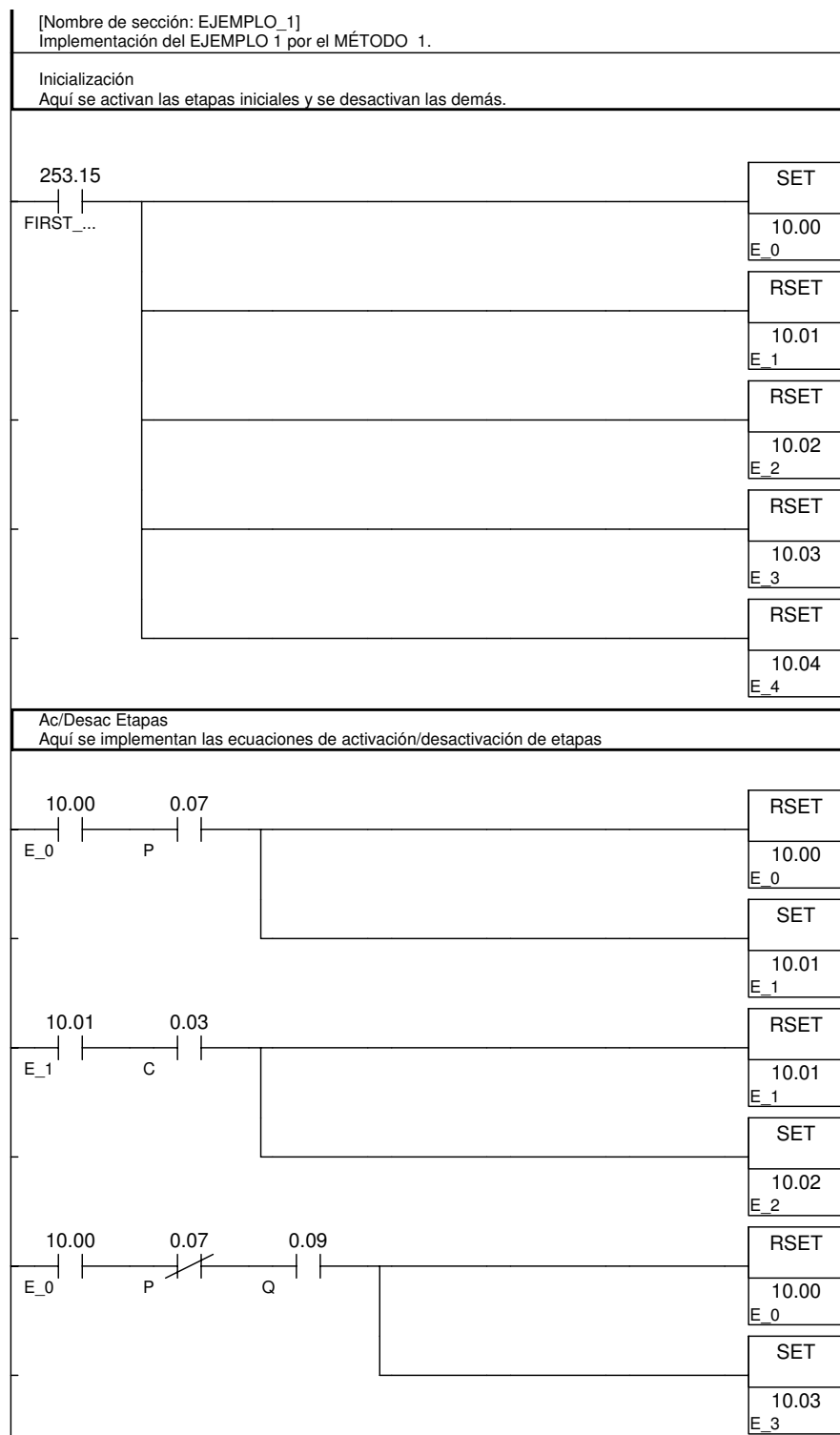


Figura 6.7: Diagrama de contactos que implementa la solución del ejemplo 1.

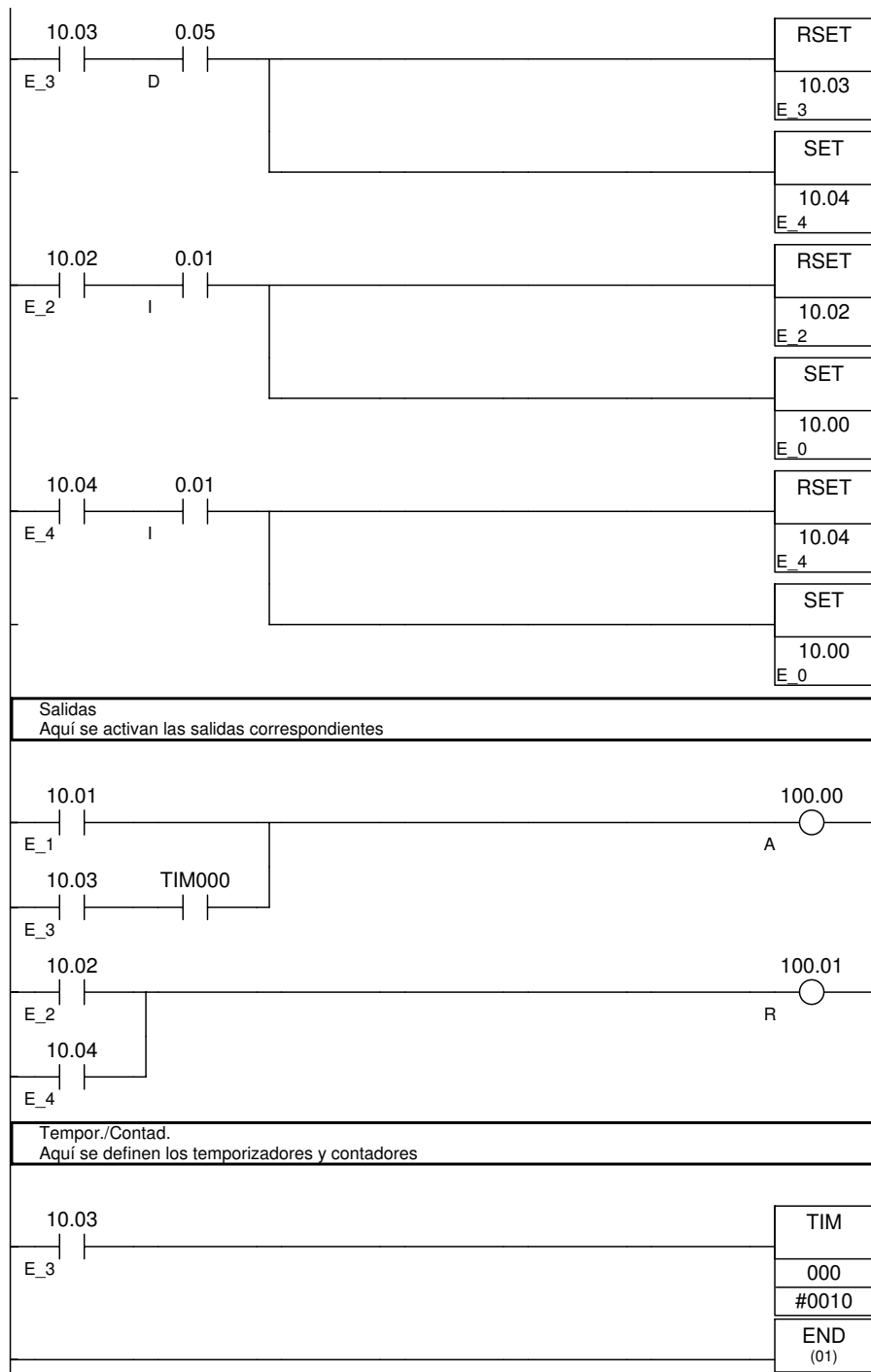


Figura 6.8: Diagrama de contactos que implementa la solución del ejemplo 1.

MÉTODO 2

Este método da lugar a un programa más largo, pero funciona bien aunque haya estados inestables o transiciones por flanco. Los bloques que contiene el algoritmo son:

1. **Inicialización** de las etapas mediante el bit de inicio. Este bloque sólo debe ejecutarse una vez, en el primer ciclo de scan. En este bloque se ponen a 1 las etapas iniciales de todos los Grafcets (las marcadas con doble cuadrado), poniéndose a 0 las demás.
2. **Copia de las variables** que definen las *etapas* en otras variables duplicadas (*copia_de_etapas*).
3. **Detección de flancos** de las entradas que dan lugar a transiciones por flanco. En este bloque, de cada variable cuyo flanco se tenga que calcular se obtiene una variable de bit que valdrá 1 únicamente durante el siguiente ciclo de scan si la variable ha pasado de 0 a 1 o de 1 a 0 (según sea un flanco ascendente o descendente).
4. **Desactivación** de las etapas anteriores a las transiciones franqueables, utilizando para las validaciones las variables *copia_de_etapas*. Se tiene que incluir una red (ecuación) lógica para cada transición.
5. **Activación** de las etapas posteriores a las transiciones franqueables, utilizando para las validaciones las variables *copia_de_etapas*. De esta forma, solo se avanza una etapa cada vez que se ejecuta el bloque completo. La activación va después de la desactivación para garantizar que sea prioritaria. Se tiene que incluir una red (ecuación) lógica para cada transición.
6. **Ejecución de las acciones impulsionales** que correspondan. Aquí hay que incluir todas las acciones que se han de ejecutar una sola vez cuando se activa la etapa. Un ejemplo podría ser incrementar un contador que cuenta el número de veces que se ha activado la etapa (es una acción impulsional asociada al paso de 0 a 1 de la etapa). La condición de ejecución de las acciones impulsionales es: $etapa \cdot copia_de_etapa$.
7. **Si ha habido cambios en alguna etapa (estado inestable) volver al bloque 2**. Es decir, se vuelven a ejecutar los bloques anteriores hasta que $etapas = copia_de_etapas$, es decir, hasta que se llegue a un estado estable.
8. **Activación de las acciones a nivel**. En este caso ya se está en un estado estable, por lo que se activan las acciones a nivel, poniendo a 1 las salidas asociadas a las etapas que están activas. Se debe incluir una red (ecuación lógica) para cada salida a nivel.

9. Definición de los **temporizadores y contadores**. Aquí se incluye las funciones que definen los contadores y temporizadores utilizados. Puede ser necesario poner algunos de los temporizadores o contadores junto con las acciones impulsionales en el bloque 6 (si se ven afectados por eventos de flanco o por etapas inestables). Por ejemplo, si un contador cuenta el número de veces que se ha activado una etapa, lo más lógico es ponerlo en el bloque de las acciones impulsionales, ya que así contará bien aunque la etapa sea inestable.

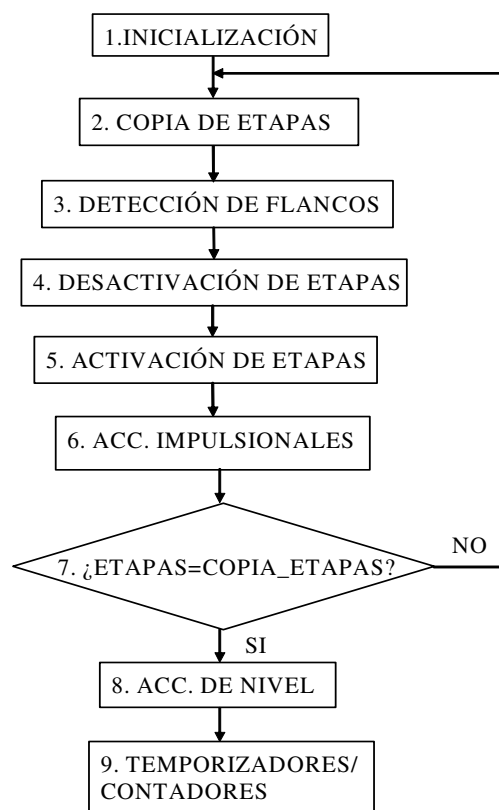


Figura 6.9: Diagrama de flujo de la estructura de programa del método 2 de implementación.

Con este método, como se observa en la figura 6.9, cada vez que se ejecutan los bloques 4 y 5 el sistema avanza como mucho una etapa. En un mismo ciclo de scan el sistema va evolucionando (ejecutando el bloque 4 y 5) hasta que se alcanza un estado estable. Si hay estados inestables intermedios, el programa activa la variable correspondiente durante un instante, activando las acciones impulsionales y memorizadas. Las acciones de nivel se ejecutan solo cuando el estado es estable. El ejemplo 2 es una muestra de automatismo implementado mediante este método.

Si hay receptividades (transiciones) que dependen de un flanco en una variable, su validez se comprueba cada vez que se ejecuta el bloque 3. De esta

forma un flanco sólo es utilizado en un paso de evolución de etapas, pues en el siguiente paso se queda inactivo. En los autómatas de Omron la detección de un flanco en una variable se implementa con la función DIFU. En el ejemplo 3 se desarrolla un caso con receptividad por flanco.

Ejemplo 2

Sea el disco de la figura 6.10. El disco tiene una placa metálica que es

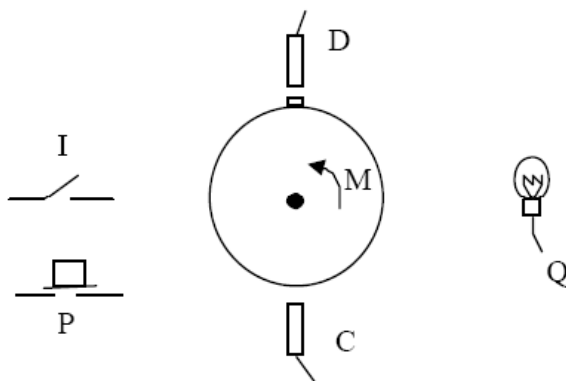


Figura 6.10: Sistema del ejemplo 2.

detectada por los detectores inductivos C y D. Inicialmente se supone el disco en el punto C. Cuando se pulsa P, el disco empieza a girar. Si el interruptor I está desactivado, el disco debe girar hasta que llegue a C. Si el interruptor I está activado cuando la placa pasa por D, el disco debe parar hasta que se desactive I. En cualquier caso, además, cada vez que se activa D se debe incrementar un contador de vueltas y activar una salida Q (luz) durante 200 ms.

La figura 6.11 muestra una posible solución. En ella, si I no está activado, la etapa 2 es inestable, por lo que se pasa de la 1 a la 3. Eso significa que la variable M (de nivel) debe permanecer activa (no se debe desactivar ni siquiera en un pequeño pulso). Por otra parte, la acción impulsional de incrementar el contador debe activarse, pues se pasa por la etapa 2.

El programa con el método 1 se representa en las figuras 6.12, 6.13, 6.14 y 6.15. En este caso, si I no está activo, cuando se active D, la etapa 2 (inestable) se activará en una red y se desactivará en la siguiente, por lo que la acción impulsional no llegará a ejecutarse. Esto se podría resolver colocando la detección del flanco de E2 y la acción impulsional inmediatamente después de la ecuación de activación de E2. El inconveniente es que el procedimiento deja de ser sistemático.

Por otra parte, el comportamiento depende del orden en que se escriban las ecuaciones, que en principio es arbitrario. En este caso, puede suceder que el sistema esté un ciclo de scan completo en la etapa 2, lo que provocaría el

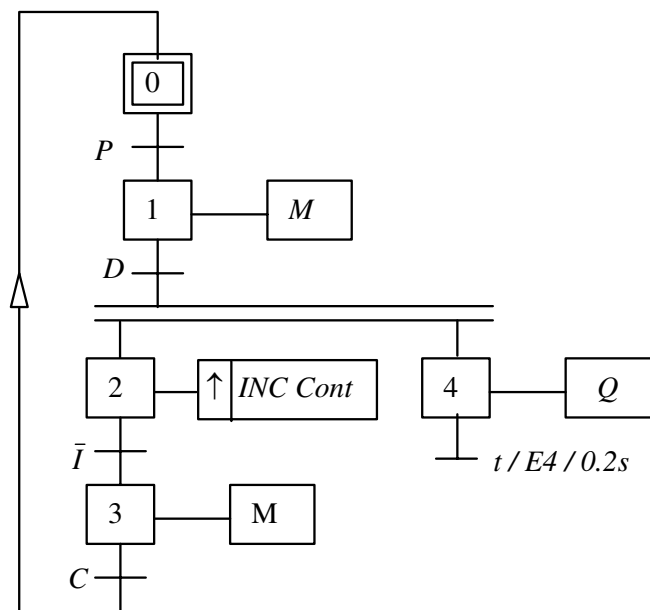


Figura 6.11: Grafcet que resuelve el ejemplo 2.

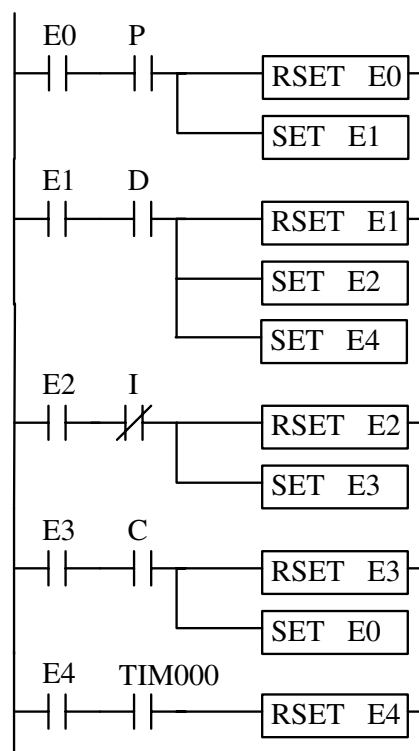


Figura 6.12: Implementación de las activaciones/desactivaciones del ejemplo 2 mediante el método 1.

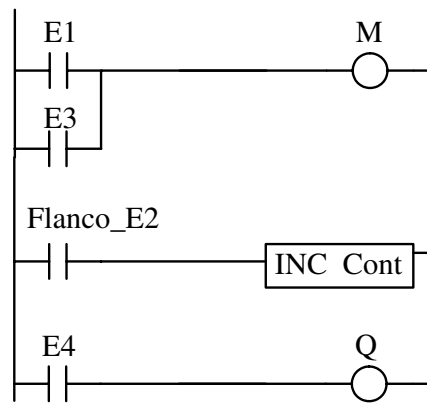


Figura 6.13: Implementación de las salidas en el ejemplo 2 mediante el método 1.



Figura 6.14: Implementación de la detección de flancos en el ejemplo 2 mediante el método 1.

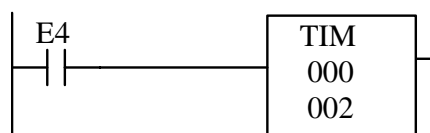


Figura 6.15: Implementación del bloque de temporizadores en el ejemplo 2 mediante el método 1.

apagado de M durante ese ciclo. Esto sucedería si el programa se escribiera en el orden representado en la figura 6.16. En este caso, si $I = 0$ cuando se

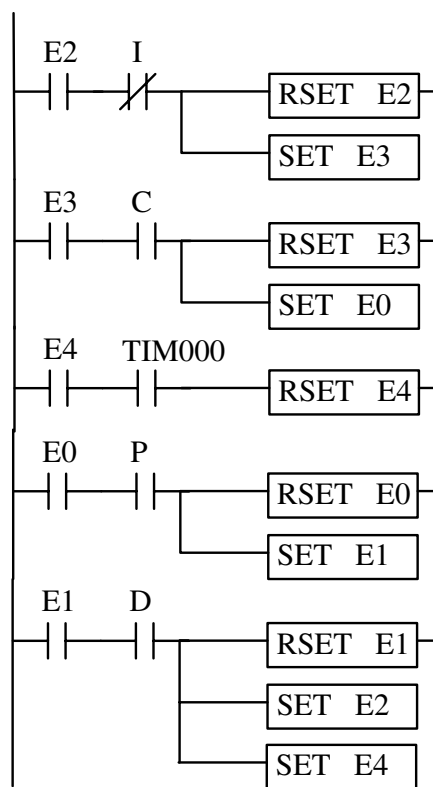


Figura 6.16: Cambio en el orden de implementación de las transiciones en el ejemplo 2 mediante el método 1.

activa D, como la última ecuación es la que activa E2, esta etapa permanecerá activa hasta el siguiente ciclo de scan. Cuando se ejecuten las ecuaciones de las salidas, el motor se desactivará durante un ciclo. En ese caso sí se ejecutaría la acción impulsional que incrementaría el contador.

La implementación por el método 2 resuelve estos problemas, tal y como se observa en las figuras 6.17 y 6.18.

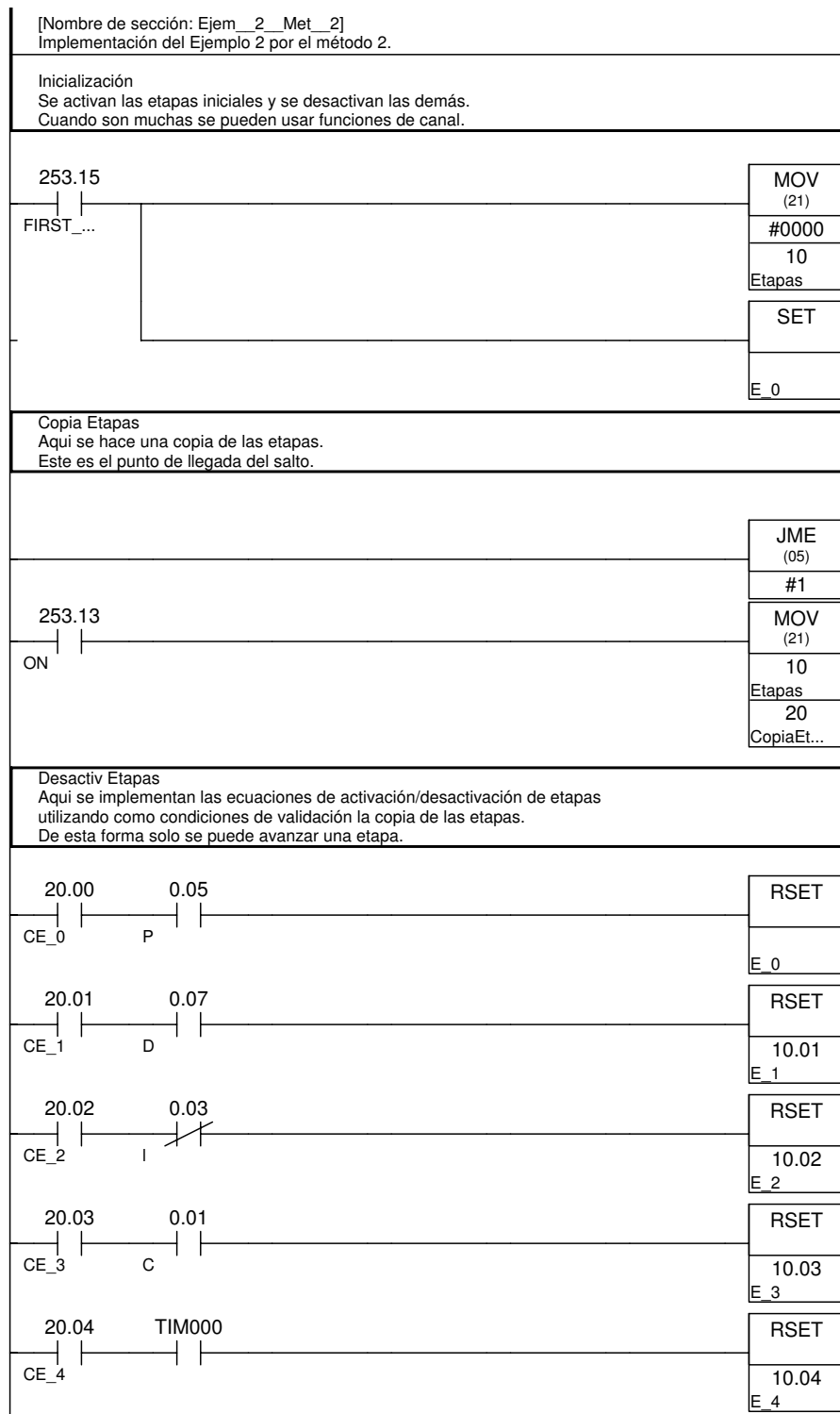


Figura 6.17: Diagrama de contactos que implementa la solución del ejemplo 2 mediante el método 2.

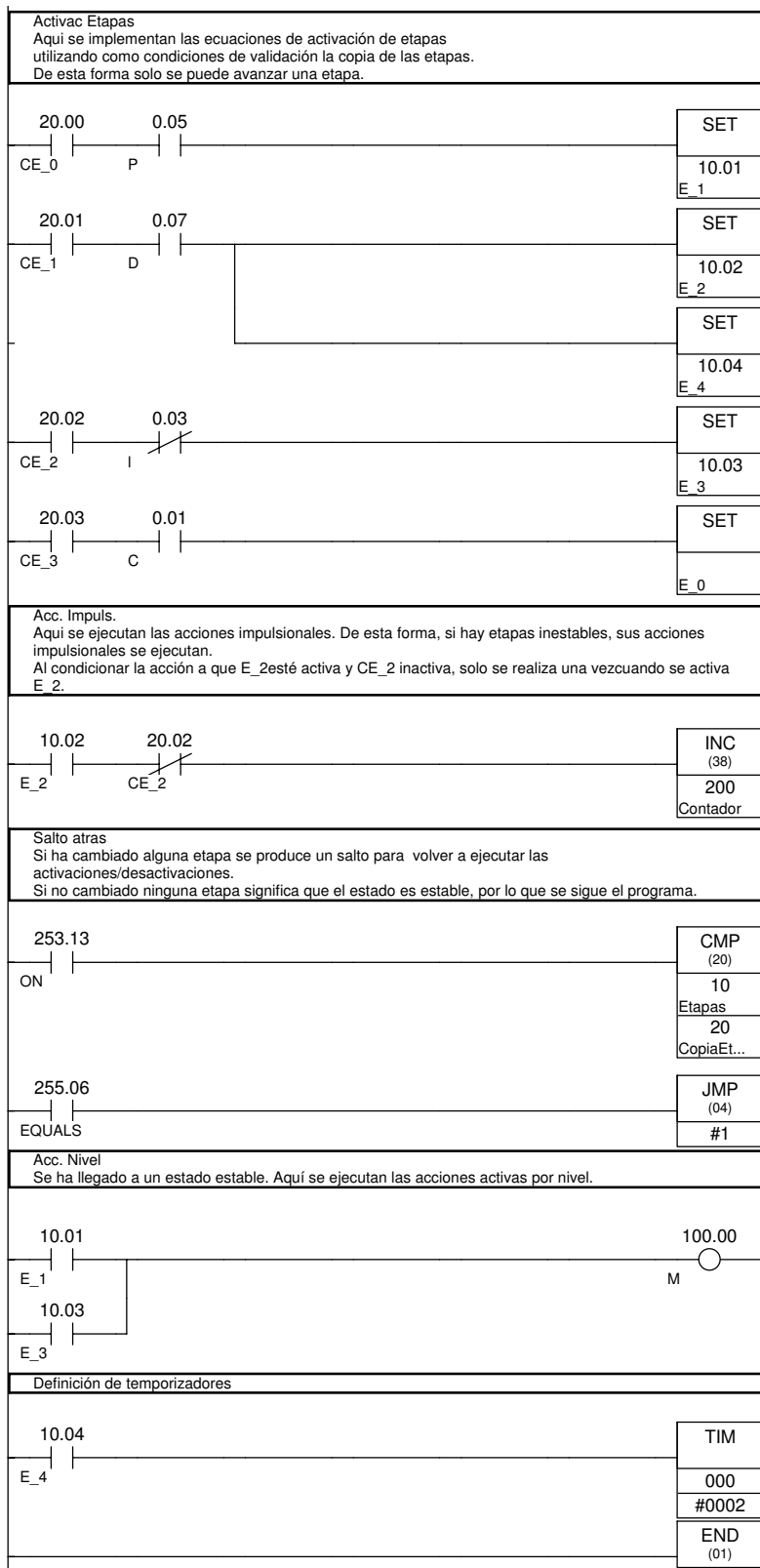


Figura 6.18: Diagrama de contactos que implementa la solución del ejemplo 2 mediante el método 2.

Ejemplo 3

Considérese el problema de poner en marcha y apagar un motor con el mismo pulsador. Cuando se pulsa P (en el flanco de subida) el motor se pone en marcha. Cuando se vuelve a pulsar P (el siguiente flanco de subida) el motor debe pararse.

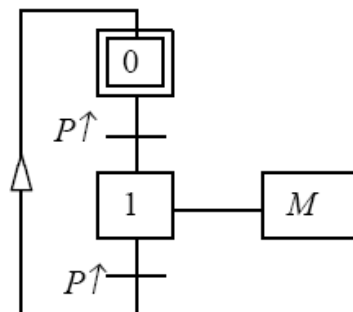


Figura 6.19: Grafset que resuelve el ejemplo 3.

Una posible solución (la más simple) sería la representada en la figura 6.19. El método 1 no sirve para implementar este Grafset, ya que sus ecuaciones (activación/desactivación más salidas) representadas en la figura 6.20 son inestables. Resulta evidente que el sistema se quedaría siempre en la etapa 0.

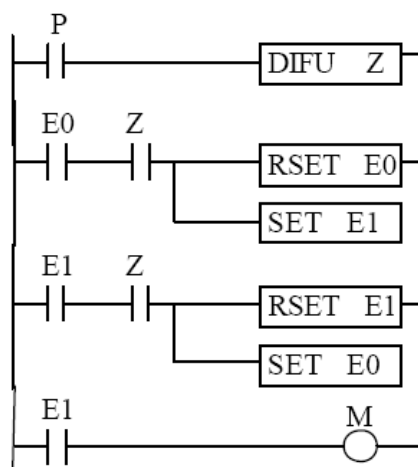


Figura 6.20: Diagrama de contactos que implementa la solución del ejemplo 3 mediante el método 1.

Con el método 2, sin embargo el funcionamiento sí sería correcto. En la figura 6.21 se muestra el programa resultante para un PLC CQM1 de Omron.

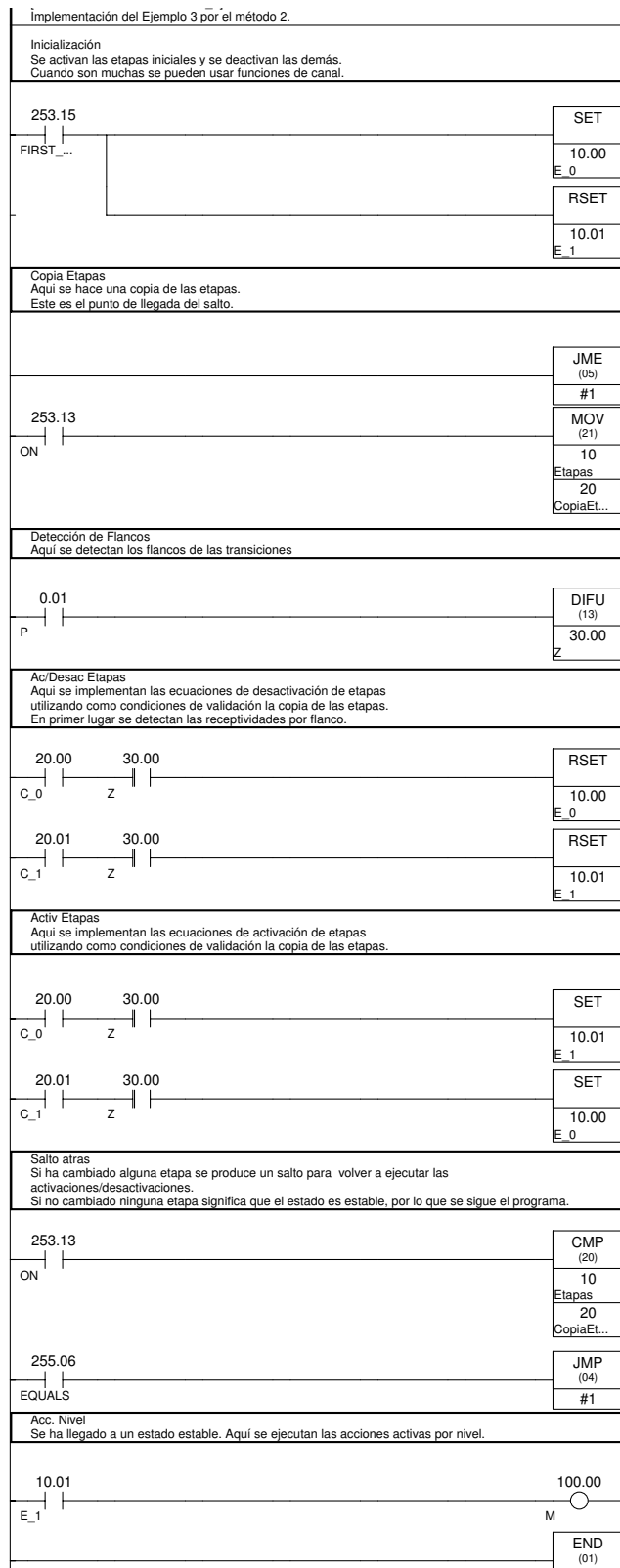


Figura 6.21: Diagrama de contactos que implementa la solución del ejemplo 3 mediante el método 2.

Ejemplo 4

Considérese el sistema de la figura 6.22.

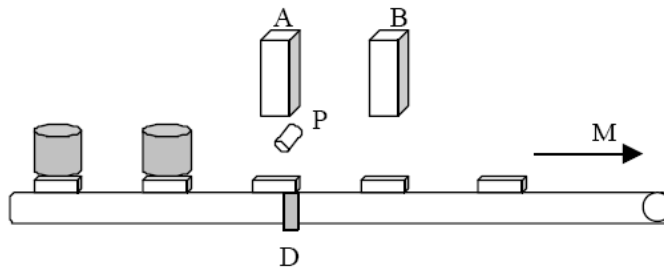


Figura 6.22: Sistema del ejemplo 4.

Por una cinta accionada por el motor M circulan palets espaciados regularmente de forma que cuando hay un palet en A, hay otro en B. Sobre estos palets puede o no haber un producto sobre el que hay que realizar dos operaciones A y B (podrían ser de llenado, taponado, etiquetado, control, etc.). El detector inductivo D detecta el paso de los palets, mientras que la fotocélula P detecta la presencia de producto (si hay producto se activa antes que D). Supondremos, en este caso, que las operaciones A y B están temporizadas (duran 1 y 2 segundos respectivamente).

Una posible solución (solución a) al problema se representa en la figura 6.23. En este caso si vienen dos productos consecutivos se activan dos etapas

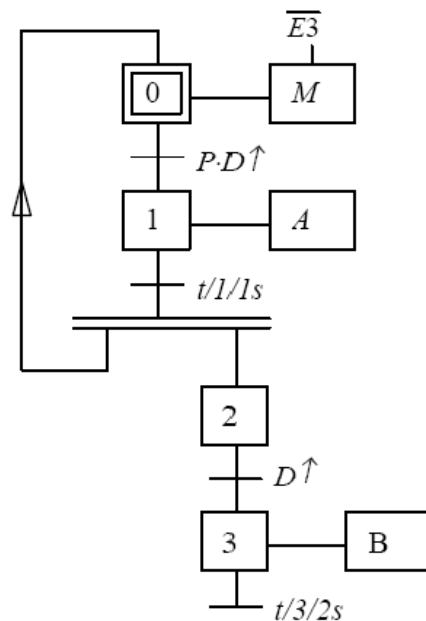


Figura 6.23: Grafcet que resuelve el ejemplo 4.

de la secuencia, realizándose las dos operaciones. El programa que implementa esta solución se realiza sin problemas con el método 1 o con el método 2, tal y como se muestra en las figuras 6.24, 6.25 y 6.26.

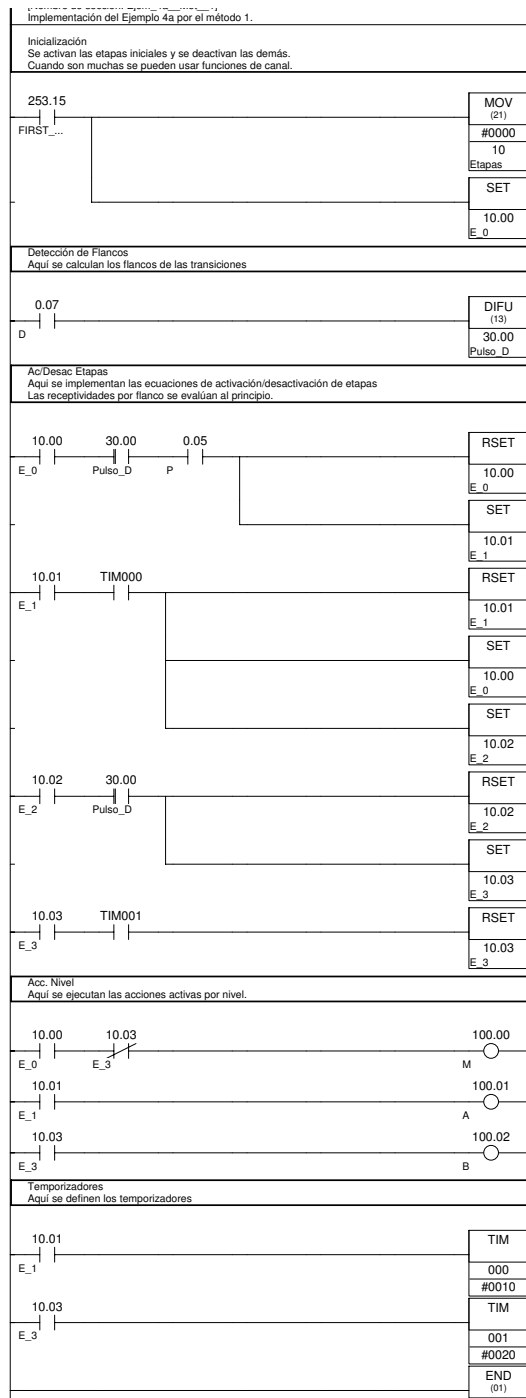


Figura 6.24: Diagrama de contactos que implementa la solución del ejemplo 4 mediante el método 1.

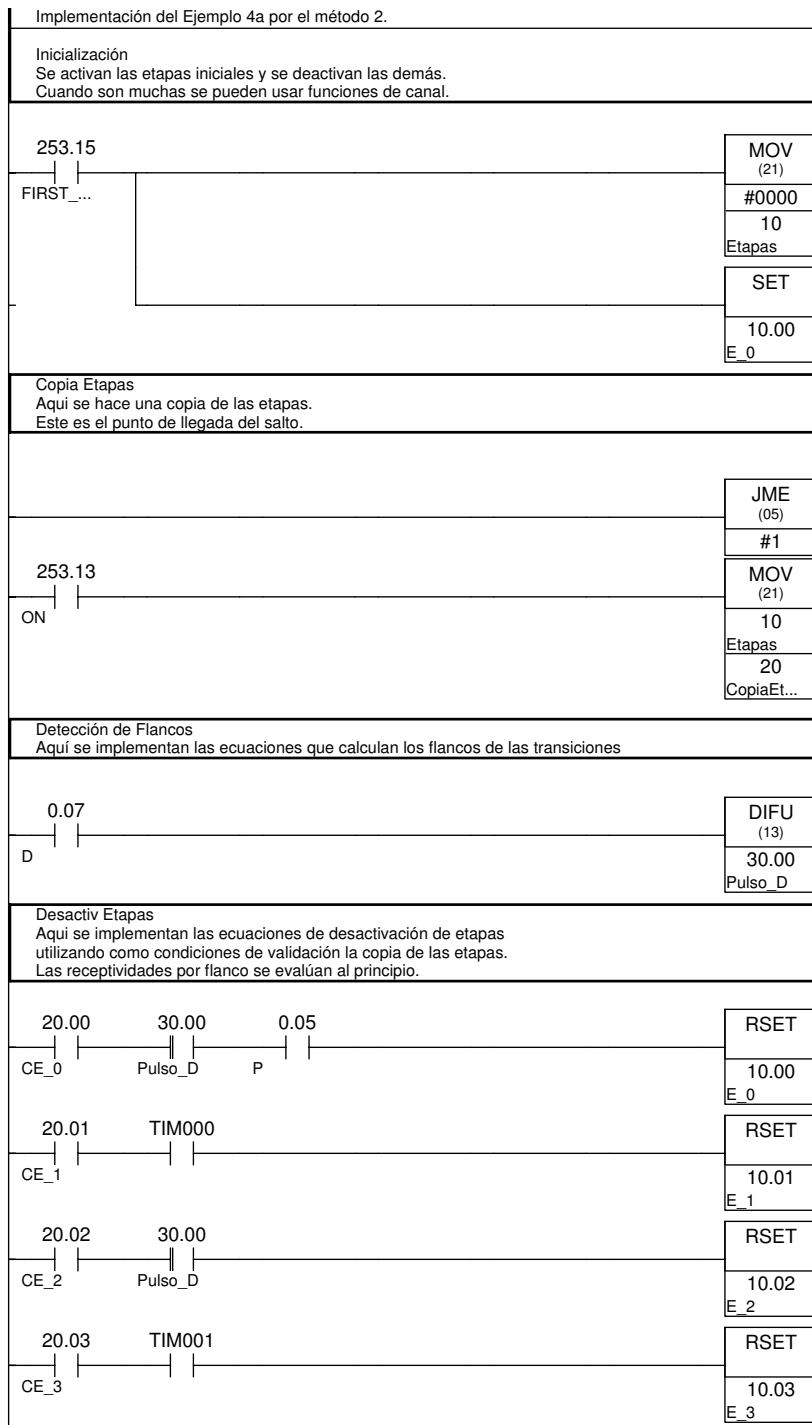


Figura 6.25: Diagrama de contactos que implementa la solución del ejemplo 4 mediante el método 2.

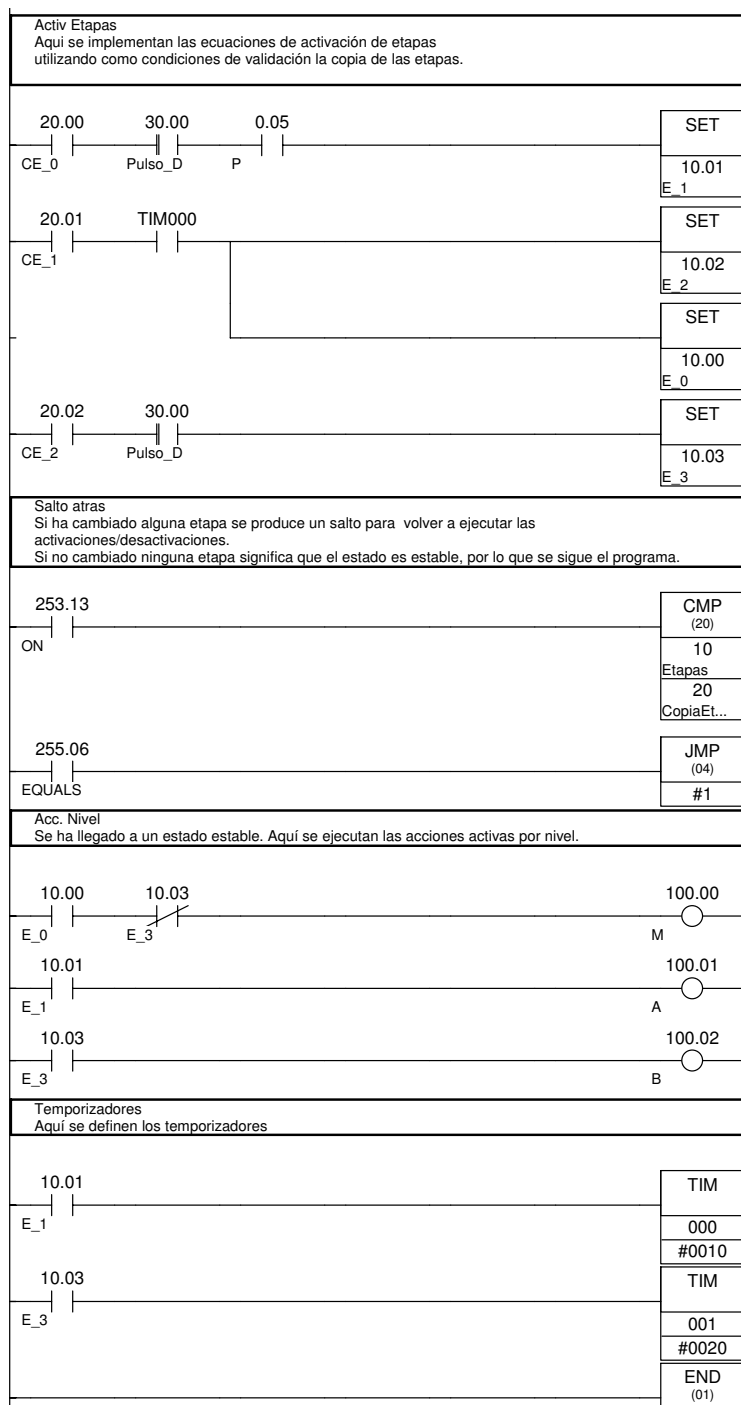


Figura 6.26: Diagrama de contactos que implementa la solución del ejemplo 4 mediante el método 2.

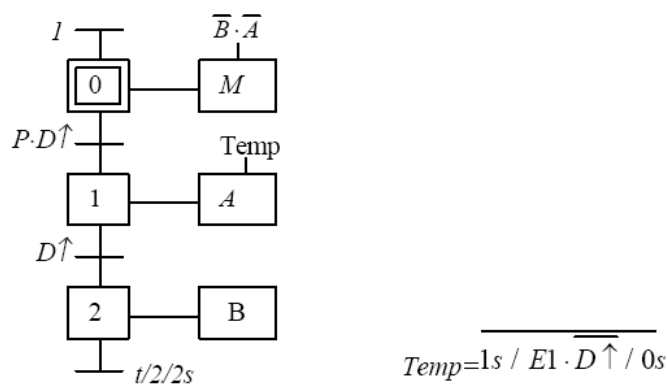


Figura 6.27: Grafcet que resuelve el ejemplo 4. Solución B.

Otra solución (solución B) podría ser la representada en la figura 6.27. En este caso la etapa 0 está siempre activa. El temporizador se ha expresado según la norma IEC848. En esta ocasión, no basta con poner la etapa 1, pues cuando pasan dos botes seguidos ésta no se desactiva, y por lo tanto el temporizador no se resetearía.

El programa que implementa esta solución no se puede realizar con el método 1. En cambio sí se puede realizar con el método 2, tal y como se muestra en las figuras 6.28 y 6.29.

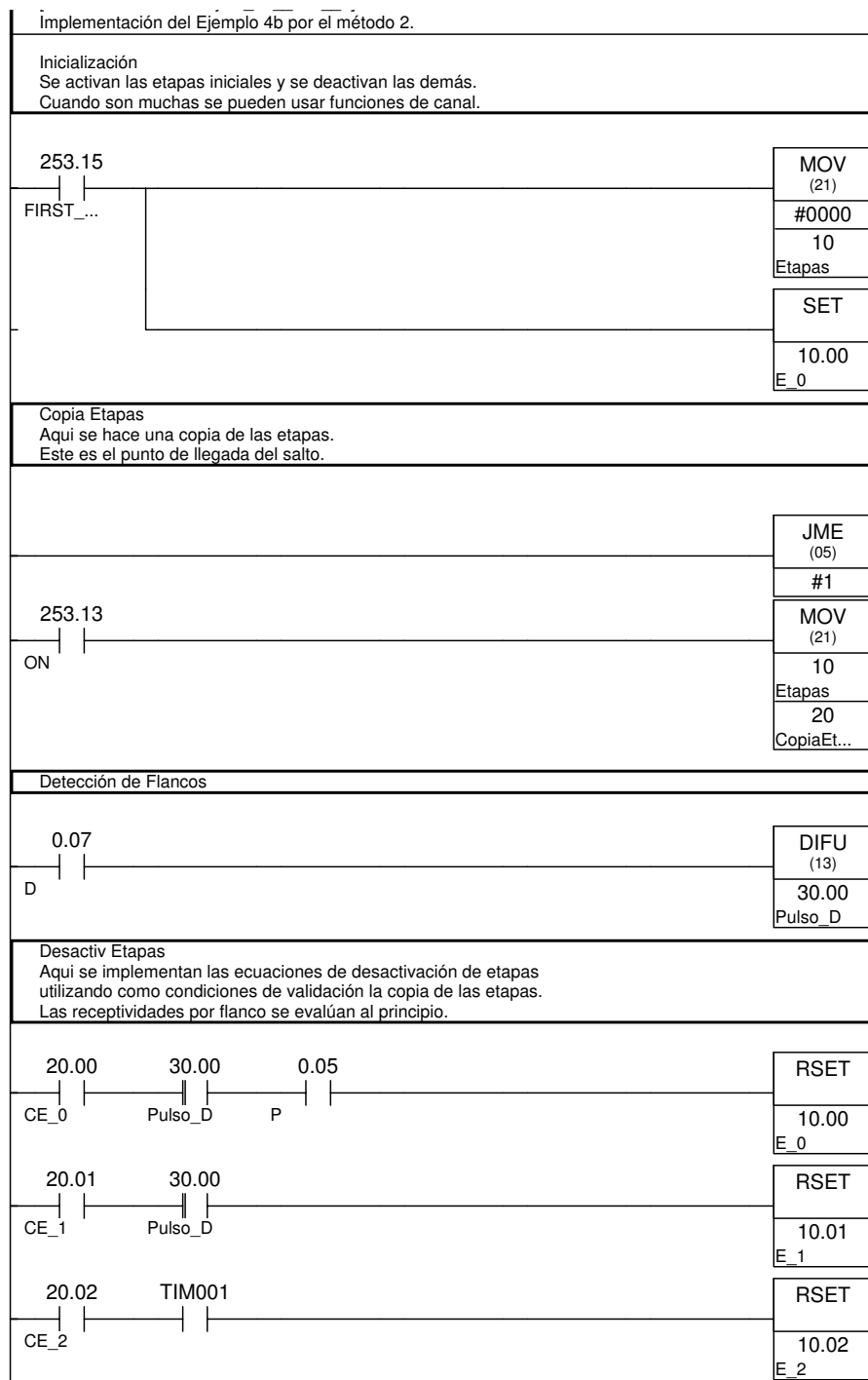
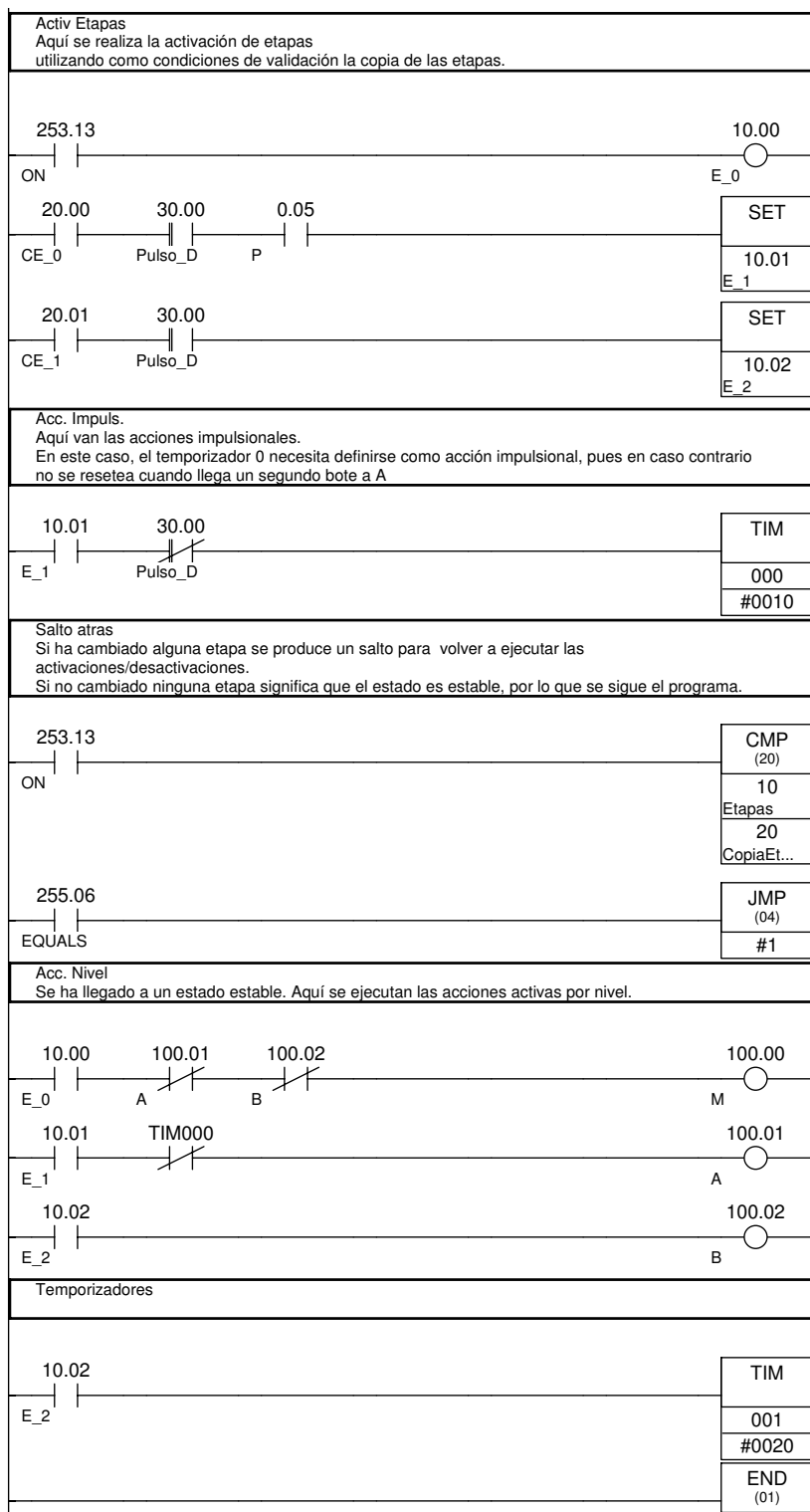


Figura 6.28: Diagrama de contactos que implementa la solución B del ejemplo 4 mediante el método 2.



20.01 30.00

SET

CE_1

Pulso_D

Figura 6.29: Diagrama de contactos que implementa la solución B del ejemplo 4 mediante el método 2.

6.3 IMPLEMENTACIÓN DEL ALGORITMO DE CONTROL CUANDO HAY FORZADOS

Cuando el proceso se modela mediante varios Graficets y existen forzados de unos sobre otros, el algoritmo de implementación cambia ligeramente para tener en cuenta las reglas del forzado. Éstas son:

1. El forzado es prioritario respecto de cualquier otra evolución (sea de activación o de desactivación).
2. Las reglas de evolución del Graficet no se aplican a un Graficet mientras éste permanezca forzado.

Cuando hay varios Graficets interrelacionados es conveniente utilizar **el método 2**, pues eso nos garantiza el funcionamiento correcto. Las reglas del forzado se tienen en cuenta en la implementación en el programa del autómatas; basta con introducir un pequeño cambio en el citado método 2. La modificación consiste en que después del bloque 5 se debe ejecutar los forzados que correspondan a las variables copia de etapas activas. Al situar los forzados después de las activaciones se garantiza que éstos son prioritarios a aquellas. Si hay forzados por flanco, debe obtenerse el flanco de la copia de la etapa donde está el forzado (esto se añade al punto 3).

El algoritmo final quedaría:

1. **Inicialización** de las etapas mediante el bit de inicio. Este bloque sólo debe ejecutarse una vez, en el primer ciclo de scan. En este bloque se ponen a 1 las etapas iniciales de todos los Graficet (las marcadas con doble cuadrado), poniéndose a 0 las demás.
2. **Copia de las variables** que definen las *etapas* en otras variables duplicadas (*copia_de_etapas*).
3. **Detección de flancos** de las entradas que dan lugar a transiciones por flanco, y de las variables *copia_de_etapas* a las que se asocian forzados por flanco. En este bloque, de cada variable cuyo flanco se tenga que calcular se obtiene una variable de bit que valdrá 1 únicamente durante el siguiente ciclo de scan si la variable ha pasado de 0 a 1 o de 1 a 0 (según sea un flanco ascendente o descendente).
4. **Desactivación** de las etapas anteriores a las transiciones franqueables, utilizando para las validaciones las variables *copia_de_etapas*. Se tiene que incluir una red (ecuación) lógica para cada transición.
5. **Activación** de las etapas posteriores a las transiciones franqueables, utilizando para las validaciones las variables *copia_de_etapas*. De esta forma, solo se avanza una etapa cada vez que se ejecuta el bloque completo. La activación va después de la desactivación para garantizar que sea prioritaria. Se tiene que incluir una red (ecuación) lógica para cada transición.

6. **Ejecución de forzados.** Para ello se utiliza como condición de validación la variable *copia_de_etapa* correspondiente. Si el forzado es por flanco la condición de validación será el flanco en la variable *copia_de_etapa* correspondiente (no en la variable *etapa*). Este flanco se calculará en el bloque 3.
7. **Ejecución de las acciones impulsionales** que correspondan. Aquí hay que incluir todas las acciones que se han de ejecutar una sola vez cuando se activa la etapa. Un ejemplo podría ser incrementar un contador que cuenta el número de veces que se ha activado la etapa (es una acción impulsional asociada al paso de 0 a 1 de la etapa). La condición de ejecución de las acciones impulsionales es: $etapa \cdot copia_de_etapa$.
8. **Si ha habido cambios en alguna etapa (estado inestable) volver al bloque 2.** Es decir, se vuelven a ejecutar los bloques anteriores hasta que $etapas = copia_de_etapas$, es decir, hasta que se llegue a un estado estable.
9. **Activación de las acciones a nivel.** En este caso ya se está en un estado estable, por lo que se activan las acciones a nivel, poniendo a 1 las salidas asociadas a las etapas que están activas. Se debe incluir una red (ecuación lógica) para cada salida a nivel.
10. Definición de los **temporizadores y contadores.** Aquí se incluyen las funciones que definen los contadores y temporizadores utilizados. Puede ser necesario poner algunos de los temporizadores o contadores junto con las acciones impulsionales en el bloque 7 (si se ven afectados por eventos de flanco o por etapas inestables). Por ejemplo, si un contador cuenta el número de veces que se ha activado una etapa, lo más lógico es ponerlo en el bloque de las acciones impulsionales, ya que así contará bien aunque la etapa sea inestable.

El bloque de ecuaciones de forzado de Grafcets (bloque 6) se implementa con instrucciones de SET y RSET o con instrucciones de ANDW y ORW, con la condición de activación igual a la copia de la etapa o al flanco de la copia de la etapa, según sea forzado por nivel o por flanco. El siguiente ejemplo ilustra la cuestión.

Ejemplo 5

Considérese el ejemplo 4, al que se desea añadir las siguientes funciones:

- Habrá un pulsador de MARCHA y otro de PARADA. Al iniciarse el sistema todo estará desactivado hasta que se pulse el pulsador de marcha. Cuando se pulse el pulsador de parada, el sistema acabará de procesar los elementos que lleguen hasta que se active D sin que haya pieza (P

= 0) y sin que quede producto en B que procesar. En ese momento se pasará a la situación de inicio.

- Además, habrá un pulsador de EMERGENCIA (pulsador con enclavamiento). Este pulsador produce un corte de la alimentación del motor de la cinta y de los procesos A y B, independientemente del automatismo. Cuando se accione el pulsador de emergencia hay que desactivar los procesos A y B, y la cinta (aunque ya estén, de hecho, apagados por falta de alimentación) y poner en marcha una sirena. Esa sirena se desactivará cuando se presione un pulsador de parada de sirena. Para rearmar el sistema (después de resolver manualmente el problema que hubiera) se deberá rearmar el pulsador de emergencia (ponerlo a off) y después apretar el pulsador de marcha.
- Además, habrá un pulsador de parada de cinta temporal, y un pulsador de reinicio de cinta, que se utilizarán para parar la cinta en cualquier momento.

Una solución podría ser la representada en la figura 6.30. El Grafcet maestro (G0) es el que controla los modos de marcha y parada y la emergencia.

El programa que implementa los Grafcets anteriores se muestra en las figuras 6.31, 6.32, 6.33, 6.34 y 6.35.

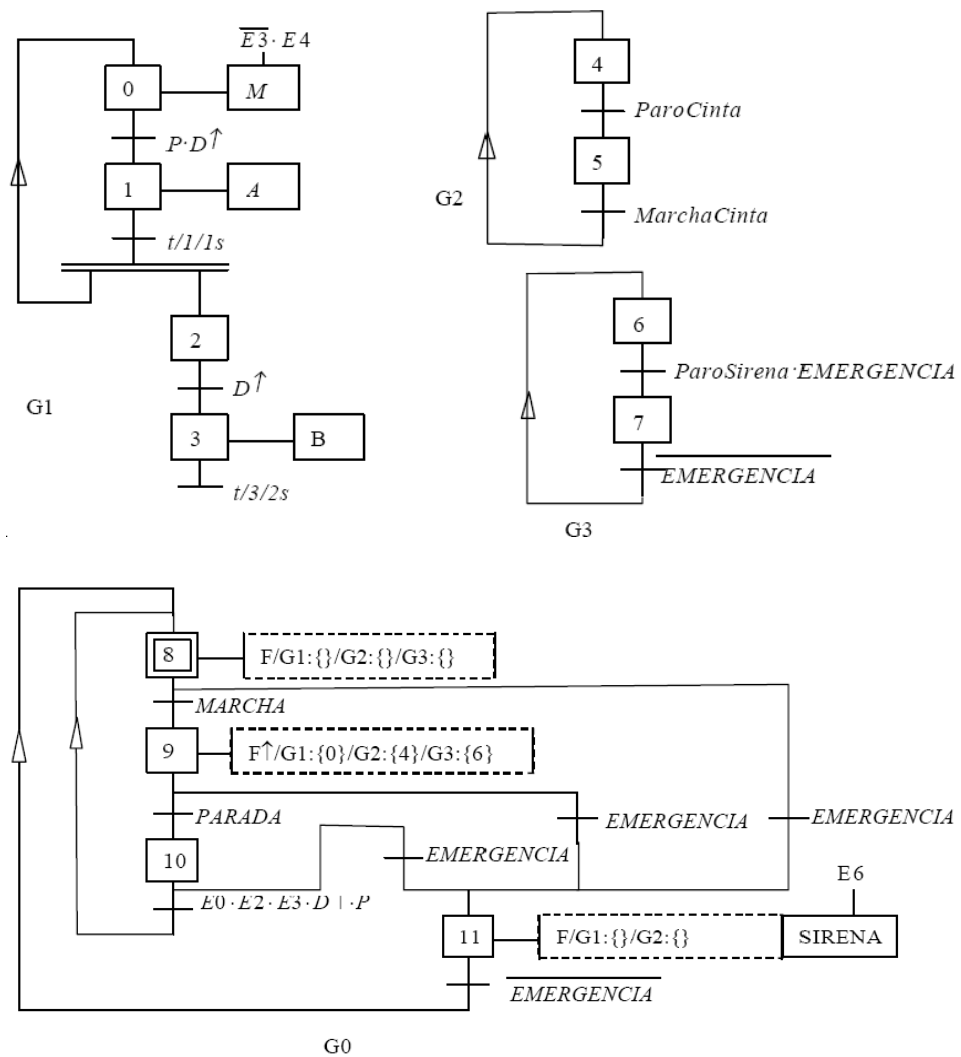


Figura 6.30: Grafset que resuelve el ejemplo 5.

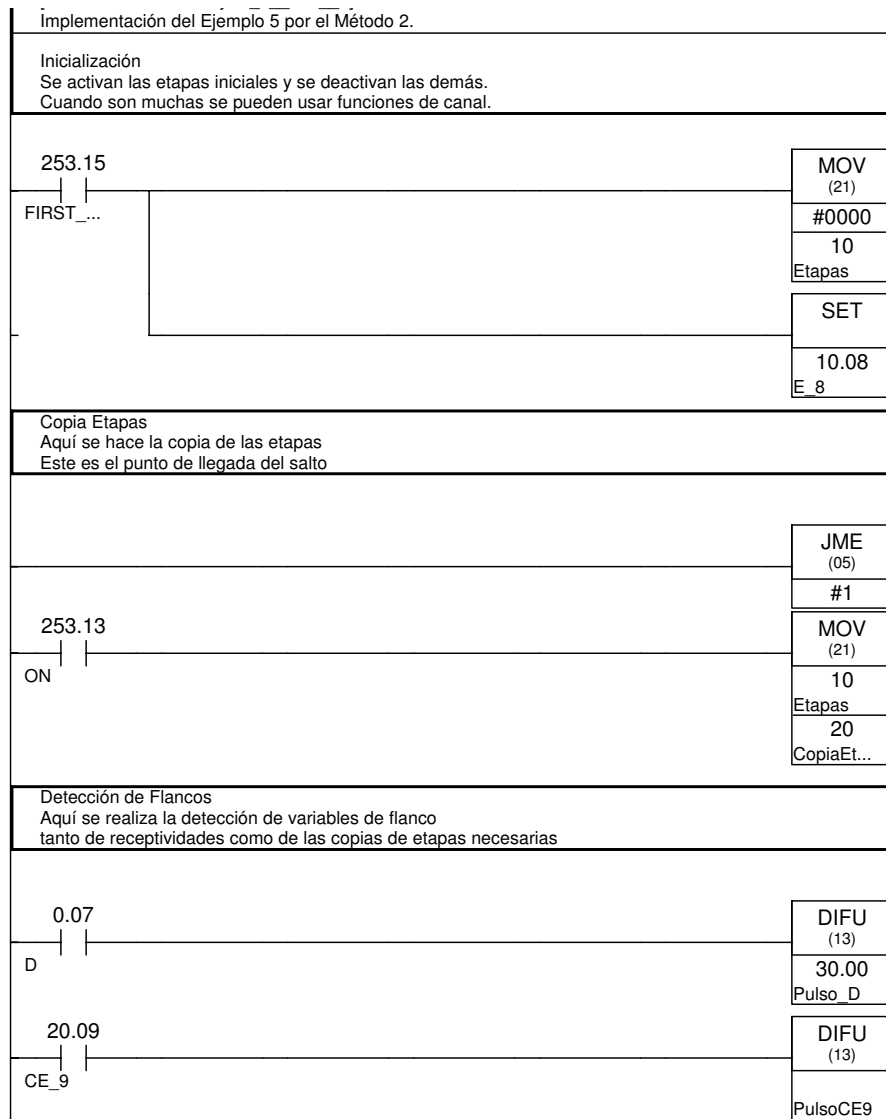


Figura 6.31: Diagrama de contactos que implementa la solución del ejemplo 5 mediante el método 2.

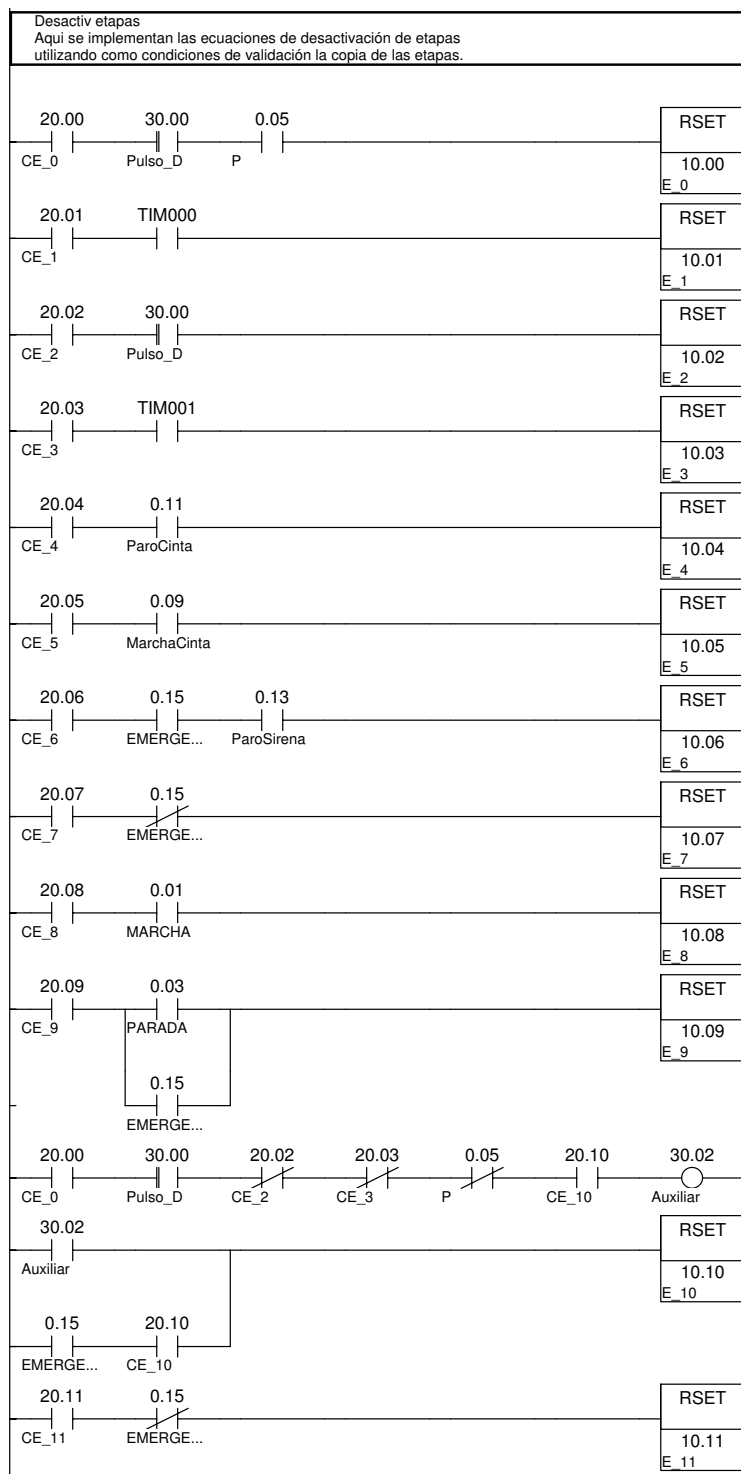


Figura 6.32: Diagrama de contactos que implementa la solución del ejemplo 5 mediante el método 2.

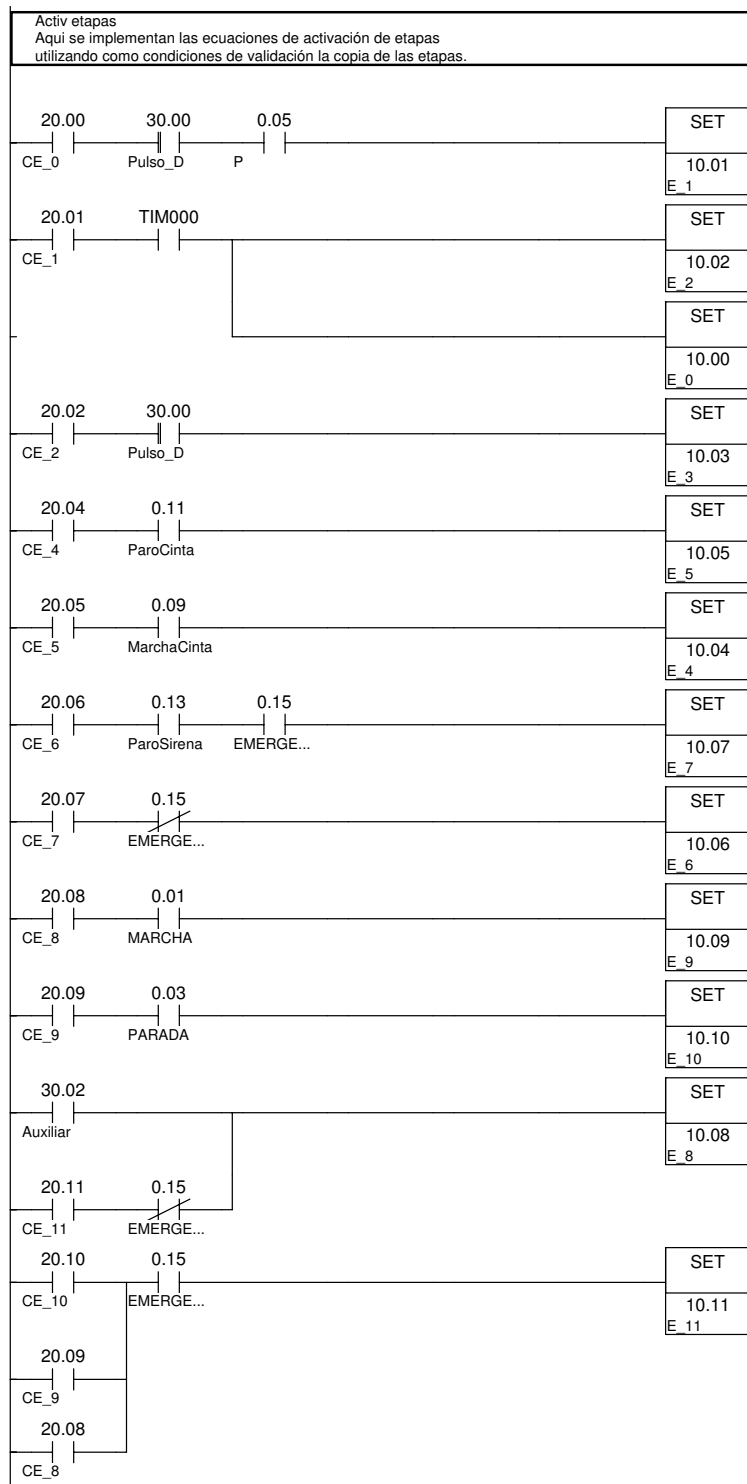


Figura 6.33: Diagrama de contactos que implementa la solución del ejemplo 5 mediante el método 2.

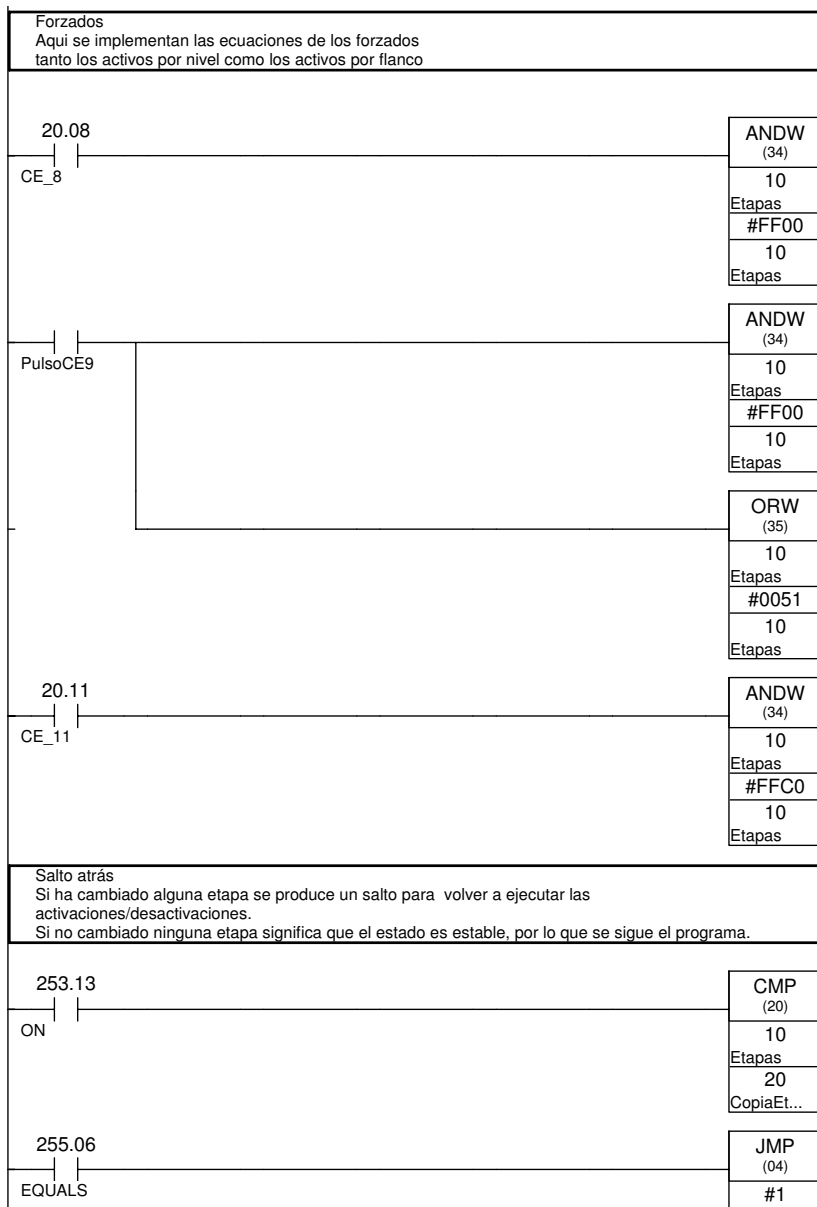


Figura 6.34: Diagrama de contactos que implementa la solución del ejemplo 5 mediante el método 2.

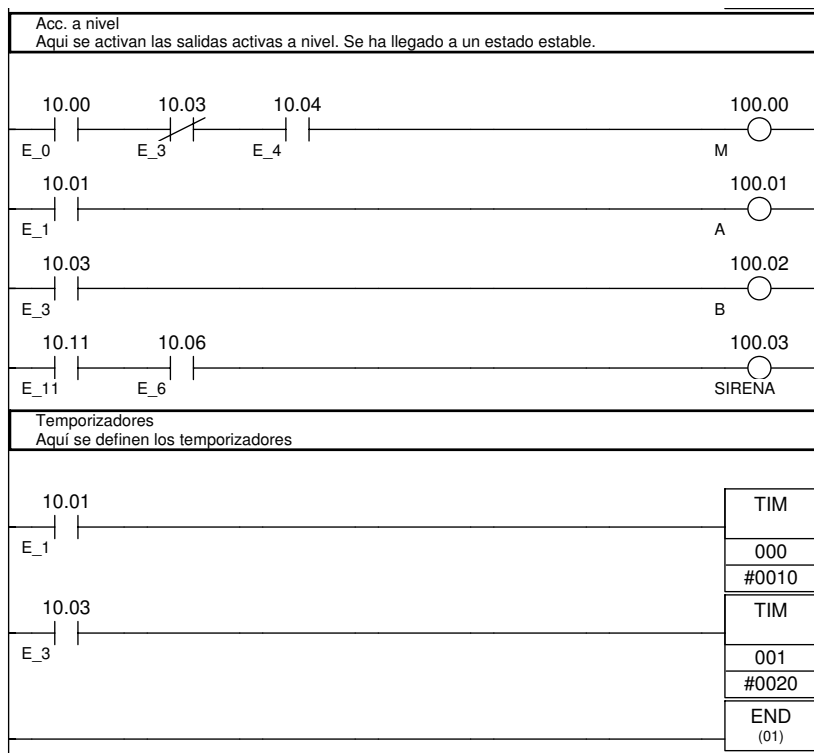


Figura 6.35: Diagrama de contactos que implementa la solución del ejemplo 5 mediante el método 2.

En este ejemplo, las ecuaciones de los forzados se implementan mediante instrucciones lógicas de canal (ANDW y ORW). La instrucción ANDW se utiliza para poner a 0 los bits deseados de un canal, mientras que la instrucción ORW se utiliza para poner a 1 los bits deseados de un canal. Por ejemplo, el forzado de la etapa 8 se implementa haciendo el AND lógico entre el número FF00 en hexadecimal (es decir, 1111111100000000 en binario) y la variable Etapas, resultando que los bits del 0 al 7 de dicha variable se ponen a 0. En cambio, el forzado de la etapa 9, después de poner a 0 mediante la instrucción ANDW las etapas de la 0 a la 7, utiliza la instrucción ORW para hacer un OR lógico entre el número 0051 en hexadecimal (0000000001010001 en binario) y la variable Etapas, resultando en que las etapas 0, 4 y 6 se ponen a 1.

6.4 IMPLEMENTACIÓN DEL ALGORITMO DE CONTROL EN OTRAS PLATAFORMAS

El autómata programable dispone de un lenguaje de programación y unas instrucciones específicamente diseñadas para implementar automatismos secuenciales. No obstante, los algoritmos descritos en la sección anterior pueden implementarse fácilmente en cualquier otro computador siempre que disponga de una interfaz adecuada de entrada y salida para obtener las señales de los sensores y dar la señal a los actuadores.

La diferencia fundamental reside en que el autómata realiza un ciclo de scan del que el usuario sólo tiene que programar una parte (según los algoritmos descritos), mientras que el bucle infinito, la actualización de entradas y salidas, la gestión de comunicaciones, la gestión de temporizadores, etc; las realiza el PLC sin que el usuario tenga que programarlas. Eso quiere decir que en otra plataforma, el algoritmo de control debe insertarse dentro de un bucle infinito, de manera que en cada ciclo de ese bucle se ejecute el algoritmo, se refresquen las señales de entrada y de salida y se realicen las tareas necesarias de comunicación o de gestión de temporizadores.

Respecto a la programación de las ecuaciones de activación y desactivación de etapas, o las ecuaciones de salidas o forzados, solo hay que tener presente el significado de las ecuaciones en diagramas de contactos, para traducirlas a ensamblador, C o cualquier otro lenguaje que se utilice. La figura 6.36 muestra la traducción a lenguaje C de dos líneas de diagrama de contactos.

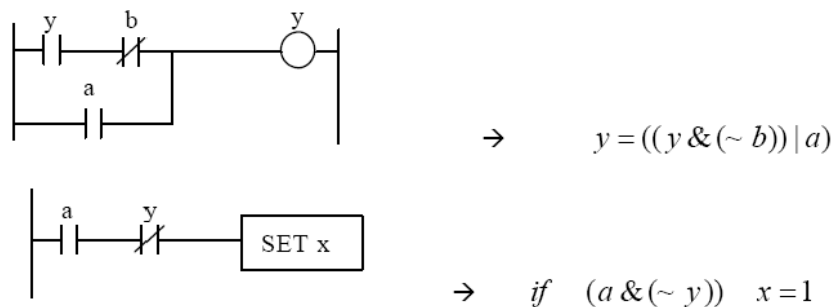


Figura 6.36: Representación en C de un diagrama de contactos.

Con respecto a las señales de entrada y salida, es conveniente tener una copia de las mismas en memoria, y trabajar con esas copias en las ecuaciones de activación, desactivación y salida, realizando una actualización de las entradas y salidas físicas cada bucle (igual que hacen los autómatas). Esto garantiza que se cumplen las reglas de evolución del Grafcet.

La detección de flancos requiere crear una variable que sea una copia de aquella cuyo flanco se quiere detectar. En C, el código equivalente a la instrucción DIFU de Omron, que calcula un flanco, podría ser:

$$\text{Flanco_A} = A \& (\sim \text{Copia_A})$$

$$Copia_A = A;$$

Una mención especial merecen los temporizadores. En los autómatas programables existen instrucciones que permiten contar un tiempo de retardo de forma muy sencilla. La implementación de estas instrucciones en otra plataforma depende mucho de cómo sea dicha plataforma y de los temporizadores hardware de que disponga.

Si el computador dispone de un temporizador hardware capaz de contar un tiempo suficientemente largo, la implementación de las instrucciones de temporizador del autómata puede consistir simplemente en comparar el valor actual del temporizador con un valor de consigna, de forma que cuando el primero sea mayor que el segundo quiere decir que ha transcurrido el tiempo especificado (lo cual llevaría a poner a 1 un bit). Esto se puede hacer, por ejemplo, en una plataforma basada en PC, pues en un PC se dispone de un reloj de tiempo real, que cuenta periodos de tiempo arbitrariamente largos.

En plataformas basadas en microcontrolador, normalmente lo que se tiene es un temporizador hardware, por ejemplo de 16 bits, que se incrementa cada cierto periodo desde 0 hasta la cuenta máxima (65535 para 16 bits), volviendo después a 0 y así indefinidamente. Este temporizador permite contar un tiempo máximo de $65536 \cdot periodo$ segundos, que normalmente es un tiempo demasiado corto para las aplicaciones de automatización habituales. En ese caso habría que programar una rutina de interrupción periódica, que se ejecutase cada cierto tiempo (contado con ese temporizador de 16 bits) para implementar un temporizador por software (un contador que se incrementa cada interrupción). La interrupción se podría programar cada 100 ms, por ejemplo, con lo que se podrían contar tiempos con precisión de décimas de segundo.

Un ejemplo en C que implementaría dos temporizadores parecidos a los del autómata de OMRON, utilizando un único temporizador capaz de contar un tiempo suficientemente largo, bien por hardware, o bien por software (mediante interrupción), sería:

```
if (E1 & (~Copia_E1))  t_final_1 = t_actual + t_a_contar_1;
Copia_E1 = E1;
if (E2 & (~Copia_E2))  t_final_2 = t_actual + t_a_contar_2;
Copia_E2 = E2;
if (E1 & (t_actual >= t_final_1) & (t_anterior < t_final_1))  bit_fin_t_1 = 1;
else  bit_fin_t_1 = 0;
if (E2 & (t_actual >= t_final_2) & (t_anterior < t_final_2))  bit_fin_t_2 = 1;
else  bit_fin_t_2 = 0;
t_anterior = t_actual;
```

Donde se ha supuesto que se quiere contar un tiempo $t_a_contar_1$ desde que se activa la etapa E1 hasta que se activa el bit $bit_fin_t_1$, que se desactiva una vez se desactiva la etapa E1. De la misma manera, se cuenta un tiempo $t_a_contar_2$ asociado a la etapa E2. El bit de salida se utilizaría para la transición de salida de la etapa, por ejemplo. Se supone que t_actual se incrementa

por hardware cada cierto periodo. La detección del flanco de la variable E1 es necesaria para que la variable *t_final1* solo se defina una vez en el momento que se activa E1. La utilización de la variable *t_anterior* es necesaria debido a que *t_actual* vuelve a 0 cuando alcanza su valor máximo, por lo que podría darse el caso de que al definir *t_final1* se produjera un desbordamiento en la suma, quedando *t_final1* menor que *t_actual*. Considérese por ejemplo, que *t_actual* es un contador de 16 bits, es decir, pasa a 0 después de llegar a 65535. Si en un momento dado, *t_actual* = 60000, y se quiere contar 10000 periodos, la suma daría 70000, es decir, desbordaría los 16 bits quedando un valor *t_final1* = 4465. Evidentemente, *t_actual* > *t_final1*, por lo que si simplemente se hace esta comparación, el bit *bit_fin_t1* se activaría inmediatamente. La condición que se propone, en cambio, es que *t_actual* haya pasado de ser menor a ser mayor que *t_final1*, por lo que el desbordamiento en la suma no da problemas.

El código anterior podría servir para tantos retardos de tiempo en un Graficet como fuera necesario, utilizando únicamente un temporizador hardware.

Ejemplo de implementación en otra plataforma

El siguiente código en C muestra cómo se podría implementar el ejemplo 1 por el método 1 en un PC, utilizando lenguaje C++ de Borland. Para implementar el temporizador se usa la función *clock()*, que devuelve el número de ciclos de reloj transcurridos desde que se puso en marcha el ordenador. El número de ciclos de reloj por segundo está en la constante del sistema *CLK_TCK*.

```
#include <time.h>
boolean C, D, I, A, R, P, Q, bitT, E0, E1, E2, E3, E4, CE3;
clock_t t_final, t_a_contar, t_anterior;
void main()
{
    /* Inicialización */
    E0=1;
    E1=0;
    E2=0;
    E3=0;
    E4=0;
    t_a_contar=1*CLK_TCK;
    while (1)
    {
        /* Activación y desactivación de etapas*/
        if (E0 & P)
        {
            E0=0;
            E1=1;
```



```

}
if (E1 & C)
{
    E1=0;
    E2=1;
}
if (E0 & (~ P) & Q)
{
    E0=0;
    E3=1;
}
if (E3 & D)
{
    E3=0;
    E4=1;
}
if (E2 & I)
{
    E2=0;
    E0=1;
}
if (E4 & I)
{
    E4=0;
    E0=1;
}

/* Activación de salidas*/
A=(E1|(E3 & bitT));
R=(E2|E4);

/* Temporizador*/
if (E3&(~ CE3))  t_final=clock()+t_a_contar;
CE3=E3;
if (E3 & (clock()>=t_final) & (t_anterior<t_final))  bitT=1;
else  bitT=0;
t_anterior=clock();
}
}

```

CAPÍTULO 7

SISTEMAS DE CONTROL DISTRIBUIDO

7.1 INTRODUCCIÓN

Se denomina sistema de control centralizado a aquél en el que hay un único equipo de control, el cual está conectado a todos los sensores y actuadores. Este equipo de control, generalmente un autómata, implementa el algoritmo lógico de control de todo el proceso.

Se denomina sistema de control distribuido o descentralizado a aquél en el que hay varios equipos de control, cada uno de los cuales controla una parte del proceso global. Los sensores y actuadores de cada parte están conectados a su correspondiente equipo de control. Para que el sistema global funcione correctamente es necesario que los distintos equipos puedan comunicarse entre sí, es decir, que exista una red de comunicación a través de la cual se puedan transmitir datos de configuración y valores de variables del proceso entre los distintos equipos. Un sistema de control distribuido está formado, por tanto, por equipos conectados mediante una red de comunicación.

Las estructuras de control distribuido tienen las siguientes ventajas:

- Posible puesta en servicio simultánea e independiente de partes concretas de la instalación.
- Programas más pequeños y sencillos.
- Procesamiento paralelo por sistemas de automatización repartidos. Tiempos de reacción más cortos.
- Estructura supervisora que puede asumir funciones de diagnóstico.
- Aumento de la disponibilidad de las distintas partes de la instalación.

Evidentemente, la concreción de las anteriores ventajas pasa por la existencia de un sistema de comunicación eficiente, completo y seguro.

7.2 NIVELES DE UN SISTEMA DE CONTROL DISTRIBUIDO

Se suele distinguir varios niveles cuando se describe el control distribuido de una planta. Una posible estructuración puede ser la que se muestra en la figura 7.1, donde se distinguen cuatro niveles, los cuales se describen a continuación.

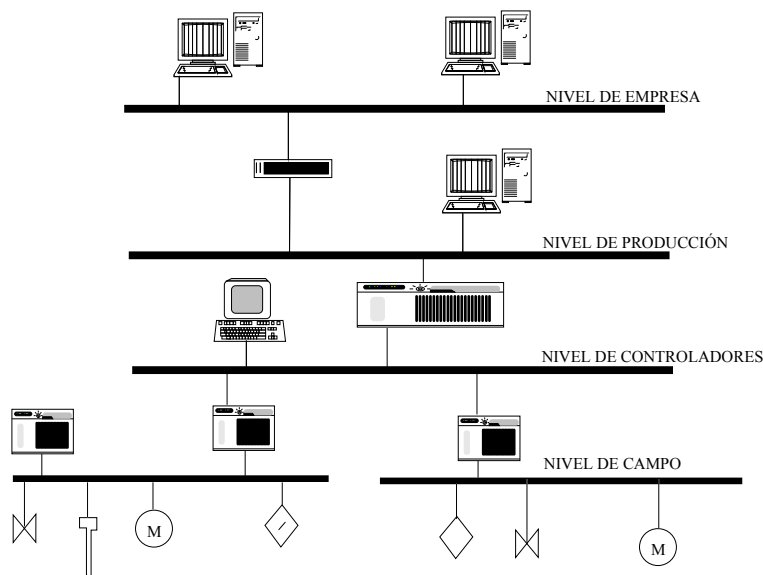


Figura 7.1: Niveles de un sistema de control distribuido.

- **Nivel de campo.** En este nivel se conecta un equipo de control, típicamente un autómatas programable, con otros elementos como variadores de frecuencia, controladores PID, módulos remotos de entradas y salidas, o sensores o actuadores “inteligentes”. Generalmente, la conexión se realiza mediante una red denominada bus de campo.
- **Nivel de controladores.** En este nivel se conectan entre sí diversos equipos de control que pueden ser autómatas programables u ordenadores, cada uno de los cuales gobierna una parte del proceso. Los distintos elementos conectados entre sí, en este nivel, constituyen las células de producción. Suele haber un equipo de control, normalmente un ordenador o autómatas de gama alta, que es el encargado de enlazar la red de este nivel con la red del nivel superior. La red que implementa este nivel suele ser también un bus de campo, aunque suele ser de más capacidad que el utilizado en el nivel de campo.
- **Nivel de producción.** En este nivel se conectan entre sí los ordenadores o autómatas principales de cada célula. Éstos se conectan a su vez a los ordenadores de gestión. La red de este nivel suele ser una red de más alto nivel, como una red Ethernet.
- **Nivel de empresa.** En este nivel se conectan los ordenadores de gestión de la empresa.

En la práctica, los 4 niveles no suelen distinguirse de forma nítida, sino que en función del proceso productivo pueden existir sólo algunos de esos niveles, o puede haber un solapamiento entre varios niveles. Por ejemplo, es habitual que

existan redes en las que se integran sensores y actuadores inteligentes, variadores de frecuencia y reguladores PID, con varios equipos de control (autómatas o PC), es decir, una mezcla de los niveles de campo y de controladores. También es frecuente que exista una única red que comunica los ordenadores de gestión y los ordenadores de producción, es decir, que integra el nivel de empresa (o de gestión) y el nivel de producción.

La figura 7.2 muestra un ejemplo de sistema de control distribuido de una industria. En ella se observa una red de nivel corporativo (red Ethernet), una red de planta y control (TransparentReady) y varias redes de campo (red CanOpen y red ModBus).

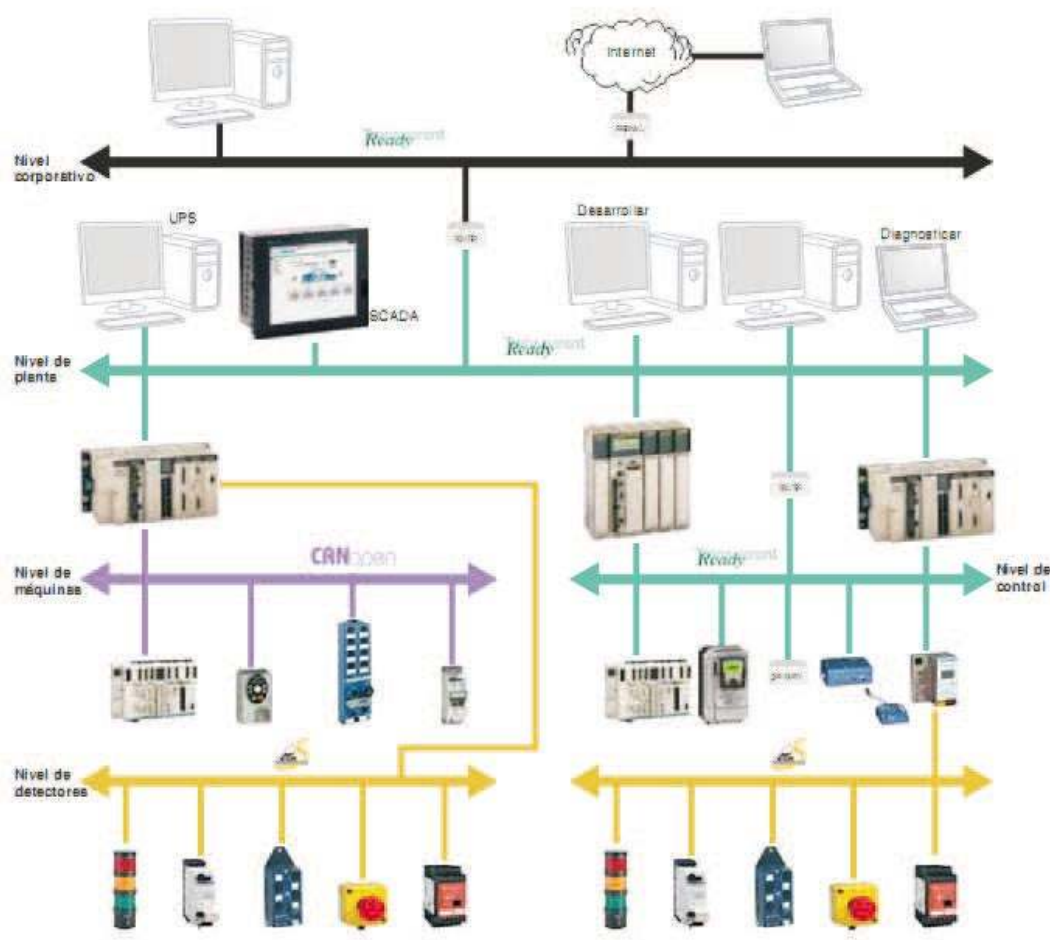


Figura 7.2: Ejemplo de sistema de control distribuido de Schneider Electric.

7.3 REDES DE COMUNICACIÓN

Un sistema de control distribuido está formado por diversos equipos conectados por medio de una o varias redes de comunicación. Las propiedades

fundamentales que definen una red de comunicación son las siguientes:

- Topología de la red.
- Medio físico.
- Características eléctricas de las señales.
- Método de arbitraje o de acceso.
- Modo de funcionamiento.
- Tamaño y estructura de los mensajes que se transmiten.
- Número de nodos.
- Distancia de transmisión.
- Velocidad de transmisión.

Topología

Se llama topología a la disposición o estructura física de los equipos, conocidos como nodos, y los cables de la red. Las topologías más habituales son:

- **Estrella:** En esta topología hay un equipo central, llamado hub, al que se conectan todos los equipos, figura 7.3. Del hub parte un cable para cada nodo.

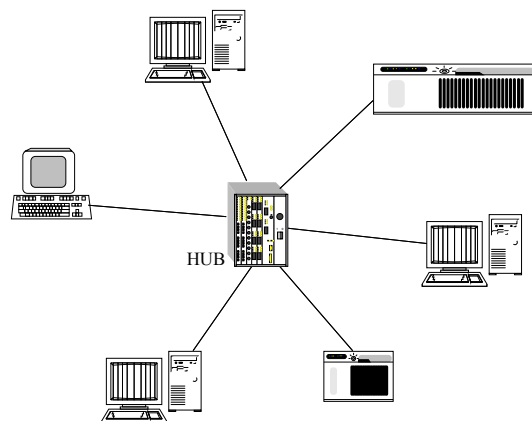


Figura 7.3: Topología en estrella.

Cualquier mensaje que se quiera transmitir se envía al hub, que lo reenvía al destinatario. Un inconveniente importante de esta topología es que si falla el hub, falla toda la red. Otro inconveniente es que la longitud total de cable es elevada, lo que encarece el sistema. La ventaja es que si hay un problema en un cable, solo se ve afectado un equipo. Esta topología se utiliza en redes de ordenadores, pero no es habitual en sistemas de control distribuido.

- **Anillo:** Cada equipo tiene un conector de entrada y uno de salida, uniéndose un equipo con otro hasta formar un anillo, figura 7.4.

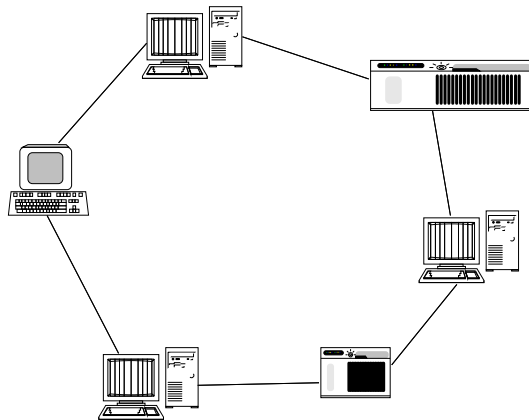


Figura 7.4: Topología en anillo.

Cuando un nodo quiere enviar un mensaje a otro, se lo envía al nodo contiguo. Éste lo recibe y comprueba si el mensaje es para él, transmitiéndolo al siguiente nodo en caso contrario. La longitud de cableado es similar a la topología en bus. Un inconveniente de esta topología es que si falla un equipo o un cable, falla la comunicación de toda la red. Otro inconveniente es que el tiempo de transmisión de un nodo a otro depende mucho de la posición relativa entre éstos. También es un inconveniente el hecho de que para transmitir un mensaje de un nodo a otro, todos los nodos intermedios tengan que intervenir. Esto hace que los nodos necesiten un dispositivo hardware más complejo (y más caro). Esta topología se utiliza en redes de ordenadores, pero no en sistemas de control distribuido.

- **Bus:** Existe un solo cable, llamado bus, al que se conectan todos los equipos, figura 7.5.

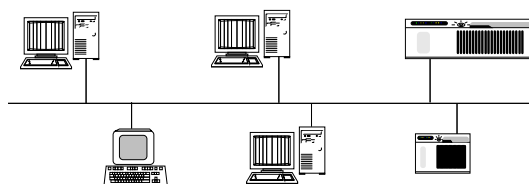


Figura 7.5: Topología en bus.

Tiene la ventaja de que la longitud total de cable es menor. Otra ventaja importante es la facilidad de conexión de nuevos equipos. El inconveniente mayor es que si hay un problema en el cable se ven afectados muchos equipos. El fallo de un equipo, sin embargo, no afecta a la red.

- **Árbol:** Es similar a la topología bus, pero del bus principal pueden salir ramas largas a las que se conectan los equipos de forma similar a como se conectan en el bus principal, figura 7.6.

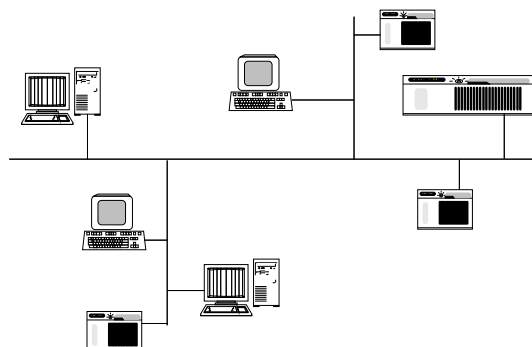


Figura 7.6: Topología en árbol.

Las características son similares a la topología de bus, pero tienen la ventaja adicional de que el cableado es más versátil.

Las redes de bus y de árbol son las que se utilizan en sistemas de control distribuido industriales. Estas redes suelen denominarse **buses de campo**, y permiten la conexión de autómatas programables con sensores, controladores PID, variadores de frecuencia, otros autómatas e incluso ordenadores.

Medio físico

Se refiere al medio a través del cual se transmiten los datos en una red de comunicaciones. En general, las comunicaciones pueden ser por cable o por radio.

La transmisión por radio es especialmente útil cuando las distancias entre los nodos son grandes o cuando hay que comunicar con un equipo móvil, como un vehículo de transporte auto-guiado, por ejemplo. La ventaja de las comunicaciones inalámbricas es que no requieren la instalación de cable. Los inconvenientes más importantes son dos: el primero es que la velocidad de transmisión no es muy alta, y el segundo es que se necesita un emisor y un receptor de radio en cada nodo de la red, lo cual encarece el equipo de comunicación respecto de los usados en la comunicación por cable. Actualmente existen algunos estándares de comunicaciones inalámbricas orientados a aplicaciones de automatización como el ZigBee, concebida para aplicaciones que requieren comunicaciones seguras con baja tasa de envío de datos y bajo consumo. Los tipos de redes inalámbricas se tratarán con más detalle en la Sección 7.7.

En cuanto a los tipos de cables que forman las redes cableadas existe una gran diversidad. Éstos pueden ser:

- **Par trenzado sin apantallar:** Son pares de conductores trenzados entre sí. La señal se transmite como la diferencia de tensión entre los dos conductores del par. En función del tipo de red el cable puede estar formado por un solo par, o por varios pares. El trenzado reduce el ruido electromagnético captado por el cable. Tiene la ventaja de ser el más barato y fácil de instalar. No todos los cables de par trenzado son iguales, sino que existen distintos tipos con distintas características en función del conductor del que están hechos. Los hay que permiten velocidades de transmisión de hasta 10 Mb/s.
- **Par trenzado apantallado:** Es igual al anterior, pero con una pantalla conductora rodeando los pares trenzados que llevan la señal. Esta pantalla se conecta a masa, mejorando mucho la inmunidad frente a ruidos electromagnéticos. Debido a esta mejora permiten una mayor velocidad de transmisión. El inconveniente es su mayor precio.
- **Cable coaxial:** Consta de un conductor central rodeado de aislante y de una pantalla metálica exterior. La señal se transmite como diferencia de tensión entre el conductor y la pantalla. Son muy inmunes al ruido electromagnético.
- **Fibra óptica:** Constan de varias fibras ópticas en el mismo cable. Éstas fibras pueden ser de plástico o de vidrio. Las de vidrio son más caras y más difíciles de instalar (menos flexibles), pero presentan una atenuación menor de la señal, por lo que sirven para distancias mayores sin repetidores. La gran ventaja de la fibra óptica es que no se ven afectadas en absoluto por el ruido electromagnético, además de permitir velocidades de transmisión muy elevadas (100 Mb/s o más). El inconveniente respecto de los cables conductores es la mayor dificultad de instalación y el mayor precio, tanto del cable como de los equipos de transmisión.

El medio físico que más se utiliza en buses de campo es el par trenzado, apantallado o sin apantallar.

Características eléctricas de las señales

Cuando el medio físico es el par trenzado o el cable coaxial, una propiedad importante de la red son las características eléctricas de las señales que permiten la transmisión de la información. La transmisión siempre se realiza en serie. Se transmite un bit detrás de otro hasta completar todos los bits que forman el mensaje. Las características eléctricas de las señales determinan cuándo el bit transmitido vale 1 y cuándo vale 0. Lo más habitual es la transmisión diferencial en la que se utilizan dos hilos, y la señal vale 0 ó 1 en función de la diferencia de tensión entre los dos hilos. La tensión absoluta de los hilos (tensión en modo común) no afecta al valor de la señal transmitida. Si los dos hilos están trenzados, el ruido electromagnético afecta por igual a los dos, por

lo que no afecta a la diferencia de tensión entre los mismos, y por tanto a la señal transmitida. De esta forma se consiguen distancias de transmisión muy largas.

En la mayoría de redes de comunicación, la alimentación de los elementos de la red va aparte de los dos hilos por los que se transmite la señal. Sin embargo, en algunos buses comerciales los dos hilos trenzados permiten transmitir la señal y dar alimentación a los elementos al mismo tiempo. Esto sucede, sobre todo, en buses de nivel bajo que comunican un autómata programable con sensores o actuadores. En ese caso, entre los dos hilos hay siempre una diferencia de tensión positiva (que alimenta a los sensores), transmitiéndose la señal digital mediante la variación de esa tensión entre dos valores. Por ejemplo, se podría tener una tensión media de 20 V, y que la tensión en el bus fuera de 22 V para transmitir un 1, y de 18 V para transmitir un 0.

Método de arbitraje

Se denomina arbitraje o método de acceso, al método por el cual se distribuye entre los nodos de la red el derecho a utilizar ésta para transmitir un mensaje. Los métodos de acceso más extendidos son:

- **Maestro único.** Uno de los nodos de la red es el maestro, siendo los demás esclavos. El maestro es el único que puede utilizar la red cuando quiera. El esclavo sólo puede transmitir después de haber recibido un mensaje del maestro, como respuesta a dicho mensaje. Este método se usa sobre todo en redes de bus o de árbol.
- **Pase de testigo.** Es similar al anterior, pero el maestro no es siempre el mismo nodo. El nodo que actúa como maestro en un momento dado es el que tiene el testigo que se va pasando de un nodo a otro de la red. El paso del testigo no es más que la transmisión de un mensaje especial. En un momento dado sólo puede transmitir el nodo que tiene el testigo (que es el maestro en ese momento). Cuando éste termina sus transmisiones pasa el testigo al siguiente nodo. El paso del testigo de un nodo a otro puede ser gestionado por un dispositivo especial de la red.
- **CSMA/CD, *Carrier Sense Multiple Acces / Collision Detection*.** Cualquier nodo puede empezar a transmitir cuando lo requiera. Simplemente debe comprobar que el bus no está siendo utilizado por otro nodo. Si está siendo utilizado, el nodo que desea transmitir espera hasta que quede libre. El inconveniente es que puede haber colisiones si dos nodos empiezan a transmitir al mismo tiempo. En este método, si hay una colisión, los dos nodos dejan de transmitir, y esperan un tiempo para intentarlo de nuevo. Evidentemente, el tiempo de espera debe ser diferente para cada nodo con el fin de evitar una nueva colisión.

- **CSMA/BA:** *Carrier Sense Multiple Acces / Bitwise Arbitration*. Es similar al anterior, salvo que cuando hay una colisión dejan de transmitir todos los nodos menos uno que es el que tiene mayor prioridad. Además, este nodo no se ve afectado por la colisión (no pierde nada de tiempo en la transmisión a pesar de la colisión). La prioridad se establece por la dirección del nodo que transmite (a menor dirección mayor prioridad).

Modo de funcionamiento

Además de la forma de arbitrar el acceso a la red descrita anteriormente, cada red de comunicación comercial utiliza un modo de funcionamiento. Los más comunes son:

- **Maestro-esclavo con maestro único.** En este caso todos los esclavos están siempre a la escucha, o sea, sin transmitir. Cada esclavo tiene un número de nodo (dirección) diferente. Cuando el maestro quiere leer una variable de un esclavo, o le quiere transmitir el valor de una variable, le envía un mensaje cuya cabecera contiene el número de nodo de ese esclavo. Cada esclavo recibe el mensaje y comprueba si el número de nodo coincide con el suyo propio. Si no coincide, el mensaje se descarta. Si coincide, el mensaje se procesa, y en el caso de ser solicitada una respuesta, el esclavo la transmite hacia el maestro. Son típicas las redes que admiten 31 esclavos y 1 maestro, que es el número de direcciones de nodo que se pueden codificar con 5 bits. Normalmente, en este tipo de redes el maestro realiza ciclos en los que comunica sucesivamente con cada uno de los esclavos.
- **Maestro-esclavo con varios maestros.** Igual que el anterior, pero puede haber varios nodos que pueden ser maestros. Si el arbitraje es por pase del testigo, los maestros se van pasando el testigo de uno a otro, por lo que en cada momento hay un único nodo que actúa como maestro. Si el arbitraje es de acceso múltiple, puede haber varios maestros actuando a la vez.
- **Punto a punto.** Todos los nodos tienen el mismo rango y pueden transmitir un mensaje a cualquier otro nodo. El arbitraje puede ser por paso de testigo, en cuyo caso cada nodo solo puede transmitir cuando tiene el testigo, o por acceso múltiple, donde cada nodo puede transmitir siempre que el bus esté libre.
- **Productor/consumidor.** Cada nodo de la red es productor de una serie de variables, o sea, genera esas variables, y es consumidor de otra serie de variables que son producidas por otros nodos. El nodo que produce una variable envía un mensaje dirigido al resto de nodos en el que transmite el valor de la variable en cuestión. Los nodos consumidores de esa variable almacenan el valor de la misma en su memoria. Normalmente,

se tiene un nodo especial que decide qué variables deben transmitirse. El procedimiento más habitual es que se transmitan cíclicamente todas las variables de la red, con lo que se garantiza que los consumidores tienen un valor actualizado de dichas variables.

- **Cliente/Servidor.** Un nodo es el servidor, que almacena la información, por ejemplo la configuración o los programas de los nodos clientes. Los clientes acceden al servidor para descargarse la información que necesitan. Este modo de funcionamiento se utiliza, no en la fase de intercambio de datos entre equipos, sino en la fase de configuración o programación, por ejemplo, para cargar un programa desde un ordenador (servidor) en un autómatas programable (cliente) a través de la red de comunicación.

Una red de comunicación comercial puede funcionar con una mezcla de varios de los métodos descritos anteriormente.

Tamaño y estructura de los mensajes

El tamaño de los mensajes es el número de bits que contienen los mensajes que se transmiten entre los nodos de la red. Los buses de nivel más bajo, los que comunican únicamente sensores y actuadores con un equipo de control, tienen mensajes muy cortos, pues se transmite muy poca información (valores de los sensores y actuadores), y es importante que se refresquen esos valores con la mayor frecuencia posible. Si el bus es de nivel más alto, que comunica autómatas entre sí y con ordenadores, por ejemplo, el tamaño de los mensajes también es mayor pues la cantidad de información que se tiene que transmitir es mayor, y no es crítico el refresco muy frecuente. En la práctica, los tamaños van desde unos pocos bytes hasta cientos o miles de bytes.

La estructura de los mensajes determina el significado de cada bit que se transmite. Por ejemplo, podría indicar que en primer lugar hay un bit de inicio, después hay 5 bits que definen la dirección del que envía el mensaje, después otros 5 bits que definen la dirección del destinatario del mensaje, después 8 bits que definen el código de la operación, después 16 bits que definen el dato que se transmite, después 5 bits para detección de errores, y por último un bit de stop.

El tamaño y estructura de los mensajes junto con el modo de funcionamiento y el método de arbitraje definen lo que se conoce como protocolo de comunicación.

Máximo número de nodos

El máximo número de nodos que admite la red depende fundamentalmente de las características eléctricas de la misma, aunque puede depender también del protocolo. Por ejemplo, si se reservan únicamente 5 bits para la dirección

de cada nodo sólo puede haber $2^5 = 32$ nodos, puesto que las direcciones de cada nodo deben ser diferentes.

Las redes de más bajo nivel suelen admitir un número de nodos menor que las de más alto nivel. El número de nodos varía entre 32 y varios cientos, incluso miles.

Distancia de transmisión

La distancia de transmisión depende únicamente de las características eléctricas de la red. Las redes de más alto nivel suelen admitir distancias mayores. Las distancias varían entre unos cien metros y varios kilómetros. En algunas redes se puede aumentar la distancia utilizando unos dispositivos repetidores.

Velocidad de transmisión

La velocidad de transmisión mide el número de bits que se transmiten por segundo a través de la red. Esta velocidad depende de las características eléctricas de la red. Para una misma red la velocidad máxima que se puede conseguir depende de la distancia de transmisión. A mayor distancia, menor velocidad. Las velocidades de las redes comerciales pueden variar entre unos 10 Kb/s hasta unos 10 Mb/s.

Es importante resaltar que la velocidad de transmisión no es una medida completa de la rapidez de funcionamiento de la red, ya que ésta depende también del número de nodos y de la cantidad de información que se tiene que transferir. Por ejemplo, los buses de sensores y actuadores con un máximo de 32 nodos no necesitan tener velocidades de transmisión muy altas, pues la cantidad de información transmitida es pequeña, y se puede realizar un ciclo completo de comunicación con todos los nodos en unos pocos milisegundos. Si en lugar de 32 nodos hay 128, y además son nodos más complejos que tienen más información que transmitir, la misma velocidad de transmisión dará lugar a un ciclo de comunicación completo mucho más largo, por lo que se necesitará una velocidad de transmisión mayor.

7.4 ENLACES ESTÁNDARES PARA COMUNICACIONES DIGITALES

Muchas de las características comentadas anteriormente dependen de manera directa de las propiedades eléctricas de la red y de la forma en que los datos (bits, bytes) son transmitidos. Por ejemplo, sólo es posible conectar más de dos nodos a una red si los circuitos electrónicos de los nodos permiten dicha interconexión.

Existen varios estándares que especifican las propiedades eléctricas y la forma de transmitir los datos. A continuación, se comentan dos de los más usados: el RS232 y el RS485, siendo este último la forma de enlace en buses de campo con un uso muy extendido.

RS232

La conexión serie RS232 permite comunicar únicamente dos elementos entre sí. Ejemplos habituales son la comunicación de un autómatas con un terminal, o un autómatas con un ordenador, o un autómatas con un variador de frecuencia.

La conexión RS232 básica dispone de 3 hilos; R (recibir), T (transmitir) y masa. Las señales R y T se cruzan: el R de un equipo con el T del otro y viceversa. El uno lógico queda definido por una tensión de -12 V respecto de masa, mientras el 0 lógico queda definido por una tensión de $+12\text{ V}$.

La velocidad de transmisión puede ser como máximo de 38.4 Kbits/s . La distancia máxima de transmisión es de unos 10 a 15 metros.

Los datos se transmiten en serie en paquetes de 1 carácter. El formato del mensaje suele tener un bit de inicio, 7 u 8 bits que definen el carácter que se transmite, 1 bit de paridad, y 1 o 2 bits de stop.

La comunicación es asíncrona. La transmisión del carácter empieza con un flanco de subida de la señal. A partir de ese flanco el receptor lee un bit de la línea cada cierto tiempo, determinado por la velocidad de transmisión. Por esta razón el equipo que transmite y el que recibe deben estar configurados a la misma velocidad de transmisión. La transmisión del carácter finaliza cuando se han transmitido todos los bits. El número de bits total del mensaje debe ser también el mismo para el equipo que transmite y para el que recibe.

El estándar RS232 sólo define las características eléctricas de la transmisión. Para que dos equipos puedan comunicarse entre sí, además deben utilizar el mismo protocolo, o sea, el significado de los mensajes (formados por un conjunto de caracteres) que se envían.

RS485

Se puede considerar como una modificación de la conexión serie RS232 que mejora sustancialmente las características de aquella, tanto en velocidad de transmisión, como en distancia, permitiendo además la conexión de hasta 32 elementos entre sí. Define únicamente las características eléctricas de la transmisión pero no el protocolo de comunicación.

La transmisión se realiza a través de 2 hilos en modo diferencial, o sea, ninguno de los dos hilos está unido a masa. El uno lógico queda definido por una tensión mayor en un hilo que en el otro, típicamente 5 V mayor, mientras el cero lógico queda definido por una tensión menor en este hilo que en el otro, típicamente -5 V . La tensión total respecto de masa de los hilos puede variar entre -7 V y $+12\text{ V}$, pero no afecta a la señal transmitida. Los dos hilos suelen ir trenzados. De esta forma, las interferencias electromagnéticas afectan por igual a los dos hilos, y no modifican la tensión diferencial entre ellos, solo modifican la tensión absoluta. Ésta es la razón por la que son muy inmunes al ruido, y permiten distancias de transmisión muy largas.

La velocidad máxima de transmisión es de 10 Mbit por segundo, aunque depende de la distancia. La distancia máxima de transmisión es de 1.2 km .

Se pueden conectar hasta 32 elementos al mismo bus de 2 hilos. En un momento dado, solo un elemento puede estar transmitiendo a través de los 2 hilos. Todos los demás deben estar escuchando. Si dos elementos intentan transmitir a la vez, se produce una colisión que no es detectada, por lo que el arbitraje debe impedir a toda costa que dos elementos intenten transmitir a la vez. Este arbitraje es el de maestro-esclavo, pudiendo haber varios maestros que se pasan un testigo entre ellos.

Cada elemento del bus tiene una dirección asociada: un número entre 0 y 31. Los mensajes que se envían incluyen la dirección del destinatario. De esta forma un nodo sólo hace caso del mensaje que contiene su dirección. Cuando solo hay un maestro, éste suele tener la dirección 0.

Diversos protocolos comerciales estándar utilizan la conexión física RS485, Modbus o ProfiBus DP, por ejemplo. Además, muchos equipos industriales (sensores, controladores PID, variadores de frecuencia, autómatas) utilizan esta conexión física con un protocolo específico propio del fabricante del equipo, que suele funcionar en modo maestro-esclavo simple.

Existen convertidores de RS232 a RS485 que permiten conectar un equipo con puerto serie RS232 a un bus RS485 a 2 hilos. El inconveniente es que la velocidad queda limitada en ese caso por el RS232, con un máximo de 38.4 Kb/s.

7.5 BUSES DE CAMPO

Un bus de campo es un sistema de transmisión de información que simplifica enormemente la instalación y operación de máquinas y equipamientos industriales utilizados en procesos de producción. Típicamente, son redes digitales, bidireccionales, multipunto, montadas sobre un bus serie, que conectan dispositivos de campo como PLC, sensores, controladores PID, variadores de frecuencia y otros actuadores. Cada dispositivo de campo incorpora cierta capacidad de proceso, que lo convierte en un dispositivo inteligente. Cada uno de estos elementos será capaz de ejecutar funciones simples de diagnóstico, control o mantenimiento, así como de comunicarse bidireccionalmente a través del bus cableado. Las principales ventajas de los buses de campo se derivan parcialmente del hecho de que la información se transmite/recibe digitalmente, lo cual implica:

- Reducción significativa del cableado. El hecho de que los buses de campo sean redes multipunto hace necesario sólo un cable central al que se conectan todos los dispositivos de campo, como sensores y actuadores.
- Facilita de forma considerable las labores de mantenimiento: el sistema de bus de campo permite al operador verificar los estados de los dispositivos de campo y de las interacciones entre éstos. De esta forma la detección de averías se simplifica, reduciéndose los tiempos de reparación. Además,

se dispone de sistemas de diagnóstico en línea que permiten predecir las necesidades de mantenimiento/calibración de los dispositivos de campo.

- **Flexibilidad:** la disponibilidad de dispositivos de campo capaces de realizar funciones de control, diagnóstico y comunicación hace más fácil la expansión/renovación del sistema que si todas estas tareas se desarrollasen en un sistema centralizado.
- **Simplificación de la recopilación de la información del proceso:** los valores de las variables del proceso están disponibles para todos los dispositivos de campo en unidades de ingeniería. Esto elimina las tareas de conversión a estas unidades y permite al sistema de control dedicarse a tareas más importantes. Es posible, además, la comunicación entre varios dispositivos de campo y de éstos a su vez con el sistema de control. Esto permite que varios dispositivos de campo relacionados trabajen de forma combinada.

Características generales de los buses de campo

Actualmente existen un gran número de aplicaciones de control distribuido que requieren diferentes características de los buses de campo que se utilizan en cada una. Por ejemplo, para algunas aplicaciones se necesitan redes con velocidades de comunicación elevadas, para otras se necesitan redes seguras, mientras que en otros casos se precisa de redes con mayor funcionalidad. Estas diferencias en los requerimientos de las aplicaciones ha dado lugar a buses de campo con diferentes prestaciones, los cuales pueden clasificarse dentro de las siguientes categorías:

Buses de alta velocidad y baja funcionalidad

Están diseñados para integrar dispositivos simples como finales de carrera, fotocélulas, relés y actuadores simples, funcionando en aplicaciones de tiempo real, y agrupados en una pequeña zona de la planta, típicamente una máquina. Algunos ejemplos son:

- **CAN:** diseñado originalmente para su aplicación en vehículos.
- **SDS:** bus para la integración de sensores y actuadores, basado en CAN.
- **ASI:** bus serie diseñado por Siemens para la integración de sensores y actuadores.

Buses de alta velocidad y funcionalidad media

Permiten el envío eficiente de mensajes de tamaño medio. Estos mensajes permiten que los dispositivos conectados tengan mayor funcionalidad de modo

que incluyan aspectos como la configuración, calibración o programación. Por tanto, son buses capaces de controlar dispositivos de campo complejos, de forma eficiente y a bajo costo. Normalmente, disponen de funciones utilizables desde programas basados en ordenador para acceder, cambiar y controlar los diversos dispositivos que constituyen el sistema. Algunos incluyen funciones estándar para distintos tipos de dispositivos que facilitan su conexión aun siendo de distintos fabricantes. Algunos ejemplos son:

- DeviceNet: desarrollado por Allen-Bradley, utiliza como base el bus CAN, e incorpora una capa de aplicación orientada a objetos.
- LONWorks: red desarrollada por Echelon.
- BitBus: Red desarrollada por INTEL.
- DIN MessBus: estándar alemán de bus de instrumentación, basado en comunicación RS-232.
- InterBus-S: bus de campo alemán de uso común en aplicaciones medias.

Buses de altas prestaciones

Suelen ser redes multi-maestro en las que las comunicaciones maestro-esclavo se realizan según el esquema pregunta-respuesta. La recuperación de datos desde el esclavo tiene un límite máximo de tiempo para favorecer la velocidad. Tienen la capacidad de hacer varios tipos de direccionamiento: único (*unicast*) y múltiple (*multicast* y *broadcast*). Soportan la petición de servicios a los esclavos basada en eventos y la comunicación de variables y bloques de datos, además permiten la descarga y ejecución remota de programas. Estos buses presentan altos niveles de seguridad, opcionalmente, con procedimientos de autenticación.

Algunos ejemplos son:

- Profibus.
- WorldFIP.
- Fieldbus Foundation.

Buses para áreas de seguridad intrínseca

Incluyen modificaciones en el medio físico de comunicación para cumplir con los requisitos específicos de seguridad intrínseca en ambientes con atmósferas explosivas. La seguridad intrínseca es un tipo de protección por la que el componente en cuestión no tiene posibilidad de provocar una explosión en la atmósfera circundante. Un circuito eléctrico o una parte de un circuito tienen seguridad intrínseca cuando alguna chispa o efecto térmico en este circuito

producidos en las condiciones de prueba establecidas por un estándar, dentro del cual figuran las condiciones de operación normal y de fallo específicas, no puede ocasionar una ignición. Algunos ejemplos son:

- HART.
- Profibus PA.
- WorldFIP.

Buses de campo comerciales

Una red de comunicación comercial estándar queda definida por una serie de características que incluyen desde el tipo de cable y las características eléctricas de las señales, hasta el arbitraje y el modo de comunicación, pasando por la estructura y tamaño de los mensajes. A continuación se describen brevemente algunas de las redes de comunicación estándar más utilizadas.

CAN

El estándar CAN define solo la parte física del bus, existiendo diversos protocolos comerciales. El bus es de 2 hilos (par trenzado). Las señales se transmiten también de modo diferencial. Un uno lógico queda definido por una tensión de 5 V entre los dos hilos, mientras un cero lógico queda definido por una tensión de 0 V entre los dos hilos.

Permite una velocidad de hasta 1 Mb/s, y una distancia de hasta 1 km.

El arbitraje es de acceso múltiple (CSMA/BA). Cualquier nodo puede transmitir si el bus no está ocupado. Si se produce una colisión porque dos nodos intentan transmitir a la vez, se impone el que tiene la dirección más baja, por lo que el otro deja de transmitir sin afectar a la transmisión del primero. Esto permite que haya varios maestros, aunque también se puede funcionar como maestro/esclavo.

Para realizar este arbitraje, la primera parte del mensaje contiene siempre la dirección del nodo que lo envía. Las características eléctricas del bus hacen que si un nodo intenta poner un uno en el bus, y otro nodo intenta poner un cero, el bus se queda en cero. De esta forma cada nodo va transmitiendo bit a bit y comprobando que efectivamente el bus va tomando los valores que él transmite. Si en un momento dado el nodo intenta poner un uno, pero el bus se queda en cero, eso significa que hay otro nodo transmitiendo, por lo que este nodo deja de transmitir y espera a que el bus quede libre para intentar de nuevo la transmisión. Como lo primero que contiene el mensaje es la dirección del nodo que transmite, el que tiene la dirección menor se impone y continúa la transmisión. Esta forma de resolver la colisión permite que no se pierda nada de tiempo pues el nodo prioritario continúa la transmisión normalmente.

Entre los protocolos comerciales estándar que utilizan el bus CAN destacan: DeviceNet, SDS y CANOpen.

ASI bus

Es un bus comercial de 2 hilos que se utiliza para conectar sensores y actuadores a un autómatas programable. Define tanto la parte física como el protocolo de comunicación. A través de los 2 hilos permite no solo transmitir la información, sino también alimentar a los sensores y actuadores siempre que su consumo sea pequeño.

El arbitraje es maestro-esclavo. El autómatas es el maestro, y los sensores, los esclavos. Hasta 31 esclavos pueden conectarse al bus. La distancia máxima es de 100 m. La topología de la red es en bus lineal o en árbol. La velocidad de transmisión no es muy alta, pero permite un tiempo de ciclo para los 31 nodos de unos 5 ms.

Modbus

Es un protocolo desarrollado por *Gould Modicon*, actualmente *Schneider Electric*. Las comunicaciones se soportan sobre enlaces de comunicación serie RS-232 ó RS-485. Esto hace que Modbus adolezca de las limitaciones de estos enlaces:

- La velocidad de comunicación es relativamente lenta: 9600-115000 baudios, o sea, entre 0.010 Mbps y 0.115 Mbps. Las velocidades de otras redes de control están entre los 5 y 10 Mbps y Ethernet llega hasta los 100 Mbps.
- Sólo soporta dos nodos si el enlace se hace sobre RS-232 y entre 20 y 30 nodos si se utiliza RS-485. La única forma de superar esta cantidad es mediante una estructura anidada de maestros y esclavos que incrementa su complejidad con el número de nodos.
- Sólo se puede tener un nodo maestro y el resto serán esclavos, de forma que sólo un dispositivo, el maestro, dispone de los datos en tiempo real de los valores de los sensores enviados por los esclavos.

A pesar de estas limitaciones, Modbus está ampliamente extendido y es aceptado por muchos fabricantes y usuarios como un estándar para comunicaciones industriales con probadas capacidades.

Modbus tiene algunas características que son fijas y otras que son elegidas por el usuario. Entre las fijas, están el formato y secuencia de las tramas, tratamiento tanto de los errores de comunicación como de las condiciones de excepción. Entre las características que puede seleccionar el usuario están el medio, las características y el modo de transmisión que puede ser RTU (formato hexadecimal) o ASCII (formato ascii).

El modo RTU, también llamado Modbus-B (Modbus-Binary), es el preferido ya que los mensajes son más cortos y por tanto la comunicación más ágil.

El modo ASCII, también llamado Modbus-A, es utilizado principalmente para realizar pruebas.

Los paquetes Modbus pueden ser encapsulados en paquetes TCP/IP, y de esta forma ser enviados a través de una red de área local o global.

Modbus-Plus. Modbus-Plus es una variante de Modbus que elimina la limitación de un sólo maestro. Como se ha comentado, para compartir información entre varios maestros en una red Modbus es necesario crear una estructura jerárquica cuya complejidad crece con el número de nodos. El protocolo Modbus-Plus fue uno de los primeros basado en paso de testigo.

Modbus-Plus permite hasta 64 nodos en un segmento de red. De estos 64 nodos, 32 pueden ser conectados a un cable de 450 m. Para conectar los otros 32 es necesario un repetidor mediante el cual el cable puede ser extendido hasta un máximo de 1800 m. A cada nodo se le asigna una dirección entre 01 y 64. Se pueden crear redes de mayor tamaño uniendo varios segmentos de red mediante puentes.

Los mensajes generados en un segmento de red pueden ser direccionados a otros segmentos a través de los puentes. El paso de un mensaje por un puente añade retrasos que no son conocidos a priori y son aleatorios. Un comportamiento determinista de la temporización de los mensajes se tiene sólo dentro de un segmento de red.

Los mensajes incluyen una ruta que permite su paso entre dispositivos. La ruta incluye la dirección del segmento de red donde está el nodo destino, así como la dirección de éste dentro de dicho segmento. Además, incluye las direcciones de los segmentos intermedios entre los que son origen y destino del mensaje.

Cada segmento de red mantiene su propia secuencia de paso del testigo, independiente de los otros segmentos. El testigo no pasa a través de los *puentes*. La secuencia del paso del testigo depende de la dirección de cada nodo: de las más bajas a las más altas, de forma cíclica. Si un nodo no está operativo, el testigo se pasa al siguiente después de una espera de aproximadamente 100 ms y su dirección se borra de la lista de direcciones. Cuando el nodo esté nuevamente activo, su dirección se incluye nuevamente pasados 5 ó 15 segundos en el peor de los casos. Este proceso es transparente al usuario.

Mientras un nodo tiene el testigo, envía sus mensajes si tiene algo que transmitir. El nodo receptor envía inmediatamente un mensaje de confirmación al nodo transmisor. Si el nodo transmisor solicita algún dato, entonces el receptor prepara la trama y la envía la próxima vez que tenga el testigo. Después de que un nodo envíe todos sus mensajes, pasa el testigo al siguiente nodo.

Cuando un nodo tiene el testigo, éste puede además escribir datos en una base de datos global al segmento de red. Los datos entonces son enviados a los nodos formando parte del mensaje del testigo que llega a todos los nodos. Este es el mecanismo que se utiliza para hacer llegar información de una forma rápida a los nodos de un segmento.

ProfiBus

ProfiBus es un estándar ampliamente establecido en la industria de procesos y manufactura, el cual soporta la conexión de bloques de sensores, de actuadores, de sub-redes de menor tamaño e interfaces de operador. La conexión al bus es soportada por un circuito integrado producido por varios fabricantes.

Una red ProfiBus puede tener hasta 127 nodos distribuidos en una distancia de hasta 24 Km, si se usan repetidores y fibra óptica. La velocidad varía desde los 9600 bps hasta los 12 Mbps. El tamaño de los mensajes puede ser de hasta 256 bytes: 12 byte de cabecera y hasta 244 byte de datos. El acceso al medio se hace por encuesta o por paso de testigo.

Existen dos tipos de nodos: maestros y esclavos. Los maestros controlan el bus y cuando tienen acceso a él pueden iniciar una comunicación enviando mensajes. Los esclavos sólo reconocen mensajes recibidos, enviados por nodos maestros, a los que responden con otros mensajes en caso necesario. Los esclavos suelen ser sensores o actuadores con sus correspondientes transmisores.

Existen tres versiones ProfiBus: Profibus DP (maestro/esclavo), ProfiBus FMS (maestro múltiple/punto a punto) y ProfiBus PA (intrínsecamente segura).

Profibus DP (*distributed peripheral*). Está orientada al intercambio de datos entre maestros (controladores) y esclavos (sensores, actuadores). Permite el uso de varios dispositivos maestros, en cuyo caso cada esclavo es asignado a un único maestro. Esto quiere decir que un único maestro puede escribir en un esclavo, aunque todos los maestros pueden leer datos de un esclavo. El intercambio de datos es generalmente cíclico. Es adecuada para aprovechar el cableado industrial para la transmisión de señales de 4-20 mA y alimentación de 24 V.

ProfiBus FMS (*Fieldbus Message Specification*). Está orientada al intercambio de datos entre maestros. Utiliza un modelo de comunicación en el que los nodos se comportan como iguales entre sí, es decir, actúan simultáneamente como esclavos o maestros respecto a los demás nodos de la red. Este modelo se conoce como punto a punto (*peer to peer*). Soporta hasta 126 nodos de los cuales todos pueden ser esclavos. La cabecera de los mensajes es más compleja y larga que la de los mensajes de ProfiBus DP.

COMBI mode. Es cuando FMS y DP son usados simultáneamente en una misma red, para lo que se utilizan dispositivos que soportan este modo. Una arquitectura típica sería un PC (maestro primario) y un PLC (maestro secundario) al cual están conectados sensores y actuadores (esclavos). Los mensajes entre el PC y el PLC se enviarían vía FMS, y entre el PLC y los sensores y actuadores, vía DP.

Tabla 7.1: Relación entre velocidad y distancia para una red ProfiBus

Velocidad	Distancia
9.6Kbps	1200m
19.2Kbps	1200m
93.75Kbps	1200m
187.5Kbps	600m
500Kbps	300m
1.5Mbps	200m
12Mbps	100m

ProfiBus PA. Es similar a ProfiBus DP, excepto que los niveles de corrientes y voltajes son reducidos para cumplir los requerimientos de seguridad para la industria de procesos. Muchas tarjetas maestras DP/FMS soportan ProfiBus PA, pero su conexión requiere el uso de un convertidor DP/PA.

Tanto la variante DP como PA se basan en el estándar RS-485, aunque en el caso de PA se introduce las variaciones para garantizar los requerimientos de seguridad. ProfiBus-DP tiene las siguientes características:

- La topología de red es un bus lineal.
- El medio es un par trenzado con o sin pantalla, dependiendo de la aplicación.
- La velocidad de transmisión puede estar entre los 9.6 Kbps y los 12 Mbps, dependiendo de la distancia según la tabla 7.1.

En ProfiBus existen dos formas de controlar el acceso al medio: paso de testigo y encuesta. El testigo circula siguiendo un anillo lógico de un maestro al siguiente en orden de dirección ascendente, al final el maestro de dirección más alta pasa el testigo al maestro de dirección más baja. El método de encuesta, por otro lado, se utiliza para la comunicación entre un maestro que tiene el testigo y los esclavos asociados a él.

A continuación se describe de forma más detallada el mecanismo de paso del testigo en una red ProfiBus.

Inicialización

El maestro con menor dirección se auto-envía dos tramas de testigo, lo cual informa al resto de maestros que el anillo lógico está formado sólo por el maestro de menor dirección. Entonces comienza a enviar un mensaje especial llamado “mensaje de GAP” a direcciones sucesivas en orden creciente. Al primer maestro que responde a estos mensajes se le pasa el testigo, y el maestro con dirección más baja actualiza su sucesor y además crea una lista llamada “lista de GAP” con las direcciones de las que no recibió respuesta.

Actualización

Cada maestro es responsable de añadir a su lista de GAP los nuevos nodos maestros que se conectan al bus y que tienen direcciones entre la suya y la siguiente dirección registrada en la lista. Una vez el maestro que está en posesión del testigo finaliza el envío de los mensajes, explora una dirección entre la suya y la siguiente que hay en la lista de direcciones. Para ello envía mensajes de GAP a esas direcciones. Cuando recibe una respuesta, entonces pasa el testigo al nuevo maestro el cual actualiza también su nuevo sucesor.

Control de acceso al medio

El tiempo que un maestro estará en posesión del testigo se calcula a partir de criterios temporales. Para ello se definen los siguientes parámetros:

- Tiempo de rotación deseado (T_{TR}): es el tiempo que se desea transcurra entre dos recepciones del testigo consecutivas por una unidad maestro. Este tiempo se debe fijar cuando se diseña la red y se considera el mismo para todos los nodos maestros.
- Tiempo real de rotación del testigo (T_{RR}): es el tiempo transcurrido entre dos recepciones del testigo consecutivas por una unidad maestro.

Cuando un maestro recibe el testigo, calcula el tiempo que dispondrá para acceder al medio, como $T_{TH} = T_{TR} - T_{RR}$, donde T_{TH} se conoce como tiempo de retención del testigo.

Si $T_{TH} > 0$, el maestro envía los mensajes según su prioridad. Para ello, ProfiBus hace la siguiente distinción entre los mensajes:

- Mensajes de ALTA prioridad.
- Mensajes de BAJA prioridad, los cuales pueden ser:
 - Mensajes cíclicos.
 - Mensajes No cíclicos.
 - Mensajes de GAP.

Primero se envía los mensajes de alta prioridad, después los de prioridad baja en el siguiente orden: primero los cíclicos, luego, si queda tiempo ($T_{TH} > 0$) los no cíclicos. Los mensajes de Gap sólo se envían cuando no queda ningún mensaje por enviar.

Si $T_{TH} < 0$, entonces sólo se envía un mensaje de alta prioridad y se libera el testigo al siguiente maestro.

Paso del testigo

El maestro activo envía el testigo al maestro siguiente con mayor dirección. Si el transmisor del testigo no detecta ninguna actividad en el bus, dentro de

un lapso de tiempo determinado, reenvía el testigo al mismo maestro hasta en dos ocasiones más. Si no detecta actividad en el bus, entonces envía el mensaje de testigo al maestro con dirección siguiente. Este proceso se repite hasta que un maestro comienza a utilizar el bus.

Recuperación del testigo

Cuando el testigo se pierde, por algún motivo, no es necesario re-inicializar el sistema. Después de un tiempo de espera sin recibir el testigo, el maestro con dirección más baja procede con sus mensajes y genera un nuevo testigo que se pasa al maestro con dirección superior.

Otros buses comerciales

En las secciones anteriores se han descrito las características generales de algunos de los buses de campo más utilizados actualmente. Existen otros muchos que han sido desarrollados por diversos fabricantes para las más diversas aplicaciones. A continuación se presenta una lista con algunos de ellos, donde se especifica la dirección en Internet para obtener información de cada uno.

DeviceNet (<http://www.odva.org>)
AS-I (<http://www.as-interface.com>)
Seriplex (<http://www.seriplex.org>)
CANOpen (<http://www.can-cia.de>)
SDS (<http://www.honeywell.com/sensing/prodinfo/sds/>)
Interbus-S (<http://www.interbusclub.com/>)
WorldFIP (<http://www.worldfip.org>)
ControlNet (<http://www.controlnet.org>)
ARCNet (<http://www.arcnet.com>)
P-Net (http://www.infoside.de/infida/pnet/p-net_uk.htm)
Hart (<http://www.thehartbook.com/default.asp>)
LonWorks (<http://osa.echelon.com/Program/LonWorksIntroPDF.htm>)
BITbus (<http://www.bitbus.org>)
Sercos (<http://www.sercos.com/technology/index.htm>)
Lightbus (<http://www.beckhoff.com/english.asp?lightbus/default.htm>)

7.6 ETHERNET COMO RED DE CAMPO: ETHERNET CONMUTADA

Ethernet se ha convertido en el estándar de comunicación en red más común en el ámbito doméstico y laboral al ser usado en la mayor red de intercambio de datos: Internet. Esto ha provocado el abaratamiento de los componentes hardware que soportan esta tecnología debido a la producción masiva de los mismos. Entre las ventajas de Ethernet está su velocidad, que puede llegar a los 100 Mbit/s para comunicaciones sobre par trenzado. Esta velocidad es muy

superior a las velocidades conseguidas con la mayoría de los buses de campo industriales, las cuales, en general, son menores que 10 Mbit/s.

El bajo precio del hardware y la alta velocidad, unido al aumento de los requerimientos de los sistemas de control, ha motivado un interés creciente por aplicar Ethernet como estándar para la automatización industrial. El principal escollo que se ha tenido que saldar en ese sentido es que Ethernet, en su versión más tradicional, tiene un comportamiento temporal estocástico, o sea que no se puede predecir con exactitud el tiempo que se tardará en recibir la respuesta a un mensaje que se envíe desde un nodo. Desde el punto de vista de la automatización esto es un problema pues en las aplicaciones de control existen restricciones temporales que son críticas para el correcto funcionamiento de los procesos, incluso para la seguridad de los mismos. El manejo de alarmas ante situaciones de peligro es un ejemplo de ello.

El comportamiento no determinista de Ethernet se debe a la conjunción de dos factores: por un lado, que el método de arbitraje usado es el CSMA/CD, en el que la colisión de acceso al medio se resuelve asignando tiempos de espera a los nodos antes de que intenten acceder nuevamente, y por otro, que el dispositivo de interconexión entre los nodos se limita a repetir la información recibida por un puerto en todos los demás. Dicho dispositivo se conoce como repetidor o *hub*, ver figura 7.3.

Algunas variaciones introducidas a Ethernet han mejorado considerablemente sus prestaciones. Concretamente, la sustitución del repetidor por un dispositivo electrónico más sofisticado llamado conmutador (*switch*) ha sido uno de los cambios más importantes. La introducción del conmutador ha dado lugar a lo que se conoce como Ethernet conmutada (*Switched Ethernet*).

El conmutador, al igual que el repetidor, tiene varios puertos a los cuales se conectan los nodos de la red, como se puede ver en la figura 7.7. Sin embargo, a diferencia del repetidor, el conmutador realiza una interconexión más inteligente entre los puertos que permite reducir la congestión de la red, causada por el comportamiento no determinista del arbitraje CSMA/CD.

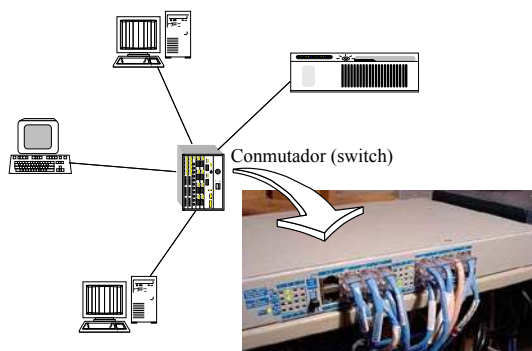


Figura 7.7: Esquema de una red Ethernet con conmutador.

Entre las mejoras que produce el uso del conmutador en una red Ethernet están las siguientes:

- Permite la comunicación full-duplex, o sea que dos dispositivos pueden transmitir datos entre ellos al mismo tiempo, sin tener que esperar a que termine uno para comenzar el otro.
- El conmutador soporta dispositivos que trabajan a velocidades diferentes: 10 Mbits/s, 100 Mbits/s, 1 Gbit/s o 100 Gbit/s. Además, de forma automática adapta la velocidad de cada puerto a la del equipo que se encuentra conectado, lo cual implica un mejor aprovechamiento del ancho de banda total de la red.
- Existen varios mecanismos para asignar prioridades a los mensajes. Los mensajes son transmitidos en función de su prioridad: aquellos con mayor prioridad son transmitidos antes. Esto permite que la información crítica sea transmitida de forma más ágil si se envía en mensajes de prioridad alta.

Estas mejoras han hecho que el comportamiento de Ethernet se adapte mejor a los requerimientos temporales de las aplicaciones de control, y por tanto, su uso se está extendiendo como una alternativa a los buses de campo, especialmente los de nivel intermedio.

7.7 REDES INALÁMBRICAS

El estado actual de la tecnología de comunicaciones por radio frecuencia (RF) ha hecho que en los últimos años se puedan aplicar sistemas de control distribuido mediante conexiones inalámbricas sin incurrir en un incremento de coste excesivo de los componentes, ni exigir excesivos conocimientos de radiofrecuencia a los ingenieros que aplican dicha tecnología al control de la planta. Además, aporta una simplificación importante en lo referente a instalación y cableado.

Clasificación de las redes inalámbricas de RF

Las comunicaciones inalámbricas por RF se pueden clasificar mediante los siguientes criterios:

- **Rango de frecuencia de trabajo y ancho de banda.** Las bandas ISM (Industrial, Scientific and Medical bands) para sistemas de comunicaciones digitales inalámbricas empleando la radiofrecuencia, son las que no necesitan licencia (siempre que no se pasen los límites de potencia) y que además son gratuitas. Las frecuencias de trabajo estandarizadas son: 314 MHz en USA (potencia máxima +30 dBm), 434 MHz (+10 dBm) y 868 MHz (+14 dBm) en Europa en AM o FM.
- **Tipo de modulación digital.** La modulación digital permite un mejor aprovechamiento del canal de comunicación, lo que posibilita transmitir

más información en forma simultánea, protegiéndola de posibles interferencias y ruidos. Básicamente, la modulación consiste en hacer que un parámetro de la onda portadora cambie de valor de acuerdo con las variaciones de la señal moduladora, que es la información que queremos transmitir. Existen tres tipos básicos de modulación:

Modulación por desplazamiento de Amplitud (ASK). En la que se varía la amplitud de la señal portadora para indicar si el dato transmitido es un 0 o un 1. Las ventajas de este tipo de modulación son el sencillo diseño (menor coste) y el bajo consumo, especialmente si se utiliza el método o modulación OOK (On/Off Keying), Modulación On/off, donde con un 0 digital no hay potencia de salida y un 1 digital se entrega toda la señal portadora. La desventaja es la fragilidad en presencia de interferencias por ruido eléctrico, que pueden provocar errores en los datos recibidos.

Modulación por desplazamiento de Frecuencia (FSK). La modulación por desplazamiento de frecuencia (FSK), donde con un 0 digital se transmite una portadora a una frecuencia y con un 1 digital se transmite la portadora a otra frecuencia distinta, con la misma amplitud. La ventaja de este tipo de modulación es la mejor robustez ante la presencia de interferencias. La desventaja es la complejidad del sistema (mayor coste) y el consumo que permanece siempre presente durante la transmisión. Se utiliza en los módems de baja velocidad. Se emplea separando el ancho de banda total en dos bandas; los módems pueden transmitir y recibir datos por el mismo canal simultáneamente.

Modulación por desplazamiento de Fase (PSK). Se codifican los valores binarios como cambios de fase de la señal portadora. Ya sea asignando una fase al valor digital 0 y 1 o un desplazamiento de fase determinado al pasar de un valor a otro (DPSK). Técnicamente, utilizando el concepto de modulación PSK, es posible aumentar la velocidad de transmisión a pesar de los límites impuestos por el canal. De aquí que este tipo de modulación haya sido el que más desarrollo ha tenido en los últimos años.

- **Protocolo de comunicación.** La importancia del protocolo de comunicación ha llevado a las redes a diferenciarse por el protocolo utilizado.

Redes inalámbricas comerciales

Entre las redes inalámbricas comerciales se pueden distinguir las que no utilizan un protocolo establecido, dejando al usuario que utilice el protocolo que considere oportuno para su aplicación y las que tienen un protocolo establecido, que realizan la comunicación entre elementos de forma transparente. Esto permite una puesta en marcha rápida y el uso de componentes más extendidos pero limita la aplicación dependiendo de la idoneidad del protocolo

Redes genéricas sin protocolo establecido

Estos sistemas no utilizan ningún protocolo estándar. Los circuitos integrados dentro de este grupo se basan en un transmisor integrado en un solo circuito, exceptuando la antena, el cristal y algunos componentes externos, sin necesidad de ajustes de RF. Normalmente, la frecuencia de trabajo, la velocidad de transmisión y la potencia de salida son programables. Están por debajo de la potencia máxima permitida sin necesidad de licencia. Son fácilmente conectables a un microcontrolador. El receptor también es un circuito integrado o puede estar integrado en el propio emisor. El receptor dispone de un sistema para dejarlo dormido y activarse rápidamente. Dentro de esta categoría podemos encontrar componentes en Analog Devices (la serie ADF7010-20), Freescale que ofrece un microcontrolador con el emisor integrado, Texas Instruments, etc.

WirelessUSB

Esta red, aunque entraría dentro de las redes sin protocolo establecido, tiene ciertas peculiaridades que hacen que la tratemos por separado. La gran ventaja de esta tecnología, es que ha entrado en el mercado de consumo USB (como ratones, teclados, joysticks... del mercado de la informática), para seguir con el mercado industrial con las ventajas de un muy bajo costo, como aplicaciones de enlace inalámbrico punto a punto o punto a multipunto que no exceda los 64kbps del ancho de banda disponible. De hecho, se trata de una interfaz SPI, que empaqueta los datos entrantes y los prepara para una transmisión sin hilos a 2,4 GHz. El usuario no tiene que preocuparse de codificar, decodificar paquetes o manejar los errores, así como de preparar el enlace de radio. WirelessUSB ofrece al usuario una variedad de opciones desde la transmisión simple entre dos dispositivos o entre un dispositivo master y varios esclavos, en comunicación bidireccional.

ZigBee IEEE 802.15.4.

Iniciado por Philips, Honeywell, Invensys y seguido por Motorola (ahora Freescale), Mitsubishi y así hasta 25 empresas para crear un sistema estándar de comunicaciones inalámbrico y bidireccional, para usarlo dentro de dispositivos de domótica, automatización de edificios (denominado inmótica), control industrial, periféricos de PC y sensores médicos. Los miembros de esta alianza justifican el desarrollo de este estándar para cubrir el vacío que se produce por debajo del Bluetooth. Puede transmitir con un simple protocolo de 20 kB/s hasta 250 Kbps trabajando a una frecuencia de 2,4 GHz con la tecnología GSSS, bajo consumo y rangos entre 10 y 75 metros, aunque las condiciones físicas ambientales son las que determinan las distancias de trabajo. La idea de ponerle el nombre ZigBee vino de una colmena de abejas pululando alrededor de su panal y comunicándose entre ellas.

Bluetooth

Bluetooth opera en una banda no licenciada ISM (Industrial Scientific Medical) de 2.4-2.5 GHz permitiendo la transmisión de voz y datos, de forma rápida y segura con un rango de hasta 10 metros con 1 miliwatio o 100 metros si se usa un amplificador con 100 miliwatios. Puede transferir datos de forma asimétrica a 721 Kbps y simétricamente a 432 Kbps. Se puede transmitir voz, datos e incluso vídeo. Dentro de una aplicación típica de Bluetooth nos podemos encontrar los siguientes elementos:

- Master: es el dispositivo Bluetooth que establece e inicializa la conexión, la secuencia de control hopping y la temporización de los demás dispositivos colocados en lo que se llama una red Piconet.
- Slave: es el dispositivo habilitado en una Piconet. Una red Piconet tiene un máximo de 7 esclavos.
- Piconet: una red de hasta 8 dispositivos conectados (1 maestro + 7 esclavos).
- Scatternet: red formada por diferentes redes Piconet.

WiFi o WLAN IEEE 802.11.

Es un sistema de comunicación sin hilos WLAN (Wireless Local Area Network) que se utilizaba en un principio para redes de PC y periféricos. La iniciaron un consorcio de diferentes compañías en 1990. La transmisión de datos trabaja en modo bidireccional con un protocolo CSMA/CD, que evita colisiones monitorizando el nivel de señal en la red. La versión más conocida actualmente es la 802.11g y se conoce con el nombre comercial de WiFi (Wireless Fidelity). La asociación WECA es la encargada de vigilar y certificar que los productos WiFi cumplen todas las normas y que, por lo tanto, son compatibles con los dispositivos comercializados hasta la fecha.

Estandares WLAN 802.11. En el estándar IEEE 802.11x se encuentran las especificaciones tanto físicas como a nivel MAC que hay que tener en cuenta a la hora de implementar una red de área local inalámbrica. Esta norma ha sufrido diferentes extensiones para su mejora desde que se creó. Así, hoy en día se pueden encontrar las siguientes especificaciones:

- 802.11: especificación para 1-2 Mbps en la banda de los 2.4 GHz, usando salto de frecuencias(FHSS) o secuencia directa (DSSS).
- 802.11b: extensión de 802.11 para proporcionar 11 Mbps usando DSSS.
- Wi-Fi (Wireless Fidelity): promulgado por el WECA para certificar productos 802.11b capaces de interoperar con los de otros fabricantes.

- 802.11a extensión de 802.11 para proporcionar 54 Mbps usando OFDM.
- 802.11g extensión de 802.11 para proporcionar 20-54 Mbps usando DSSS y OFDM. Es compatible hacia atrás con 802.11b. Tiene mayor alcance y menor consumo de potencia que 802.11a.

Como todos los estándares 802 para redes locales del IEEE, en el caso de las WLAN, también se centran en los dos niveles inferiores del modelo OSI, el físico y el de enlace, por lo que es posible correr por encima cualquier protocolo (TCP/IP o cualquier otro) o aplicación, soportando los sistemas operativos de red habituales, lo que supone una gran ventaja para los usuarios, que pueden seguir utilizando sus aplicaciones habituales, con independencia del medio empleado, sea por red de cable o por radio.

Arquitectura WLAN IEEE 802.11. Uno de los requisitos del estándar IEEE 802.11 es que se puede utilizar en las redes cableadas existentes. El estándar 802.11 ha resuelto este problema con el uso de un punto de acceso. El punto de acceso es la integración lógica entre redes cableadas LAN y 802.11.

Las funciones del punto de acceso son integrar las estaciones inalámbricas con la red local existente. Dentro de la misma red local pueden existir diferentes puntos de acceso, con diferentes estaciones de trabajo inalámbricas asociadas. El protocolo de red permite que una estación, dependiendo de la calidad de la comunicación con los diferentes puntos de acceso, cambie el enlace de uno a otro de forma transparente para la estación de trabajo. Esto permite que la estación tenga movilidad por toda la zona dentro del rango de acción de los diferentes puntos de acceso.

Por otro lado, la red inalámbrica debe garantizar la seguridad del acceso a la misma. Para ello, hay dos tipos de servicios de autenticación que ofrece 802.11. La primera es de sistema abierto de autenticación. Esto significa que cualquier persona que intenta autenticarse recibirá autenticación. El segundo tipo es Shared Key Authentication. Para llegar a ser autenticados, los usuarios deben estar en posesión de una clave de red.

Por último, el protocolo también asegura privacidad de los datos transmitidos. La clave de red permite la privacidad mediante el uso de encriptación utilizando Wired Equivalent Privacy (WEP) algoritmo de privacidad. La clave de red se entrega a todas las estaciones antes de iniciar la conexión.

Ejemplos de componentes industriales. Aunque inicialmente este tipo de protocolos tanto ethernet como IEEE 802.11 se crearon para satisfacer las necesidades de comunicación entre ordenadores, actualmente existen multitud de componentes de enlace para autómatas, así como puntos de enlace para redes RS-485 u otras de ámbito industrial. En este apartado se describen algunos de esos componentes.

MOXA AWK-1100. MOXA AirWorks AWK-1100 permite a usuarios wireless acceder a los recursos de una red sin cables, actuando como un punto de acceso. AWK-1100 puede autentificar y autorizar a los usuarios wireless por IEEE 802.1X y RADIUS. Está diseñado para operar en rangos de temperatura que van de 0 a 60°C, y está preparado para las duras condiciones industriales.



Figura 7.8: Imagen del componente MOXA AWR-1100 de OMRON.

MOXA NPort W2150. Es una puerta de acceso que permite interaccionar no solo con redes Ethernet e inalámbricas sino también con las redes industriales más extendidas: RS-232/422/485.



Figura 7.9: Imagen del componente MOXA NPort W2150 de OMRON.

7.8 SISTEMAS SCADA

Se define un software SCADA (*Supervisory Control And Data Acquisition*) como un programa que comunica el ordenador con los equipos que controlan

un proceso, típicamente autómatas programables, con el objetivo de que el operador pueda supervisar desde el ordenador el funcionamiento de todo el proceso.

A través de la red de comunicación, el SCADA puede leer valores de la memoria de los equipos o escribir valores en ella. De esta forma, el programa puede mostrar en el monitor del ordenador de forma gráfica el estado de las distintas variables del proceso controlado. Por otra parte, el operador puede introducir órdenes de marcha y paro o consignas de funcionamiento para los distintos equipos de control a través del programa SCADA, que se encarga de transmitirlos a los equipos.

El ordenador en el que está el SCADA puede estar conectado mediante una red a los ordenadores de la empresa, por lo que se pueden transmitir los datos de producción recogidos de los equipos a los ordenadores de gestión. Esta transferencia se puede hacer incluso a través de Internet a ordenadores situados en cualquier parte del mundo. Esto facilita la integración de toda la información de la empresa y la coordinación de los distintos departamentos.

Es importante resaltar que el objetivo del SCADA no es el control del proceso en tiempo real, sino únicamente la supervisión del mismo, la recogida de datos y la transmisión esporádica de consignas de funcionamiento. El control en tiempo real lo realizan los equipos de control, como autómatas programables, que pueden reaccionar de forma muy rápida a los cambios del proceso. El programa SCADA va refrescando los valores que lee de los equipos a una frecuencia que no suele ser suficiente para el control en tiempo real. De hecho, cuando se desarrolla la aplicación se definen las frecuencias con las que el SCADA debe refrescar cada una de las variables que se leen o escriben en los equipos conectados. Es típico un refresco cada segundo. El mínimo tiempo de refresco admisible depende del SCADA y del tipo de red de comunicación, pero nunca puede ser demasiado bajo: un valor típico es 100 ms.

La red de comunicación utilizada puede ser cualquiera de las existentes en el mercado, si se conecta en el ordenador la tarjeta de comunicaciones necesaria y los equipos de control soportan dicha red. Sin embargo, es habitual utilizar directamente uno de los puertos serie RS232 del ordenador junto con un convertidor RS232-RS485 para conectar los equipos al bus RS485 de 2 hilos. No obstante, cada vez es más habitual utilizar una red Ethernet para este propósito. Para que la comunicación funcione de forma transparente para el usuario, se necesita un driver adecuado a la red y a los equipos que se conectan.

Una dificultad de las operaciones de control industrial es la de compartir información entre dispositivos inteligentes de campo, y con el resto de la empresa. El problema hasta ahora se ha resuelto escribiendo un sinnúmero de protocolos, que definen de qué manera se estructuran los datos que transmite cada dispositivo. Esta diversificación obliga a los desarrolladores de software SCADA a incorporar muchos drivers para cada fabricante. Para evitar esto se ha desarrollado una norma de intercambio de datos para el nivel de planta basada en la tecnología OLE denominada OPC (OLE for Process Control),

que permite un método para el flujo transparente de datos entre aplicaciones corriendo bajo sistemas operativos basados en Microsoft Windows. OPC es un primer paso concreto que permite crear fácilmente una red para compartir los datos de los dispositivos a nivel de proceso. Mediante OPC, es el fabricante el que debe proporcionar su servidor OPC, mientras que los software SCADA son clientes OPC.

Un paquete SCADA comercial suele tener dos programas diferentes: un programa de desarrollo que permite desarrollar la aplicación a la medida del proceso que se quiere supervisar, y un programa de ejecución que sirve para ejecutar la aplicación desarrollada. Una vez desarrollada la aplicación solo se necesita para funcionar, el programa de ejecución.

El programa de desarrollo dispone de librerías gráficas con objetos que permiten desarrollar la aplicación de forma muy simple. Esos objetos incluyen indicadores digitales y analógicos, gráficos de tendencia, gráficos de barras, botones, ventanas de diálogo con el usuario, elementos físicos como bombas o depósitos, etc.

En general, el desarrollo de una aplicación SCADA requiere los siguientes pasos:

1. Instalar los drivers adecuados para los equipos conectados y el tipo de red utilizada.
2. Configurar cuáles son los equipos conectados y su dirección en la red.
3. Definir las variables que se van a leer o escribir en esos equipos y la frecuencia de refresco de las mismas.
4. Definir las variables internas al ordenador que se van a utilizar.
5. Estructurar la aplicación en distintas ventanas, que facilitan el acceso a la información. Si la aplicación es muy simple puede requerir únicamente una ventana.
6. Diseñar cada ventana arrastrando los objetos necesarios de las librerías y configurando las propiedades (acciones asociadas) de dichos objetos.

Ejemplo de sistema SCADA

Con el objetivo de mostrar las funcionalidades básicas que debe satisfacer un sistema SCADA, en esta sección se presenta un ejemplo desarrollado para un proceso industrial real.

Descripción del proceso

El proceso para el cual se desarrolla el sistema SCADA es una planta de regeneración de condensado. Este tipo de plantas se usa para realizar un trata-

miento al agua que sale de la etapa de condensación en centrales térmicas. Dicho tratamiento tiene como objetivo eliminar los residuos sólidos y des-ionizar el agua.

En la figura 7.10 se muestra un esquema general de la planta, la cual consta de una batería de tres intercambiadores de iones que utilizan resina para su funcionamiento. Cuenta además con un sistema de regeneración externo de las resinas agotadas durante el intercambio iónico. También forman parte de la instalación un sistema de recirculación del condensado y otro para la producción de aire a presión, necesario en algunas fases de la regeneración.

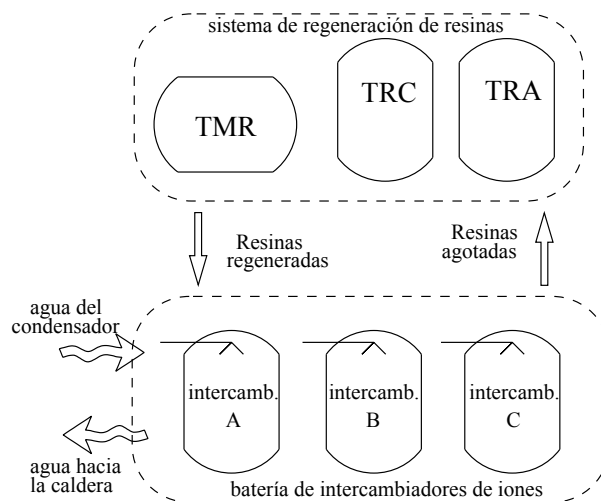


Figura 7.10: Esquema de la planta de regeneración de condensado.

De los tres intercambiadores, dos están funcionando a la vez, mientras el tercero está de reserva con una carga de resina regenerada. El intercambiador de reserva entra en funcionamiento tan pronto como uno de los que está funcionando agota su resina y por tanto no es capaz de completar la des-ionización del condensado. Esta situación se detecta midiendo la conductividad del agua a la salida del intercambiador. El intercambiador que ha dejado de funcionar se recarga con resina regenerada y se queda de reserva. Este ciclo se repetirá tan pronto como otro intercambiador agote su resina.

El sistema de regeneración externo de resina cuenta con un tanque de regeneración aniónica (TRA), otro de regeneración catiónica (TRC) y otro en el que se almacena la resina regenerada (TMR). Para la regeneración de la resina en el TRA se utiliza sosa, mientras que en la regeneración en el TRC se utiliza ácido. Tanto la sosa como el ácido se añaden con sus respectivas bombas de trasiego. El caudal de estas sustancias depende de las concentraciones en los tanques de regeneración.

Un aspecto importante que determina la eficiencia de este proceso es el caudal de agua que pasa por los intercambiadores. Ese caudal se controla mediante una bomba de regeneración. El caudal de agua debe estar en un rango entre F_{min} y F_{max} . Si el caudal de agua hacia la planta de regeneración es mayor

que F_{max} , se ponen en marcha los dos intercambiadores. Si para algunos de los intercambiadores el caudal es menor que F_{min} , será necesario mantener el caudal por encima de F_{min} lo cual se consigue recirculando agua desde la salida del intercambiador a través del sistema de recirculación de condensado. Para ello se utiliza una bomba de recirculación.

Sistema SCADA

Las diferentes funciones que debe realizar la aplicación SCADA para implementar el control automatizado de la planta son las siguientes:

- Debe reflejar el estado de la planta en tiempo real.
- Debe mostrar los parámetros de control más importantes, como el caudal de agua de condensado, los valores de conductividad a la salida de los intercambiadores, etc.
- Debe permitir la configuración de los tiempos de las etapas y del número de trasvases de resinas.
- Debe mostrar gráficas de la evolución, parámetros como el caudal de condensado, valores de conductividad, etc.
- Debe mostrar un panel de alarmas de la planta, generar un fichero de registro de alarmas diario e imprimirlas de forma continua, tal y como se van produciendo.
- Debe permitir la selección del modo de control: manual o automático.
- En modo manual y en condición de paro de toda la planta debe ser posible la puesta en marcha y parada de la bomba de regeneración, bomba de recirculación y las bombas de trasiego de ácido y de sosa.

La aplicación consta de varias ventanas que se clasifican en los siguientes tipos:

- Ventana de inicio y presentación.
- Ventanas de descripción de la planta: proceso
- Ventanas de descripción de elementos del proceso.
- Ventanas de configuración de parámetros.
- Ventana de gráficas.
- Ventana de alarmas.

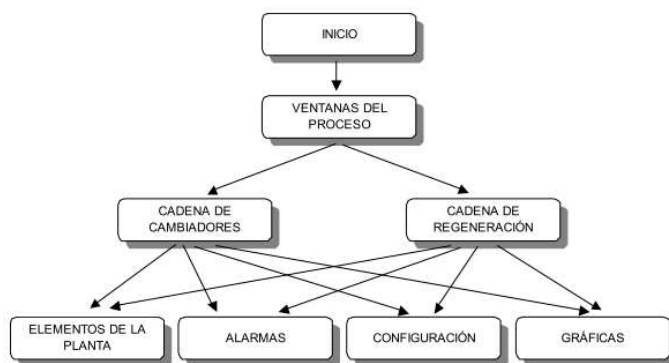


Figura 7.11: Esquema de distribución de ventanas.

En la figura 7.11 se muestra un diagrama jerárquico de la distribución de ventanas de la aplicación SCADA de visualización y control.

En las ventanas del proceso aparece un menú de opciones como el de la figura 7.12. Las opciones son las siguientes:

- Cadena de intercambiadores: pulsando sobre este botón se abre la ventana que contiene la cadena de intercambiadores. Permanecerá inactivo si la ventana abierta es precisamente la de los intercambiadores.
- Cadena de regeneración de resinas: pulsando este botón se abre la ventana que contiene la cadena de regeneración de resinas. Análogamente al caso anterior, este botón permanecerá inactivo si la ventana abierta es la que contiene la cadena de regeneración.
- Configuración: este botón da paso a la ventana de opciones de configuración, que se describirá más adelante.
- Gráficas: este botón abre la ventana de gráficas, en las que, como se verá más adelante, se representa el caudal de condensado y los valores de conductividad a la salida de los intercambiadores.
- Inicio: este botón abre la ventana de inicio de la aplicación.
- Alarmas: este botón abre la ventana donde se representa el panel de alarmas de la plana.



Figura 7.12: Menú de opciones.

A continuación se describirán cada una de las ventanas de la aplicación. Para ello se han agrupado en ventanas del proceso, ventanas de elementos

particulares del proceso (intercambiadores y bombas dosificadoras de ácido y sosa), ventanas de configuración y ventanas de gráficas.

Ventana de inicio

Esta ventana, que se muestra en la figura 7.13, es la que aparecerá inicialmente en el funcionamiento de la aplicación de control. En ella se describe el título de la aplicación y se muestra la barra de menús para poder acceder a cualquier opción de control y visualización.



Figura 7.13: Ventana de inicio.

Ventanas del proceso

Para definir completamente la planta de regeneración de condensado se han creado dos ventanas, la primera donde se muestra el grupo de cambiadores donde se realiza la eliminación de los iones del agua de condensado, y la segunda donde se muestra la cadena de regeneración de condensado.

Ventana de grupo de intercambiadores

En esta ventana, la cual se muestra en la figura 7.14, aparece el grupo de cambiadores iónicos junto a la bomba de regeneración, bomba de recirculación y el compresor que suministra aire a presión necesario para el funcionamiento de la planta. Además de la representación física de esta parte de la planta, se incluyen dos indicadores: caudal de condensado y valores de conductividad a la salida de los cambiadores.

Esta ventana también dispone de un indicador del modo de control de la planta, manual o automático, de manera que se conozca el modo de funcionamiento sin necesidad de abrir la ventana de configuración.

También se pueden observar unos pulsadores con el nombre "Prueba" al lado de la bomba de recirculación y de la bomba de regeneración. Estos pulsadores

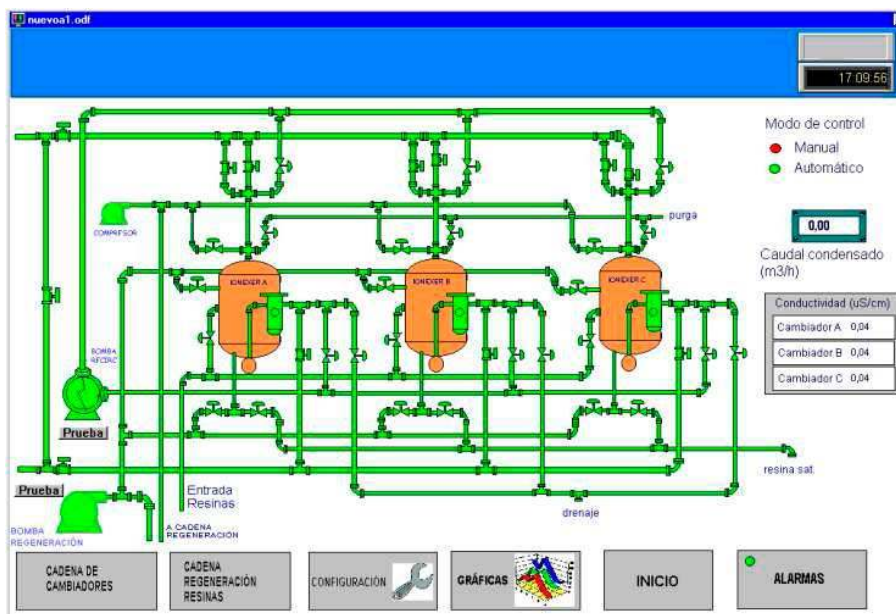


Figura 7.14: Ventana del grupo de intercambiadores.

ponen en marcha y detienen la bomba correspondiente, en caso de que el control de la planta esté en modo manual y toda ella se encuentre en reposo.

La representación de la cadena de intercambiadores en la ventana ha reproducido el esquema físico de la planta real con la finalidad de que los operarios reconozcan cada elemento de ésta de forma inmediata.

En el botón de Alarmas se puede observar un indicador verde. Este indicador cambia el color verde al rojo cuando se activa una o más de las alarmas definidas. De esta forma no es necesario tener abierta la ventana del panel de alarmas para cerciorarse de que alguna o algunas se han activado. En caso de que así fuera el operario de la planta lo reconocería de forma inmediata gracias a este indicador.

Cadena de regeneración

En esta ventana, figura 7.15, se pueden observar los cuatro tanques, el tanque de regeneración catiónica (TRC), el tanque de regeneración aniónica (TRA), el tanque de mezcla de resinas (TMR) y el tanque de agua caliente (TAC). Complementa a los tanques toda la instalación de valvulería y conducciones.

Se muestran las bombas dosificadoras de ácido y de sosa junto a sus depósitos. En éstos pueden observarse los indicadores de nivel alto y bajo. En caso de activarse alguno de ellos, cambiarían de color verde a color rojo y se activarían las alarmas correspondientes. También se muestran los valores de turbidez a la salida del drenaje del TRC y en un recuadro a la izquierda, los valores de temperatura del agua de dilución mezclada con el ácido y la sosa para las regeneraciones catiónica y aniónica, respectivamente (en $^{\circ}C$). Estos indicadores

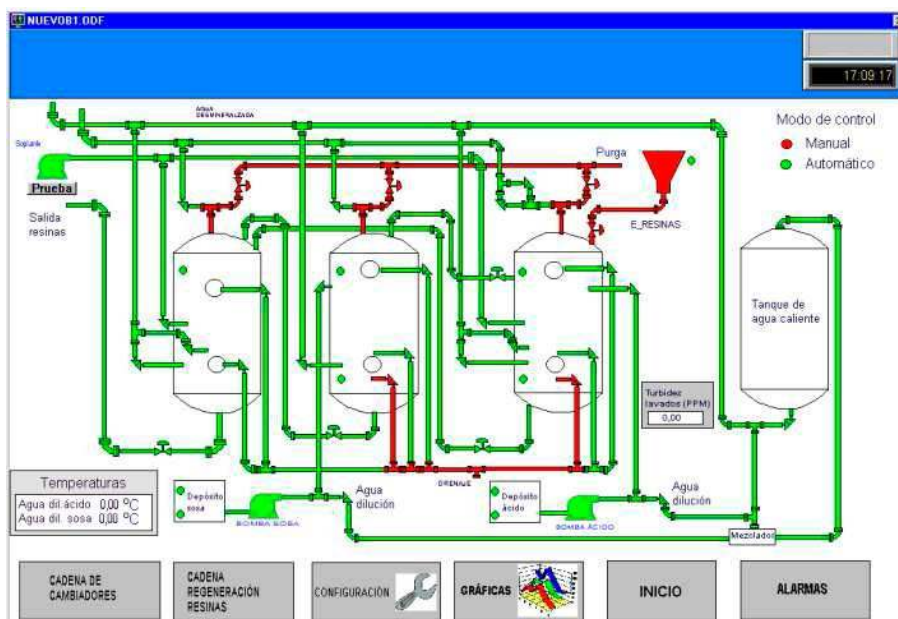


Figura 7.15: Ventana del sistema de regeneración de resinas.

proporcionan el valor en tiempo real de las magnitudes que representan.

En esta ventana se pueden observar los indicadores de nivel de los tanques TRC, TRA y TMR; de la tolva de recogida de resinas agotadas y de los depósitos de ácido y de sosa. Estos indicadores estarán de color verde en estado de reposo, pero cuando se alcancen estos niveles se cambian a color rojo, activándose además la alarma correspondiente.

Ventanas de elementos del proceso

En este apartado se van a describir las ventanas que muestran las características más importantes de los elementos siguientes: cambiadores, bomba dosificadora de ácido y bomba dosificadora de sosa.

Ventana de estado de los cambiadores

Esta ventana es la misma para los tres cambiadores. Como se ha indicado en el apartado anterior, aparte de mostrar el estado del cambiador, muestra el valor de conductividad de agua a su salida, el tiempo acumulado de funcionamiento y alarmas de presión diferencial, tanto entre la entrada y salida del cambiador como del filtro strainer usado para evitar la salida de resina al circuito de agua.

Para abrir la ventana de estado de un cambiador basta con pulsar sobre éste. En ese instante se abrirá una ventana como la mostrada en la figura 7.16. En este caso, se indica que el intercambiador está en servicio: no está ni parado ni saturado.

Ventana de estado de las bombas dosificadoras de ácido y sosa

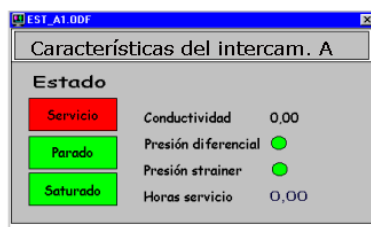


Figura 7.16: Ventana de características de un intercambiador.

En estas ventanas se representa el control de la concentración de ácido y sosa en los tanques de regeneración catiónica y aniónica respectivamente, indicando el valor de la concentración actual y el de la concentración de referencia. Para abrir estas ventanas se debe pulsar sobre los símbolos de las bombas de ácido y sosa respectivamente, apareciendo ventanas como las mostradas en las figuras 7.17 y 7.18.

En las ventanas se observa cómo los depósitos de ácido y de sosa tienen dos indicadores. Estos indicadores son los de nivel alto y bajo, y están configurados como alarmas, por lo que en caso de quedarse el depósito sin ácido o sosa, el indicador de nivel bajo pasaría de color verde a rojo y además se activaría la alarma de nivel bajo.

Se representa además el conexionado entre los controladores PID, los sensores de concentraciones y las bombas dosificadoras.

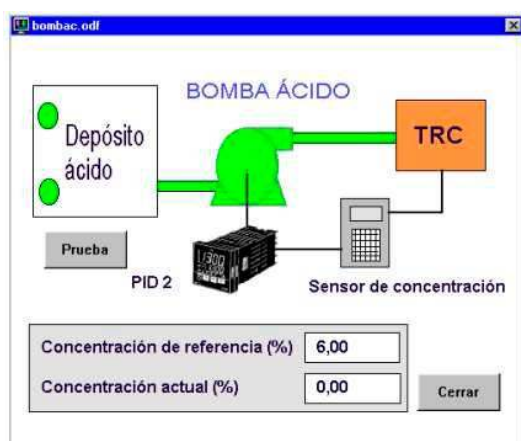


Figura 7.17: Ventana de estado de la bomba dosificadora de ácido.

Ventana de alarmas de la planta

Esta ventana, que se muestra en la figura 7.19 se abre cuando desde cualquier ventana de la aplicación SCADA se pulsa sobre el botón del menú de opciones de Alarmas. Las alarmas representadas estarán de color amarillo mientras no estén activadas. En caso de darse alguna de ellas, cambiarían de color, pasando del amarillo al rojo.

Como ya se ha comentado, en caso de que al menos una señal configurada como alarma se active, el indicador que aparece en el botón de alarmas de la



Figura 7.18: Ventana de estado de la bomba dosificadora de sosa.

barra de menús cambiará de color verde a rojo. De esta forma, el operador conoce la existencia de la activación de una alarma sin necesidad de tener abierta la ventana de alarmas.



Figura 7.19: Ventana de alarmas.

Ventana de configuración del sistema de control

Esta ventana se abrirá siempre que desde el menú de opciones se pulse el botón Configuración. La ventana que se abrirá será como la que se muestra en la figura 7.20.

En esta ventana se encuentran agrupados por diferentes botones las características de configuración del control. A continuación comentaremos algunas de ellas.

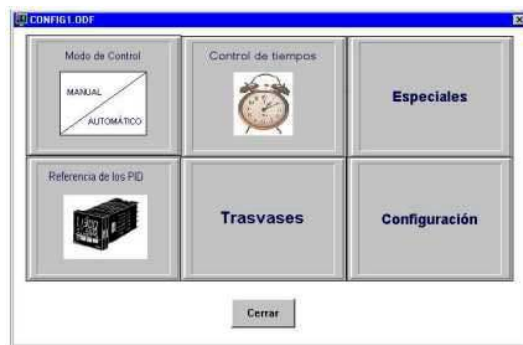


Figura 7.20: Ventana de configuración del control.

Configuración del modo de control

Si se pulsa el botón Modo de control de la ventana de configuración aparece la ventana mostrada en la figura 7.21. En esta ventana se puede determinar el modo de control de la planta de regeneración de condensado: manual o automático.



Figura 7.21: Ventana de configuración del modo de control.

Cuando el control está definido como automático, se realizan todas las secuencias de control sin solicitar la confirmación del operador. Mientras que si el control es seleccionado como manual, la secuencia de control se detiene en varias posiciones a la espera de confirmación del operador de control.

Configuración de las referencias de los controladores PID

Pulsando sobre el botón Referencia de los PID, el operario accede a una ventana en la que puede modificar el valor de referencia a asignar a cada uno de los controladores PID que interviene en el control de la planta. Estos son los controladores del compresor, y de las concentraciones de los tanques de des-ionización catiónica y aniónica.

Configuración de número de trasvases

Si desde la ventana de configuración se pulsa el botón Trasvases, se accede a la ventana de configuración del número de trasvases de resinas desde el intercambiador saturado, al TRC y del TRC y TRA, al TMR. En la figura 7.23

Etapas	Referencia actual	Nuevo valor
Cambiador - TRC:		
Vaciado	60,00	60,00
Transporte	42,00	42,00
Vaciado final	42,00	42,00
TMR - cambiador:		
Transp. y vaciado	44,00	44,00
Transp. y llenado	44,00	44,00
Lav. y sep.		
Drenaje forzado	50,00	50,00
Transp. resina	27,00	27,00
Reg. catiónica		
Drenaje forzado	73,00	73,00
Transp. resina	17,00	17,00
Reg. aniónica		
Drenaje forzado	17,00	17,00
Transp. resina	17,00	17,00
Prueba	10,00	10,00

PID 2: Control de la concentración de ácido	
Valor actual de concentración (%)	0,00
Valor actual de referencia (%)	6,00
Nuevo Valor de referencia (%)	6,00
Referencia de prueba (%)	0,00
Nueva Referencia de prueba (%)	3,00

PID 3: Control de la concentración de sosa	
Valor actual de concentración (%)	0,00
Valor actual de referencia (%)	8,00
Nuevo Valor de referencia (%)	8,00
Referencia de prueba (%)	0,00
Nueva Referencia de prueba (%)	5,00

Figura 7.22: Ventana de configuración de las referencias de los controladores PID.

se muestra la ventana de configuración de estas variables.

Trasvase	Valor actual	Nuevo valor
Cambiador A - TRC	3,00	3,00
Cambiador B - T	2,00	2,00
Cambiador C - T	2,00	2,00
TRC - TMR	2,00	2,00
TRA - TMR	2,00	2,00

Figura 7.23: Ventana de configuración del número de trasvases. Entrada de nuevo valor.

Ventanas gráficas

Una de las ventajas de uso de sistemas SCADA consiste en disponer de una visualización de la evolución de las variables más importantes del proceso como son en este caso el caudal de condensado y los valores de conductividad a la salida de los cambiadores. Las gráficas de estas variables se muestran en la ventana de gráficas, figura 7.24. La ventana de gráficas se puede visualizar pulsando el botón de Gráficas del menú de opciones.

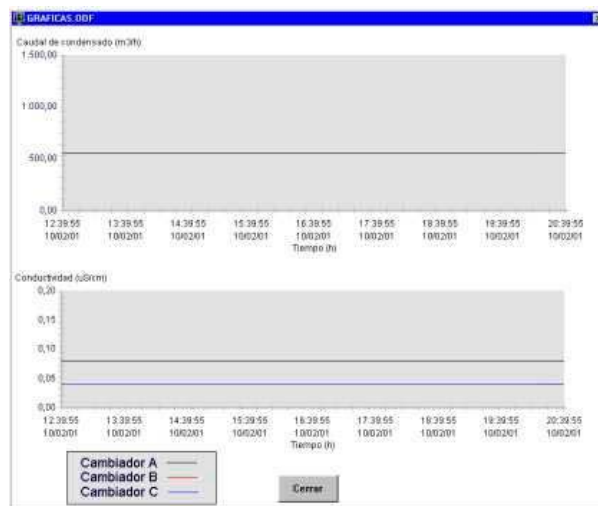


Figura 7.24: Ventana de gráficas.

Anexo A

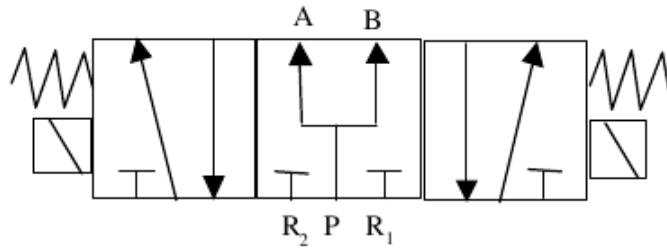
EJERCICIOS SOBRE SENSORES Y ACTUADORES

Ejercicio 1

Selecciona la opción correcta en las siguientes afirmaciones:

1. El retardo a la disponibilidad de un sensor es:
 - a) El tiempo durante el cual un sensor es capaz de detectar un objeto desde que se desconecta de la alimentación.
 - b) El tiempo que transcurre desde que se detecta un objeto hasta que se activa la salida.
 - c) El tiempo que tarda un sensor desde que se conecta a la alimentación hasta que es capaz de detectar un objeto.
 - d) El tiempo que tarda un sensor en desactivar su salida después de haber dejado de detectar un objeto.
2. Se llama corriente residual a:
 - a) La corriente que consume el detector estando desactivado.
 - b) La corriente que circula por un detector cuando está detectando un objeto.
 - c) La corriente que circula por la carga cuando el detector se activa.
 - d) La corriente que consume un detector con salida a transistor.
3. Un detector inductivo detecta:
 - a) Campos magnéticos.
 - b) Inductancia.
 - c) Cualquier tipo de material.
 - d) Objetos metálicos.
4. Un detector reflex polarizado es aquél que:
 - a) Detecta objetos con polarización + y -.
 - b) Sólo detecta objetos que reflejan la luz polarizada.
 - c) Sólo detecta objetos brillantes.
 - d) Evita que los objetos brillantes no sean detectados.
5. La curva de detección de un detector representa:
 - a) La zona del espacio en la que se produce la detección.
 - b) La ganancia del detector en función de la distancia.
 - c) El margen de operación respecto de la distancia.

- d) La tensión de salida del detector en función del tamaño y la distancia del objeto.
6. La siguiente figura representa una electroválvula de:

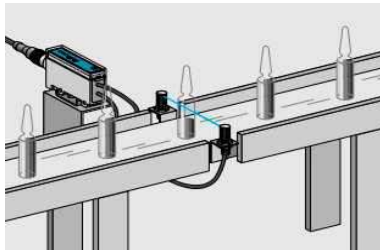


- a) 3 vías, 2 posiciones y permite controlar un cilindro de simple efecto.
- b) 2 vías, 3 posiciones y permite controlar un cilindro de doble efecto.
- c) 5 vías, 3 posiciones y permite controlar un cilindro de doble efecto.
- d) 5 vías, 3 posiciones y permite controlar un cilindro de simple efecto.
7. El uso de las fases A y B en un codificador (encoder) incremental permite:
- a) Conocer el sentido de giro.
- b) Aumentar la resolución.
- c) Conocer el sentido de giro y aumentar la resolución.
- d) Aumentar la fiabilidad del dispositivo.

Ejercicio 2

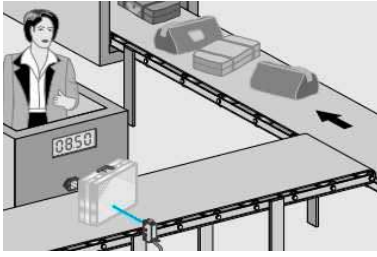
Las siguientes figuras¹ muestran algunas aplicaciones de los detectores.

1. ¿Qué tipo de detector fotoeléctrico se puede usar en cada una de ellas? Explica.
2. ¿En qué casos puede usarse detectores que no sean fotoeléctricos sino inductivos o capacitivos?
- a) Detección de líquido en una ampolla de vidrio.

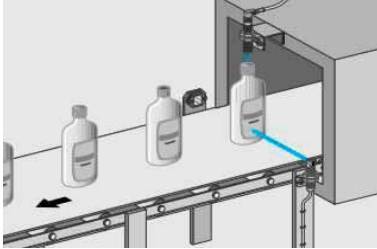


- b) Detección de objetos brillantes.

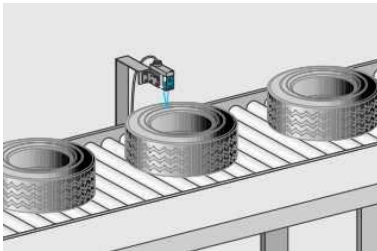
¹Cortesía de Schneider Electric



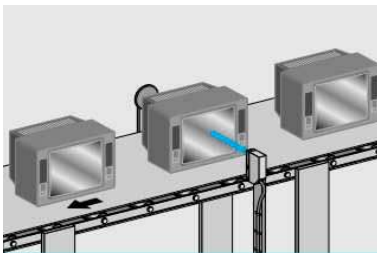
c) Detección de tapones de plástico.



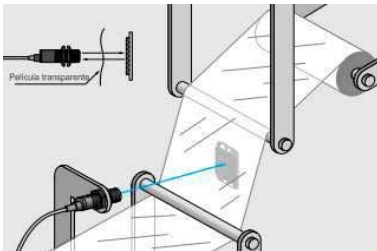
d) Detección de objetos de color oscuro.



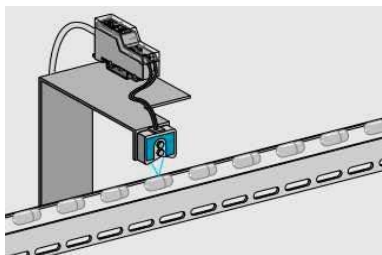
e) Conteo de televisores.



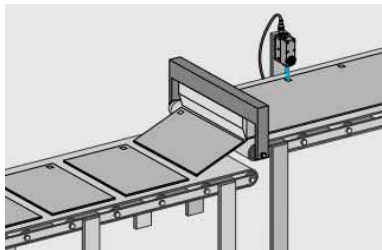
f) Detección de película de plástico transparente.



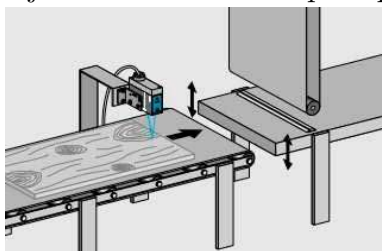
g) Conteo de comprimidos.



h) Detección de marcas para corte.



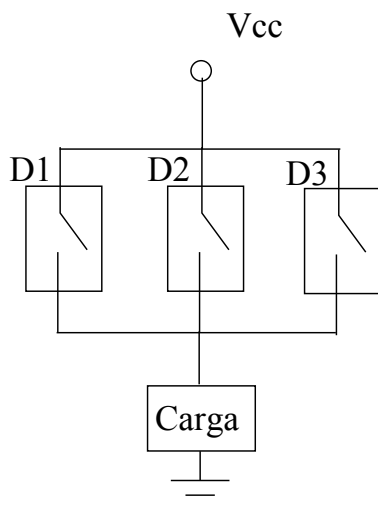
i) Ajuste de altura de cepillo para placa de madera.



Ejercicio 3

La siguiente figura representa la conexión en paralelo de los detectores D1, D2 y D3.

1. Supón que la tensión residual de los tres detectores es de 1 V. ¿Cuál debe ser la mínima tensión de alimentación V_{cc} si la tensión de activación de la carga es de 18 V?



2. ¿Qué corriente residual deben tener los detectores si la corriente de activación de la carga es de 5 mA? Explica.
3. Resuelve los apartados 1 y 2 considerando que los tres detectores están conectados en serie en vez de en paralelo.
4. ¿Qué precaución debe tenerse en cuenta respecto a la tensión de trabajo de los detectores cuando están conectados en serie?

Ejercicio 4

El margen de operación o ganancia de un detector fotoeléctrico se define como el cociente entre la cantidad de luz recibida y la cantidad de luz mínima necesaria para producir la activación de la salida.

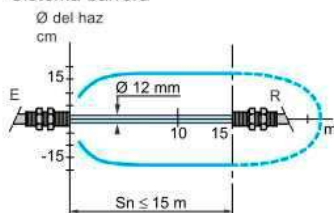
1. ¿Depende el margen de operación de la distancia del objeto que se detecta? Explica.
2. ¿Depende el margen de operación de las propiedades del objeto que se detecta? Explica.
3. ¿Qué valor de margen de operación se tiene sobre la curva de detección? Explica.
4. ¿En qué rango de valores está el margen de operación a la distancia nominal de detección?
5. ¿Qué representa la curva de respuesta de un detector?

Ejercicio 5

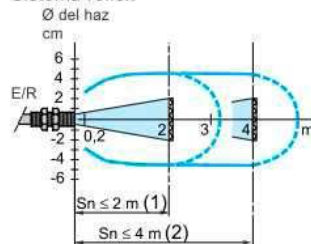
En la siguiente figura se muestran las curvas de detección y ganancia de tres detectores de Schneider Electric: detectores de barrera, reflex y difuso.

Curvas de detección

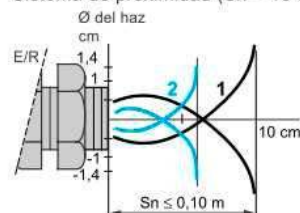
Sistema barrera



Sistema reflex

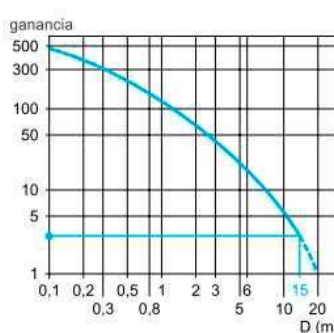


Sistema de proximidad (Sn = 10 cm)

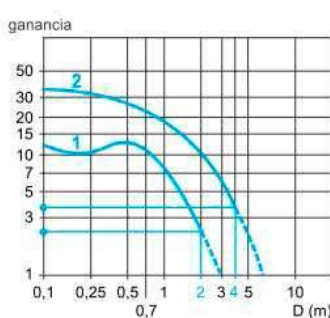


Curvas de ganancia (temperatura ambiente: + 25 °C)

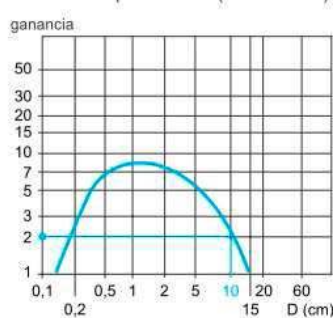
Sistema barrera



Sistema reflex



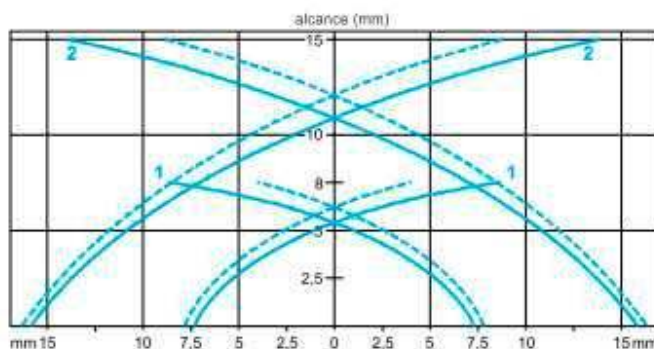
Sistema de proximidad (Sn = 10 cm)



1. ¿Cuál es la mayor distancia a la que se puede detectar objetos con cada uno de estos detectores?
2. ¿Qué detector será más fiable o robusto para detectar objetos ubicados a 10 cm aproximadamente, el reflex con $S_n < 2m$ o el difuso? Explica basándote en la curva de ganancia.
3. ¿Qué significa la forma cóncava de la curva de ganancia del detector difuso? ¿Detectará este detector objetos ubicados a menos de 0,1 cm?

Ejercicio 6

En la siguiente figura se representan las curvas de detección de dos detectores de proximidad inductivos de Schneider Electric.

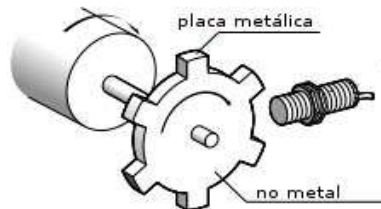


1. ¿Qué representan las líneas continuas y discontinuas?
2. ¿Cómo se pueden determinar de forma experimental estas curvas?
3. Supón que las curvas mostradas en la figura han sido obtenidas para un cuerpo de metal no ferromagnético. ¿Si se usara un cuerpo con las mismas dimensiones pero de metal ferromagnético, las zonas de detección serían mayores o menores? Explica.

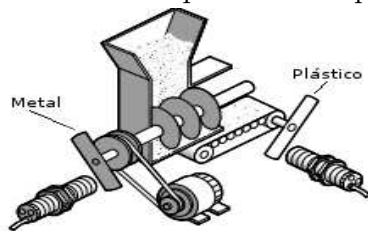
Ejercicio 7

Las siguientes figuras, extraídas de los manuales de Schneider Electric, muestran algunas aplicaciones de los detectores.

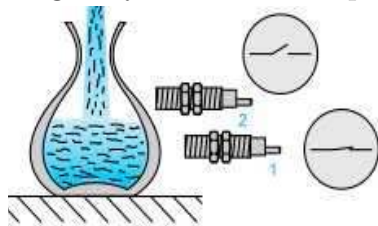
1. ¿Qué tipo de detector, inductivo o capacitivo, se puede usar en cada una de ellas? Explica.
 2. ¿En qué casos pueden usarse detectores que no sean capacitivos o inductivos sino fotoeléctricos?
- a) Medición de velocidad de giro.



- b) Control de ruptura de acoplamiento.



- c) Llegada y llenado de recipiente de vidrio.

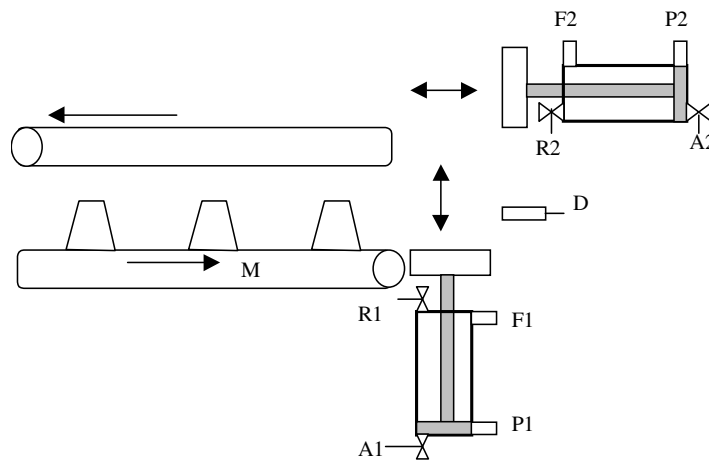


Anexo B

PROBLEMAS DE DISEÑO DE AUTOMATISMOS

Ejercicio 8 Manipulación con cilindros neumáticos I

Considérese el sistema:



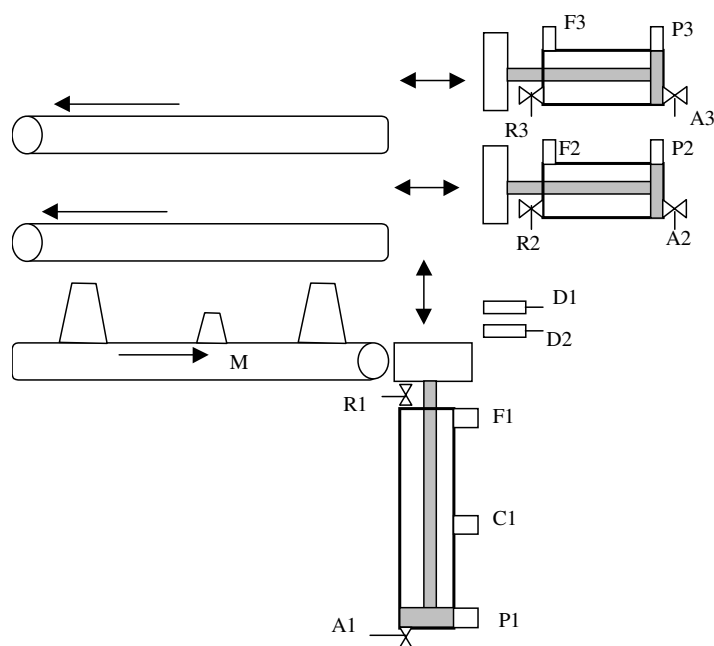
Por la cinta inferior llegan piezas. Cuando una pieza se sitúa encima de la superficie de elevación el detector óptico D, se activa. En ese momento hay que parar la cinta, subir el cilindro hasta su posición superior (hasta que se activa el detector magnético F1), mover el cilindro horizontal hasta F2 y volver los dos cilindros a su posición inicial (P1 y P2 activos). Las señales A1 y A2 hacen avanzar los cilindros, y las señales R1 y R2 los hacen retroceder. Si se desactivan las dos señales, el cilindro se queda parado. La cinta inferior se mueve mediante la señal M, mientras que la superior está siempre en marcha (la controla otro proceso). Se tienen, además, un pulsador de marcha y otro de parada. Cuando se inicia el automatismo, el sistema debe estar en reposo. Al pulsar la marcha se pondrá en funcionamiento. Cuando se pulse la parada el sistema acabará de trasladar la última pieza (si hay alguna en la plataforma de subida) antes de volver al inicio y quedar en reposo.

1. Obtener el Grafcet que resuelve el problema y su implementación en autómatas programables.

- Obtener los Graficets que resuelven el problema y su implementación en autómata programable si además se tiene un pulsador de emergencia con enclavamiento que al pulsarse desactiva todo el proceso. Una vez solucionado el problema y rearmado el pulsador de emergencia se volverá al inicio.

Ejercicio 9 Manipulación con cilindros neumáticos II

Considérese el sistema:



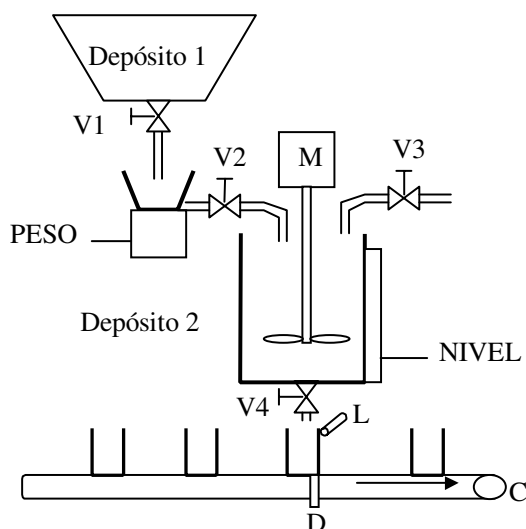
Por la cinta inferior llegan piezas de dos tamaños. Cuando una pieza se sitúa encima de la superficie de elevación, el detector óptico D2 se activa. Si la pieza es grande, además se activa D1. En ese momento hay que parar la cinta, subir el cilindro hasta la posición C1 (si la pieza es pequeña) o hasta la posición F1 (si la pieza es grande). Después hay que mover el cilindro horizontal correspondiente (el 2 o el 3) hasta F2 (o F3), y volver los dos cilindros a su posición inicial. Las señales A1, A2 y A3 hacen avanzar los cilindros, y las señales R1, R2 y R3 los hacen retroceder. Si se desactivan las dos señales, el cilindro se queda parado. La cinta inferior se mueve mediante la señal M, mientras que las superiores están siempre en marcha (las controla otro proceso). Se tiene, además, un pulsador de marcha y otro de parada. Cuando se inicia el automatismo, el sistema debe estar en reposo. Al pulsar la marcha se pondrá en funcionamiento. Cuando se pulse la parada el sistema acabará de trasladar la última pieza (si hay alguna en la plataforma de subida) antes de volver al inicio y quedar en reposo.

- Obtener el Graficet que resuelve el problema y su implementación en autómata programable.

2. Obtener los Graficets que resuelven el problema y su implementación en autómatas programables si además se tiene un pulsador de emergencia con enclavamiento que al pulsarse deje congelado todo el proceso (todas las salidas inactivas), pero de forma que una vez solucionado el problema y rearmado el pulsador de emergencia el sistema continúe desde la posición en la que estaba.

Ejercicio 10 Mezcladora-embotelladora I

El sistema de la figura sirve para producir y embotellar una disolución de un componente en agua.

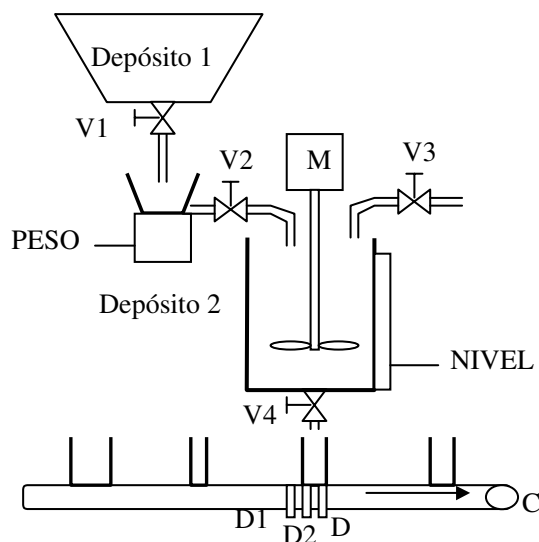


El funcionamiento es como sigue: en un principio se supone el depósito 2 vacío y la báscula también vacía. El procedimiento para hacer la disolución es abrir la válvula V1 hasta que la señal analógica PESO alcance un valor determinado (Peso_Deseado). Después se abrirá la válvula V2 durante 5 segundos para pasar el contenido al depósito 2. Después se abrirá la válvula V3 hasta que el NIVEL (variable analógica) alcance el valor Nivel_Deseado. A continuación, el motor M se conectará durante 10 segundos para disolver el compuesto. Una vez realizada la mezcla, se puede empezar a embotellar. Cuando las botellas llegan al punto de llenado (detector D), se para la cinta (C) y la válvula V4 se abre hasta que se activa L (bote lleno), volviéndose a poner en marcha la cinta hasta el siguiente bote. Cada 5 minutos hay que activar el agitador M durante 10 segundos, para evitar que se formen posos, aunque esto no debe interrumpir el proceso de llenado. Cuando el NIVEL desciende hasta un valor Nivel_Mínimo hay que interrumpir el llenado y proceder a hacer una nueva mezcla. En esta nueva mezcla se añadirá agua hasta un NIVEL igual a Nivel_Mínimo+Nivel_Deseado, para garantizar que la concentración de la mezcla es la misma.

Para el control del proceso habrá un pulsador de MARCHA que iniciará todo el proceso y un pulsador de PARADA. Cuando se pulse PARADA el sistema continuará el proceso, llenando botes hasta que se vacíe por completo el depósito 2 (dejando abierta la válvula V4 durante 10 segundos desde que el NIVEL se hace cero). El último bote (a medio llenar) será evacuado hasta que se vuelva a activar D, y el sistema quede en reposo. Diseñar el automatismo mediante Grafcet y obtener el programa del autómeta que lo implementa.

Ejercicio 11 Mezcladora-embotelladora II

Hacer el ejercicio anterior con las siguientes modificaciones. En el proceso de mezcla: la báscula se llena al principio (con V1) hasta el Peso_Máximo. Después se abre la válvula V2 hasta que el PESO se ha decrementado en la cantidad Peso_Deseado. Cuando (tras varios procesos de mezcla) la báscula baja por debajo de Peso_Mínimo, la báscula se recarga hasta Peso_Máximo. En el proceso de embotellado: ahora no hay detector de bote lleno L. En su lugar, se abre la válvula V4 hasta que el NIVEL se decrementa en una cantidad que depende del tamaño del bote a llenar: Nivel_Pequeño, Nivel_Mediano o Nivel_Grande. El tamaño del bote se obtiene de los detectores D1, D2 y D. Si es pequeño sólo queda activo D. Si es mediano queda activo D2 y D, y si es grande, D1, D2 y D. El nuevo sistema se muestra en la siguiente figura.

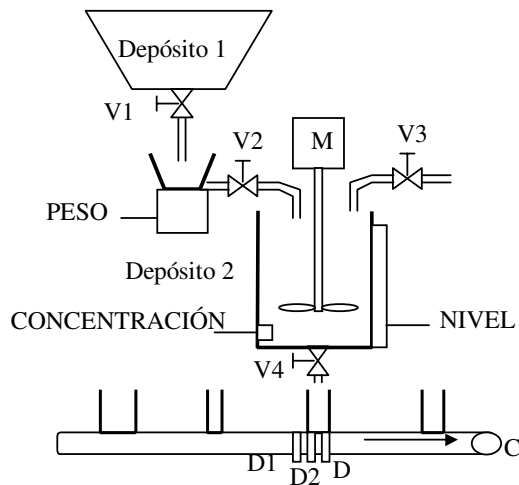


Obtener el Grafcet que resuelve el problema y su implementación en autómeta programable.

Ejercicio 12 Mezcladora-embotelladora III

Considerar el proceso anterior al que, por otra parte, se quiere incorporar un sensor de concentración para supervisar que la mezcla es correcta antes de

embotellarla. Ese sensor da una salida analógica: CONCENTRACIÓN, tal y como muestra la figura.



Después de acabar la mezcla y agitar, y cada vez que se agita de nuevo la mezcla, se tendrá que comprobar que CONCENTRACIÓN está entre dos límites, Con_Máxima y Con_Mínima. Si no lo está, se encenderá una luz de alarma y se parará el proceso. En el modo de alarma, un operario conectará manualmente una manguera a la salida de V4. Después, accionará un pulsador de EVACUACIÓN y se vaciará por completo el depósito 2. Cuando termine se apretará un pulsador de FIN_DE_ALARMA y el sistema volverá al inicio. Obtener el Grafcet que resuelve el problema y su implementación en autómatas programable.

Ejercicio 13 Montacargas

Diseñar un automatismo que controle un montacargas de 3 plantas. Las señales de entrada serán:

- FOTOCÉLULA. Situada en la puerta de la jaula. Está a 1 si no se interpone ningún objeto, y a cero si se interpone algún objeto.
- PUERTA_CERRADA. Es un detector que se pone a 1 cuando la puerta se ha cerrado del todo.
- P1, P2 y P3 son los detectores que indican que la jaula está en un piso. Están a 1 si el ascensor está situado en ese piso.
- L1, L2 y L3 son los pulsadores para indicarle al montacargas que se desplace hasta un piso.
- PARO es un pulsador que produce la parada inmediata del montacargas en el lugar donde esté.

Las señales de salida son:

- SUBIR. Mueve el montacargas hacia arriba.
- BAJAR. Mueve el montacargas hacia abajo.
- CERRAR_PUERTA. Cuando está a 1 se cierra la puerta del montacargas. Cuando está a cero, se abre.

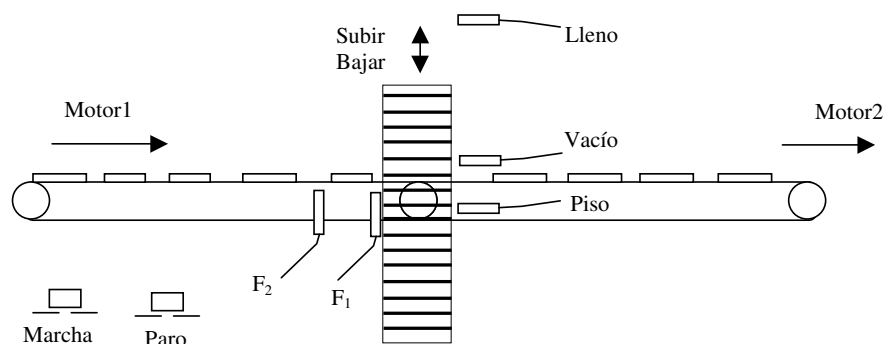
El funcionamiento es como sigue:

- Estando en reposo el montacargas, la puerta estará abierta.
- Cuando se pulse un piso, se debe cerrar la puerta y a continuación moverse hasta que se llegue a ese piso. Si se corta la fotocélula antes de que se haya cerrado la puerta, ésta se abrirá, esperará 2 segundos, y volverá a intentar cerrarla para después moverse hacia el piso.
- El pulsador de PARO debe hacer que se detenga el ascensor, volviéndose a poner en marcha cuando se pulse L1, L2 o L3 hasta llegar al piso elegido en ese momento.
- Cada vez que llega a un piso debe abrir la puerta y mantenerla abierta al menos 2 segundos antes de atender una llamada. Si no hay otra llamada, se quedará la puerta abierta.
- El ascensor sólo atenderá una orden. Cuando acabe el movimiento (cerrar puerta, moverse hasta el piso destino, abrir puerta, esperar 2 segundos) podrá atender otra llamada. La orden solo queda cancelada cuando se pulsa PARO, ya que después el montacargas irá al nuevo piso elegido.

Obtener el Grafcet que resuelve el problema y su implementación en autó-mata programable.

Ejercicio 14 Separador de azulejos

El sistema de la figura representa un sistema de separación de azulejos.



Las entradas del sistema de control son:

- F1, F2. Son dos detectores que detectan las baldosas.
- Piso. Es un detector que se activa cuando el compenser está exactamente en un piso.
- Lleno. Es un detector que se activa cuando el compenser está lleno (en la posición más elevada).
- Vacío. Es un detector que se activa cuando el compenser está vacío (en la posición más baja).
- Motor2. Es una señal que indica si el motor 2 está activo o no (tramo posterior en marcha o parado).
- Marcha. Es un pulsador simple.
- Paro. Es un pulsador simple.

Las salidas del sistema de control son:

- Motor1. Pone en marcha el motor de la cinta anterior al compenser.
- Subir. Mueve el compenser hacia arriba.
- Bajar. Mueve el compenser hacia abajo.

Descripción del funcionamiento del proceso.

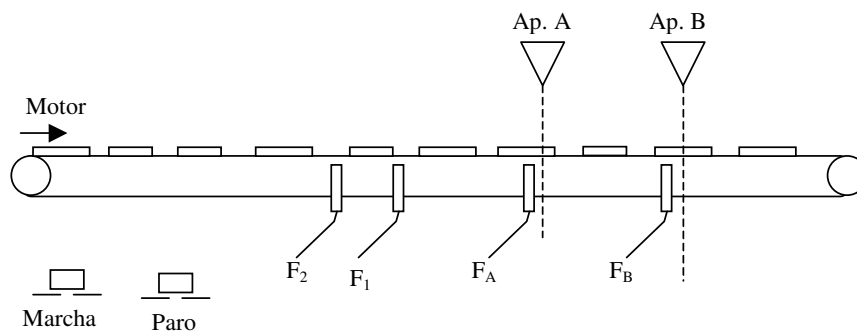
- Por la cinta llegan, al azar, azulejos de dos tamaños (de 30 cm y de 60 cm). El compenser (buffer) se utiliza para separar los azulejos según su tamaño. Para ello, el compenser debe cargar los azulejos de 30 cm, dejando pasar los de 60 cm. Cuando el compenser se llena, se debe parar el motor de la cinta 1, y vaciar todos los azulejos de 30 cm almacenados, empezando después un nuevo ciclo. Se supone que aunque el motor 1 esté parado, los azulejos salen por la cinta 2 sin problemas.
- La detección del tamaño del azulejo se realiza mediante los detectores F1 y F2, (si el azulejo es de 30 cm cuando se activa F1 ya no está activo F2, mientras que si es de 60 sí están activos a la vez).
- El compenser debe coger un azulejo justo cuando se desactiva el sensor F1 tras el paso del azulejo. Para cogerlo simplemente subirá un piso.
- Para descargar cada azulejo se debe bajar un piso y esperar medio segundo.
- Si el motor 2 se para, el motor 1 debe pararse también hasta que el motor 2 vuelva a funcionar.

- Se supone que las baldosas llegan lo bastante separadas unas de otras (cuando el compenser coge una baldosa puede subir un piso antes de que llegue la siguiente al detector F1).
- Estando el sistema en reposo, al pulsar Marcha debe ponerse en funcionamiento. Cuando se pulse Paro, el sistema debe detenerse en cuanto compruebe que el compenser ha terminado la operación de subir piso o bajar piso (es decir, evitando que el compenser se quede entre dos pisos).

Realizar el (o los) diagrama GRAFCET que resuelve el problema de automatización anterior, así como el programa de autómatas que implementa el automatismo anterior.

Ejercicio 15 Línea de tratamiento de azulejos

El sistema de la figura representa una línea de tratamiento de azulejos.



Las entradas del sistema de control son:

- F1, F2, FA, FB. Son detectores que detectan las baldosas.
- Marcha. Es un pulsador simple.
- Paro. Es un pulsador simple.

Las salidas del sistema de control son:

- Motor. Pone en marcha el motor de la cinta.
- A. Abre la válvula que activa el aspersor de la aplicación A.
- B. Abre la válvula que activa el aspersor de la aplicación B.

Descripción del funcionamiento del proceso.

- Por la cinta llegan, al azar (no separados uniformemente), azulejos de dos tamaños (de 30 cm y de 60 cm). Los azulejos de 30 cm deben recibir únicamente el tratamiento A, mientras los azulejos de 60 cm deben recibir únicamente el tratamiento B.

- La detección del tamaño del azulejo se realiza mediante los detectores F1 y F2, (si el azulejo es de 30 cm, cuando se activa F1 ya no está activo F2, mientras que si es de 60, sí están activos a la vez).
- La activación de la aplicación (A o B) debe hacerse durante 1 segundo desde que se activa la fotocélula.
- Estando el sistema en reposo, al pulsar Marcha debe ponerse en funcionamiento. Cuando se pulse Paro el sistema debe detenerse, pero sin dejar un azulejo a medias (si un azulejo está en medio de una aplicación ésta debe finalizar antes de parar el sistema).

Realizar el (o los) diagrama GRAFCET que resuelve el problema de automatización anterior, así como el programa de autómatas que lo implementa.

Ejercicio 16 Aparcamiento automático

Diseñar el automatismo que controla un aparcamiento automático con 30 plazas disponibles. Se tiene las siguientes variables de salida:

- SubirE, SubirS, BajarE, BajarS para subir y bajar la barrera de entrada y la de salida.
- COMPLETO y LIBRE para indicar si está completo o hay plazas libres.
- TicketEntrada, que produce la salida de un ticket de entrada para que lo coja el usuario.

Las variables de entrada serán:

- E_Arriba, E_Abajo, S_Arriba y S_Abajo, que son detectores que indican que las barreras están arriba o abajo.
- Fotocélula_E y Fotocélula_S, que son dos fotocélulas situadas exactamente bajo las barreras para detectar un vehículo que está pasando.
- Coche_Presente, que es un detector que se activa si hay un vehículo en la zona de entrada (antes de la barrera de entrada).
- Pagado es una variable producida por el elemento que acepta los tickets de salida. Da un pulso de 200 ms de duración si se ha introducido un ticket de salida.

El funcionamiento es como sigue: cuando hay sitio libre y un coche se sitúa en la zona de entrada (se activa CochePresente), se activa TicketEntrada y se abre automáticamente la barrera de entrada, cerrándose en cuanto la Fotocélula_E deja de detectar al coche. Cada vez que entra un coche se debe decrementar un contador que contiene el número de plazas libres.

Para salir del aparcamiento el usuario debe ir a pagar a la caja, donde recibe un ticket de salida. Cuando introduce ese ticket, la barrera de salida debe levantarse, bajándose en cuanto la Fotocélula_S deja de detectar al vehículo. Cada vez que sale un vehículo, el contador de plazas libres debe incrementarse. Cuando el número de plazas libres sea cero, se debe encender el letrero COMPLETO, mientras que si no es cero debe estar activado el letrero LIBRE.

Diseñar el Grafcet que resuelve el problema y el programa que lo implementa.

Ejercicio 17 Ascensor

Diseñar un automatismo que controle un ascensor de tres plantas.

En la jaula hay un pulsador para cada planta y un pulsador de paro de emergencia, así como una fotocélula en la puerta corredera. La puerta corredera tiene un detector de puerta abierta y otro de puerta cerrada. Además se tiene un detector inductivo en la jaula, que se activa únicamente cuando la jaula está situada exactamente en una planta. En cada planta hay un pulsador de llamada y un final de carrera de roldana que está activo mientras la jaula esté entre 50 cm por encima y 50 cm por debajo de la planta.

Como salidas se tiene SubirJaula, BajarJaula, para subir y bajar el ascensor; Velocidad (si vale 1, el motor gira a velocidad alta y si vale 0 a velocidad baja); CerrarPuerta y AbrirPuerta para abrir y cerrar la puerta.

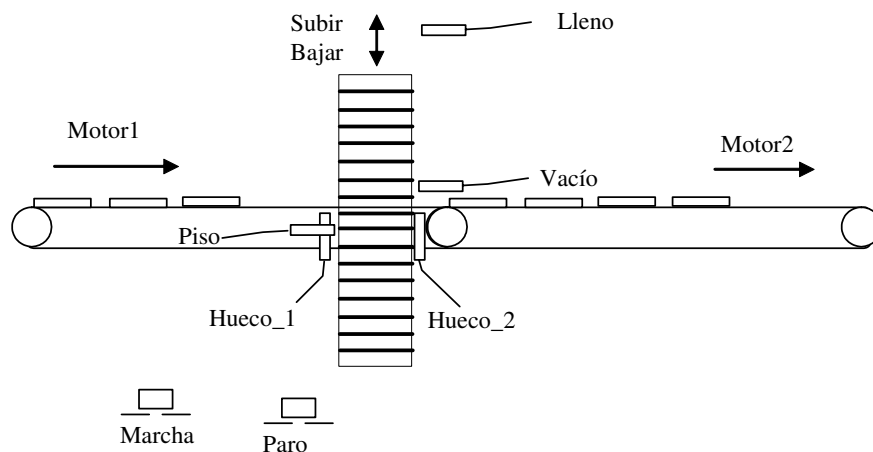
El funcionamiento es como sigue:

- El funcionamiento del movimiento del ascensor cuando va de un piso a otro debe ser tal que se mueva a velocidad baja los primeros 50 cm y los últimos 50 cm del recorrido y a velocidad alta el resto.
- El pulsador de paro de emergencia debe hacer que se detenga el ascensor, volviéndose a poner en marcha cuando se pulsa un piso (desde dentro, no de llamada externa).
- Cada vez que llega a un piso debe abrir la puerta y mantenerla abierta al menos 5 segundos antes de atender una llamada. Si no hay otra llamada, se quedará la puerta abierta.
- Si la puerta se está cerrando y se activa la fotocélula, se debe abrir y esperar 1 segundo para volver a cerrarse.
- El ascensor sólo atenderá una orden (de llamada externa o de pulsador interior). Cuando acabe el movimiento (cerrar puerta, moverse hasta el piso destino, abrir puerta, esperar 5 segundos) podrá atender otra llamada.

Obtener el Grafcet que resuelve el problema y su implementación en autó-mata programable.

Ejercicio 18 Buffer en línea de fabricación de azulejos

El sistema de la figura representa un compenser (buffer) habitual en una línea de azulejos. El compenser se sitúa al final de un tramo de transporte y su objetivo es acumular azulejos cuando el tramo siguiente se para, para evitar que se pare el tramo anterior.



Las entradas del sistema de control son:

- Hueco_1. Es un detector que detecta la baldosa.
- Hueco_2. Es otro detector para detectar las baldosas.
- Piso. Es un detector que se activa cuando el compenser está exactamente en un piso.
- Lleno. Es un detector que se activa cuando el compenser está lleno (en la posición más elevada).
- Vacío. Es un detector que se activa cuando el compenser está vacío (en la posición más baja).
- Motor2. Es una señal que indica si el motor 2 está activo o no (tramo posterior en marcha o parado).
- Marcha. Es un pulsador simple.
- Paro. Es un pulsador simple.

Las salidas del sistema de control son:

- Motor1. Pone en marcha el motor de la cinta anterior al compenser.
- Subir. Mueve el compenser hacia arriba.
- Bajar. Mueve el compenser hacia abajo.

Descripción del funcionamiento del proceso. Estando el sistema en reposo, al pulsar Marcha debe ponerse en funcionamiento. En funcionamiento normal, Motor2 está activo, y Motor1, también. Si se desactiva Motor2 (debido a un fallo en el tramo posterior de la línea), el compenser debe empezar a cargar azulejos. Para ello, cuando se active Hueco_2 el compenser debe subir un piso (sin parar Motor1). Esta carga continuará hasta que se active Lleno (en cuyo caso se debe desactivar Motor1) o hasta que se active de nuevo Motor2.

Si el compenser se ha llenado, el sistema permanecerá parado hasta que se active Motor2. Cuando se vuelva a activar el Motor2, el Motor1 debe activarse, y el compenser debe empezar a descargar.

Para descargar un azulejo tiene que haber detectado un hueco suficiente. Habrá un hueco suficiente si cuando se desactiva Hueco_2 (al salir el azulejo), Hueco_1 no está activo (no hay un azulejo llegando). Para descargar, simplemente bajará un piso.

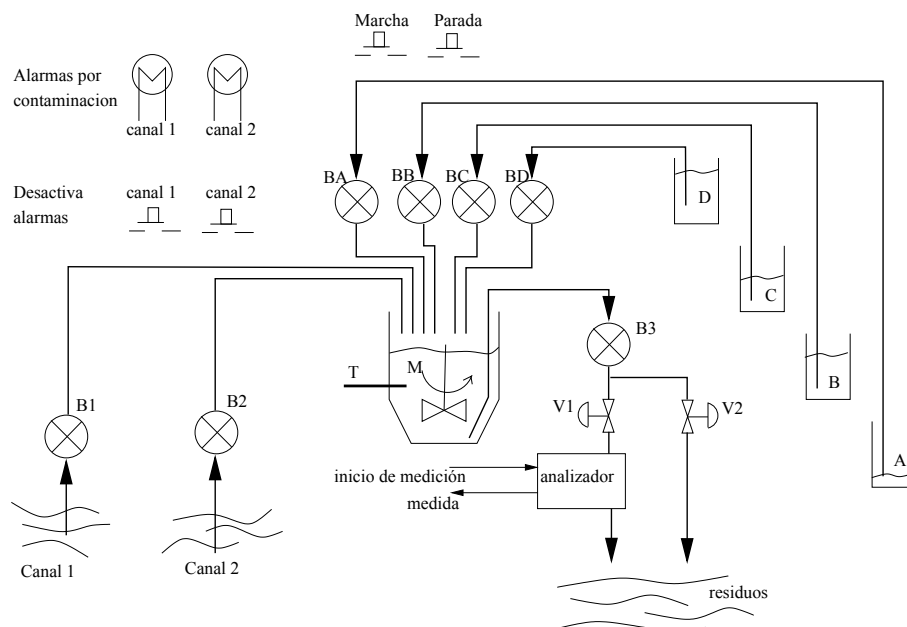
La descarga seguirá hasta que se active Vacío o hasta que se desactive Motor2, en cuyo caso se procederá a cargar azulejos tal y como se ha explicado arriba.

Si en funcionamiento normal se pulsa Paro, el sistema debe detenerse hasta que se pulse Marcha.

Realizar el (o los) diagrama GRAFCET que resuelve el problema de automatización anterior.

Ejercicio 19 Medidor automático de contaminantes

El siguiente esquema representa un sistema para medir la concentración de contaminantes en el agua que fluye por dos canales.



El funcionamiento es como a continuación se describe. Las bombas B1 y B2 toman agua de los canales 1 y 2 respectivamente. El análisis de cada canal se hace por separado, alternándose. Entre los análisis de los canales deben transcurrir 30 minutos.

Cuando el sistema se pone en marcha, mediante el pulsador Marcha, la bomba B1 extrae agua del canal 1 hasta el recipiente. Esta bomba tiene que estar activa durante 30 s, entonces comienzan a añadirse los reactivos A, B, C y D. Primero se activa la bomba BA durante 5 s al cabo de los cuales se activa la bomba BB, manteniéndose activa la bomba BA. A partir de ese momento las bombas BA y BB se mantienen activas durante 3 s, entonces se desactivan y se activa la bomba BC durante 10 s. Cuando se desactiva la bomba BC, se activa la BD durante 5 segundos. En el momento en que se detiene la bomba BC y se activa la BD hay que poner en marcha el motor M del agitador, que debe permanecer activado hasta que la mezcla formada por el agua y los reactivos alcance una temperatura de 30 grados centígrados. La temperatura se mide con el sensor T. En ese momento la mezcla estará lista para pasar al analizador, lo cual se consigue activando la bomba B3 y abriendo la válvula V1, manteniendo la V2 cerrada. La bomba B3 debe estar activa durante 25 s, tiempo suficiente para que la mezcla llegue al analizador. Entonces se activa la señal digital inicio para que se comience el análisis. Tres segundos después de la activación de inicio de medición, estará disponible el resultado del análisis a través de la señal analógica medida. Si el valor de medida es mayor que medida_máxima entonces se debe activar una alarma por contaminación del canal 1. Una vez finalizada la medición se debe vaciar por completo el recipiente, para lo cual se debe abrir la válvula V2, cerrar la V1 y activar la bomba B3 durante 40 s, después de lo cual el sistema quedará listo para realizar el análisis del agua del canal 2. Dicho análisis se realiza pasados 30 minutos y siguiendo el mismo procedimiento que el utilizado para el canal 1.

Si durante cualquiera de los análisis se acciona el pulsador de Parada, el análisis debe finalizar y el sistema se detendrá, quedando listo para realizar el siguiente análisis cuando se pulse el botón Marcha. Cuando se active la alarma de uno de los dos canales, no se realizarán más análisis de ese canal y la alarma se mantendrá activa. En este caso, el sistema continuará realizando análisis del canal no contaminado a intervalos de 60 minutos. Las alarmas por contaminación sólo se pueden desactivar de forma manual mediante los pulsadores correspondientes.

Obtener el Grafcet que resuelve el problema y su implementación en autó-mata programable.

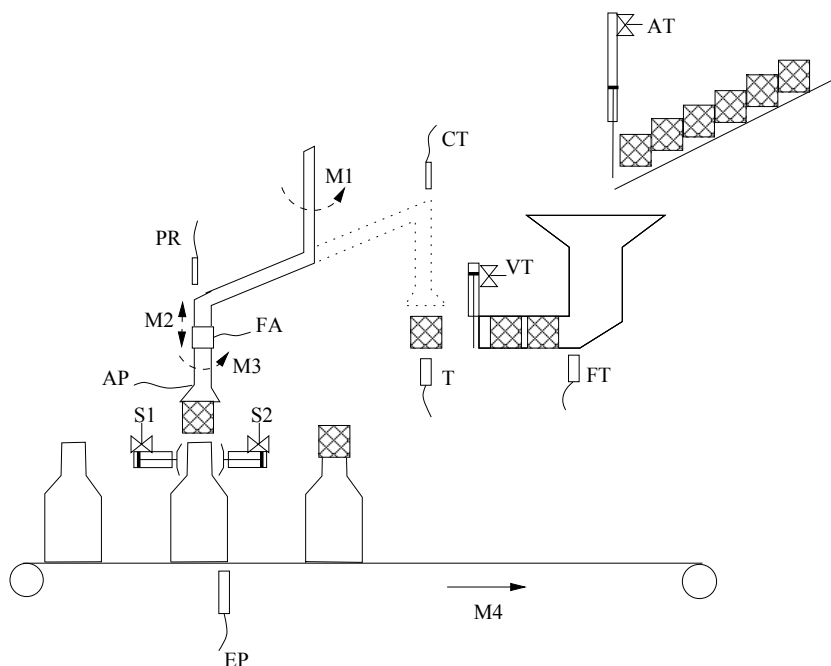
Ejercicio 20 Máquina roscadora de tapones

El esquema representa una máquina roscadora de tapones en una línea de envasado de detergentes. El funcionamiento de la máquina es como a continuación se describe.

El detector EP se activa cuando un envase llega por la cinta hasta la posición donde se realizará el roscado del tapón. En ese momento se debe detener la cinta, desactivar el motor M4, y activar los cilindros de sujeción S1 y S2 para comenzar el proceso de roscado, el cual se realiza por un cabezal diseñado para este fin.

El cabezal de roscado tiene tres actuadores: actuador de giro neumático M1, que realiza el movimiento de giro horizontal del cabezal (si M1 está activo, el cabezal gira hasta la posición CT, mientras que si está desactivado, el cabezal gira hasta la posición PR); cilindro neumático M2 para realizar el movimiento vertical (si M2 está activo el cabezal baja, si está desactivado, el cabezal sube) y M3 para realizar el roscado. El cabezal incorpora además una pinza con accionamiento neumático para la captura del tapón que se mantiene cerrada mientras la señal AP está activa, así como un sensor de fuerza FA cuya salida se activa cuando en la operación de roscado se alcanza la fuerza de apriete deseada. El funcionamiento del cabezal es como sigue:

Activando el actuador M1, se desplaza el brazo del cabezal hasta la posición CT. Para coger un tapón se activa el cilindro M2 durante 3 segundos después de los cuales se activa AP para agarrar el tapón, desactivándose M2, lo que provoca la subida del cabezal con el tapón en la pinza. Desactivando el actuador M1, se desplaza el cabezal a la posición de roscado PR. Si hay una botella presente se comienza la operación de roscado, activando el cilindro M2 durante 2 segundos después de los cuales se activa el motor M3 (manteniendo M2 Activo) hasta que se activa el sensor FA indicando que se ha alcanzado la fuerza de apriete deseada. Entonces se desactivan AP, M2 y M3, lo que provoca la subida del cabezal durante 1 segundo, al cabo del cual se activa el actuador M1 para llevar nuevamente el cabezal hasta la posición CT y coger un nuevo tapón, repitiéndose el ciclo antes descrito.

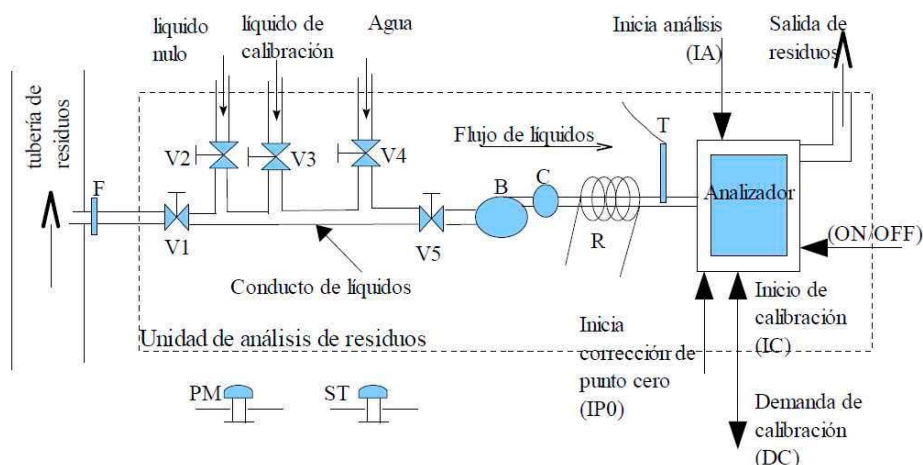


La ubicación de tapón para que sea agarrado por el cabezal se lleva a cabo mediante un cilindro de simple efecto con electroválvula VT que debe activarse para dejar pasar un tapón cada vez que se desactive el detector T, indicando que un tapón ha sido retirado por el cabezal. El sistema cuenta, además, con un detector FT que se activa cuando sólo quedan dos tapones en el depósito. La caída de tapones al depósito es controlada por un cilindro de simple efecto accionado mediante la electroválvula AT. En el depósito siempre debe haber tapones disponibles.

Obtener el Grafcet que resuelve el problema y su implementación en autó-mata programable.

Ejercicio 21 Analizador de residuos líquidos

En la planta química se ha instalado una unidad analizadora de residuos líquidos. El corazón de esta instalación es el analizador, un dispositivo que permite medir las concentraciones de distintos componentes en un líquido que se hace circular por él. La unidad de análisis está compuesta además por una serie de dispositivos (válvulas, bombas, sensores, resistencia calefactora, etc.) cuyo objetivo es garantizar un correcto funcionamiento del analizador. La siguiente figura muestra un esquema de la unidad de análisis.



El funcionamiento de la unidad de análisis de residuos es como a continuación se describe: la unidad de análisis tiene seis modos de funcionamiento excluyentes (sólo puede estar en un modo a la vez):

Modo 0; Puesta en marcha: la puesta en marcha de la unidad de análisis se inicia con el pulsador PM. Se enciende al analizador activando la señal ON/OFF, que permanecerá activa durante todo el funcionamiento de la unidad. Con las válvulas V1 y V5 abiertas y V2, V3 y V4 cerradas se hace pasar la corriente de residuos por el conducto de líquidos activando la bomba B. Desde la puesta en marcha de la bomba hasta que la corriente de residuos comienza a pasar por el analizador transcurren 30 s. Sólo después de este tiempo y verificando que la temperatura T sea mayor de 200°C, se indica al analizador que comience el análisis mediante la activación de la señal IA. Si transcurridos los 30 s después del encendido de la bomba $T < 200^{\circ}\text{C}$, no se inicia el análisis hasta que $T > 200^{\circ}\text{C}$. Este modo sólo se ejecuta al inicio del funcionamiento de la unidad de análisis.

Modo 1; Análisis: en este modo las válvulas V1 y V5 están abiertas mientras que V2, V3 y V4 están cerradas. El líquido a ser analizado es extraído de la tubería de residuos por succión mediante la bomba B. Es filtrado por el filtro F y calentado hasta una temperatura en el rango de 200°C a 250°C por la resistencia calefactora R. Este es el modo en que, por defecto, se encuentra la unidad. En este modo, la señal IA está activa y el resto de señales del analizador (excepto ON/OFF) están desactivadas.

Modo 2; Corrección del punto cero: este procedimiento se realizará cada 12 horas. Para ello se hace pasar líquido nulo por el analizador indicándole que inicie la corrección mediante la activación de la señal IP0 (el resto de señales del analizador, excepto ON/OFF, desactivadas). Antes de comenzar la corrección del punto cero es necesario purgar el conducto de los líquidos de residuos; esto se consigue haciendo pasar líquido nulo durante 15 s, sólo entonces se indica al analizador que comience la corrección del punto cero. El procedimiento, que se realiza por el propio analizador, tiene una duración de 50 s, tiempo durante el

cual ha de estar circulando líquido nulo. Una vez finalizado el ajuste de punto cero se vuelve al modo análisis.

Modo 3; Calibración: el procedimiento de calibración se inicia por demanda de analizador mediante la activación de la señal DC. Para la calibración se hace pasar líquido de calibración por el analizador. En este modo también es necesario purgar el conducto de líquidos haciendo pasar líquido de calibración durante 15 s, sólo entonces se indica a analizador, por activación de la señal IC (resto de señales, excepto ON/OFF, desactivadas), que inicie el procedimiento de calibración, el cual dura 100 s, tiempo durante el cual estará pasando líquido de calibración por el analizador. Una vez finalizada la calibración se vuelve al modo análisis.

En estos cuatro modos, el caudal de los líquidos hacia el analizador se garantiza mediante la bomba B. Así mismo, la temperatura de entrada de los líquidos (residuos, líquido nulo o líquido de calibración, según el modo) se ha de mantener entre los 200°C y los 250°C mediante la resistencia calefactora R. Dicha temperatura se mide con la termo-resistencia T. El control de la temperatura se realizará mediante un controlador todo/nada: si $T < 210^{\circ}\text{C}$ activar R hasta que $T = 240^{\circ}\text{C}$, si $T > 240^{\circ}$ desactivar R hasta que $T = 210^{\circ}\text{C}$. La temperatura de los residuos que se extraen de la tubería es aproximadamente de 100°C, mientras que el líquido nulo y líquido de calibración que se suministra a la unidad de análisis están a la temperatura ambiente.

Además de los cuatro modos de funcionamiento antes descritos existen dos modos más que se describen a continuación:

Parada forzada por temperatura baja: independiente del modo (de los cuatro anteriores) en que se encuentre la unidad de análisis, si en algún momento la temperatura del líquido a la entrada del analizador desciende por debajo de los 200°C se ha de activar una alarma; además, desactivar IA, desactivar la resistencia calefactora y detener la bomba. El sistema debe permanecer en esas condiciones hasta que se pulse el botón PM, momento en el cual se activaría nuevamente la señal IA, la resistencia calefactora y se pondría en marcha la bomba, retornado así, la unidad de análisis al modo 1 (análisis).

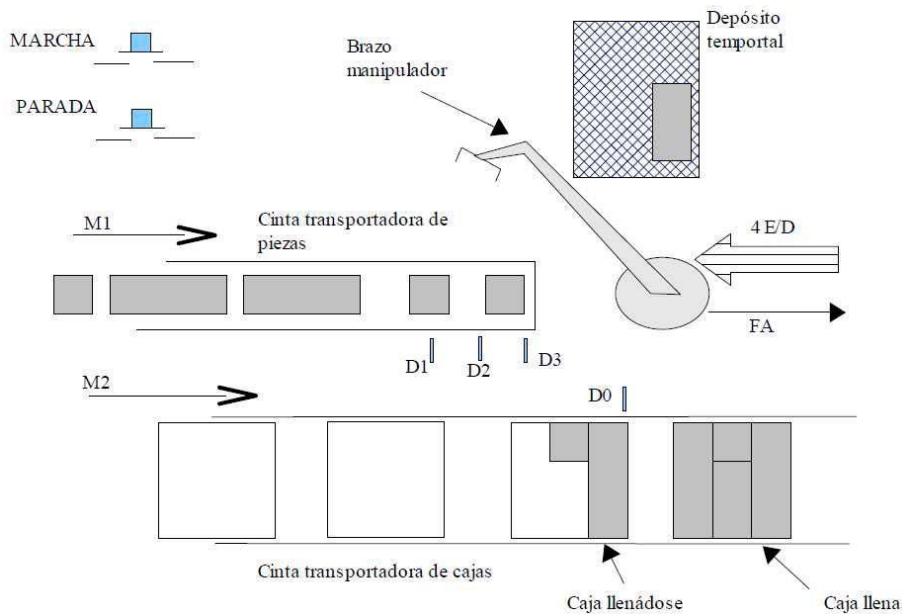
Modo 4; Parada normal: se inicia con el pulsador ST. Se desactiva la resistencia calefactora. Se apaga el analizador desactivando la señal ON/OFF. Con las válvulas V1, V2 y V3 cerradas y V4, V5 abiertas se hace pasar aguas por el conducto de líquidos hacia el analizador durante 50 s. Trascurrido ese tiempo se detiene la bomba y se cierran las válvulas V4 y V5. La unidad queda lista para una nueva puesta en marcha (modo 0).

El tiempo de apertura y cierre de las válvulas es de 2 s. Esta consideración sobre el tiempo de apertura y cierre de las válvulas se ha de tener en cuenta durante los cambios entre los distintos modos de funcionamiento.

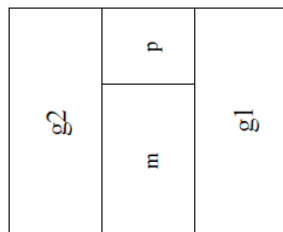
Obtener el Grafcet que resuelve el problema y su implementación en autó-mata programable.

Ejercicio 22 Sistema de empaquetado automático

La siguiente figura muestra la vista de planta de un sistema de empaquetado automático.



Las piezas grandes, medianas y pequeñas, a empaquetar, llegan hasta el brazo manipulador por una cinta transportadora movida por el motor M1. La distribución del tipo de piezas en la cinta es totalmente aleatoria. Cuando una pieza llega al final de la cinta (detector D3), el motor debe detenerse, parando la cinta a la espera de que el brazo manipulador agarre la pieza para depositarla en la caja que se está llenando (si es posible) o en el depósito temporal. Si la pieza es grande, estarán activos los detectores D1, D2 y D3, si es mediana estarán activos D2 y D3 y si es pequeña sólo estará activo D3. En cada caja deben ir dos piezas grandes, una mediana y una pequeña con la siguiente distribución:



Debido a que las piezas llegan por la cinta hasta el brazo manipulador de forma aleatoria, se utiliza un depósito temporal para almacenar las piezas que no puedan ser introducidas en la caja que se esté llenando. Para llenar una caja se deben usar, de forma prioritaria (siempre que sea posible), las piezas que estén en el depósito temporal. Una vez la caja está llena se activa el motor M2 para mover la cinta hasta que se desactive D0 y se vuelva a activar, indicando

la presencia de una nueva caja vacía. El proceso de llenado de cajas se repite continuamente. A modo de ejemplo, considérese la situación representada en la figura anterior. La secuencia de movimiento de las piezas debe ser la siguiente en ese caso:

- Pieza mediana en el depósito a la caja en la posición m.
- Pieza pequeña en cinta al depósito temporal.
- Avanza cinta de piezas activando M1 hasta activar D3
- Pieza pequeña en cinta al depósito temporal.
- Avanza cinta de piezas activando M1 hasta activar D3.
- Pieza grande de la cinta a la caja en la posición g2.
- Avanza la cinta de piezas activando M1 hasta que se active D3 y la cinta de cajas activando M2 hasta que una caja vacía sea detectada por la activación de D0 (tras su desactivación).

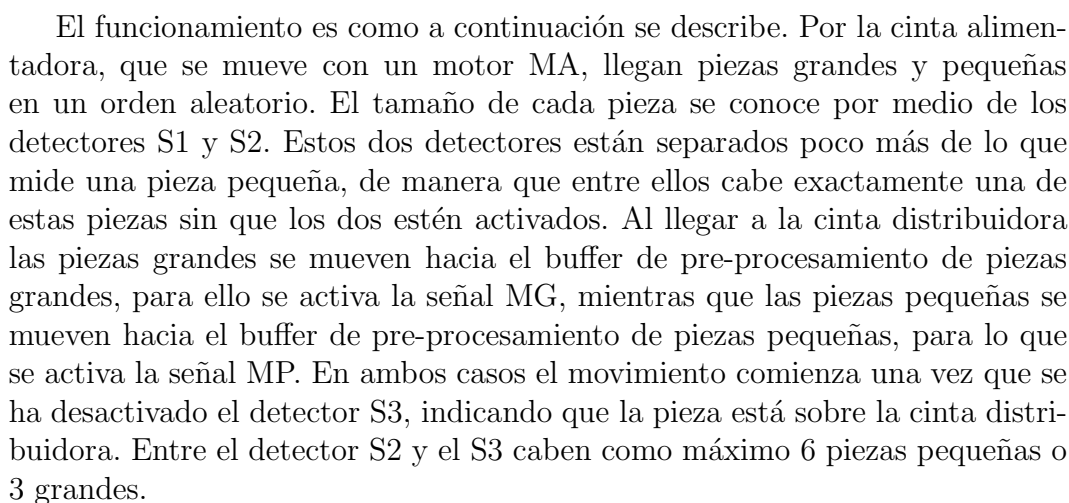
El brazo manipulador tiene cuatro entradas digitales mediante las cuales se le envían, desde el Autómata, las órdenes de movimiento y captura de piezas usando la siguiente codificación:

<i>Acción</i>	<i>Código</i>
Coger pieza grande de cinta	0000
Coger pieza mediana de cinta	0001
Coger pieza pequeña de cinta	0010
Coger pieza grande de depósito temporal	0011
Coger pieza mediana de depósito temporal	0100
Coger pieza pequeña de depósito temporal	0101
Poner pieza grande en depósito temporal	0110
Poner pieza mediana en depósito temporal	0111
Poner pieza pequeña en depósito temporal	1000
Poner pieza grande en caja, posición g1	1001
Poner pieza grande en caja, posición g2	1010
Poner pieza pequeña en caja, posición p	1011
Poner pieza mediana en caja, posición m	1100

Además, tiene una salida digital (FA) que se activa cuando el manipulador ha completado cualquiera de las acciones anteriores. Para realizar las acciones, el brazo parte de una posición inicial, realiza la acción y regresa a la posición inicial activando FA, que se desactiva cuando el brazo recibe el código de una nueva acción. Si no se está realizando ninguna acción, el brazo permanece en la posición inicial. Habrá un pulsador de MARCHA y otro de PARADA. Al iniciarse el sistema todo estará desactivado hasta que se pulse el pulsador de

Obtener el Grafcet que resuelve el problema y su implementación en autó-mata programable.

El siguiente esquema representa una línea de tratamiento y embalaje de piezas.



Mientras la cinta distribuidora está moviendo la pieza al buffer correspondiente, la cinta alimentadora se mantiene parada hasta que la pieza llegue al buffer, lo cual se indica mediante los detectores SG para las piezas grandes, o SP para las piezas pequeñas. Cada uno de los buffers de pre-procesamiento tiene un sistema de control propio que es capaz de hacer las operaciones necesarias para almacenar o dispensar las piezas. Cada buffer tiene una señal de entrada: ALMG y ALMP. Si esta señal (ALMG o ALMP) está activa, el buffer debe almacenar las piezas que le lleguen y no dispensar, si está desactivada, el buffer debe dispensar piezas y no almacenarlas; en ese caso, las piezas que entran al buffer salen directamente de éste.

De los buffer de pre-procesamiento las piezas pasan a las cintas de procesamiento, que se mueven por la activación de las señales MPG para las piezas grandes y MPQ para las piezas pequeñas. A las piezas grandes se les aplica una operación A que dura 5 s, mientras que a las piezas pequeñas se les aplica la operación B que dura 15 segundos. Estas operaciones comienzan cuando se activan los detectores SPG y SPQ respectivamente. Las velocidades de las cintas son las adecuadas en cada caso para que las operaciones se realicen con las piezas en movimiento. Una vez finalizadas las operaciones, las piezas continúan hasta los buffer de post-procesamiento. Desde estos buffers un brazo robot coge las piezas y las coloca en cajas para su posterior comercialización.

Las cajas para el embalaje de las piezas vienen por la cinta de embalaje. Existen dos formas de realizar el embalaje que se distinguen mediante etiquetas pegadas en las cajas:

- Tres piezas pequeñas: las cajas tienen una etiqueta negra.
- Una pieza pequeña y una grande: las cajas tienen una etiqueta blanca.

La llegada de una caja a la posición de embalaje se detecta con el detector SEM, mientras que el tipo de embalaje se detecta con un detector STE: si está activo, tres piezas pequeñas y si no está activo una pieza grande y una pequeña. Durante el llenado de una caja, la cinta de embalaje se mantiene parada y se pone en marcha cuando finaliza éste, hasta que llegue otra caja vacía a la posición de llenado.

El llenado de una caja sólo comenzará si el número de piezas en los buffer de post-procesamiento es suficiente para completarla. En caso contrario, se encenderá una alarma de buffer vacío (APP o APG) y se esperará hasta que estén las piezas disponibles en los buffers. Sólo entonces se apagará la alarma y se procederá al llenado de la caja.

El brazo robot tiene dos entradas (E1, E2) mediante las cuales se le da las órdenes según se muestra en el cuadro siguiente:

Código (E1, E2)	Orden
00	Permanecer en la posición inicial de reposo.
01	Llenar caja con tres piezas pequeñas.
10	Llenar caja con una piezas pequeña y una grande.
11	Abortar cualquier operación y regresar a la posición de reposo.

El brazo cuenta con un sistema de control propio que se encarga de controlar todos los movimientos para que se ejecute la orden dada. Además tiene una señal de salida FO la cual se activa cada vez que el brazo finaliza alguna de las órdenes recibidas.

Dado que la demanda de piezas es variante en el tiempo debido a las diferentes formas de embalaje, la función de los buffers es almacenar las piezas que van llegando en exceso. En un inicio todos los buffers están vacíos, y los buffers de post-procesamiento están listos para almacenar las piezas que lleguen en exceso. Estos buffers tienen un sistema de control propio que hace que funcionen de forma autónoma. Por otra parte, los buffers de pre-procesamiento sólo deben comenzar a almacenar piezas cuando los buffers de post-procesamiento estén llenos, lo cual ocurre en el caso de las piezas grandes cuando el número de piezas en el de post-procesamiento sea de 20 piezas y en el caso de las pequeñas, cuando el número de piezas en el buffer de post-procesamiento sea 30. En ese caso, la cinta de procesamiento se debe parar y se pondrá en marcha cuando nuevamente haya espacio en los buffers de post-procesamiento. La capacidad de los buffers de pre-procesamiento es de 30 piezas, tanto de piezas grandes como de pequeñas. Si alguno de los buffers de pre-procesamiento se llena, hay que detener el motor de la cinta alimentadora cuando la pieza que llegue a la cinta distribuidora sea del tipo del que está lleno el buffer de pre-procesamiento.

Las piezas que salen de los buffers de pre-procesamiento se detectan por los detectores SG1 y SP1, y las que entran a los buffers de post-procesamiento se detectan por los detectores SG2 y SP2, para las piezas grandes y pequeñas respectivamente.

El sistema tiene un botón de MARCHA que inicia el procesamiento de las piezas desde las siguientes condiciones iniciales: 1) todos los buffers vacíos; 2) cintas de procesamiento y distribución sin piezas; 3) no hay piezas entre los detectores S1 y S3.

El botón de PARADA permite detener el procesamiento y embalaje de las piezas. Cuando se pulsa este botón se detiene la cinta alimentadora. Si hay alguna pieza en la cinta distribuidora, se lleva hasta el buffer correspondiente. Se procesan todas las piezas que haya en los buffers de pre-procesamiento hasta que éstos queden vacíos. Tampoco deben quedar piezas en los buffers de post-procesamiento. El sistema se detendrá cuando haya una caja vacía o a medio llenar pero no queden piezas para finalizar la operación.

También hay un botón de EMERGENCIA, de forma que cuando se pulsa, todas las cintas deben detenerse y el brazo robot, ir a la posición de reposo inicial. Además debe encenderse una ALARMA, la cual se apagará después del rearme del botón de EMERGENCIA, quedando el sistema listo para comenzar el funcionamiento mediante el botón de MARCHA.

Obtener el Grafcet que resuelve el problema y su implementación en autó-mata programable.

Bibliografía

- [1] *Catalogo General, Telemecanique 2009.*
- [2] *Componentes neumáticos. Parker Pneumatic, 1996.*
- [3] *Comunicación Wi-Fi con MOXA AirWorks AWK-1100, Omron.*
- [4] *Comunicaciones Industriales Omron.*
- [5] *Manual de automatismos de mando neumático. Grupo Schneider, 2009.*
- [6] *Manual de programación CQM1/CPM1/CPM1A/SRM1, Omron.*
- [7] *Modem WireLess DeviceNet WD30, Omron.*
- [8] Josep Balcells and José Luis Romeral. *Autómatas programables.* Marcombo, 1997.
- [9] B. R. Bannister. *Instrumentación transductores e interfaz.* Addison-Wesley Iberoamericana, 1994.
- [10] Antonio Barrientos, Luis Felipe Peñin, and Jesús Carrera. *Automatización de la fabricación: autómatas programables, actuadores, transductores.* Universidad Politécnica de Madrid, 1999.
- [11] Oriol Boix. *Curso de GRAFCET y GEMMA.* Universitat Politècnica de Catalunya, 2002.
- [12] Oriol Boix Aragonés, Antoni Sudrià Andreu, and Joan Bergas Jané. *Automatització industrial con Grafcet.* Ediciones UPC, 1997.
- [13] W. Bolton. *Instrumentación y control industrial.* Paraninfo, 1996.
- [14] Jean-Calude Bossy, Paul Brard, Patrice Faugère, and Christian Merlaud. *GRAFCET. Prácticas y aplicaciones.* Universitat Politècnica de Catalunya, 1995.
- [15] K. Clements-Jewery and W. Jeffcoat. *The PLC Workbook.* Prentice Hall, 1996.

- [16] Antonio Creus Solé. *Instrumentación industrial*. Marcombo Boixareu Editores, 2009.
- [17] Alterto G. Davie and Mario B. Villar. *Introduccion a la automatizacion industrial*. Editorial Universitaria de Buenos Aires, 1965.
- [18] Clarence W. de Silva. *Sensors and Actuators Control system instrumentation*. CRC Press, 2007.
- [19] Joan Domingo Peña. *Comunicaciones en el entorno industrial*. Editorial UOC, 2003.
- [20] Andrés García Higuera. *El control automático en la industria*. Universidad de Castilla la Mancha, 2005.
- [21] Emilio García Moreno. *Automatización de procesos industriales: robótica y automática*. Centro de Formación de Postgrado UPV, 1999.
- [22] S.S. Iyengar and L. Min. *Advances in distributed sensor technology*. Prentice - Hall, 1995.
- [23] J.R. Jordan. *Serial networked field instrumentation*. John Wiley & Sons, 1995.
- [24] Enrique Mandado, Jorge Marcos Acevedo, Celso Fernández, and Serafín Pérez José I. Armesto. *Automatas Programables. Entornos y aplicaciones*. Thomson Editores, 2005.
- [25] Jordi Mayné. *Estado actual de las comunicaciones inalámbricas*. SILICA, 2005.
- [26] Albert Mayol. *Autómatas programables*. Marcombo Boixareu editoras, 1992.
- [27] G. Michel. *Autómatas programables Industriales*. Marcombo, 1990.
- [28] H. Norton. *Sensores y analizadores*. Gustavo Gili, 1984.
- [29] R. Payas. *Transductores y acondicionadores de señal*. Marcombo, 1994.
- [30] Manuel Pineda Sánchez. *Automatización de maniobras industriales mediante autómatas programables*. Universidad Politécnica de Valencia, 2006.
- [31] Pere Ponsa Asensio and Ramon Vilanova. *Automatización de procesos mediante la guía GEMMA*. Universitat Politècnica de Catalunya, 2005.
- [32] A. Simon. *Autómatas programables*. Paraninfo, 1989.
- [33] Jon Stenerson. *Fundamentals of programmable Logic Controllers, sensors and communications*. Prentice - Hall Career & Technology, 1998.