

• FSD-FMTC •

Fundamentos de Sistemas Dinámicos

• Física, Matemáticas y Teoría de Control •



FSD-FMTC

Fundamentos de Sistemas Dinámicos

Física, Matemáticas y
Teoría de Control

Arcenio Brito H.

Versión 2.0

2026

FSD-FMTC
Fundamentos de Sistemas Dinámicos
Física, Matemáticas y Teoría de Control

Autor
Arcenio Brito H.

Versión
2.0

Obra original
Mayo de 2007

Edición actual
Agosto de 2026

Lugar de edición
MX

Contacto
arcbritoh@gmail.com

© 2026 Arcenio Brito H.
Todos los derechos reservados.

El autor autoriza la lectura, consulta, estudio e impresión del contenido de esta obra para formación personal y fines educativos.

Quedan prohibidas su alteración, distorsión o uso indebido que desvirtúe su sentido original.

Para uso comercial se requiere autorización expresa del autor.

Prólogo 2026

¿Por qué escribir un libro sobre los principios que subyacen a disciplinas aparentemente inconexas?

Este trabajo surgió de una serie de interrogantes planteadas mientras cursaba la licenciatura en Ingeniería Eléctrica acerca de la relación entre la física, las matemáticas y el comportamiento de los sistemas. Específicamente, sobre cómo disciplinas aparentemente distintas terminan describiendo fenómenos mediante estructuras matemáticas formalmente equivalentes y si, desde una perspectiva ingenieril, las matemáticas pueden entenderse más como un lenguaje de representación que como un mero conjunto de procedimientos.

A lo largo de la licenciatura, la variedad de asignaturas del mapa curricular producía una sensación de fragmentación: **cada materia parecía introducir una terminología propia y abordar problemas inconexos**. Sin embargo, conforme avanzaba en la carrera, comencé a percibir algo que rara vez se enseñaba de forma explícita: muchas de estas disciplinas no eran sino manifestaciones de principios generales que se repetían consistentemente. Parte de las bases en Física que sustentaron esta perspectiva surgieron durante el bachillerato, etapa en la cual el autor participó a nivel nacional en la competencia de Física de CONACIBA en 2001.

Hacia el año 2006, al finalizar la etapa de licenciatura, aquella inquietud intelectual se consolidó en una práctica constante. La programación, como recurso de aprendizaje desde el bachillerato, pasó a utilizarse de forma continua durante la carrera para resolver problemas específicos. Esto incluyó el desarrollo de una herramienta en VB6 para procesar parámetros de transformadores, generadores y carga eléctrica en sistemas eléctricos distribuidos, generando dinámicamente diagramas unifilares y topologías en anillo dentro de AutoCAD.

Poco después, en el Centro de Investigación en Ingeniería y Ciencias Aplicadas (CIICAp), se desarrolló una aplicación en VB6 para la ingesta, parseo y conversión de datos no estructurados a formatos estructurados, con la cual se procesó en grandes volúmenes información científica de la APS (American Physical Society) y otras fuentes, como parte de un proyecto de investigación de ese instituto; trabajo con el cual la universidad acreditó las horas de formación práctica del autor.

Posteriormente me incorporé a un programa de formación en C++ dentro de una empresa tecnológica dedicada al desarrollo de software para dispositivos PocketPC. Esta experiencia resultó determinante: la rigurosidad sintáctica y el control de estructuras en C++ facilitaron la posterior transición hacia otros lenguajes como Java y C# con relativa

facilidad, ya que estos heredaban gran parte de los principios, paradigmas y modelos de programación presentes en C++.

A lo largo de los años he tenido la oportunidad de aplicar estos conocimientos en distintos sectores (*energético, manufacturero, financiero y retail*), participando en proyectos relacionados con el *desarrollo de software, la ingeniería de datos y la ingeniería de soporte a sistemas críticos*.

Lejos de representar ámbitos independientes, estas experiencias reforzaron la idea que había comenzado a tomar forma durante la licenciatura: **independientemente de si se trata de sistemas físicos, procesos industriales, software, sistemas empresariales en producción o sistemas de ingeniería de datos a gran escala, los principios fundamentales para comprender, modelar, controlar y asegurar su estabilidad conservan una notable unidad conceptual.**

Este es, en gran medida, el libro que me hubiera gustado leer cuando inicié mis estudios universitarios. No un texto dedicado a presentar cientos de técnicas o artificios matemáticos para resolver ecuaciones aisladas, sino un marco orientado a comprender los principios que vinculan estos campos del conocimiento y a responder muchas de las preguntas que naturalmente surgen durante los primeros semestres de una carrera de ingeniería o de disciplinas afines.

Preguntas tales como: ¿Por qué el estudio de la física, del cálculo y las ecuaciones diferenciales resulta indispensable para comprender el comportamiento de los sistemas? ¿Para qué sirven realmente las transformadas integrales en el análisis de ingeniería? ¿Por qué la ecuación que describe un sistema mecánico masa-resorte presenta una estructura matemática idéntica a la de un circuito eléctrico RLC? ¿Y por qué, a pesar de la multiplicidad de sistemas, las herramientas matemáticas para analizarlos parecen ser siempre las mismas?

El razonamiento que dio origen a este trabajo se derivó de la observación de que muchas de esas preguntas no se responden estudiando cada disciplina de manera aislada, sino integrando los principios fundamentales que las distintas ramas de la ingeniería y las ciencias aplicadas tienen en común. **Entre esos principios, el concepto de transformación ocupa un lugar central, al proporcionar un marco unificado para representar un mismo sistema o fenómeno físico mediante diferentes representaciones, dominios matemáticos o marcos de referencia, preservando las relaciones fundamentales que describen su comportamiento.**

En esta obra, el término transformación se utiliza en su sentido matemático, físico e ingenieril. Transformar un sistema no implica necesariamente mejorarlo, optimizarlo o corregirlo; significa modificar su estado, su representación o el dominio en el que se analiza mediante una operación que preserve las relaciones fundamentales de interés. La valoración de dicha transformación pertenece al ámbito de los objetivos perseguidos y de la responsabilidad con la que el conocimiento se aplica, y no al concepto de transformación en sí mismo.

Las transformadas dejan entonces de entenderse únicamente como técnicas matemáticas para convertirse en manifestaciones particulares de ese principio general de transformación. Las leyes físicas dejan de verse como expresiones independientes y comienzan a revelar similitudes estructurales que trascienden cada disciplina. Los modelos matemáticos se despojan de su carácter exclusivo para entenderse como distintas representaciones de un mismo comportamiento dinámico. La posible formalización de este enfoque dentro de un marco matemático más general representa una línea de estudio cuyo desarrollo excede el alcance y objetivos de este libro.

Ese enfoque se refleja en la estructura de este trabajo. Los primeros capítulos establecen el lenguaje matemático y los principios físicos sobre los que se fundamenta el análisis de los sistemas dinámicos. A partir de ellos se desarrolla el modelado de sistemas pertenecientes a distintos dominios de la ingeniería, mostrando que, aunque cambien los fenómenos físicos subyacentes, la formulación utilizada para describir su comportamiento conserva una notable unidad. Los capítulos dedicados a teoría de control e instrumentación se presentan como una consecuencia natural de esos fundamentos y no como el propósito central del texto.

Con el paso del tiempo, casi dos décadas después, aquella intuición inicial no hizo sino confirmarse. Las herramientas evolucionan, aparecen nuevos lenguajes, nuevas arquitecturas y nuevas tecnologías; sin embargo, **los principios fundamentales que rigen el comportamiento, la interacción, el sostenimiento y la evolución de los sistemas complejos permanecen esencialmente invariables.**

La presente edición conserva prácticamente intacto el contenido desarrollado en la versión original. Las modificaciones realizadas se limitaron principalmente a los capítulos 4 y 5, específicamente a la unificación de ecuaciones, la actualización de algunos diagramas y, en determinados casos, a mejorar la claridad expositiva del texto.

Este trabajo no pretende responder todas las preguntas relacionadas con el análisis de sistemas lineales y la teoría de control, ni está exento de errores u omisiones. Como toda creación humana, es susceptible de revisión, ampliación y perfeccionamiento. La organización y el desarrollo del contenido privilegian la continuidad del razonamiento, introduciendo los conceptos en el momento en que resultan necesarios para construir una comprensión gradual e incremental de los temas, favoreciendo la integración sistémica de las ideas frente a la rigurosa formalización aislada.

Más que presentar una colección de técnicas o procedimientos, este trabajo busca desarrollar una forma de razonar acerca de los sistemas, partiendo de sus fundamentos matemáticos y físicos para comprender cómo pueden describirse, modelarse y analizarse. Desde esta perspectiva, las herramientas y métodos utilizados constituyen una consecuencia natural de esa comprensión y no su punto de partida.

Si estas páginas logran que, al estudiar una nueva disciplina, el lector deje de ver únicamente fórmulas aisladas y comience a reconocer principios que se manifiestan bajo diferentes formulaciones matemáticas y físicas, entonces este trabajo habrá cumplido plenamente el propósito con el que fue concebido.

Porque, cuando diferentes sistemas terminan describiéndose mediante las mismas estructuras matemáticas, quizá no sea la matemática la que se repite, sino la naturaleza de los sistemas que buscamos comprender.

Arcenio Brito H.

México, agosto de 2026

Historial de versiones

Versión	Fecha	Descripción
1.0	Mayo 2007	Desarrollo del trabajo original.
1.1	Julio 2026	Actualización y revisión. Nuevo título, prólogo, mejoras de redacción en capítulos IV y V, uniformización de fórmulas y actualización de figuras.
2.0	Agosto 2026	Se agregan anexos: Ecuaciones y Leyes Matemáticas y Físicas. Se incrementa la versión a 2.0 para denotar el esfuerzo en la revisión y consistencia de la obra.

Entorno de desarrollo

La revisión de 2026 se realizó empleando Microsoft Word 2003 y un entorno de sistema operativo equivalente al de la edición original, conservando la continuidad tecnológica del proyecto.

Herramienta	Versión
Procesador de texto	Microsoft Word 2003
Editor de fórmulas	Microsoft Equation Editor 3.1
Dibujos técnicos, diagramas y figuras 2D y 3D	AutoCAD 2004
Programación en C++	Microsoft Visual C++ 6.0
Métodos numéricos	MATLAB 4.2c.1
LaTeX-2e	Actualización de ecuaciones y diagramas técnicos (edición 2026)
LaTeX-2e	Implementación de los anexos (edición 2026)

Nota de la edición 2026

Esta edición corresponde a la **versión 2.0** del trabajo original desarrollado entre **2006 y 2007**. Con el propósito de preservar el contenido original y mejorar su claridad expositiva, se realizaron únicamente las siguientes revisiones:

- actualización del título de la publicación;
- incorporación del prólogo correspondiente a la edición 2026;
- incorporación del historial de cambios;
- revisión de redacción en los capítulos IV y V para mejorar la claridad del texto;
- actualización de algunas ecuaciones, figuras y diagramas;
- uniformización de la notación matemática utilizando el editor de ecuaciones de Microsoft Word;
- corrección de ortografía y ajustes de formato para mejorar la legibilidad del documento;
- actualización de la tabla de contenido derivada de los ajustes realizados.
- implementación de anexos "Ecuaciones y Leyes Matemáticas y Físicas", desarrollado en LaTeX2e, que amplía y complementa el material original con nuevos temas y herramientas de consulta.

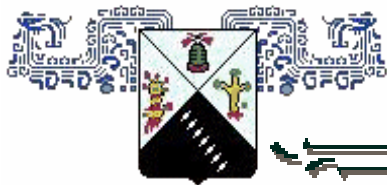
Los capítulos **I, II y III**, que constituyen el núcleo conceptual del trabajo, conservan íntegramente su contenido, redacción y estructura originales.

Prefacio de la edición original 2007

El presente trabajo tiene por objeto presentar los fundamentos matemáticos y físicos que permitan el análisis de sistemas físicos dinámicos. Se repasan las técnicas de notación y transformación matemática más comunes y se presentan las características de los bloques funcionales que componen los sistemas físicos complejos. Se emplea la función de transferencia por transformada de Laplace en dominio complejo, y el espacio de estados para la representación de ecuaciones diferenciales lineales que resultan del análisis de un sistema físico. Se modelan sistemas que involucran diferentes tipos de energía y se destacan las equivalencias estructurales entre sus modelos matemáticos. Se presenta una introducción teórica a los sistemas de control y finalmente se estudian algunos instrumentos para su medición y corrección.

Extensión del prefacio edición 2026

En la versión v2.0-r2-2026 de la obra se incorpora la sección de anexos «Ecuaciones y Leyes Matemáticas y Físicas». Este apartado está concebido como un compendio de consulta inmediata para lectores de todos los niveles, desde el preuniversitario hasta el avanzado, que deseen acceder con rapidez a la notación, las definiciones, las fórmulas y las leyes empleadas a lo largo del libro. Mientras que el cuerpo de la obra adopta un enfoque sistémico, construyendo el conocimiento de manera progresiva para mostrar la convergencia de las herramientas matemáticas y físicas en el análisis de sistemas dinámicos, los anexos desglosan y organizan estos recursos de forma independiente, tal como se estudian en las aulas, facilitando así su consulta y memorización. De este modo, el lector cuenta con dos vías complementarias: el razonamiento y el prontuario.



UNIVERSIDAD AUTÓNOMA DEL
ESTADO DE MORELOS

"POR UNA HUMANIDAD CULTA"

FACULTAD DE CIENCIAS QUÍMICAS E INGENIERÍA

*PRINCIPIOS FÍSICOS Y MATEMÁTICOS
PARA EL ANÁLISIS DE SISTEMAS
DINÁMICOS. INTRODUCCIÓN AL
CONTROL*

*Trabajo de
desarrollo
profesional
por etapas*

P R E S E N T A:
ARCENIO BRITO H.

MX, MAYO 2007

PORTADA EDICION ORIGINAL

Nota de la portada 2026: La portada de la edición 2026 se diseñó con la figura de Quetzalcóatl en su versión v1.1-r1-2026. Posteriormente, al revisar el lema de la UAEM en la versión v2.0-r2-2026, se encontró que Quetzalcóatl es parte de su logotipo, como se describe brevemente a continuación:

<https://www.uaem.mx/vida-universitaria/identidad-universitaria/lema-y-logotipo-universitaria.php>

El logotipo de la UAEM integra elementos de la cosmovisión nahua: Tamoanchan, Huaxtepec, Cuauhnáhuac y la serpiente emplumada, Quetzalcóatl, junto con el lema "Por una humanidad culta", que orienta el quehacer académico hacia la difusión del conocimiento y el desarrollo integral del espíritu humano.

Convención de notación

1. Los **variables** (*magnitudes de valor dinámico*) se denotan por letras minúsculas ($x, v, a, f, \dots$).
2. Las **constantes** (*parámetros de valor fijo*) por mayúsculas ($G, M, R, \dots$).
3. Los **escalares** (*magnitudes puramente numéricas*) por letras minúsculas ($x, v, a, f, \dots$).
4. Los **vectores** (*entidades con dirección y sentido*) indistintamente por letras con flecha superior ($\vec{x}, \vec{v}, \vec{a}, \vec{f}, \dots$) o bien en negritas ($\mathbf{x}, \mathbf{v}, \mathbf{a}, \mathbf{f}, \dots$).
5. Los **vectores unitarios** (*indicadores de dirección adimensionales*) con circunflejo superior ($\hat{r}, \hat{n}$).
6. Los **módulos de vectores** (*magnitudes escalares no negativas*) se denota como el escalar correspondiente o usando símbolos de *valor absoluto* ($|\vec{x}|, |\vec{v}|, |\vec{a}|, |\vec{f}|, \dots$).
7. Los **números complejos** (*estructuras de parte real e imaginaria*) por letras con circunflejo ($\hat{Z}, \hat{Y}$) o bien en algunos casos por letras mayúsculas negritas ($\mathbf{Z}, \mathbf{Y}$).
8. Los **fasores** (*representaciones complejas de señales armónicas*) se denotan por letras con tilde superior ($\tilde{V}, \tilde{I}$).
9. Las **matrices** (*arreglos bidimensionales de elementos*) se denotan por negritas en mayúsculas ($\mathbf{M}, \mathbf{X}$).
10. Los **elementos de una matriz** (*componentes indexadas por fila y columna*) se denotan con la letra minúscula correspondiente y dos subíndices (*fila, columna*) a_{ij}, m_{ij} .

Notas

1. A menos que se indique lo contrario, todos los **fasores** representan valores RMS (*root mean square*).
2. Para el análisis trifásico se adopta la secuencia de fase positiva (ABC), donde $V_b = a^2 V_a$ y $V_c = a V_a$, con el operador de rotación $a = e^{j120^\circ}$. La matriz de transformación de componentes simétricas $\mathbf{A}$ y su inversa $\mathbf{A}^{-1}$ se definen en consecuencia.
3. Por razones técnicas, algunas figuras y diagramas utilizan términos en inglés, de acuerdo con la terminología comúnmente utilizada en la literatura técnica y la práctica profesional.
4. Salvo indicación expresa en contrario, la notación matemática y la simbología se mantienen uniformes en la obra para representar un mismo concepto o magnitud física.

TABLA DE CONTENIDO

I. MÉTODOS MATEMÁTICOS DE TRANSFORMACIÓN1

NOTACIÓN SIGMA, INDUCCIÓN MATEMÁTICA.....	3
<i>Notación sigma</i>	4
NOTACIÓN REAL Y COMPLEJA	7
NOTACIÓN DE RAZONES DE CAMBIO Y SUMAS INFINITESIMALES	14
NOTACIÓN INTEGRODIFERENCIAL	25
<i>Ecuaciones integrodiferenciales</i>	25
<i>ED de primer orden y primer grado</i>	25
<i>ED de primer orden y grado superior</i>	27
<i>ED lineal de orden superior</i>	27
NOTACIÓN MATRICIAL Y SISTEMAS DE ECUACIONES LINEALES	29
NOTACIÓN VECTORIAL.....	40
<i>Algebra vectorial</i>	41
<i>Geometría vectorial</i>	42
<i>Cálculo vectorial</i>	43
<i>Electromagnetismo y las ecuaciones de Maxwell</i>	48
TRANSFORMACIÓN DE SISTEMAS COORDENADOS	52
TRANSFORMACIÓN FASORIAL.....	53
<i>Transformación fasorial en circuitos eléctricos de corriente alterna</i>	56
<i>Transformación de fuentes de ca</i>	61
<i>Transformación delta y estrella</i>	64
<i>Transformación unidad</i>	70
TRANSFORMACIÓN EN COMPONENTES SIMÉTRICAS	73
TRANSFORMACIÓN DE LAPLACE	76
<i>Transformación de funciones</i>	78
<i>Transformación inversa de Laplace</i>	80
<i>Transformación de ecuaciones diferenciales lineales invariantes en el tiempo</i>	83
TRANSFORMACIÓN EN SERIES DE POTENCIAS	85
<i>Convergencia uniforme</i>	87
<i>Serie de potencias de convergencia uniforme</i>	88
<i>Solución de ecuaciones diferenciales por series de potencias</i>	92
TRANSFORMACIÓN DE FOURIER	96
<i>Serie de Fourier</i>	96
<i>La transformación de Fourier</i>	100
TRANSFORMACIÓN ZETA.....	104

II. BLOQUES FUNCIONALES DE SISTEMAS FÍSICOS109

INTRODUCCIÓN.....	113
<i>Sistema General de Unidades de Medida.....</i>	<i>113</i>
SISTEMAS MECÁNICOS	120
<i>Fundamentos de cinemática y dinámica traslacional.....</i>	<i>120</i>
<i>Bloques funcionales de sistemas mecánicos con movimientos traslacionales</i>	<i>123</i>
<i>Fundamentos de cinemática y dinámica rotacional</i>	<i>124</i>
<i>Bloques funcionales de sistemas mecánicos con movimientos rotacionales.....</i>	<i>127</i>
SISTEMAS ELÉCTRICOS.....	129
<i>Fundamentos de electrostática y electrodinámica.....</i>	<i>129</i>
<i>Bloques funcionales de sistemas eléctricos</i>	<i>130</i>
SISTEMAS FLUÍDICOS	133
<i>Fundamentos de hidrostática e hidrodinámica</i>	<i>133</i>
<i>Bloques funcionales de sistemas fluidicos hidráulicos</i>	<i>135</i>
<i>Fundamentos de fisicoquímica y termodinámica.....</i>	<i>137</i>
<i>Bloques funcionales de sistemas fluidicos neumáticos</i>	<i>141</i>
SISTEMAS TÉRMICOS.....	143
<i>Fundamentos de transferencia de calor y calorimetría.....</i>	<i>143</i>
<i>Bloques funcionales de sistemas térmicos.....</i>	<i>145</i>
EQUIVALENCIAS ENTRE BLOQUES FUNCIONALES DE SISTEMAS FÍSICOS	146

III. MODELADO DE SISTEMAS DINÁMICOS145

MODELADO DE SISTEMAS DINÁMICOS	149
<i>Modelado mediante función de transferencia.....</i>	<i>150</i>
<i>Representación del sistema mediante diagramas de bloques.....</i>	<i>151</i>
<i>Modelado mediante espacio de estados.....</i>	<i>155</i>
<i>Relación entre funciones de transferencia y espacio de estados</i>	<i>157</i>
SISTEMAS MECÁNICOS	158
SISTEMAS ELÉCTRICOS.....	161
SISTEMAS FLUÍDICOS	164
SISTEMAS TÉRMICOS.....	166
SISTEMAS ELECTROMECAÑICOS	167

IV. SISTEMAS DE CONTROL.....171

CONTROL AUTOMÁTICO.....	175
<i>Sistema de lazo abierto.....</i>	<i>175</i>
<i>Sistema de lazo cerrado.....</i>	<i>176</i>
<i>Control de un proceso industrial.....</i>	<i>178</i>
<i>Señales de transmisión.....</i>	<i>180</i>
<i>Unidades de uso común en sistemas de control.....</i>	<i>181</i>
COMPONENTES DE UN SISTEMA DE CONTROL.....	182
<i>Sensores y transmisores.....</i>	<i>182</i>
<i>Válvulas de control.....</i>	<i>183</i>
<i>Controladores.....</i>	<i>186</i>
CONTROLADORES LÓGICOS PROGRAMABLES.....	192
<i>Componentes del PLC.....</i>	<i>193</i>
<i>Operación del PLC.....</i>	<i>195</i>

V. INSTRUMENTOS DE CONTROL195

INSTRUMENTOS DE CONTROL INDUSTRIAL.....	199
<i>Instrumentos de medición.....</i>	<i>199</i>
<i>Instrumentos de corrección.....</i>	<i>210</i>
INSTRUMENTOS DE MEDICIÓN UTILIZADOS EN PROCESOS INDUSTRIALES.....	216
<i>Instrumentos que miden cantidades mecánicas o termodinámicas.....</i>	<i>216</i>
<i>Instrumentos que miden cantidades eléctricas.....</i>	<i>223</i>

VI. ANEXOS

A. NOTACIÓN MATEMÁTICA, LÓGICA Y RELACIONAL

- A.1 Alfabeto griego*
- A.2 Símbolos lógicos fundamentales*
- A.3 Conjuntos y operaciones entre conjuntos*
- A.4 Relaciones de orden y comparación*
- A.5 Operadores matemáticos comunes*
- A.6 Notación de funciones y mapeos*
- A.7 Notación para variables dependientes del tiempo y señales*

B. PRINCIPIOS Y AXIOMAS MATEMÁTICOS FUNDAMENTALES

- B.1 Propiedades de la igualdad*
- B.2 Propiedades de los números reales*
- B.3 Propiedades de orden*

C. ÁLGEBRA

- C.1 Operaciones elementales*
- C.2 Productos notables*
- C.3 Factorización*
- C.4 Leyes de exponentes*
- C.5 Leyes de radicales*
- C.6 Ecuación lineal*
- C.7 Ecuación cuadrática*
- C.8 Sistema de dos ecuaciones lineales*
- C.9 Logaritmos*
- C.10 Introducción a límites*
- C.11 Álgebra booleana y tablas de verdad*

D. GEOMETRÍA PLANA Y DEL ESPACIO

- D.1 Principios fundamentales de la geometría*
- D.2 Figuras planas: perímetros, áreas y elementos*
- D.3 Figuras del espacio: áreas, volúmenes y elementos*
- D.4 Teoremas geométricos fundamentales*
- D.5 Geometría analítica*

E. TRIGONOMETRÍA

- E.1 Razones trigonométricas fundamentales*
- E.2 Identidades pitagóricas*
- E.3 Paridad (simetría)*
- E.4 Fórmulas de suma y diferencia de ángulos*
- E.5 Fórmulas del ángulo doble*
- E.6 Fórmulas del ángulo medio*
- E.7 Ley de senos*
- E.8 Ley de cosenos*
- E.9 Identidades trigonométricas adicionales*

F. VECTORES

- F.1 Notación*
- F.2 Operaciones fundamentales*
- F.3 Propiedades importantes*
- F.4 Aplicaciones físicas de vectores*

G. NÚMEROS COMPLEJOS

G.1 Representaciones

G.2 Operaciones básicas

G.3 Fórmulas fundamentales

G.4 Ejemplos de operaciones con números complejos

H. MATRICES Y SISTEMAS DE ECUACIONES LINEALES

H.1 Definición y tipos de matrices

H.2 Operaciones con matrices

H.3 Determinante de una matriz cuadrada

H.4 Matriz inversa

H.5 Rango de una matriz

H.6 Sistemas de ecuaciones lineales

H.7 Métodos de solución

H.8 Clasificación de sistemas (Teorema de Rouché-Frobenius)

H.9 Ejemplos ilustrativos

H.10 Aplicaciones en ingeniería y sistemas dinámicos

I. LEYES Y PRINCIPIOS FÍSICOS FUNDAMENTALES

I.1 Mecánica clásica

I.2 Oscilaciones y dinámica

I.3 Electricidad y circuitos

I.4 Fluidos y termodinámica

I.5 Elasticidad y materiales

J. FUNDAMENTOS DE CÁLCULO

J.1 ¿Qué es el cálculo?

J.2 La derivada: cambio instantáneo

J.3 La integral: acumulación

J.4 El Teorema Fundamental del Cálculo

J.5 ¿Qué es una ecuación diferencial?

J.6 La relación entre derivada, integral y ecuación diferencial

J.7 ¿Qué es la transformada de Laplace?

J.8 Teoremas importantes del cálculo

J.9 Resumen conceptual integrado

J.10 Conexión con el resto de la obra

K. CONCEPTOS ESENCIALES DE CONTROL

K.1 Función de transferencia

K.2 Sistemas elementales

K.3 Conexiones de bloques

K.4 Lazo cerrado y error

K.5 Controladores PID

K.6 Ecuación característica y estabilidad

K.7 Comportamiento dinámico en el tiempo

L. ECUACIONES DE MAXWELL

L.1 Las cuatro ecuaciones de Maxwell

L.2 Forma vectorial (diferencial) de las ecuaciones

L.3 Ecuación de onda electromagnética

L.4 Explicación de las ecuaciones de Maxwell

L.5 Conexión con el resto de la obra

Capítulo 1

Métodos Matemáticos de Transformación

Con el objeto de tener una sólida base matemática que nos permita el correcto análisis de los *sistemas físicos dinámicos*, en este capítulo se presentan los siguientes temas de *notación y transformación matemática*:

- Notación *sigma*, *inducción matemática*
- Notación *real y compleja*
- Notación de *razones de cambio y sumas infinitesimales*
- Notación *integrodiferencial*
- Notación *matricial y sistemas de ecuaciones lineales*
- Notación *vectorial y campos vectoriales*
- Transformación de *sistemas coordenados*
- Transformación *fasorial*
- Transformación *unidad*
- Transformación en *componentes simétricas*
- Transformación de *Laplace*
- Transformación en *series de potencias*
- Transformación de *Fourier*
- Transformación *zeta*

Notación sigma, inducción matemática

Número entero. El conjunto de números enteros es $Z = \{\dots, -2, -1, 0, 1, 2, \dots\}$, el conjunto de enteros positivos es $Z^+ = \{1, 2, 3, \dots\} = \{Z\} \mid Z > 0$; el conjunto de enteros negativos es $Z^- = \{\dots, -3, -2, -1\} = \{Z\} \mid Z < 0$. El conjunto de números enteros no negativos es: $Z^{0+} = \{0, 1, 2, 3, \dots\} = \{Z\} \cup \{0\}$. El factorial de un número (véase función Gamma) $n \mid n \in Z^+$ se define por $n! = 1 \cdot 2 \cdot 3 \cdot \dots \cdot n = n \cdot (n-1)!$, sus propiedades son: (1) $0! = 1$; (2) $(n+1)! = (n+1) \cdot n!$, en general el factorial es una operación recursiva:

$$n! = n \cdot (n-1)! = n \cdot (n-1) \cdot (n-2) \cdot (n-3) \cdot \dots \cdot (1) = \frac{(n+1)!}{n+1} \quad (1)$$

Número booleano. El álgebra booleana fue desarrollada por George Boole en 1854 (publicado en *An Investigation of the Laws of Thought*). Sea $x \in \{0, 1\}$ una variable booleana; se definen 3 operaciones fundamentales: complemento de x (x'), producto y suma y obedecen las siguientes identidades:

$$\begin{aligned} x \in \{0, 1\} \quad & x' = 1 \Leftrightarrow x = 0; x' = 0 \Leftrightarrow x = 1 \\ x + 0 = x \quad & x + 1 = 1 \quad x + x = x \quad x + x' = 1 \\ x \cdot 0 = x \quad & x \cdot 1 = x \quad x \cdot x = x \quad x \cdot x' = 0 \\ x + y = y + x \quad & x + (y + z) = (x + y) + z \quad x(y + z) = xy + xz \\ x \cdot y = y \cdot x \quad & x \cdot (y \cdot z) = (x \cdot y) \cdot z \quad x + y \cdot z = (x + y) \cdot (x + z) \\ (x + y)' = x' y' \quad & (x')' = x \quad (xy)' = x' + y' \end{aligned} \quad (2)$$

Principios de conteo. La regla de la suma establece que: “*sean A y B dos tareas mutuamente excluyentes (no pueden ocurrir simultáneamente) que pueden realizarse de m y n maneras; entonces la realización de una u otra tarea puede llevarse a cabo de m+n maneras*”.

La regla del producto establece que: “*si un procedimiento puede ser descompuesto en dos etapas consecutivas A y B, de manera que para m formas de realizar A, existan n formas de realizar B; entonces el procedimiento completo total puede ser efectuado en m×n formas*”. De los principios de conteo expuestos se deducen los conceptos de permutación, arreglo, combinación y combinación con repetición.

Orden relevante	Repeticiones	Tipo	Fórmula
Si	No	Permutación	$P(n, r) = n! / (n-r)!$
Si	Si	Arreglo	n^r
No	No	Combinación	$C(n, r) = \binom{n}{r} = n! / [r!(n-r)!]$
No	Si	Combinación c/repetición	$\binom{n+r-1}{r}$

(3)

Notación sigma

De todas las operaciones existentes, *la suma es la operación lineal más elemental* y sobre la cual se pueden desarrollar cualesquiera otras (inclusive el producto). De manera ordinaria la *suma finita* se representa por $S_n = a_0 + a_1 + \dots + a_n$; mientras que la *suma infinita* se expresa como $S = S_\infty = a_0 + a_1 + a_2 + \dots$. Existe una manera reducida conocida como *notación sigma* y se representa por:

$$S_n = \sum_{k=0}^n a_k = a_0 + a_1 + \dots + a_n, \quad S = S_\infty = \sum_{k=0}^{\infty} a_k = a_0 + a_1 + a_n + \dots \quad (4)$$

Las *propiedades fundamentales de la notación sigma* son:

$$\begin{aligned} (1) \quad & \sum_{k=0}^n (a_k \pm a_k) = \sum_{k=0}^n a_k \pm \sum_{k=0}^n b_k \leftarrow \text{Linealidad} \\ (2) \quad & \sum_{k=0}^n c a_k = c \sum_{k=0}^n a_k \\ (3) \quad & \sum_{k=0}^n c = n \cdot c \leftarrow \text{El producto como suma} \end{aligned} \quad (5)$$

Con frecuencia es necesario realizar un *corrimiento de índices* para aplicar la propiedad (1), como se muestra a continuación:

$$\begin{aligned} S &= S_1 + S_2 = \sum_{k=2}^{\infty} 6k C_k X^{k-1} + \sum_{k=0}^{\infty} 6C_k X^{k+1}, \quad \begin{aligned} &\text{si } i = k-1 \text{ en } S_1 \rightarrow k = i+1 \\ &\text{si } i = k+1 \text{ en } S_2 \rightarrow k = i-1 \end{aligned} \\ S &= \sum_{i=1}^{\infty} 2(i+1)C_{i+1} X^i + \sum_{i=-1}^{\infty} 6C_{(i-1)} X^i = \sum_{i=1}^{\infty} [2(i+1)C_{i+1} + 6C_{(i-1)}] X^i \end{aligned} \quad (6)$$

En general si aumenta en k el índice de la sumatoria, disminuyen en k los límites de la sumatoria:

$$\sum_{i=A}^N x(i) = \sum_{i=A-k}^{N-k} x(i+k) \quad (7)$$

Principio de identidad para sumas. Si $\sum_{k=0}^{\infty} a_k x^k = \sum_{k=0}^{\infty} b_k x^k$ entonces $a_k = b_k$ para toda $k \geq 0$; en particular si $\sum_{k=0}^{\infty} a_k x^k = 0$ entonces $a_k = 0$.

Principio de inducción matemática

Sea $S(n)$ una expresión matemática que involucre una o más ocurrencias de la variable $n \mid n \in \mathbb{Z}^+$, si $S(1)$ es cierta y $S(k)$ es cierta siendo $k \mid k \in \mathbb{Z}^+$; entonces $S(n)$ es cierta para toda $n \in \mathbb{Z}^+$. *El principio de inducción tiene su aplicación en la deducción de formulas y demostración de teoremas.*

 **Ejemplo 1.1.** *Demostrar por inducción matemática que:* $\sum_{i=1}^n i = n(n+1)/2$.

Solución. La *aplicación del principio de inducción* se realiza como sigue: (1) se muestra la certeza de $S(1)$; (2) se plantea una hipótesis para $S(k)$, y se demuestra que $S(k+1)$ es cierto, entonces (3) $S(n)$ queda demostrado:

$$S(n) = \sum_{i=1}^n i = 1 + 2 + 3 + \dots + n = n(n+1)/2$$

$$(1) \quad S(1) = \sum_{i=1}^1 i = 1 \Leftrightarrow 1(1+1)/2 = 2/2 = 1$$

(2) Sea $S(k)$ cierto por hipótesis, entonces $S(k+1)$ también debe ser cierto, basta con demostrar que:

$$S(k+1) = \frac{(k+1)(k+2)}{2}$$

$$\begin{aligned} S(k+1) &= \sum_{i=1}^{k+1} i = 1 + 2 + 3 + \dots + k + (k+1) = \sum_{i=1}^k i + (k+1) \\ &= \frac{k(k+1)}{2} + \frac{2(k+1)}{2} = \frac{k(k+1) + 2(k+1)}{2} = \frac{(k+2)(k+1)}{2} \end{aligned}$$

$$(3) \quad \therefore S(n) = \frac{n(n+1)}{2} \quad \clubsuit$$

 **Ejemplo 1.2.** Demostrar por inducción matemática que:

$$\sum_{i=1}^n i^2 = n(n+1)(2n+1)/6.$$

Solución. Se aplican los pasos (1) a (3) y se tiene que:

$$\sum_{k=1}^n k^2 = n(n+1)(2n+1)/6$$

$$S(n) = \sum_{k=1}^n k^2 = 1 + 4 + 9 + \dots + n^2 = n(n+1)(2n+1)/6$$

$$(1) \quad S(1) = \sum_{k=1}^1 k^2 = 1 \Leftrightarrow 1(1+1)(2+1)/6 = 6/6 = 1$$

(2) Sea $S(k)$ cierto por hipótesis, entonces $S(k+1)$ también debe ser cierto, basta con demostrar que:

$$S(k+1) = \frac{(k+1)(k+2)(2k+3)}{6}$$

$$\begin{aligned} S(k+1) &= \sum_{i=1}^{k+1} i^2 = 1 + 4 + 9 + \dots + k^2 + (k+1)^2 = \sum_{i=1}^k i^2 + (k+1)^2 \\ &= \frac{k(k+1)(2k+1)}{6} + \frac{6(k+1)^2}{6} = \frac{k(k+1)(2k+1) + 6(k+1)^2}{6} = \frac{[k(2k+1) + 6(k+1)](k+1)}{6} \\ &= \frac{[2k^2 + 7k + 6](k+1)}{6} = \frac{(k+2)(2k+3)(k+1)}{6} \end{aligned}$$

$$(3) \quad \therefore S(n) = \frac{n(n+1)(2n+1)}{6} \quad \square$$

Se han desarrollado *formulas para sumatorias de uso frecuente* cuya demostración se lleva a cabo mediante inducción matemática:

Tabla 1.1 Formulas de sumatorias frecuentes

$\sum_{k=1}^n x^k = \frac{x - x^{n+1}}{1 - x}$	$\sum_{k=1}^n k^2 = \frac{n(n+1)(2n+1)}{2}$	$\sum_{k=1}^{\infty} \frac{1}{k^6} = \frac{\pi^6}{945}$
$\sum_{k=a}^b k = \frac{(a+b)(b-a+1)}{2}$	$\sum_{k=1}^{\infty} r^k = \frac{r}{1-r}, \quad r < 1$	$\sum_{k=1}^{\infty} \frac{(-1)^{k+1}}{k} = \log(2)$
$\sum_{k=1}^n k = \frac{n(n+1)}{2}$	$\sum_{k=1}^{\infty} \frac{1}{k^2} = \frac{\pi^2}{6}$	$\sum_{k=1}^{\infty} \frac{(-1)^{k+1}}{k^2} = \frac{\pi^2}{12}$
$\sum_{k=a}^b k^2 = \frac{(b^2+b)(2b+1) - (a^2-a)(2a-1)}{6}$	$\sum_{k=1}^{\infty} \frac{1}{k^4} = \frac{\pi^4}{90}$	$\sum_{k=1}^{\infty} \frac{(-1)^{k+1}}{k^4} = \frac{7\pi^2}{720}$

Notación real y compleja

Número real. Un *intervalo abierto* I se representa por (a,b) , y un *intervalo cerrado* por $[a,b]$, siendo a y b ($a < b$) puntos de la *recta real*. El *conjunto* $\{x \mid x \in \mathfrak{R}\}$ de los *números reales* o *escalares* ($\mathfrak{R}$) que incluye los *enteros* $Z = \{\dots, -2, -1, 0, 1, 2, \dots\}$ y *naturales* $N = \{1, 2, 3, \dots\}$, se clasifica en *racionales* $Q = \{p/q \mid p, q \in Z, q \neq 0\}$ e *irracionales* ($\pi, \sqrt{2}, e$), según puedan o no expresarse como un cociente de 2 enteros.

Geométricamente el *conjunto de escalares* $\mathfrak{R}$ representa la *recta*, $\mathfrak{R}^2$ el *plano* y $\mathfrak{R}^3$ el *espacio*. Dados $a, b \in \mathfrak{R}$ se efectúan dos *operaciones algebraicas fundamentales* *suma* ($c = a + b$) y *producto* ($d = a * b$) y cumplen con las *propiedades*: *cerradura* ($c, d \in \mathfrak{R}$), *conmutativa* ($a * b = b * a$), *asociativa* ($a + b + c = a + (b + c)$), *identidad* ($a + 0 = a$, $a * 1 = a$), *aditiva inversa* ($a - a = 0$), *multiplicativa inversa* ($a * (1/a) = 1 \mid a \neq 0$), *distributiva* ($a * (b + c) = a * b + a * c$).

Límites. Un *valor límite* es el valor al que se *acercas infinitesimalmente* una expresión cuando una de sus *variables* tiende a un *valor constante*:

$$\begin{aligned} \lim_{x \rightarrow 0} c/x = \infty, \quad \lim_{x \rightarrow \infty} cx = \infty, \quad \lim_{x \rightarrow \infty} x/c = \infty, \quad \lim_{x \rightarrow \infty} c/x = 0 \\ \lim_{x \rightarrow 0} c^x = 1, \quad \lim_{x \rightarrow -\infty} c^x = 0, \quad \lim_{x \rightarrow +\infty} c^x = \infty, \quad \lim_{x \rightarrow +\infty, |c| < 1} c^x = 0, \quad \lim_{x \rightarrow 0} \frac{\sin x}{x} = 1 \end{aligned} \quad (8)$$

Los siguientes son los *teoremas fundamentales de los límites*:

$$\begin{aligned} \text{Si } \lim_{x \rightarrow a} u(x) = A, \quad \lim_{x \rightarrow a} v(x) = B, \quad \lim_{x \rightarrow a} w(x) = C \text{ entonces :} \\ (1) \quad \lim_{x \rightarrow a} (u \pm v \pm w) = A \pm B \pm C \\ (2) \quad \lim_{x \rightarrow a} (uvw) = ABC \\ (3) \quad \lim_{x \rightarrow a} \left(\frac{u}{v} \right) = \frac{A}{B} \end{aligned} \quad (9)$$

Si $f(x)$ no esta definida para $x=a$; pero $\lim_{x \rightarrow a} f(x) = B$, entonces $f(x)$ será continua para $x=a$, si se toma como valor de $f(x)$ para $x=a$ el valor B:

$$\lim_{x \rightarrow 2} \frac{x^2 - 4}{x - 2} = \frac{(x-2)(x+2)}{x-2} = x+2 \therefore \lim_{x \rightarrow 2} \frac{x^2 - 4}{x - 2} = 2+2 = 4 \quad (10)$$

Algebra. Cualquier *binomio* elevado *potencia entera* n , puede desarrollarse mediante la regla del *binomio de Newton*, donde los coeficientes se obtienen por *combinación*, en símbolos:

$$(a+b)^n = \sum_{k=0}^n \binom{n}{k} a^{n-k} b^k = a^n + na^{n-1}b + \frac{n(n-1)}{2} na^{n-2}b^2 + \dots + nab^{n-1} + b^n$$

$$\binom{n}{k} = \frac{n!}{k!(n-k)!} \leftarrow \text{combinación de } n \text{ objetos tomados } k \text{ a la vez} \quad (11)$$

Fórmulas de *productos notables* y *factorización*:

$$(a \pm b)^2 = a^2 \pm 2ab + b^2 \quad (a \pm b)^3 = a^3 \pm 3a^2b + 3ab^2 + b^3 \quad (a+b)(a-b) = a^2 - b^2$$

$$a^2 - b^2 = (a+b)(a-b) \quad a^3 - b^3 = (a-b)(a^2 + ab + b^2) \quad a^3 + b^3 = (a+b)(a^2 - ab + b^2) \quad (12)$$

Propiedades de *exponentes* y *radicales*:

$$b^n b^m = b^{n+m} \quad (b^n)^m = b^{nm} \quad (ab)^n = a^n b^n \quad \left(\frac{a}{b}\right)^n = \frac{a^n}{b^n} \quad \frac{b^n}{b^m} = b^{n-m}$$

$$\sqrt[n]{a} = a^{\frac{1}{n}} \quad \sqrt[n]{ab} = \sqrt[n]{a} * \sqrt[n]{b} \quad \sqrt[n]{a/b} = \sqrt[n]{a} / \sqrt[n]{b} \quad (13)$$

Propiedades de los *logaritmos*:

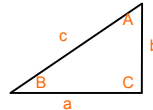
$$\log_b b = 0 \quad \log_b 1 = 0 \quad \log_b rs = \log_b r + \log_b s$$

$$\log_b \frac{r}{s} = \log_b r - \log_b s \quad \log_b r^p = p \log_b r \quad (14)$$

Trigonometría. Una *circunferencia* es una curva cerrada cuyos puntos equidistan a uno fijo llamado *centro* C , su *longitud* es $2\pi r$, siendo r su *radio*; la relación entre la *longitud del arco* s , el *ángulo subtendido* θ , y el *radio* r , esta dado por $s=r\theta$, donde θ se expresa en *radianes* ($\pi \text{ rad}=180^\circ$). Un *triángulo* (equilátero, isósceles o escaleno) de *lados* a, b, c y *ángulos* A, B, C de *área* S , obedece a la *ley de senos* y *cosenos*; si es un *triángulo rectángulo* obedece al *teorema de Pitágoras* $c^2 = a^2 + b^2$ (el cual se deduce de la ley de cosenos ya que $\cos(90^\circ)=0$:

$$S = \sqrt{t(t-a)(t-b)(t-c)} \leftarrow t = (a+b+c)/2$$

$$\frac{a}{\sin A} = \frac{b}{\sin B} = \frac{c}{\sin C}, \quad c^2 = a^2 + b^2 - 2ab \cos C \quad (15)$$



Para un *triángulo rectángulo* se definen las siguientes *razones de sus lados*:

$$\sin B = b/c, \quad \cos B = a/c, \quad \tan B = \sin B / \cos B = b/a, \quad \cot B = \cos B / \sin B = a/b$$

$$\sec B = 1/\cos B = c/a, \quad \csc B = 1/\sin B = c/b$$

$$\sin^2 B + \cos^2 B = 1, \quad 1 + \tan^2 B = \sec^2 B, \quad 1 + \cot^2 B = \csc^2 B, \quad \lim_{B \rightarrow \pm\pi/2} \tan B = \infty \quad (16)$$

Una *línea recta* es una figura que solo tiene extensión, si se sitúa el plano cartesiano XY, y forma un ángulo θ con el eje X y pasa a través de dos puntos (x_1, y_1) , (x_2, y_2) , se define su *pendiente* como $m = \tan \theta = (y_2 - y_1) / (x_2 - x_1)$, se denomina *abscisa* a la *coordenada x*, y *ordenada* a la *coordenada y*; las siguientes son ecuaciones de la recta:

$$\begin{array}{l} \text{Ecuación : } y - y_1 = m(x - x_1), \quad y = mx + y_0, \quad \frac{y - y_1}{x - x_1} = \frac{y_2 - y_1}{x_2 - x_1}, \quad \frac{x}{x_0} + \frac{y}{y_0} = 1 \\ \text{Datos : } \quad m, (x_1, y_1), \quad m, (0, y_0), \quad (x_1, y_1), (x_2, y_2), \quad (x_0, 0), (0, y_0) \end{array} \quad (17)$$

Función. Una *función, transformación o aplicación escalar de una variable* $y=f(x)$ es una *regla de correspondencia* que asigna a cada $x \in X$ un elemento específico $y \in Y$, el conjunto X se llama *dominio*, y al conjunto Y *contradominio*. El conjunto $\{f(x) | x \in X\}$ se denomina *rango* de f , y es un *subconjunto del contradominio*, en símbolos se escribe como: $f: X \rightarrow Y$. Una *función escalar* f , de 3 variables con dominio U , y valores escalares $f=f(x,y,z)$ se denota por: $f: U \subset \mathbb{R}^3 \rightarrow \mathbb{R}$, y asigna un escalar a cada punto del espacio (x,y,z) . Una *función vectorial* f , de 3 variables con dominio U , y valores vectoriales $\mathbf{C}=\mathbf{C}(x,y,z)$ se denota por $f: U \subset \mathbb{R}^3 \rightarrow \mathbb{R}^3$, y asigna un vector a cada punto del espacio (x,y,z) .

En la notación $y=f(x)$, x se denomina *variable dependiente* e y *variable independiente*. Una *función* es *inyectiva* si a elementos distintos del rango le corresponden elementos distintos del dominio $x_1 \neq x_2 \rightarrow f(x_1) \neq f(x_2)$, es *surjectiva* si para cada y del rango existe una x del dominio tal que $y=f(x)$. Si la función es inyectiva y surjectiva entonces es *biyectiva*, *biunívoca* o *función uno a uno*. Una *función par* cumple que $f(-x)=f(x)$ (ejemplos: $f(x)=x^2$, $f(x)=\cos(x)$); de lo contrario es *función impar* $f(-x)=-f(x)$. En general, las funciones pueden ser: *algebraicas* (*polinomial, racional*) o *transcendentales* (*exponencial, logarítmica, trigonométrica*).

$$\begin{array}{l} \text{algebraica} \left\{ \begin{array}{l} f(x) = a_n x^n + a_{n-1} x^{n-1} + \dots + a_1 x + a_0, \quad f(x) = \frac{g(x)}{h(x)} \Leftrightarrow g \text{ y } h \text{ son polinomios} \\ \text{transcendental} \left\{ \begin{array}{l} f(x) = b^x, \quad f(x) = \log_b x, \quad f(x) = \sin x \quad f(x) = \cos x \quad f(x) = \tan x \end{array} \right. \end{array} \right. \quad (18)$$

El *recíproco* de una función $f(x)$ es $1/f(x)$. La *inversa* de una función *biyectiva* $y=f(x)$ se denota por $x=f^{-1}(y)$, la *función y su inversa* satisfacen que $x = f(f^{-1}(x)) = f^{-1}(f(x))$. La inversa se obtiene por dos pasos: (1) despejar x de $y=f(x)$, (2) red denominar x por y , la función resultante es $f^{-1}(x)$. La función *logaritmo* y *exponencial* son inversas $y = b^x \rightarrow x = \log_b y$.

Una función $f(x)$ es *continua* (sin rupturas ni saltos) en $c \in [a, b]$ si cumple las siguientes tres condiciones: (1) $f(c)$ existe; (2) $\lim_{x \rightarrow c} f(x)$ existe; (3) $\lim_{x \rightarrow c} f(x) = f(c)$. Una función $f(t)$ es *periódica* y de *periodo* p , si se satisface que $f(t) = f(t + np) \rightarrow p \in \mathbb{R} \wedge n \in \mathbb{Z}^+$, (las funciones seno y coseno son periódicas con $p=2\pi$).

Numero complejo. Un *número imaginario puro* tiene la forma jb , $b \in \mathbb{R}$, $j = i = \sqrt{-1}$, la constante i (introducida por Euler; en ingeniería eléctrica se prefiere usar j , por representar i la corriente eléctrica) tiene la propiedad de que $i^2 = -1$. Un *número complejo* $A = a + jb$, se compone de dos partes una *real* $\text{Re}[A] = a$ y otra *imaginaria* $\text{Im}[A] = b$ ($a, b \in \mathbb{R}$), separadas mediante el *símbolo imaginario*; el conjunto de los números complejos se representa por: $C = \{a + jb \mid a, b \in \mathbb{R}, j^2 = -1\}$. Puede ser expresado en cualquiera de las siguientes formas (*rectangular*, *exponencial*, *polar* (abreviación de la forma exponencial), *trigonométrica* y en el *plano complejo*):

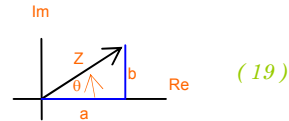
$$\hat{A} = a + jb = |\hat{A}|e^{j\theta} = |\hat{A}| \angle \theta = |\hat{A}| \cos(\theta) + j|\hat{A}| \sin(\theta)$$

$$\theta = \arctan(b/a) \leftarrow \text{argumento o ángulo de } \hat{A}$$

$$|\hat{A}| = \sqrt{a^2 + b^2} \leftarrow \text{módulo, valor absoluto o magnitud de } \hat{A}$$

$$\hat{A}^* = a - jb \leftarrow \text{conjugado de } \hat{A}$$

$$|\hat{A}^*| = |\hat{A}|$$



El ángulo θ apropiado se elige basado en el cuadrante dado por (a, b) siendo el ángulo positivo (+) en sentido horario y negativo (-) en sentido antihorario; en ingeniería eléctrica se expresa en grados en física en radianes distinguiéndose uno del otro por el uso de $^\circ$ o múltiplos de π . La *notación rectangular* es idónea para la *suma y resta* (partes reales se suman con partes reales, partes imaginarias con partes imaginarias), mientras que la *polar* es idónea para la *multiplicación y división* (por cumplir con las leyes de los exponentes), sean **A** y **B** dos números complejos, entonces:

$$\hat{A} = a + jb = |\hat{A}| \angle \theta_A \quad \hat{B} = c + jd = |\hat{B}| \angle \theta_B$$

$$\hat{A} + \hat{B} = (a + c) + j(b + d) \quad \hat{A} - \hat{B} = (a - c) + j(b - d) \quad (20)$$

$$\hat{A} * \hat{B} = |\hat{A}| |\hat{B}| \angle (\theta_A + \theta_B) \quad \frac{\hat{A}}{\hat{B}} = \frac{|\hat{A}|}{|\hat{B}|} \angle (\theta_A - \theta_B)$$

Las siguientes identidades son de gran importancia:

$$1 = e^{j0} = 1 \angle 0$$

$$\lim_{\theta \rightarrow 0} \arctan(1/\theta) = \lim_{\theta \rightarrow 0} \arctan(+\infty) = \pi/2 \quad \text{y} \quad \lim_{\theta \rightarrow 0} \arctan(-1/\theta) = \lim_{\theta \rightarrow 0} \arctan(-\infty) = -\pi/2 \quad (21)$$

$$j = 0 + j = 1 \angle \arctan(1/0) = 1 \angle \arctan(+\infty) = 1 \angle (\pi/2) = 1 \angle 90^\circ$$

$$-j = 0 - j = 1 \angle \arctan(-1/0) = 1 \angle \arctan(-\infty) = 1 \angle (-\pi/2) = 1 \angle -90^\circ$$

$$j^2 = \sqrt{-1}^2 = -1$$

$$\frac{1}{j} = \frac{1 \angle 0}{1 \angle (\pi/2)} = e^{0-\pi/2} = e^{-\pi/2} = 1 \angle (-\pi/2) = -j \quad (22)$$

$$(j^1, j^2, j^3, j^4 \dots j^9, j^{10}, j^{11}, j^{12} \dots) = (j, -1, -j, 1 \dots j, -1, -j, 1 \dots)$$

De lo anterior se concluyen la siguiente regla para el *producto unitario complejo*: “el producto de un número complejo por la constante imaginaria (j) ocasiona un incremento de 90° ($\pi/2$ rad) en el ángulo del número complejo”, es decir:

$$\begin{aligned} \dot{A} &= |\dot{A}| \angle \phi \\ j\dot{A} &= |\dot{A}| \angle (\phi + \pi/2) \\ -j\dot{A} &= |\dot{A}| \angle (\phi - \pi/2) \end{aligned} \quad (23)$$

Sea $r \in \Re$, entonces r puede representarse como un número complejo en forma polar como: $r = r \angle 0$. Sea $\mathbf{Z} = a + jb$, $\mathbf{Z} \in \mathbb{C}$, su *inverso* $\mathbf{Y} \in \mathbb{C}$, es igual a su *recíproco* y está dado por: $\hat{Y} = 1/\hat{Z} = (1/\sqrt{a^2 + b^2}) \angle (-\tan^{-1} b/a)$.

La siguiente identidad se conoce como *teorema de Demoiivre* :

$$\begin{aligned} (Ze^{j\theta})^n &= Z^n (e^{j\theta})^n = Z^n (\cos \theta + j \sin \theta)^n = Z^n (\cos(n\theta) + j \sin(n\theta)) \\ z = Ze^{j\theta} &\longrightarrow z^n = Z^n (\cos(n\theta) + j \sin(n\theta)) \longleftrightarrow Z = |z| = \sqrt{a^2 + b^2} \end{aligned} \quad (24)$$

Función polinomial. Es una función algebraica cuyos términos son monomios; un *monomio* es una agrupación de variables cuyos exponentes son enteros (xy^3 , xy^2z). Sea $P(x)=0$ un *polinomio* en x , de *grado* n , el *teorema fundamental del algebra* (demostrado por Gauss) afirma que un *polinomio* de *grado* $n \geq 1$ tiene exactamente n *raíces* no necesariamente distintas, que pueden ser *reales* o *complejas*, obtenidas por *factorización*, *división sintética* o mediante *aproximaciones numéricas*. Para $n=2$ se emplea la *fórmula cuadrática*.

$$ax^2 + bx + c = 0 \rightarrow x_{1,2} = \frac{-b}{2a} \pm \frac{\sqrt{b^2 - 4ac}}{2a} \quad (25)$$

Un *polinomio* $P(x) = a_0x^n + a_1x^{n-1} + \dots + a_{n-2}x + a_{n-1}$, de *grado* n , para el cual $n \in \mathbb{Z}^+$ puede representarse como un *producto de binomios* de la forma $(x-r_1)^{m_1} \dots (x-r_n)^{m_n} = 0$, $m_i \in \mathbb{Z}^+$, los valores r_i se denominan *raíces del polinomio* y son números reales o complejos. Sea $k \in \mathbb{Z}$, el *teorema del residuo* afirma que si $P(x)$ es divisible entre $(x-k)$, el residuo es $P(k)$; el *teorema del factor* afirma que si $P(k)=0$, entonces $(x-k)$ es un factor de $P(x)$. El *método de división sintética*, permite obtener sistemáticamente el *cociente* $C(x)$ y el *residuo* $R(x)$ de un polinomio $P(x)$ de *grado* n , dividido por un *factor* de la forma $(x-k)$, siempre que $k \in \mathbb{Z}$; $P(x)/(x-k)$:

$$\begin{array}{cccccccc|c} a_0 & a_1 & a_2 & \cdots & a_{n-2} & a_{n-1} & & \\ \downarrow & kb_0 & kb_1 & \cdots & kb_{n-3} & kb_{n-2} & & \\ b_0 = a_0 & b_1 = a_1 + k \cdot b_0 & b_2 = a_2 + k \cdot b_1 & \cdots & b_{n-2} = a_{n-2} + k \cdot b_{n-3} & b_{n-1} = a_{n-1} + k \cdot b_{n-2} & & +k \end{array} \quad (26)$$

$$C(x) = b_0x^{n-1} + b_1x^{n-2} + b_2x^{n-3} + \dots + b_{n-2} \quad R(x) = b_{n-1}/(x-k)$$

Un *método numérico eficiente para el cálculo de raíces* es el de *Newton*, que requiere un *valor inicial para x* (método abierto) y cuya *ecuación de recurrencia* esta dada por:

$$x_{i+1} = x_i - \frac{f(x_i)}{f'(x_i)} \quad |f(x_i)| < \varepsilon \quad (27)$$

Función racional. Es aquella que tiene la forma $F(s)=P(s)/Q(s)$. Siendo P y Q *polinomios* de $s | s \in \mathfrak{R} \vee \mathbb{C}$. F es una *fracción propia* si el grado de Q es mayor que el grado de P , es *fracción impropia* en caso contrario. Las *fracciones racionales* a menudo se representan por conveniencia como una suma finita de *fracciones parciales* $F_i(x)$, (como ocurre en algunos casos de *integración* y en la determinación de la *transformada inversa de Laplace*) esto es:

$$F(s) = \frac{P(s)}{Q(s)} = F_1(s) + F_2(s) + \dots + F_n(s) \quad \leftrightarrow \quad \text{Grado}(Q) > \text{Grado}(P) \quad (28)$$

La expresión anterior puede escribirse como:

$$\frac{P(s)}{Q(s)} = \frac{P(s)}{(s-\alpha)^n R(s)} = \left(\frac{a_0}{(s-\alpha)^n} + \frac{a_1}{(s-\alpha)^{n-1}} + \dots + \frac{a_{n-1}}{(s-\alpha)} \right) + (\dots) \quad (29)$$

Donde a_k son *coeficientes*, $a_k \in \mathfrak{R} \vee \mathbb{C}$ y las elipses representan la expansión en fracciones parciales de las raíces de $R(s)$; *las raíces del denominador se denominan*

polos ($s-\alpha$), las raíces del numerador se denominan ceros. Los coeficientes están dado por:

$$a_k = \frac{1}{k!} \frac{d^k}{ds^k} \left(\frac{P(s)}{R(s)} \right) \Big|_{s=\alpha} \quad (30)$$

En el empleo de esta ecuación considérese que es necesario derivar solo en el caso de existir *raíces repetidas* ya que $D^0 F(s) = F(s)$. Observe que en la evaluación α cambia su signo.

 **Ejemplo 1.3.** Descomponer en fracciones parciales la siguiente *función racional* que involucra *polos distintos*: $F(s) = \frac{s+3}{(s+1)(s+2)}$

Solución. Aplicando la *descomposición en fracciones parciales* se tiene:

$$\begin{aligned} \frac{P(s)}{Q(s)} &= \frac{s+3}{(s+1)(s+2)} = \frac{a_0}{s+1} + \frac{b_0}{s+2} \\ a_0 &= \frac{1}{0!} \frac{d^0}{ds^0} \left(\frac{s+3}{s+2} \right) \Big|_{s=-1} = \left(\frac{s+3}{s+2} \right) \Big|_{s=-1} = \frac{-1+3}{-1+2} = 2 \\ b_0 &= \frac{1}{0!} \frac{d^0}{ds^0} \left(\frac{s+3}{s+1} \right) \Big|_{s=-2} = \left(\frac{s+3}{s+1} \right) \Big|_{s=-2} = \frac{-2+3}{-2+1} = -1 \\ \therefore F(s) &= \frac{2}{s+1} + \frac{-1}{s+2} \quad \clubsuit \end{aligned}$$

 **Ejemplo 1.4.** Descomponer en fracciones parciales la siguiente *función racional* que involucra *polos complejos conjugados*: $F(s) = \frac{2s+12}{s^2+2s+5}$

Solución. Aplicando la *descomposición en fracciones parciales* se tiene:

$$\begin{aligned} \frac{P(s)}{Q(s)} &= \frac{2s+12}{s^2+2s+5} = \frac{2s+12}{[s+(1+j2)][s+(1-j2)]} = \frac{a_0}{[s+(1+j2)]} + \frac{a_1}{[s+(1-j2)]} \\ a_0 &= \frac{1}{0!} \frac{d^0}{ds^0} \left(\frac{2s+12}{s+(1-j2)} \right) \Big|_{s=-1-j2} = \left(\frac{2s+12}{s+(1-j2)} \right) \Big|_{s=-1-j2} = \frac{-2-j4+12}{-1-j2+1-j2} = 1+j2.5 \\ b_0 &= \frac{1}{0!} \frac{d^0}{ds^0} \left(\frac{2s+12}{s+(1+j2)} \right) \Big|_{s=-1+j2} = \left(\frac{2s+12}{s+(1+j2)} \right) \Big|_{s=-1+j2} = \frac{-2+j4+12}{-1+j2+1+j2} = 1-j2.5 \\ \therefore F(s) &= \frac{1+j2.5}{[s+(1+j2)]} + \frac{1-j2.5}{[s+(1-j2)]} = \frac{2.693\angle 68.198}{[s+(2.236\angle 63.43)]} + \frac{2.693\angle -68.198}{[s+(2.236\angle -63.43)]} \quad \clubsuit \end{aligned}$$

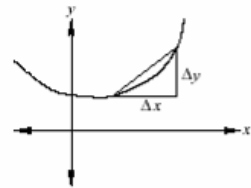
 **Ejemplo 1.5.** Descomponer en fracciones parciales la siguiente *función racional* que involucra *polos múltiples*: $F(s) = \frac{s^2 + 2s + 3}{(s+1)^3}$

Solución. Aplicando la *descomposición en fracciones parciales* se tiene:

$$\begin{aligned} \frac{P(s)}{Q(s)} &= \frac{s^2 + 2s + 3}{(s+1)^3} = \frac{a_0}{(s+1)^3} + \frac{a_1}{(s+1)^2} + \frac{a_2}{(s+1)} \\ a_0 &= \frac{1}{0!} \frac{d^0}{ds^0} (s^2 + 2s + 3) \Big|_{s=-1} = (s^2 + 2s + 3) \Big|_{s=-1} = 1 - 2 + 3 = 2 \\ a_1 &= \frac{1}{1!} \frac{d^1}{ds^1} (s^2 + 2s + 3) \Big|_{s=-1} = (2s + 2) \Big|_{s=-1} = -2 + 2 = 0 \\ a_2 &= \frac{1}{2!} \frac{d^2}{ds^2} (s^2 + 2s + 3) \Big|_{s=-1} = \frac{1}{2} (2) \Big|_{s=-1} = 1 \\ \therefore F(s) &= \frac{2}{(s+1)^3} + \frac{0}{(s+1)^2} + \frac{1}{(s+1)} = \frac{2}{(s+1)^3} + \frac{1}{(s+1)} \quad \clubsuit \end{aligned}$$

Notación de razones de cambio y sumas infinitesimales

Derivada. La *derivada* de una función $y=f(x)$ es la *razón de cambio* de la variable dependiente (y) cuando la variable independiente (x) cambia infinitesimalmente. Sea $P(a, f(a))$ un punto a lo largo de la *gráfica*, *curva*, *imagen* o *traza* de $y=f(x)$, la *derivada* es la *pendiente* de la *línea recta tangente* a la *curva* en tal punto, se simboliza y define por:



$$y' = \dot{y} = y^{(1)} = Dy = f'(x) = \frac{dy}{dx} = \lim_{\Delta x \rightarrow 0} \frac{f(x + \Delta x) - f(x)}{\Delta x} = \lim_{\Delta x \rightarrow 0} \frac{\Delta y}{\Delta x} \quad (31)$$

El *teorema de Rolle* afirma que si una función $f(x)$ es continua en $[a, b]$ y diferenciable en (a, b) y $f(a)=f(b)$, entonces existe un punto c en (a, b) tal que $f'(c)=0$. Una consecuencia es el *teorema del valor medio para la derivada* y establece que:

$$f'(c) = \frac{f(b) - f(a)}{b - a} \quad (32)$$

Sea $f(x)$ una función continua en (a,b) ; $f(x)$ es *creciente* donde $f'(x)>0$, es *decreciente* donde $f'(x)<0$ y *estacionaria* donde $f'(x)=0$ (si es $f(x)$ creciente o decreciente en $[a,b]$ es *monótona* en $[a,b]$). Los *puntos críticos* (x_c) de $f(x)$ son aquellos *valores de x para los cuales la pendiente es igual a cero* $f'(x_c)=0$, es un *máximo* si $f'(x_c)$ cambia de $+$ a $-$ o bien si $f''(x_c)<0$; es un *mínimo* si $f'(x_c)$ cambia de $-$ a $+$ o bien si $f''(x_c)>0$. Una función $f(x)$ puede ser aproximada en $x=a$ por un *polinomio de grado n* , mediante la *serie de Taylor*:

$$f(x) = f(a) + (x-a)f'(a) + \frac{(x-a)^2}{2!}f''(a) + \dots + \frac{(x-a)^n}{n!}f^{(n)}(a) + \dots = \sum_{k=0}^n \frac{(x-a)^k}{k!}f^{(k)}(a) \quad (33)$$

En $a=0$ la función se aproxima a un polinomio cuyos coeficientes están dados por la siguiente, conocida como *serie de Maclaurin*:

$$f(x) = f(0) + f'(0)x + f''(0)\frac{x^2}{2!} + \dots + f^{(n)}(0)\frac{x^n}{n!} + \dots = \sum_{k=0}^n f^{(k)}(0)\frac{x^k}{k!} \quad (34)$$

Sea $f(x)$ y $g(x)$ funciones diferenciables y a, b constantes entonces se tienen las siguientes *propiedades de diferenciación*:

$$\begin{aligned} u &= u(x), & v &= v(x) & a, b &= \text{constantes} \\ \frac{da}{dx} &= 0, & \frac{dx}{dx} &= 1, & \frac{d(au)}{dx} &= a \frac{du}{dx}, \\ \frac{d}{dx}(au + bv) &= a \frac{du}{dx} + b \frac{dv}{dx} \leftarrow \text{Linealidad} \\ \frac{d}{dx}(uv) &= \frac{du}{dx}v + u \frac{dv}{dx} \leftarrow \text{Regla del producto} \\ \frac{d}{dx}\left(\frac{u}{v}\right) &= \frac{v \frac{du}{dx} - u \frac{dv}{dx}}{v^2} \leftarrow \text{Regla del cociente} \\ \frac{d}{dx}(u^a) &= au^{a-1} \frac{du}{dx} \leftarrow \text{Regla de la potencia} \\ \frac{d}{dx}u(f(x)) &= \frac{du}{dv} * \frac{dv}{dx} \leftarrow \text{Regla de la cadena} \end{aligned} \quad (35)$$

Aplicando las *propiedades de diferenciación* y la *definición de derivada* se han desarrollado *fórmulas de derivación* para las principales funciones, como las mostradas en la **tabla 1.2**:

Tabla 1.2 Derivadas de funciones elementales

$\frac{d}{dx} \ln u = \frac{1}{u} \frac{du}{dx}$	$\frac{d}{dx} \sin u = \cos u \frac{du}{dx}$	$\frac{d}{dx} \arcsin u = \frac{1}{\sqrt{1-u^2}} \frac{du}{dx}$
$\frac{d}{dx} \log u = \frac{\log e}{u} \frac{du}{dx}$	$\frac{d}{dx} \cos u = -\sin u \frac{du}{dx}$	$\frac{d}{dx} \arccos u = -\frac{1}{\sqrt{1-u^2}} \frac{du}{dx}$
$\frac{d}{dx} a^u = a^u \ln a \frac{du}{dx}$	$\frac{d}{dx} \tan u = \sec^2 u \frac{du}{dx}$	$\frac{d}{dx} \arctan u = \frac{1}{1+u^2} \frac{du}{dx}$
$\frac{d}{dx} e^u = e^u \frac{du}{dx}$	$\frac{d}{dx} \cot u = -\csc^2 u \frac{du}{dx}$	$\frac{d}{dx} \operatorname{arc cot} u = -\frac{1}{1+u^2} \frac{du}{dx}$
$\frac{d}{dx} u^v = v u^{v-1} \frac{du}{dx} + u^v \ln u \frac{dv}{dx}$	$\frac{d}{dx} \sec u = \sec u \tan u \frac{du}{dx}$	$\frac{d}{dx} \operatorname{arc sec} u = \frac{1}{u\sqrt{u^2-1}} \frac{du}{dx}$
	$\frac{d}{dx} \csc u = -\csc u \cot u \frac{du}{dx}$	$\frac{d}{dx} \operatorname{arc csc} u = -\frac{1}{u\sqrt{u^2-1}} \frac{du}{dx}$

Si $y=f(x)$ una *función continua y diferenciable* en $[a,b]$ entonces la *derivada de la función inversa* $x=f^{-1}(y)$ esta dada por, $1/f'(x)$, en símbolos:

si $y = f(x)$ tiene como inversa $x = f^{-1}(y)$ entonces :

$$D(f^{-1}) = \frac{dx}{dy} = \frac{1}{\frac{dy}{dx}} = \frac{1}{f'(x)} \quad (36)$$

Una *función explícita* tiene la forma $y=f(x)$, una *función implícita* tiene la forma $f(x,y)=0$, cuando no es posible expresar explícitamente una función se *deriva implícitamente* y luego se factoriza para $f'(x)$ si es posible.

Un *límite indeterminado* tiene la forma $0/0$ o bien ∞/∞ , la *regla de L'Hospital* (debida a J. Bernoulli) se emplea para el *cálculo de límites de formas indeterminadas* y establece que:

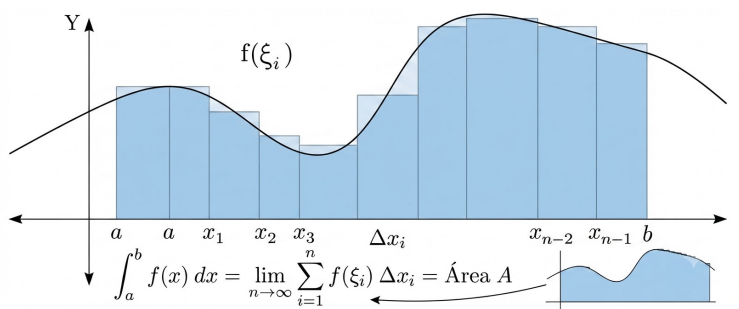
$$\lim_{x \rightarrow a} \frac{f(x)}{g(x)} = \lim_{x \rightarrow a} \frac{f'(x)}{g'(x)} \quad (37)$$

Integral. La operación inversa de la derivada es la antiderivada o integral y consiste en hallar la *primitiva* $f(x)$ dada su derivada $f'(x)$ o diferencial $f'(x)dx$ y se denota por

$$\int f'(x)dx.$$

La *integral indefinida* de $f'(x)dx$ es $\int f'(x)dx = f(x) + C$, C es la *constante de integración*, cuyo origen es cualquier constante anulada al ser derivada en la primitiva $D_x(C)=0$.

Sea $y=f(x)$ una función continua en el intervalo $[a,b]$, entonces la *integral definida* $A = \int_a^b f(x)dx = \int_a^b ydx$ representa el área A , bajo la curva engendrada por $f(x)$ y el eje de las x desde $x=a$ hasta $x=b$.



Geométricamente es la suma (conocida como suma de Riemann) de rectángulos de anchura infinitesimal dx (Δx) y altura $f(\xi_i)$ desde $x=a$ hasta $x=b$, $\xi_i \in x_i, x_{i+1}$, en símbolos:

$$\int_a^b f(x)dx = \lim_{\Delta x \rightarrow 0} \sum_{i=0}^{n-1} f(\xi_i) \Delta x \quad (38)$$

La *integral es una suma infinitesimal*, la *suma es una operación lineal*, por lo tanto la *integral es una operación lineal*, una integral definida es la suma de integrales definidas en *rangos o límites de integración*, cambiando el límite de integración cambia el *signo de la integral*, la integral de $x=a$ hasta $x=a$ es cero:

$$\begin{aligned}
 a, b, c, d &\in R \\
 \int_a^b (cf(x) + dg(x))dx &= c \int_a^b f(x)dx + d \int_a^b g(x)dx \\
 \int_a^b f(x)dx &= c \int_a^c f(x)dx + \int_c^b f(x)dx \Leftrightarrow a < c < b \\
 \int_a^b f(x)dx &= -\int_b^a f(x)dx, \quad \int_a^a f(x)dx = 0
 \end{aligned} \tag{39}$$

Sea $f(x)$ una función continua en $[a, b]$, consideremos que $x \in [a, b]$, llamemos $m = \min[f(x)]$ y llamemos $M = \max[f(x)]$ entonces, el *teorema del valor medio para la integral* establece que, existe un valor c en el intervalo $[a, b]$ para el cual:

$$\begin{aligned}
 a, b &\in R \\
 (b-a)m &\leq \int_a^b f(x)dx \leq (b-a)M \\
 c &\in [a, b] \\
 \int_a^b f(x)dx &= (b-a)f(c) \rightarrow f(c) = \frac{1}{(b-a)} \int_a^b f(x)dx
 \end{aligned} \tag{40}$$

Sea $F(x)$ la *antiderivada* de $f(x)$, es decir $F'(x) = f(x)$, entonces el *teorema fundamental del cálculo* (el cual reúne la noción de derivada con la noción de integral) establece que:

$$\int_a^b f(x)dx = F(b) - F(a) \Leftrightarrow F'(x) = f(x) \Leftrightarrow F(x) = \int f(x)dx \tag{41}$$

La *longitud del arco s*, engendrado por la gráfica de una función $y=f(x)$, de x_1 hasta x_2 , está dado por:

$$s = \int_{x_1}^{x_2} [1 + (y')^2]^{1/2} dx = \int_{y_1}^{y_2} [1 + (x')^2]^{1/2} dy \tag{42}$$

Cálculo de integrales. Para simplificar el *cálculo de integrales* se han desarrollado extensas tablas cuyas fórmulas se deducen del *teorema fundamental del cálculo* (las *integrales de funciones elementales* se muestran en la **tabla 1.3**). Cuando la función no tiene semejanza con una ninguna integral de la tabla se emplean *artificios de integración* entre cuales figuran los siguientes:

1. **Integración por partes:** Dada la función a evaluar $f(x)$ descompóngase de manera tal que sea posible aplicar la siguiente *fórmula de integración por partes*: $\int u dv = uv - \int v du$: 1) descomponer $f(x)$ en u y dv , 2) integrar dv , 3) derivar u , 4) evaluar $\int u dv = uv - \int v du$. Se emplea en los casos de diferenciales que contienen: (a) productos; (b) logaritmos; (c) funciones trigonométricas inversas.

 **Ejemplo 1.6.** Aplicar integración por partes para (a) integrar $\int x \cos x dx$, (b) demostrar que $\int \sec^3 z dz = \frac{1}{2} \sec z \tan z + \frac{1}{2} \ln|\sec z + \tan z| + C$

Solución. Se aplican los tres pasos antes numerados y se tiene que:

$$\begin{aligned}
 \text{(a)} \quad u &= x \rightarrow du = dx & dv &= \cos x dx \rightarrow v = \int \cos x dx = \sin x \\
 \int u dv &= uv - \int v du = x \sin x - \int \sin x dx = x \sin x + \cos x + C \\
 \text{(b)} \quad u &= \sec z \rightarrow du = \sec z \tan z dz & dv &= \sec^2 z dz \rightarrow v = \int \sec^2 z dz = \tan z \\
 \int u dv &= uv - \int v du = \sec z \tan z - \int \tan z \cdot \sec z \tan z dz \\
 \int \tan^2 z \cdot \sec z dz &= \int (\sec^2 z - 1) \cdot \sec z dz = \int \sec^3 z dz - \int \sec z dz \\
 &\rightarrow \int \sec^3 z dz = \sec z \tan z - \int \sec^3 z dz + \int \sec z dz \\
 &\rightarrow \int \sec^3 z dz = \frac{1}{2} \sec z \tan z + \frac{1}{2} \ln|\sec z + \tan z| + C \quad \clubsuit
 \end{aligned}$$

2. **Descomposición en fracciones parciales:** Para una *función racional propia* es posible la descomposición en fracciones parciales y la integración separada de cada fracción resultante.

3. **Sustitución conveniente:** Ciertas integrales requieren de emplear identidades o sustituciones que conduzcan a formas integrables elementales; la **tabla 1.4** muestra un conjunto de integrales de este tipo.

4. **Formas variadas:** A veces es necesario reexpresar la integral de manera que resulte una forma conocida o fácilmente integrable:

 **Ejemplo 1.7.** Integrar $\int \sec x dx$.

Solución. Multiplicando y dividiendo por $\sec x + \tan x$ resulta:

$$\int \sec x dx = \int \sec x \frac{\sec x + \tan x}{\sec x + \tan x} dx = \int \frac{\sec^2 x + \sec x \tan x dx}{\sec x + \tan x} = \ln|\sec x + \tan x| + C \quad \clubsuit$$

Método de integración por descomposición en fracciones parciales. Sea una función $H(x)$ que puede expresarse como una fracción de dos funciones cuyo numerador es la función $P(x)$ y cuyo denominador es la función $Q(x)$, entonces se puede aplicar este método para hallar la integral $\int [P(x)/Q(x)] dx$, si $H(x)=P(x)/Q(x)$ es una *fracción racional propia* ($\deg(P) < \deg(Q)$). Si la fracción es *impropia* ($\deg(P) > \deg(Q)$), primero se efectúa la *división polinomial* para obtener un polinomio, que se integra directamente, más una fracción propia que se integra vía este método.

Reglas de descomposición en fracciones parciales

- (1) Si el factor en $Q(x)$ es lineal distinto $(ax+b)$ la fracción parcial correspondiente es: $A/(ax+b)$
- (2) Si el factor en $Q(x)$ es lineal repetido $(ax+b)^k$ las fracciones parcial correspondientes son: $A_1/(ax+b) + A_2/(ax+b)^2 + \dots + A_k/(ax+b)^k$
- (3) Si el factor en $Q(x)$ es cuadrático irreducible distinto $(ax^2 + bx + c)$ la fracción parcial correspondiente es: $(Ax + B)/(ax^2 + bx + c)$. Un cuadrático es irreducible si su discriminante $\Delta=(b^2-4ac) < 0$
- (4) Si el factor en $Q(x)$ es cuadrático repetido $(ax^2 + bx + c)^k$ las fracciones parciales correspondientes son:
 $(A_1x+B_1)/(ax^2+bx+c) + (A_2x+B_2)/(ax^2+bx+c)^2 + \dots + (A_kx+B_k)/(ax^2+bx+c)^k$

Aplicación del método por fracciones parciales para integrar

- (1) **Factorizar** completamente $Q(x)$ en los factores descritos anteriormente.
- (2) **Plantear la suma de fracciones parciales** con coeficientes indeterminados (A, B, C, ...) utilizando estrictamente las reglas de la tabla anterior.
- (3) **Multiplicar toda la igualdad** por $Q(x)$ para eliminar denominadores.
Desarrollar el polinomio resultante, igualar los coeficientes de las potencias de x (o evaluar en raíces reales) y **resolver el sistema de ecuaciones lineales para obtener los valores numéricos de A, B, C, ...**
- (4) **Sustituir los coeficientes hallados** directamente dentro del signo de la integral.
- (5) **Integrar cada fracción resultante** aplicando las siguientes consideraciones:

(5a) **Si es lineal** $\frac{D}{ax+b} \rightarrow IL = \int \frac{D}{ax+b} dx = \left(\frac{D}{a}\right) \ln|ax+b|$

(5b) **Si es cuadrática** $\frac{Dx+E}{ax^2+bx+c}$ se resuelve en dos partes

Se calcula la derivada del denominador

$R' = 2ax+b$ y se reescribe el numerador como

$$(Dx+E) = \alpha(2ax+b) + \beta, \text{ donde } \alpha = D/(2a) \text{ y } \beta = E - (D \cdot b)/(2a)$$

Esto divide a la integral en dos partes:

$$Ic = Ic1 + Ic2 = \alpha \int \frac{2ax+b}{ax^2+bx+c} dx + \int \frac{1}{ax^2+bx+c} dx$$

(5b1) La solución de la primera integral es:

$$Ic1 = \alpha \int \frac{2ax+b}{ax^2+bx+c} dx = \alpha \ln|ax^2+bx+c|$$

(5b2) Para la segunda integral, se completa el cuadrado en el denominador:

$(x-h)^2 + k^2$, y se aplica:

$$Ic2 = \int \frac{1}{(x-h)^2 + k^2} dx = \frac{1}{k} \arctan\left(\frac{x-h}{k}\right) \xrightarrow{\text{donde}} k = \left(\frac{h = -b/2a}{\sqrt{4ac - b^2}}\right) / (2a)$$

- (6) **Finalmente** se suman todos los términos y se añade la constante C

 **Ejemplo 1.8.** Calcular la integral: $\int \frac{x^2 + 1}{(x - 2)(x^2 + 2x + 2)} dx$

Solución. Aplicamos el método de integración por fracciones parciales:

$$I = \int \frac{x^2 + 1}{(x - 2)(x^2 + 2x + 2)} dx = \int \frac{P(x)}{Q(x)} dx$$

(1) Factores de $Q(x)$: $(x - 2) \leftarrow$ lineal, $(x^2 + 2x + 2) \leftarrow$ cuadrático irreducible ya que $\Delta = -4$

$$(2) \quad \frac{P(x)}{Q(x)} = \frac{A}{(x - 2)} + \frac{Bx + C}{(x^2 + 2x + 2)} = \frac{x^2 + 1}{(x - 2)(x^2 + 2x + 2)}$$

$$(3) \xrightarrow{\text{multiplicando cada igualdad por el denominador } Q(x)} \frac{A(x^2 + 2x + 2)(x - 2)}{(x - 2)} + \frac{(Bx + C)(x^2 + 2x + 2)(x - 2)}{(x^2 + 2x + 2)} = (x^2 + 1)$$

$$\longrightarrow (x^2 + 1) = A(x^2 + 2x + 2) + (Bx + C)(x - 2) = (A + B)x^2 + (2A - 2B + C)x + (2A - 2C)$$

$$1x^2 \rightarrow A + B = 1$$

$$\xrightarrow{\text{igualando los coeficientes}} 0x \rightarrow 2A - 2B + C = 0 \xrightarrow{\text{resolviendo el sistema de ecuaciones}} A = \frac{1}{2} \quad B = \frac{1}{2} \quad C = 0$$

$$1 \rightarrow 2A - 2C = 1$$

$$(4) \quad I = \int \left[\frac{A}{(x - 2)} + \frac{Bx + C}{(x^2 + 2x + 2)} \right] dx = \int \frac{\frac{1}{2}}{(x - 2)} dx + \int \frac{\frac{1}{2}x + 0}{(x^2 + 2x + 2)} dx \rightarrow I = \frac{1}{2} \int \frac{1}{(x - 2)} dx + \frac{1}{2} \int \frac{x}{(x^2 + 2x + 2)} dx$$

(5) integrar cada FP resultante:

$$\text{caso 5(a), fracción lineal: } IL1 = \frac{1}{2} \int \frac{1}{(x - 2)} dx \rightarrow IL1 = \frac{1}{2} \ln|x - 2|$$

$$\text{caso 5(b), fracción cuadrática: } Ic = \frac{1}{2} \left[\int \frac{x}{(x^2 + 2x + 2)} dx \right] \longrightarrow 2Ic = \int \frac{x}{(x^2 + 2x + 2)} dx,$$

tenemos numerador: $D = 1, E = 0$; denominador: $a = 1, b = 2, c = 2$

$$R' = 2ax + b = 2x + 2, \quad \alpha = D/(2a) = 1/(2 \cdot 1) = 1/2 \quad \text{y} \quad \beta = E - (D \cdot b)/(2a) = 0 - (1 \cdot 2)/(2 \cdot 1) = -1$$

$$\text{la integral se compone de 2 partes: } Ic = Ic1 + Ic2 = \alpha \int \frac{2ax + b}{ax^2 + bx + c} dx + \beta \int \frac{1}{ax^2 + bx + c} dx$$

calcular $Ic1$ e $Ic2$:

$$5b1 \rightarrow Ic1 = \alpha \int \frac{2ax + b}{ax^2 + bx + c} dx = \alpha \ln|ax^2 + bx + c| = \frac{1}{2} \ln|1 \cdot x^2 + 2x + 2|$$

$$5b2 \rightarrow Ic2 = - \int \frac{1}{ax^2 + bx + c} dx = - \int \frac{1}{(x - h)^2 + k^2} dx = - \frac{1}{k} \arctan\left(\frac{x - h}{k}\right) \xrightarrow{\text{donde}} \frac{h = -b/2a = -2/(2 \cdot 1) = -1}{k = (\sqrt{4ac - b^2})/(2a) = 1}$$

$$\longrightarrow Ic2 = - \int \frac{1}{(x - (-1))^2 + 1^2} dx = - \frac{1}{1} \arctan\left(\frac{x - (-1)}{1}\right) = - \arctan(x + 1)$$

$$2Ic = Ic1 + Ic2 \longrightarrow Ic = \frac{1}{2} [Ic1 + Ic2] = \frac{1}{2} \left[\frac{1}{2} \ln|x^2 + 2x + 2| - \arctan(x + 1) \right] = \frac{1}{4} \ln|x^2 + 2x + 2| - \frac{1}{2} \arctan(x + 1)$$

(6) Finalmente se suman todos los términos y se agrega la constante C

$$I = IL1 + Ic = \frac{1}{2} \ln|x - 2| + \frac{1}{4} \ln|x^2 + 2x + 2| - \frac{1}{2} \arctan(x + 1) + C \quad \clubsuit$$

Tabla 1.3 Integrales de funciones elementales

$\int u dv = uv - \int v du$	$\int \sin u du = -\cos u + C$	$\int \frac{du}{\sqrt{a^2 - u^2}} = \arcsin \frac{u}{a} + C$
$\int u^n du = \frac{1}{n+1} u^{n+1} + C$	$\int \cos u du = \sin u + C$	$\int \frac{du}{a^2 + u^2} = \frac{1}{a} \arctan \frac{u}{a} + C$
$\int \frac{du}{u} = \ln u + C$	$\int \tan u du = \ln \sec u + C$	$\int \frac{du}{a^2 - u^2} = \frac{1}{2a} \ln \left \frac{a+u}{a-u} \right + C$
$\int e^u du = e^u + C$	$\int \cot u du = \ln \sin u + C$	$\int \frac{du}{u\sqrt{a^2 - u^2}} = \frac{1}{a} \operatorname{arccsc} \left \frac{u}{a} \right + C$
$\int b^u du = \frac{b^u}{\ln b} + C$	$\int \sec u du = \ln \sec u + \tan u + C$	$\int \arcsin u du = u \arcsin u + \sqrt{1-u^2} + C$
$\int u e^{au} du = \frac{e^{au}}{a^2} (au - 1) + C$	$\int \csc u du = \ln \csc u - \cot u + C$	$\int \arccos u du = u \arccos u - \sqrt{1-u^2} + C$
$\int u^n e^{au} du = \frac{u^n e^{au}}{a} - \frac{n}{a} \int u^{n-1} e^{au} du + C$	$\int \sec u \tan u du = \sec u + C$	$\int \arctan u du = u \arctan u - \ln \sqrt{1+u^2} + C$
$\int \ln u du = u \ln u - u + C$	$\int \csc u \cot u du = -\csc u + C$	$\int \operatorname{arccot} u du = u \operatorname{arccot} u + \ln \sqrt{1+u^2} + C$
$\int u^n \ln u du = u^{n+1} \left[\frac{\ln u}{n+1} - \frac{1}{(n+1)^2} \right] + C$	$\int \sec^2 u du = \tan u + C$	$\int \sec^{-1} u du = u \sec^{-1} u - \ln(u + \sqrt{u^2 - 1}) + C$
$\int e^{au} \ln u du = \frac{1}{a} e^{au} \ln u - \frac{1}{a} \int \frac{e^{au}}{u} du$	$\int \csc^2 u du = -\cot u + C$	$\int \csc^{-1} u du = u \csc^{-1} u + \ln(u + \sqrt{u^2 - 1}) + C$
$\int \frac{du}{u \ln u} = \ln(\ln u) + C$		

Tabla 1.4 Sustituciones convenientes para integración de funciones

Forma	condición	sustitución	integral resultante
$\int \sin^m u \cos^n u du$	$m \in \{Z^+ \% 2 \neq 0\}$	$\sin^2 u = 1 - \cos^2 u$	$\int [\cos^n u (1 - \cos^2 u)^{m-1}] \sin u du$
$\int \sin^m u \cos^n u du$	$n \in \{Z^+ \% 2 \neq 0\}$	$\cos^2 u = 1 - \sin^2 u$	$\int [\sin^m u (1 - \sin^2 u)^{n-1}] \cos u du$
$\int \sin^m u \cos^n u du$	$m, n \in \{Z^+ \% 2 = 0\}$	$\sin u \cos u = \frac{1}{2} \sin 2u$ $\sin^2 u = \frac{1}{2} - \frac{1}{2} \cos 2u$ $\cos^2 u = \frac{1}{2} + \frac{1}{2} \cos 2u$	varias formas elementales
$\int \sin mu \cos n u du$ $\int \sin mu \sin n u du$ $\int \cos mu \cos n u du$	$m, n \in \{Z^+\}$ $m \neq n$	$\sin mu \cos nu = \frac{1}{2} \sin(m+n)u$ $\quad + \frac{1}{2} \sin(m-n)u$ $I_1 = \frac{\cos(m+n)u}{2(m+n)}$ $I_2 = \frac{\cos(m-n)u}{2(m-n)}$	$\int \sin mu \cos n u du = -I_1 - I_2 + C$ $\int \sin mu \sin n u du = -I_1 + I_2 + C$ $\int \cos mu \cos n u du = I_1 + I_2 + C$
$\int \tan^n u du$	$n \in \{Z^+\}$	$\tan^n u = \tan^{n-2} u \tan^2 u$ $\tan^2 u = (\sec^2 u - 1)$	$\int \tan^{n-2} u (\sec^2 u - 1) du$
$\int \cot^n u du$	$n \in \{Z^+\}$	$\cot^n u = \cot^{n-2} u \cot^2 u$ $\cot^2 u = (\csc^2 u - 1)$	$\int \cot^{n-2} u (\csc^2 u - 1) du$
$\int \tan^m u \sec^n u du$	$n \in \{Z^+ \% 2 = 0\}$	$\sec^2 u = \tan^2 u + 1$	$\int [\tan^m u (\tan^2 u - 1)^{n-2}] \sec^2 u du$
$\int \tan^m u \sec^n u du$	$m \in \{Z^+ \% 2 \neq 0\}$	$\tan^m u = \tan u \tan^{m-1} u$ $\tan^2 u = \sec^2 u - 1$	$\int [\sec^{n-1} u (\sec^2 u - 1)^{m-3}] \sec u \tan u du$
$\int (\sqrt{a^2 - u^2})^k du$	$k \in \{Q \subseteq \Re\}$	$u = a \sin z \rightarrow du = a \cos z dz$ $\sqrt{a^2 - u^2} = a \cos z$	$\int (a \cos z)^k a \cos z dz$
$\int (\sqrt{a^2 + u^2})^k du$	$k \in \{Q \subseteq \Re\}$	$u = a \tan z \rightarrow du = a \sec^2 z dz$ $\sqrt{a^2 + u^2} = a \sec z$	$\int (a \sec z)^k a \sec^2 z dz$
$\int (\sqrt{u^2 - a^2})^k du$	$k \in \{Q \subseteq \Re\}$	$u = a \sec z$ $du = a \sec z \tan z dz$ $\sqrt{u^2 - a^2} = a \tan z$	$\int (a \tan z)^k a \sec z \tan z dz$

El operador modulo se designa por % y devuelve el residuo de una división entera: $1=5\%2$.

Notación integrodiferencial

Un *sistema físico* es un ente con características propias que se aísla para su estudio basado en las *leyes experimentales de la física*. Los principios del álgebra básica son suficientes para la modelación de *sistemas estáticos*. Sin embargo las máquinas y procesos naturales o artificiales son en alto grado dinámicos, tales *sistemas dinámicos* se modelan mediante *ecuaciones integrodiferenciales*.

Ecuaciones integrodiferenciales

Una *ecuación* es una igualdad que se cumple para ciertos valores de la variable o función que interviene en ella. Una *ecuación diferencial* (ED) involucra una *función*, sus *derivadas* y las *variables independientes*. Una *ecuación integrodiferencial* (EID) incluye además *integrales* de la función, de sus derivadas o de las variables independientes.

El *orden* (O) se determina por la derivada mayor, el *grado* (G) se determina por el exponente de la *derivada de mayor orden*. La ED puede ser *ordinaria* (EDO) o *parcial* (EDP) según incluya *derivadas ordinarias o parciales*. La ED es *lineal* (grado 1) si lo es en la variable dependiente y todas sus derivadas, es *homogénea* ($y'' + xy' + y = 0$) si no hay términos que sean funciones sólo de la variable independiente.

La *primitiva* o *solución* es la función que satisface la ED y puede ser *implícita* o *explícita*. La *solución general* de la ED de orden n tendrá n *constantes arbitrarias*, la *solución particular* es la determinación de los valores de las n constantes arbitrarias y se obtiene de $n+1$ *condiciones iniciales* de la función y sus n derivadas.

ED de primer orden y primer grado

La ED de grado y orden 1 tiene la forma $F(x, y, y') = 0$, donde $y = y(x)$ y su solución es $f(x, y, C) = 0 \Leftrightarrow f(x, y) = C$. Los métodos de solución dependen de su forma, los siguientes son los más comunes: *formas integrables* (*variables separables*, *ecuaciones exactas*, *ecuaciones homogéneas*), *ecuaciones lineales*.

Formas integrables

* Una ED de *variables separables*, de la forma $P(x) + Q(y)y' = 0$ se resuelve separando variables dependientes e independientes e integrando:

$$P(x) + Q(y)y' = 0 \rightarrow \int P(x)dx + \int Q(y)dy = C \quad (43)$$

* Una ED de la forma $P(x, y)dx + Q(x, y)dy = 0$, es *exacta* si $\partial P / \partial y = \partial Q / \partial x$. Siendo $P(x, y) = \partial f(x, y) / \partial x$ y $Q(x, y) = \partial f(x, y) / \partial y$. Entonces su primitiva es $f(x, y) = C$ y se obtiene sistemáticamente por *integración parcial*:

$$\begin{aligned} \frac{\partial f}{\partial x} dx &= P dx \longrightarrow f(x, y) = \int^x P dx + \phi(y) \\ \frac{\partial f}{\partial y} &= \frac{\partial}{\partial y} \int^x P dx + \frac{d\phi}{dy} = Q \longrightarrow \phi = \int \left\{ Q - \frac{\partial}{\partial y} \int^x P dx \right\} dy \\ \therefore f &= \int^x P dx + \int^y \left(Q - \frac{\partial}{\partial y} \int^x P dx \right) dy \end{aligned} \quad (44)$$

* Una ecuación de la forma $f(x, y)$ se dice que es *homogénea* y de *grado n* si al sustituir x por λx e y por λy se cumple que $f(\lambda x, \lambda y) = \lambda^n f(x, y)$. Una ED de la forma $P(x, y)dx + Q(x, y)dy = 0$ es *homogénea*, si tanto P como Q son homogéneas y del mismo grado. La transformación $y = vx \rightarrow dy = vdx + xdv$, reduce la ED homogénea a la forma:

$$P(x, v)dx + Q(x, v)dv = 0 \quad (45)$$

Y se resuelve por *separación de variables*. Después de integrar se sustituye v por y/x para recuperar las variables originales.

Ecuación lineal

Una ecuación lineal de grado 1 de la forma $y' + yP(x) = Q(x)$ se resuelve considerando el *factor integrante* $\mu(x) = e^{\int P(x)dx}$ y su primitiva es:

$$ye^{\int P(x)dx} = \int Q(x) e^{\int P(x)dx} dx + C \longrightarrow y = \left(e^{\int P(x)dx} \cdot \int \left[Q(x) \cdot e^{\int P(x)dx} dx \right] \right) + \left(C e^{\int P(x)dx} \right) = y_p + y_h \quad (46)$$

y_p se denomina *solución particular* o *respuesta en estado estable*, y_h *solución homogénea* o *respuesta transitoria* (proviene de la ED homogénea asociada $y' + yP(x) = 0$).

Una ecuación de la forma $y' + yP(x) = y^n Q(x)$ se conoce como *ecuación de Bernoulli* y se reduce a lineal mediante la transformación $v = y^{1-n}$:

$$\begin{aligned} v &= y^{1-n} = y^{-n+1} \\ \frac{dv}{dx} &= (-n+1)y^{-n} \frac{dy}{dx} \rightarrow y^{-n} \frac{dy}{dx} = \frac{1}{1-n} \frac{dv}{dx} \\ \frac{dy}{dx} + yP(x) &= y^n Q(x) \rightarrow y^{-n} \frac{dy}{dx} + y^{-n+1} P(x) = Q(x) \\ \frac{1}{1-n} \frac{dv}{dx} + vP(x) &= Q(x) \longrightarrow \frac{dv}{dx} + v\{(1-n)P(x)\} = (1-n)Q(x) \end{aligned} \quad (47)$$

ED de primer orden y grado superior

La EDO de *orden 1* y *grado n*, tiene la forma: $y'^n + P_1(x, y)y'^{n-1} + \dots + P_{n-1}(x, y)y' + P_n(x, y) = 0$.

Si expresamos $p = y'$, entonces una ED orden 1 y grado n puede escribirse como un polinomio de p:

$$p^n + P_1(x, y)p^{n-1} + \dots + P_{n-1}(x, y)p + P_n(x, y) = 0 \quad (48)$$

ED que se pueden resolver respecto a p. Considerar la ED como un polinomio de p, factorizar $(p-F_1)(p-F_2)\dots(p-F_n)=0$ y resolver cada ED resultante de primer grado $y'=F_i(x, y)$ cuya solución es $f_i(x, y, C)=0$, entonces la solución general se expresa como el producto de las soluciones particulares:

$$f_1(x, y, C) \cdot f_2(x, y, C) \dots f_n(x, y, C) = 0 \quad (49)$$

Otros casos especiales de poca aplicación práctica incluyen solución respecto x e y. La ED de Clairaut es $y = rx + f(r)$ y su solución es $y = Cx + f(C)$.

ED lineal de orden superior

La ED lineal de orden superior tiene la forma $P_0 y^n + P_1 y^{n-1} + \dots + P_{n-1} y' + P_n y = Q$, donde y^k es la *k-ésima derivada* (de grado 1), los *coeficientes* P_k son *funciones de la variable independiente* $P_k(x)$, o *constantes* y el *término independiente* Q una *función* $Q(x)$, una *constante* o *cero*. Una EDO lineal de coeficientes constantes de grado n se puede tratar como un *polinomio en la derivada de grado n* y resolver por varios

métodos, tal ecuación se denomina *invariante en el tiempo* y tiene gran aplicación en los sistemas de control y en general en la descripción física de gran variedad de fenómenos naturales.

La ED lineal de coeficientes constantes de orden n , se resuelve por varios métodos bien definidos: *coeficientes indeterminados, variación de parámetros, método del operador D y transformada de Laplace*. Los primeros dos son de ensayo y error, la transformada de Laplace es una formalización de los métodos de operador y es directo y es el método tradicional para el estudio de los sistemas de control. A continuación se muestra la ED lineal invariante de orden superior así como su *solución general* que consta de la suma de n *soluciones linealmente independientes*:

$$\begin{aligned} a_n y^{(n)} + a_{n-1} y^{(n-1)} + \dots + a_2 y^{(2)} + a_1 y^{(1)} + a_0 y^{(0)} &= 0 \\ y &= c_1 e^{r_1 x} + c_2 e^{r_2 x} + \dots + c_n e^{r_n x} \end{aligned} \quad (50)$$

A manera de ejemplo se expondrá la solución de la ED de orden 2:

$$a_2 y^{(2)} + a_1 y^{(1)} + a_0 y^{(0)} = 0 \quad (51)$$

La *ecuación característica* o *auxiliar* asociada se puede resolver por la fórmula cuadrática y se tienen 3 casos según sus raíces sean reales o complejas:

$$a_2 r^2 + a_1 r + a_0 = 0 \quad (52)$$

Caso I - **raíces reales distintas** $r_1 \neq r_2$:

$$y = c_1 e^{r_1 x} + c_2 e^{r_2 x} \quad (53)$$

Caso II - **raíces reales e iguales** $r_1 = r_2$:

$$y = c_1 e^{r_1 x} + c_2 x e^{r_2 x} \quad (54)$$

Caso III - **raíces complejas conjugadas** $r_1 = a + jb$, $r_2 = a - jb$:

$$\begin{aligned} y &= c_1 e^{(a+jb)x} + c_2 e^{(a-jb)x} \leftarrow \text{Forma } \cdot \text{ compleja} \\ y &= e^{ax} (C_1 \cos(bx) + C_2 \sin(bx)) \leftarrow \text{Forma } \cdot \text{ Real} \\ \text{donde } : C_1 &= c_1 + c_2 \quad y \quad C_2 = j(c_1 - c_2) \end{aligned} \quad (55)$$

Notación matricial y sistemas de ecuaciones lineales

Una **matriz bidimensional** es un arreglo rectangular $m \times n$ de celdas dispuestas en m filas y n columnas. Sea $\mathbf{A}$ una matriz $m \times n$, sea i el índice de fila y j el índice de columna, se denota por a_{ij} cualquier elemento de la matriz y la matriz por (a_{ij}) . La traspuesta de una matriz se forma cambiando las filas por las columnas y se denota por $\mathbf{A}^T$ o $\mathbf{A}'$, es decir: $(a_{ij})^T = (a_{ji})$.

Las **matrices de escalares** se clasifican en: *n-cuadrada*, de orden n , *n-matriz*, $n \times n$ o n^2 ; *identidad* $\mathbf{I}$ unos en su *diagonal principal* (DP) ceros en otra parte; *triangular superior* MTS o *triangular inferior* MTI según tenga ceros debajo o sobre la DP; *diagonal*, ceros en cualquier parte que no sea la DP; *simétrica* si $\mathbf{A}^T = \mathbf{A}$; *antisimétrica* si $\mathbf{A}^T = -\mathbf{A}$; *ortogonal* si $\mathbf{A}\mathbf{A}^T = \mathbf{A}^T\mathbf{A} = \mathbf{I} \rightarrow \mathbf{A}^{-1} = \mathbf{A}^T$, donde $\mathbf{A}^{-1}$ es la *matriz inversa* de $\mathbf{A}$ y cumple que $\mathbf{A}\mathbf{A}^{-1} = \mathbf{I}$; *normal* aquella que conmuta con su traspuesta $\mathbf{A}\mathbf{A}^T = \mathbf{A}^T\mathbf{A}$ (ejemplos: *simétrica*, *antisimétrica* y *ortogonal*). Sean $\mathbf{A}$ y $\mathbf{B}$ dos matrices $m \times n$ y $p \times q$.

La *suma* y *resta* solo está definida para matrices de igual dimensión $m=p$ y $n=q$, se definen por $(a_{ij}) + (b_{ij})$ y $(a_{ij}) - (b_{ij})$; el *producto por un escalar* k , se define por $k(a_{ij})$. Un *vector fila* $(\mathbf{a}_i)$, es una matriz $1 \times n$, un *vector columna* $(\mathbf{b}_j)$, es una matriz $m \times 1$ y se define su *producto escalar* cuando $n=m$ por $(\mathbf{a}_i) \cdot (\mathbf{b}_j) = (a_i) + (b_j)$.

La *multiplicación* solo esta definida si el numero de columnas de $\mathbf{A}$ es igual al numero de filas de $\mathbf{B}$, es decir $n=p$, la matriz resultante $\mathbf{C}$ será $m \times q$, y se efectúa como el *producto escalar* de cada *vector fila* de $\mathbf{A}$ denotado por $[\mathbf{A}_i]$ por cada *vector columna* de $\mathbf{B}$ denotado por $[\mathbf{B}_j]$, es decir $(c_{ij}) = [\mathbf{A}_i] \cdot [\mathbf{B}_j]$. La *división* de $\mathbf{A}$ entre $\mathbf{B}$ se define por $\mathbf{A}/\mathbf{B} = \mathbf{A}\mathbf{B}^{-1}$.

El **determinante** de una matriz *n-cuadrada* $\mathbf{A}$, se denota por $\det(\mathbf{A}) = |\mathbf{A}| = \Delta_G$, es un *escalar* (no un valor absoluto) y se obtiene por *recursividad*. A cada elemento a_{ij} de la *n-matriz* le corresponde un *signo* dado por $(-1)^{i+j}$ es decir: $(+ - + \dots)$; la *submatriz* de a_{ij} es aquella *n-matriz* que se obtiene eliminando la *i-ésima fila* y la *j-ésima columna*; el *menor* M_{ij} de a_{ij} es el *determinante* de su *submatriz* $M_{ij} = |a_{ij}|$; el *cofactor* o

adjunto de a_{ij} se define por $(-1)^{i+j} M_{ij}$. El determinante de una n -matriz es la suma de los productos de los elementos de cualquier fila o columna por sus correspondientes cofactores, en símbolos:

$$K \in \{1 \dots n\}$$

$$|A| = \sum_{j=1}^n a_{ij} * (-1)^{i+j} * |a_{ij}| \quad \Leftrightarrow \quad |A| = \sum_{j=K}^n a_{ij} * (-1)^{i+j} * |a_{ij}| \quad (56)$$

$$\text{cof}(a_{ij}) = (-1)^{i+j} * |a_{ij}| \leftarrow \text{cofactor. Ejemplo: } \begin{vmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{vmatrix} = a_{11}a_{22} - a_{12}a_{21}$$

Otra forma de expresar recursivamente el determinante es como sigue: Si $\mathbf{A}$ es una n -matriz, sea $\mathbf{A}_{ij}$ la $(n-1)$ -matriz obtenida a partir de $\mathbf{A}$ suprimiendo su i -ésimo renglón y su j -ésima columna. Entonces el determinante $|\mathbf{A}|$ a lo largo de su i -ésimo renglón está dado por:

$$|A| = \sum_{j=1}^n a_{ij} * (-1)^{i+j} * |A_{ij}| \quad (\text{i fija}) \quad (57)$$

Para una matriz de 3×3 existe un método conocido como *regla de Sarrus* y consiste en duplicar las filas 1 y 2 en la parte inferior de la matriz, el determinante es la suma de los productos de los elementos diagonales, hacia abajo son positivos y hacia arriba son negativos, es decir:

$$\begin{array}{ccc} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \\ a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \end{array} \quad (58)$$

$$= a_{11}a_{22}a_{33} + a_{21}a_{32}a_{13} + a_{31}a_{12}a_{23} - a_{31}a_{22}a_{13} - a_{11}a_{32}a_{23} - a_{21}a_{12}a_{33}$$

Si una n -matriz $\mathbf{A}$ tiene como *inversa* $\mathbf{A}^{-1}$ se dice que es *regular* o *invertible* (de lo contrario será *singular* o no invertible) y satisface que $\mathbf{A}\mathbf{A}^{-1} = \mathbf{A}^{-1}\mathbf{A} = \mathbf{I}$, y se obtiene por transformaciones elementales de renglón TER (*método de Gauss-Jordan*) sobre la matriz aumentada de $\mathbf{A}$ con la matriz identidad $\mathbf{I}$, $\mathbf{A}:\mathbf{I}$, de manera tal que se obtenga $\mathbf{I}:\mathbf{A}^{-1}$, o bien por determinantes. Sea $a \neq 0$ un escalar y k un índice $\{1 \dots n\}$, las TER son: (1) $[\mathbf{A}_i] = [\mathbf{A}_k]$, $[\mathbf{A}_j] = [\mathbf{A}_k]$; (2) $a[\mathbf{A}_i]$, $a[\mathbf{A}_j]$; (3) $a[\mathbf{A}_i] \pm [\mathbf{A}_{i+k}]$, $a[\mathbf{A}_j] \pm [\mathbf{A}_{j+k}]$; en general este procedimiento es *iterativo*. Empleando determinantes la inversa de $\mathbf{A}$ esta dada por:

$$A^{-1} = \frac{(\text{cof}(a_{ij}))^T}{|A|} = \frac{\text{adj } A}{|A|} \Leftrightarrow |A| \neq 0 \quad (59)$$

Para una *matriz* de 2-cuadrada se tiene que:

$$A = \begin{bmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{bmatrix}, \quad \det A = a_{11}a_{22} - a_{21}a_{12}, \quad \text{cof } A = \begin{bmatrix} a_{22} & -a_{21} \\ -a_{12} & a_{11} \end{bmatrix},$$

$$\text{adj } A = (\text{cof } A)^T = \begin{bmatrix} a_{22} & -a_{12} \\ -a_{21} & a_{11} \end{bmatrix}, \quad A^{-1} = \frac{\text{adj } A}{\det A} = \frac{1}{a_{11}a_{22} - a_{21}a_{12}} \begin{bmatrix} a_{22} & -a_{12} \\ -a_{21} & a_{11} \end{bmatrix} \quad (60)$$

 **Ejemplo 1.9.** Determinar la inversa de **A**, empleando determinantes.

$$A = \begin{bmatrix} 2 & 2 & 0 \\ -2 & 1 & 1 \\ 3 & 0 & 1 \end{bmatrix}$$

Solución. (1) Se determina el determinante de **A**; (2) se encuentra la matriz de cofactores y la matriz adjunta; (3) se aplica la formula:

$$(1) \quad \det A = 2 \begin{vmatrix} 1 & 1 \\ 0 & 1 \end{vmatrix} - 2 \begin{vmatrix} -2 & 1 \\ 3 & 1 \end{vmatrix} = 2 - 2(-2 - 3) = 2 - 2(-5) = 12 \neq 0 \quad \therefore \text{tiene inversa}$$

(2) Los cofactores de cada elemento a_{ij} son los determinantes de las submatrices obtenidas al eliminar la fila y columna ij :

$$\text{Cof } A = \begin{bmatrix} \begin{vmatrix} 1 & 1 \\ 0 & 1 \end{vmatrix} & -\begin{vmatrix} -2 & 1 \\ 3 & 1 \end{vmatrix} & \begin{vmatrix} -2 & 1 \\ 3 & 0 \end{vmatrix} \\ -\begin{vmatrix} 2 & 0 \\ 0 & 1 \end{vmatrix} & \begin{vmatrix} 2 & 0 \\ 3 & 1 \end{vmatrix} & -\begin{vmatrix} 2 & 2 \\ 3 & 0 \end{vmatrix} \\ \begin{vmatrix} 2 & 0 \\ 1 & 1 \end{vmatrix} & -\begin{vmatrix} 2 & 0 \\ -2 & 1 \end{vmatrix} & \begin{vmatrix} 2 & 2 \\ -2 & 1 \end{vmatrix} \end{bmatrix} = \begin{bmatrix} 1 & 5 & -3 \\ -2 & 2 & 6 \\ 2 & -2 & 6 \end{bmatrix} \rightarrow \text{adj } A = (\text{Cof } A)^T = \begin{bmatrix} 1 & -2 & 2 \\ 5 & 2 & -2 \\ -3 & 6 & 6 \end{bmatrix}$$

$$(3) \quad A^{-1} = \frac{\text{adj}(A)}{|A|} = \frac{1}{12} \begin{bmatrix} 1 & -2 & 2 \\ 5 & 2 & -2 \\ -3 & 6 & 6 \end{bmatrix} = \begin{bmatrix} \frac{1}{12} & -\frac{1}{6} & \frac{1}{6} \\ \frac{5}{12} & \frac{1}{6} & -\frac{1}{6} \\ -\frac{1}{4} & \frac{1}{2} & \frac{1}{2} \end{bmatrix} \quad \clubsuit$$

Un sistemas de n ecuaciones lineales en el que los *coeficientes* se denotan por a_{ij} , las variables por x_{ij} y los *términos independientes* por Y_i , puede transformarse matricialmente en: una *matriz de coeficientes* n -cuadrada **A**, un *vector columna de n -incógnitas* **X** y un *vector columna de n -términos independientes* **Y**; cuya relación entre sí es: **AX=Y**, la *matriz aumentada del sistema* MAS se representa por **A:Y**.

$$\begin{array}{l}
 a_{11}x_1 + a_{12}x_2 + a_{13}x_3 + \dots + a_{1n}x_n = y_1 \\
 a_{21}x_1 + a_{22}x_2 + a_{23}x_3 + \dots + a_{2n}x_n = y_2 \\
 a_{31}x_1 + a_{32}x_2 + a_{33}x_3 + \dots + a_{3n}x_n = y_3 \\
 \vdots \quad \quad \quad \vdots \quad \quad \quad \vdots \quad \quad \quad \vdots \quad \quad \quad \vdots \\
 a_{n1}x_1 + a_{n2}x_2 + a_{n3}x_3 + \dots + a_{nn}x_n = y_n
 \end{array}
 \rightarrow
 \begin{bmatrix}
 a_{11} & a_{12} & a_{13} & \dots & a_{1n} \\
 a_{21} & a_{22} & a_{23} & \dots & a_{2n} \\
 a_{31} & a_{32} & a_{33} & \dots & a_{3n} \\
 \dots & \dots & \dots & \dots & \dots \\
 a_{n1} & a_{n2} & a_{n3} & \dots & a_{nn}
 \end{bmatrix}
 \begin{bmatrix}
 x_1 \\
 x_2 \\
 x_3 \\
 \dots \\
 x_n
 \end{bmatrix}
 =
 \begin{bmatrix}
 y_1 \\
 y_2 \\
 y_3 \\
 \dots \\
 y_n
 \end{bmatrix}
 \rightarrow AX = Y \quad (61)$$

La solución $\mathbf{X}$, en forma matricial esta dada por $\mathbf{X}=\mathbf{A}^{-1}\mathbf{Y}$. A si mismo *para hallar la j -ésima incógnita x_j* se emplea la *regla de Cramer* que consiste en sustituir la *j -ésima columna de la matriz de coeficientes* (cuyo determinante es Δ_G) por el *vector columna de términos independientes* y hallar su determinante denotado por Δ_i , de manera que $x_j=\Delta_i/\Delta_G$.

Los *métodos recursivos* expuestos, son útiles para sistemas pequeños; sin embargo para *grandes sistemas* (como las matrices de admitancias de una red eléctrica) se emplean métodos cuya base son las TER y son en esencia *iterativos*. El *método de Gauss* MG consiste en tomar la MAS, aplicar TER y formar una MTS y realizar *sustitución regresiva*. El *método de Gauss-Jordan* consiste en tomar la MAS, aplicar TER y formar una *matriz identidad*, entonces el *vector de aumento* tendrá los valores de las incógnitas.

 **Ejemplo 1.10.** Resolver el sistema de ecuaciones lineales por *transformaciones elementales de renglón*.

$$\begin{bmatrix}
 2 & 6 & 1 \\
 1 & 2 & -1 \\
 5 & 7 & -4
 \end{bmatrix}
 \begin{bmatrix}
 x_1 \\
 x_2 \\
 x_3
 \end{bmatrix}
 =
 \begin{bmatrix}
 7 \\
 -1 \\
 9
 \end{bmatrix}$$

Solución. Aplicando las TER a la *matriz aumentada del sistema*, se tiene que:

$$\begin{array}{l}
 \xrightarrow{\quad} \left[\begin{array}{ccc|c} 2 & 6 & 1 & 7 \\ 1 & 2 & -1 & -1 \\ 5 & 7 & -4 & 9 \end{array} \right] \xrightarrow{R_{12}} \left[\begin{array}{ccc|c} 1 & 2 & -1 & -1 \\ 2 & 6 & 1 & 7 \\ 5 & 7 & -4 & 9 \end{array} \right] \xrightarrow{\begin{smallmatrix} -2R_1+R_2 \\ -5R_1+R_3 \end{smallmatrix}} \left[\begin{array}{ccc|c} 1 & 2 & -1 & -1 \\ 0 & 2 & 3 & 9 \\ 0 & -3 & 1 & 14 \end{array} \right] \\
 \xrightarrow{\frac{1}{2}R_2} \left[\begin{array}{ccc|c} 1 & 2 & -1 & -1 \\ 0 & 1 & \frac{3}{2} & \frac{9}{2} \\ 0 & -3 & 1 & 14 \end{array} \right] \xrightarrow{3R_2+R_3} \left[\begin{array}{ccc|c} 1 & 2 & -1 & -1 \\ 0 & 1 & \frac{3}{2} & \frac{9}{2} \\ 0 & 0 & \frac{11}{2} & \frac{55}{2} \end{array} \right] \xrightarrow{\frac{2}{11}R_3} \left[\begin{array}{ccc|c} 1 & 2 & -1 & -1 \\ 0 & 1 & \frac{3}{2} & \frac{9}{2} \\ 0 & 0 & 1 & 5 \end{array} \right] \\
 \xrightarrow{-2R_2+R_1} \left[\begin{array}{ccc|c} 1 & 0 & -4 & -10 \\ 0 & 1 & \frac{3}{2} & \frac{9}{2} \\ 0 & 0 & 1 & 5 \end{array} \right] \xrightarrow{\begin{smallmatrix} 4R_3+R_1 \\ -\frac{3}{2}R_3+R_2 \end{smallmatrix}} \left[\begin{array}{ccc|c} 1 & 0 & 0 & 10 \\ 0 & 1 & 0 & -3 \\ 0 & 0 & 1 & 5 \end{array} \right] \quad \therefore \quad \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix} = \begin{bmatrix} 10 \\ -3 \\ 5 \end{bmatrix} \quad \clubsuit
 \end{array}$$

Los métodos de aproximaciones como Jacobi y Gauss-Seidel consideran ecuaciones de recurrencia para cada x_i , cuya aproximación x_k se satisface cuando la diferencia entre valores inmediatos es un *infinitésimo* ε .

$$y_k = a_{k1}x_1 + a_{k2}y_k + \dots + a_{kk}x_k + \dots + a_{kn}x_n, \quad \left| \frac{x_k(i+1) - x_k(i)}{x_k} \right| < \varepsilon$$

$$x_k(i+1) = \frac{1}{a_{kk}} \left[y_k - \sum_{r=1}^{k-1} a_{kr}x_r(i) - \sum_{r=k+1}^n a_{kr}x_r(i) \right] \leftarrow \text{Jacobi} \quad (62)$$

$$x_k(i+1) = \frac{1}{a_{kk}} \left[y_k - \sum_{r=1}^{k-1} a_{kr}x_r(i+1) - \sum_{r=k+1}^n a_{kr}x_r(i) \right] \leftarrow \text{Gauss - Seidel}$$

Valores y vectores característicos de matrices cuadradas

Se dice que un número $\lambda | \lambda \notin C$ (*puede ser real o complejo*), es un *valor propio*, *valor característico* o *eigenvalor* de una *matriz* n -cuadrada $\mathbf{A}$, si existe un *n-vector propio*, *n-vector característico* o *n-eigenvector* $\mathbf{K}$ correspondiente al valor propio λ y satisfacen que: $\mathbf{AK} = \lambda \mathbf{K}$. Esta expresión se representa comúnmente por:

$$(\mathbf{A} - \lambda \mathbf{I})\mathbf{K} = 0 \quad (63)$$

Siendo $\mathbf{I}$, la *n-matriz identidad*. En forma extendida la ecuación anterior se expresa por:

$$\begin{array}{ccccccc} (a_{11} - \lambda)k_1 & + a_{12}k_2 & & + \dots & + a_{1n}k_n & & = 0 \\ a_{21}k_1 & & + (a_{22} - \lambda)k_2 & + \dots & + a_{2n}k_n & & = 0 \\ \dots & & \dots & & \dots & & = \vdots \\ a_{n1}k_1 & + a_{n2}k_2 & & + \dots & + (a_{nn} - \lambda)k_n & & = 0 \end{array} \quad (64)$$

Un sistema homogéneo de ecuaciones lineales con n incógnitas tiene una solución no trivial si y sólo si el determinante de la matriz de coeficientes es igual a cero. Este principio requiere que:

$$\det(\mathbf{A} - \lambda \mathbf{I}) = 0 \quad (65)$$

El desarrollo de este determinante produce un polinomio de grado n , en λ , conocido como *ecuación característica*, las raíces del polinomio son los valores propios; y la sustitución de un valor particular de λ en el sistema $(\mathbf{A} - \lambda \mathbf{I})\mathbf{K} = 0$, permite hallar el vector característico; por solución del sistema singular

resultante, lo cual requiere la asignación de un valor arbitrario para un componente del vector característico por ejemplo $v_n = k$.

 **Ejemplo 1.11.** Determinar los *valores y vectores característicos* de **A**.

$$A = \begin{bmatrix} 1 & 2 & 1 \\ 6 & -1 & 0 \\ -1 & -2 & -1 \end{bmatrix}$$

Solución. (1) Se determina el polinomio característico y sus raíces; (2) se emplean los valores característicos y se resuelve el sistema singular para hallar cada vector característico, es decir:

$$(1) \quad \det(A - \lambda I) = \begin{vmatrix} 1-\lambda & 2 & 1 \\ 6 & -1-\lambda & 0 \\ -1 & -2 & -1-\lambda \end{vmatrix} = -\lambda^3 - \lambda^2 + 12\lambda = -\lambda(\lambda + 4)(\lambda - 3) = 0 \rightarrow \lambda = \begin{bmatrix} 0 \\ -4 \\ 3 \end{bmatrix}$$

(2)

$$\text{Para } \lambda = 0 : (A - 0I | 0) = \left[\begin{array}{ccc|c} 1 & 2 & 1 & 0 \\ 6 & -1 & 0 & 0 \\ -1 & -2 & -1 & 0 \end{array} \right] \xrightarrow{\substack{-6R_1 + R_2 \\ R_1 + R_3}} \left[\begin{array}{ccc|c} 1 & 2 & 1 & 0 \\ 0 & -13 & -6 & 0 \\ -1 & -2 & -1 & 0 \end{array} \right] \xrightarrow{-\frac{1}{13}R_2} \left[\begin{array}{ccc|c} 1 & 2 & 1 & 0 \\ 0 & 1 & \frac{6}{13} & 0 \\ 0 & 0 & 0 & 0 \end{array} \right]$$

$$\xrightarrow{-2R_2 + R_1} \left[\begin{array}{ccc|c} 1 & 0 & \frac{1}{13} & 0 \\ 0 & 1 & \frac{6}{13} & 0 \\ 0 & 0 & 0 & 0 \end{array} \right] \quad \begin{array}{l} k_2 + \frac{6}{13}k_3 = 0 \\ k_1 + \frac{1}{13}k_3 = 0 \end{array} \quad \text{si } k_3 = -13 \rightarrow \begin{array}{l} k_2 = 6 \\ k_1 = 1 \end{array} \quad \therefore K_1 = \begin{bmatrix} 1 \\ 6 \\ -13 \end{bmatrix}$$

$$\text{Para } \lambda = -4 : (A + 4I | 0) = \left[\begin{array}{ccc|c} 5 & 2 & 1 & 0 \\ 6 & 3 & 0 & 0 \\ -1 & -2 & -3 & 0 \end{array} \right] \xrightarrow{\text{operaciones elementales de renglon}} \left[\begin{array}{ccc|c} 1 & 0 & 1 & 0 \\ 0 & 1 & -2 & 0 \\ 0 & 0 & 0 & 0 \end{array} \right] \quad \text{si } k_3 = -13 \quad K_2 = \begin{bmatrix} -1 \\ 2 \\ 1 \end{bmatrix}$$

$$\text{Para } \lambda = 3 : (A - 3I | 0) = \left[\begin{array}{ccc|c} -2 & 2 & 1 & 0 \\ 6 & -4 & 0 & 0 \\ -1 & -2 & -4 & 0 \end{array} \right] \xrightarrow{\text{operaciones elementales de renglon}} \left[\begin{array}{ccc|c} 1 & 0 & 1 & 0 \\ 0 & 1 & \frac{3}{2} & 0 \\ 0 & 0 & 0 & 0 \end{array} \right] \quad \text{si } k_3 = -2 \quad K_3 = \begin{bmatrix} 2 \\ 3 \\ -2 \end{bmatrix} \quad \clubsuit$$

Clases en C++ para el manejo de *vectores*, *matrices* y *determinantes* con las operaciones elementales.

```
% Clase para el manejo de vectores.
arcbritoh@gmail.com [Ms Vc++ 6]
% Archivo de declaración: Vector.h
```

```
class Vector
{
public:
    Vector ();
    Vector (const Vector& rVector);
    explicit Vector (int iSize);

    virtual ~Vector ();

    Vector& operator= (const Vector& rVector);
    Vector& operator+= (const Vector& rVector);
    Vector& operator-= (const Vector& rVector);

    Vector& operator+= (double dNum);
    Vector& operator-= (double dNum);
    Vector& operator*= (double dNum);
    Vector& operator/= (double dNum);

    bool operator== (const Vector& rVector);
    bool operator!= (const Vector& rVector);

    Vector operator- () const;
    Vector operator+ (const Vector& rVector) const;
    Vector operator- (const Vector& rVector) const;
    double operator* (const Vector& rVector) const;

    friend Vector operator+ (const Vector& rVector, double dNum);
    friend Vector operator+ (double dNum, const Vector& rVector);
    friend Vector operator- (const Vector& rVector, double dNum);
    friend Vector operator- (double dNum, const Vector& rVector);
    friend Vector operator* (double dNum, const Vector& rVector);
    friend Vector operator/ (const Vector& rVector, double dNum);

    double operator[] (int iIndex) const;
    double& operator[] (int iIndex);
    int GetSize (void) const { return _iSize; }

private:
    int _iSize;
    double* _pdArray;
};
```

```
% Clase para el manejo de matrices, basado en la clase
Vector.
arcbritoh@gmail.com [Ms Vc++ 6]
% Archivo de declaración: Matrix.h
```

```
#include "Vector.h"

class Matrix
{
public:
    Matrix ();
    Matrix (const Matrix& rMatrix);
    Matrix (int iRows, int iColumns);

    virtual ~Matrix ();

    Matrix& operator= (const Matrix& rMatrix);
    Matrix& operator+= (const Matrix& rMatrix);
    Matrix& operator-= (const Matrix& rMatrix);
    Matrix& operator*= (const Matrix& rMatrix);

    bool operator== (const Matrix& rMatrix);
    bool operator!= (const Matrix& rMatrix);

    Matrix operator+ (const Matrix& rMatrix) const;
    friend Matrix operator+ (const Matrix& rMatrix, double dNum);
    friend Matrix operator+ (double dNum, const Matrix& rMatrix);

    Matrix operator- () const;
    Matrix operator- (const Matrix& rMatrix) const;
    friend Matrix operator- (const Matrix& rMatrix, double dNum);
    friend Matrix operator- (double dNum, const Matrix& rMatrix);

    Matrix operator* (const Matrix& rMatrix) const;
    friend Matrix operator* (const Matrix& rMatrix, double dNum);
    friend Matrix operator* (double dNum, const Matrix& rMatrix);

    friend Matrix operator/ (const Matrix& rMatrix, double dNum);

    double operator! ();
    Matrix operator~ ();

    Vector operator[] (int iIndex) const;
    Vector& operator[] (int iIndex);

    void Print(void);

    Vector GetRow (int iIndex) const;
    Vector GetColumn (int iIndex) const;

    int GetRows (void) const { return _iRows; }
    int GetColumns (void) const { return _iColumns; }

    void RemoveRow (int iIndex);
    void RemoveColumn (int iIndex);

private:
    bool ValidateSizes (const Matrix& rMatrix) const;

    int _iRows;
    int _iColumns;

    Vector* _pVector;
};
```

```

% Clase para el manejo de vectores.
arcbritoh@gmail.com [Ms Vc++ 6]
% Archivo de implementación: Vector.cpp

#include "Vector.h"

Vector::Vector() :
    _iSize(0),
    _pdArray(0)
{
}

Vector::Vector(const Vector& rVector) :
    _iSize(rVector._iSize),
    _pdArray(0)
{
    _pdArray = new double[_iSize];
    memcpy(_pdArray,
           rVector._pdArray, _iSize*sizeof(double));
}

Vector::Vector(int iSize) :
    _iSize(iSize),
    _pdArray(0)
{
    _pdArray = new double[iSize];
    memset(_pdArray, 0, iSize*sizeof(double));
}

Vector::~Vector()
{
    if (0 != _pdArray)
    {
        delete[] _pdArray;
        _pdArray = 0;
    }
}

Vector& Vector::operator=(const Vector& rVector)
{
    if (this != &rVector)
    {
        if (0 != _pdArray)
        {
            delete[] _pdArray;
            _pdArray = 0;
        }
        _iSize = rVector._iSize;
        _pdArray = new double[_iSize];
        memcpy(_pdArray,
               rVector._pdArray, _iSize*sizeof(double));
    }
    return *this;
}

Vector& Vector::operator+=(const Vector& rVector)
{
    *this = *this + rVector;
    return *this;
}

Vector& Vector::operator-=(const Vector& rVector)
{
    *this = *this - rVector;
    return *this;
}

Vector& Vector::operator+=(double dNum)
{
    *this = *this + dNum;
    return *this;
}

Vector& Vector::operator-=(double dNum)
{
    *this = *this - dNum;
    return *this;
}

Vector& Vector::operator*(double dNum)
{
    *this = *this * dNum;
    return *this;
}

Vector& Vector::operator/=(double dNum)
{
    *this = *this / dNum;
    return *this;
}

bool Vector::operator==(const Vector& rVector)
{
    bool bEqual = false;
    if (rVector._iSize != _iSize)
    {
        // THROW EXCEPTION
    }
    if (0 == memcmp(_pdArray,
                    rVector._pdArray, _iSize*sizeof(double)))
    {
        bEqual = true;
    }
    return bEqual;
}

bool Vector::operator!=(const Vector& rVector)
{
    return !operator==(rVector);
}

Vector Vector::operator-() const
{
    Vector vector(_iSize);
    for (int i = 0; i < _iSize; i++)

```

1

```

{
    vector._pdArray[i] = -_pdArray[i];
}
return vector;
}

Vector Vector::operator+(const Vector& rVector) const
{
    Vector vector(_iSize);
    for (int i = 0; i < _iSize; i++)
    {
        vector._pdArray[i] = _pdArray[i] + rVector._pdArray[i];
    }
    return vector;
}

Vector Vector::operator-(const Vector& rVector) const
{
    return operator+(-rVector);
}

double Vector::operator*(const Vector& rVector) const
{
    double dMag = 0;
    if (_iSize != rVector._iSize)
    {
        // THROW EXCEPTION
    }
    for (int i = 0; i < _iSize; i++)
    {
        dMag += _pdArray[i] * rVector._pdArray[i];
    }
    return dMag;
}

Vector operator+(const Vector& rVector, double dNum)
{
    Vector vector(rVector._iSize);
    for (int i = 0; i < rVector._iSize; i++)
    {
        vector._pdArray[i] = rVector._pdArray[i] + dNum;
    }
    return vector;
}

Vector operator+(double dNum, const Vector& rVector)
{
    return operator+(rVector, dNum);
}

Vector operator-(const Vector& rVector, double dNum)
{
    Vector vector(rVector._iSize);
    for (int i = 0; i < rVector._iSize; i++)
    {
        vector._pdArray[i] = rVector._pdArray[i] - dNum;
    }
    return vector;
}

Vector operator-(double dNum, const Vector& rVector)
{
    return -operator-(rVector, dNum);
}

Vector operator*(const Vector& rVector, double dNum)
{
    Vector vector(rVector._iSize);
    for (int i = 0; i < rVector._iSize; i++)
    {
        vector._pdArray[i] = rVector._pdArray[i] * dNum;
    }
    return vector;
}

Vector operator*(double dNum, const Vector& rVector)
{
    return operator*(rVector, dNum);
}

Vector operator/(const Vector& rVector, double dNum)
{
    return operator*(rVector, 1/dNum);
}

double Vector::operator[](int iIndex) const
{
    if ((iIndex >= _iSize) || (iIndex < 0))
    {
        // THROW EXCEPTION
    }
    return _pdArray[iIndex];
}

double& Vector::operator[](int iIndex)
{
    if ((iIndex >= _iSize) || (iIndex < 0))
    {
        // THROW EXCEPTION
    }
    return _pdArray[iIndex];
}
}

```

2

```

% Clase para el manejo de matrices.
arcbtrh@gmail.com [Ms Vc++ 6]
% Archivo de implementación: Matrix.cpp

#include "Matrix.h"
Matrix::Matrix() :
    _iRows(0),
    _iColumns(0),
    _pVector(0)
{
}

Matrix::Matrix(const Matrix& rMatrix) :
    _iRows(rMatrix._iRows),
    _iColumns(rMatrix._iColumns),
    _pVector(0)
{
    _pVector = new Vector[_iRows];
    for (int i = 0; i < _iRows; i++)
    {
        _pVector[i] = rMatrix._pVector[i];
    }
}

Matrix::Matrix(int iRows, int iColumns) :
    _iRows(iRows),
    _iColumns(iColumns),
    _pVector(0)
{
    Vector vector(_iColumns);
    _pVector = new Vector[_iRows];
    for (int i = 0; i < _iRows; i++)
    {
        _pVector[i] = vector;
    }
}

Matrix::~Matrix()
{
    if (0 != _pVector)
    {
        delete[] _pVector;
        _pVector = 0;
    }
}

Matrix& Matrix::operator=(const Matrix& rMatrix)
{
    if (this != &rMatrix)
    {
        if (0 != _pVector)
        {
            delete[] _pVector;
            _pVector = 0;
        }
        _iRows = rMatrix._iRows;
        _iColumns = rMatrix._iColumns;
        _pVector = new Vector[_iRows];
        for (int i = 0; i < _iRows; i++)
        {
            _pVector[i] = rMatrix._pVector[i];
        }
        return *this;
    }
}

Matrix& Matrix::operator+=(const Matrix& rMatrix)
{
    *this = *this + rMatrix;
    return *this;
}

Matrix& Matrix::operator-=(const Matrix& rMatrix)
{
    *this = *this - rMatrix;
    return *this;
}

Matrix& Matrix::operator*=(const Matrix& rMatrix)
{
    *this = *this * rMatrix;
    return *this;
}

bool Matrix::operator==(const Matrix& rMatrix)
{
    bool bEqual = true;
    for (int i = 0; i < _iRows; i++)
    {
        if (_pVector[i] != rMatrix._pVector[i])
        {
            bEqual = false;
            break;
        }
    }
    return bEqual;
}

bool Matrix::operator!=(const Matrix& rMatrix)
{
    return !operator==(rMatrix);
}

Matrix Matrix::operator+(const Matrix& rMatrix) const
{
    Matrix matrix(_iRows, _iColumns);
    ValidateSizes(rMatrix);

    for (int i = 0; i < _iRows; i++)
    {
        matrix._pVector[i] = _pVector[i] + rMatrix._pVector[i];
    }
    return matrix;
}

Matrix operator+(const Matrix& rMatrix, double dNum)

```

```

{
    Matrix matrix(rMatrix._iRows, rMatrix._iColumns);
    matrix.ValidateSizes(rMatrix);
    for (int i = 0; i < matrix._iRows; i++)
    {
        matrix._pVector[i] = dNum + rMatrix._pVector[i];
    }
    return matrix;
}

Matrix operator+(double dNum, const Matrix& rMatrix)
{
    return operator+(rMatrix, dNum);
}

Matrix Matrix::operator-() const
{
    Matrix matrix(*this);
    for (int i = 0; i < _iRows; i++)
    {
        matrix._pVector[i] = -_pVector[i];
    }
    return matrix;
}

Matrix Matrix::operator-(const Matrix& rMatrix) const
{
    return *this + (-rMatrix);
}

Matrix operator-(const Matrix& rMatrix, double dNum)
{
    return operator+(rMatrix, -dNum);
}

Matrix operator-(double dNum, const Matrix& rMatrix)
{
    return -operator-(dNum, rMatrix);
}

Matrix Matrix::operator*(const Matrix& rMatrix) const
{
    Matrix matrix(_iRows, rMatrix._iColumns);
    if (_iColumns != rMatrix._iRows)
    {
        // THROW EXCEPTION
    }
    for (int i = 0; i < _iRows; i++)
    {
        for (int j = 0; j < rMatrix._iColumns; j++)
        {
            matrix[i][j] = _pVector[i] * rMatrix.GetColumn(j);
        }
    }
    return matrix;
}

Matrix operator*(const Matrix& rMatrix, double dNum)
{
    Matrix matrix(rMatrix._iRows, rMatrix._iColumns);
    matrix.ValidateSizes(rMatrix);
    for (int i = 0; i < matrix._iRows; i++)
    {
        matrix._pVector[i] = rMatrix._pVector[i] * dNum;
    }
    return matrix;
}

Matrix operator*(double dNum, const Matrix& rMatrix)
{
    return operator*(rMatrix, dNum);
}

Matrix operator/(const Matrix& rMatrix, double dNum)
{
    return operator*(rMatrix, 1/dNum);
}

double Matrix::operator!()
{
    double dDet = 0;
    int iSign = 0;
    if ((_iRows != _iColumns) || (_iRows == 0))
    {
        // THROW EXCEPTION
    }
    if (_iRows == 1)
    {
        dDet = _pVector[0][0];
    }
    else
    {
        Matrix matrixSub;
        for (int j = 0; j < _iColumns; j++)
        {
            matrixSub = *this;
            matrixSub.RemoveRow(0);
            matrixSub.RemoveColumn(j);
            iSign = (j % 2 == 0) ? 1 : -1;
            dDet += iSign * _pVector[0][j] * !matrixSub;
        }
    }
    return dDet;
}

Matrix Matrix::operator~()
{
    Matrix matrix(_iColumns, _iRows);
    for (int i = 0; i < _iRows; i++)
    {
        for (int j = 0; j < _iColumns; j++)
        {
            matrix[j][i] = _pVector[i][j];
        }
    }
}

```

```

        }
        return matrix;
    }
    Vector Matrix::operator[](int iIndex) const
    {
        if (0 == _pVector)
        {
            // TROW EXCEPTION
        }
        if (iIndex >= _iRows)
        {
            // TROW EXCEPTION
        }
        return _pVector[iIndex];
    }
    Vector& Matrix::operator[](int iIndex)
    {
        if (0 == _pVector)
        {
            // TROW EXCEPTION
        }
        if (iIndex >= _iRows)
        {
            // TROW EXCEPTION
        }
        return _pVector[iIndex];
    }
    Vector Matrix::GetRow(int iIndex) const
    {
        if (0 == _pVector)
        {
            // TROW EXCEPTION
        }

        if (iIndex >= _iRows)
        {
            // TROW EXCEPTION
        }
        return _pVector[iIndex];
    }
    Vector Matrix::GetColumn(int iIndex) const
    {
        Vector vector(_iRows);
        for (int i = 0; i < _iRows; i++)
        {
            vector[i] = _pVector[i][iIndex];
        }

        return vector;
    }

    bool Matrix::ValidateSizes(const Matrix& rMatrix) const
    {
        if (_iRows != rMatrix._iRows)
        {
            // THROW EXCEPTION
        }
        if (_iColumns != rMatrix._iColumns)
        {
            // THROW EXCEPTION
        }
        return true;
    }

    void Matrix::RemoveRow(int iIndex)
    {
        Matrix matrixCopy(_iRows-1, _iColumns);
        for (int i = 0, j = 0; i < _iRows; i++)
        {
            if (iIndex == i)
            {
                continue;
            }
            matrixCopy[j++] = _pVector[i];
        }
        operator=(matrixCopy);
    }
    void Matrix::RemoveColumn(int iIndex)
    {
        Matrix matrixCopy(~(*this));
        matrixCopy.RemoveRow(iIndex);
        operator=~(matrixCopy);
    }

```

```
% Transformación de una matriz cuadrada en triangular superior. arcbrth@gmail.com [Matlab 5]
% Método de Gauss. Archivo mTransfTS.m
% Entrada.: matriz mA(Y (NxN+1)
% Salida...: la matriz transformada en Triangular Superior: mTS
% Uso.....: mTransfTS(mAY)

function [mTS] = mTransfTS(mAY)
mTS=mAY; N=size(mAY,1);
if size(mAY,2)~=N+1 then
disp('¡Error! la matriz mAY debe ser NxN+1.')
return;
end
%eliminación del triangulo superior
for c=1:N-1
for f=c+1:N
%Nota mTS(f,:) hace referencia a la fila completa f de mTS
mTS(f,:)=mTS(f,:)-mTS(c,:)*mTS(f,c)/mTS(c,c)
end
end
return;
```

```
% Solución del sistema de ecuaciones por sustitución hacia atrás. arcbritoh@gmail.com [Matlab 5]
% Archivo: mSolveTS.m
% Entrada.: matriz triangular superior mTS (NxN+1)
% Salida...: un vector solución vS(N)
% Uso.....: mSolveTS(mTS)

function [vS] = mSolveTS(mTS)
N=size(mTS,1);
vS=zeros(N,1);
if size(mTS,2)~=N+1 then
disp('¡Error! la matriz mTS debe ser (NxN+1)')
return;
end
%sustitución hacia atras
f=N;
for k=1:N
c=N; sAcumulado=0;
for m=1:N-f
sAcumulado = sAcumulado + mTS(f,c)*vS(c);
c=c-1;
end
vS(f)=(mTS(f,N+1)-sAcumulado) / mTS(f,f);
f=f-1;
end
return;
```

```
% Solución del sistema de ecuaciones por método iterativo de Jacobi. arcbritoh@gmail.com [Matlab 5]
% Archivo: mJacobi.m
% Entrada.: matriz de coeficientes mA(Y (NxN),
%           vector de términos independientes vY (Nx1)
% Salida...: un vector solución vS(N)
% Uso.....: mJacobi(mAY)

function [Xs] = mJacobi(mA, vY)
N=size(mA,1);
if size(vY)~=N then
disp('¡Error! la matriz mTS debe ser (NxN+1)')
return;
end
%la matrices tienen unos como valores iniciales
Xs=ones(N,1);
XsOld=ones(N,1);
%numero máximo de iteraciones
MAXITER=100
%Error máximo permitido
MaxErr=0.0001
deltaE=0;
%formacion de la matriz D
D=zeros(N,N);
vdiagA=diag(mA);
for k=1:N
D(k,k)=vdiagA(k);
end
M=inv(D)*(D-mA);
cuentaX=0;
h=0;
while cuentaX<N
cuentaX=0;
%obtiene los valores incognitas iterativamente
Xs=M*XsOld + inv(D)*vY;
for k=1:N
deltaE=abs(Xs(k)-XsOld(k))/XsOld(k);
if deltaE<MaxErr
cuentaX=cuentaX+1;
end
end
XsOld=Xs;
h=h+1;
if h>=MAXITER
disp('¡Error! El método no converge');
break;
end
end
return;
```

Notación vectorial

En la *física clásica* (Newtoniana, en la que $v \ll c$) *todo evento en el espacio* puede ser descrito por una 4-tupla (x,y,z,t) ; el *espacio* (x,y,z) no se deforma con la *velocidad* (v) y el *tiempo* (t) es independiente del *marco de referencia* elegido.

En el *espacio* un *punto* es una triada ordenada de escalares $P(x,y,z)$, un *escalar* es cualquier número real, geoméricamente un *vector* es un segmento de recta dirigido que une dos puntos $\vec{r} = \overline{PQ}$, que tiene como *propiedades*: *magnitud* y *dirección* y se puede escribir como una triada ordenada $\mathbf{E}=(E_1,E_2,E_3)$, siendo (E_1,E_2,E_3) sus *componentes escalares* a lo largo de ejes mutuamente perpendiculares x,y,z en un *sistema coordenado ortogonal cartesiano (rectangular)*.

La *longitud*, *magnitud*, *valor absoluto*, *norma* o *módulo* del vector es $E=|\mathbf{E}|=+\sqrt{E_1^2+E_2^2+E_3^2}$, su *dirección* está dada por el *vector normalizado* o *unitario* (de longitud unidad) $\hat{\mathbf{E}}=\mathbf{E}/E$, sus *ángulos directores* respecto a los ejes x, y, z son $\theta_x, \theta_y, \theta_z$ y se relacionan con sus componentes por $E_1 = E \cos \theta_x, E_2 = E \cos \theta_y, E_3 = E \cos \theta_z$.

Se utilizan *vectores ortonormales* o de la *base canónica* (*unitarios a lo largo de los ejes coordenados*) $\mathbf{i}=(1,0,0)$, $\mathbf{j}=(0,1,0)$, $\mathbf{k}=(0,0,1)$ para representar un vector como $\mathbf{E}=E_1\mathbf{i}+E_2\mathbf{j}+E_3\mathbf{k}$. “Un punto $P(x,y,z)$ puede ser descrito por un vector de posición $\mathbf{r}=\mathbf{r}(x,y,z)$ fijo o ligado al origen $O(0,0,0)$ y se expresa como $P(\mathbf{r})$ ”, dados dos puntos P y Q el vector que va de P a Q es $\overline{PQ}=\mathbf{Q}-\mathbf{P}$, siendo su *magnitud* la *distancia* $|\overline{PQ}|$.

Un *vector libre* puede ser trasladado en el espacio conservando su *magnitud* y *dirección*, en virtud de lo cual los vectores con sus mismas propiedades se dice que son *equipolentes*.

Algebra vectorial

Sean $\mathbf{A}(A_1, A_2, A_3)$ y $\mathbf{B}(B_1, B_2, B_3)$ dos vectores; su *suma* es $\mathbf{A} + \mathbf{B}$, el *producto por un escalar* k es $k\mathbf{B}$, el *negativo* o *vector de sentido opuesto* es $-\mathbf{B}$, su *diferencia* es $\mathbf{A} - \mathbf{B}$, el *producto escalar*, *producto interno* o *producto punto* se define por $\mathbf{A} \cdot \mathbf{B}$, el *producto vectorial*, *producto externo* o *producto cruz* se define por $\mathbf{A} \times \mathbf{B}$ (el vector resultante es *perpendicular*, *ortogonal* o *normal* al plano AB , su dirección se determina por la **regla de la mano derecha**: cierre la mano derecha en dirección de $\mathbf{A}$ hacia $\mathbf{B}$ y el pulgar indicara la dirección de $\mathbf{A} \times \mathbf{B}$).

Sea un vector $\mathbf{C}(C_1, C_2, C_3)$, se define el *triple producto escalar* como $\mathbf{A} \cdot (\mathbf{B} \times \mathbf{C})$ y el *triple producto vectorial* por $\mathbf{A} \times (\mathbf{B} \times \mathbf{C})$, es decir:

$$\begin{aligned}
 \vec{A} &= (A_1, A_2, A_3), & \vec{B} &= (B_1, B_2, B_3) \\
 \vec{A} + \vec{B} &= (A_1 + B_1, A_2 + B_2, A_3 + B_3) \\
 k\vec{B} &= (kB_1, kB_2, kB_3) \\
 -\vec{B} &= (-B_1, -B_2, -B_3) \\
 \vec{A} - \vec{B} &= \vec{A} + (-\vec{B}) = (A_1 - B_1, A_2 - B_2, A_3 - B_3) \\
 \vec{A} \cdot \vec{B} &= \vec{B} \cdot \vec{A} = A_1B_1 + A_2B_2 + A_3B_3 = |\vec{A}| |\vec{B}| \cos \theta_{AB} \\
 \vec{A} \times \vec{B} &= -\vec{B} \times \vec{A} = \begin{vmatrix} i & j & k \\ A_1 & A_2 & A_3 \\ B_1 & B_2 & B_3 \end{vmatrix} = (|\vec{A}| |\vec{B}| \sin \theta_{AB}) \hat{n} \\
 \vec{A} \cdot (\vec{B} \times \vec{C}) &= (\vec{A} \times \vec{B}) \cdot \vec{C} = \begin{vmatrix} A_1 & A_2 & A_3 \\ B_1 & B_2 & B_3 \\ C_1 & C_2 & C_3 \end{vmatrix} \\
 \vec{A} \times (\vec{B} \times \vec{C}) &= (\vec{A} \cdot \vec{C}) \vec{B} - (\vec{A} \cdot \vec{B}) \vec{C} \\
 \text{el vector } \hat{n} &\text{ es normal al plano } AB \text{ hacia fuera}
 \end{aligned} \tag{66}$$

La *magnitud* de $\mathbf{A}$ puede determinarse a partir del *producto punto*, la *componente de A sobre B* $\text{comp}_B \mathbf{A}$ es un escalar, la *proyección de A sobre B* $\text{proy}_B \mathbf{A}$ es un vector (cuya magnitud es la componente de $\mathbf{A}$ sobre $\mathbf{B}$ y cuya dirección es la de $\mathbf{B}$).

$$\begin{aligned}
 A &= |\vec{A}| = \sqrt{\vec{A} \cdot \vec{A}} = \sqrt{A_1^2 + A_2^2 + A_3^2} \\
 \text{comp}_B A &= \frac{\vec{A} \cdot \vec{B}}{|\vec{B}|} = |\vec{A}| \cos \theta_{AB} \\
 \text{proy}_B A &= (\text{comp}_B A) \frac{\vec{B}}{|\vec{B}|} = (\text{comp}_B A) \hat{B} = \left(\frac{\vec{A} \cdot \vec{B}}{|\vec{B}|^2} \right) \vec{B}
 \end{aligned} \tag{67}$$

Geometría vectorial

La ecuación de la recta $\vec{r} = \vec{A} + t\vec{M}$ queda determinada por un vector $\vec{A}(A_1, A_2, A_3)$ cuya punta toca la recta y otro vector $\vec{M}(M_1, M_2, M_3)$ en la dirección de la misma, siendo t un escalar; o bien por dos puntos $P(P_1, P_2, P_3)$, $Q(Q_1, Q_2, Q_3)$ a lo largo de la misma.

$$\begin{aligned}\vec{A} &= (A_1, A_2, A_3), \quad \vec{M} = (M_1, M_2, M_3), \quad \vec{r} = (x, y, z) \\ x &= A_1 + M_1 t \\ \vec{r} = \vec{A} + t\vec{M} &\longrightarrow y = A_2 + M_2 t \\ z &= A_3 + M_3 t\end{aligned}\tag{68}$$

$$\begin{aligned}P &= (P_1, P_2, P_3), \quad Q = (Q_1, Q_2, Q_3) \\ x &= P_1 + (Q_1 - P_1)t \\ \vec{r} = \vec{P} + t(\vec{Q} - \vec{P}) &\longrightarrow y = P_2 + (Q_2 - P_2)t \\ z &= P_3 + (Q_3 - P_3)t\end{aligned}$$

La ecuación del plano $n_1x + n_2y + n_3z + k = 0$ queda completamente determinada por un punto $P(P_1, P_2, P_3)$ y un vector normal al plano $\vec{n}(n_1, n_2, n_3)$; o bien por tres puntos A, B, C no *colineales* (colineal = a lo largo de la misma línea). La *distancia* del punto $Q(Q_1, Q_2, Q_3)$ al plano $n_1x + n_2y + n_3z + k = 0$, se denota por d .

$$\begin{aligned}P &= (P_1, P_2, P_3), \quad \vec{n} = (n_1, n_2, n_3) \\ n_1x + n_2y + n_3z + k &= 0 \\ k &= -n_1P_1 - n_2P_2 - n_3P_3\end{aligned}$$

$$A = (A_1, A_2, A_3), \quad B = (B_1, B_2, B_3), \quad C = (C_1, C_2, C_3)\tag{69}$$

$$\vec{n} = \overrightarrow{BA} \times \overrightarrow{CA} = \begin{vmatrix} i & j & k \\ A_1 - B_1 & A_2 - B_2 & A_3 - B_3 \\ A_1 - C_1 & A_2 - C_2 & A_3 - C_3 \end{vmatrix}$$

$$d = \frac{|n_1Q_1 + n_2Q_2 + n_3Q_3 + k|}{|\vec{n}|} = \frac{|n_1Q_1 + n_2Q_2 + n_3Q_3 + k|}{\sqrt{n_1^2 + n_2^2 + n_3^2}}$$

Cálculo vectorial

Curvas, superficies, sólidos y regiones. La *gráfica* de una función $y=f(x)$, es el conjunto de todos los *puntos en el plano* (x,y) y engendra una *curva*, y para su análisis completo basta el *cálculo de una variable*. La *gráfica* de una función $z=f(x,y)$ es el conjunto de los *puntos en el espacio* (x,y,z) y engendra una *superficie*, y para su análisis completo se requiere el *cálculo multivariable*. Un *sólido* es un cuerpo que se genera haciendo *girar una curva en torno a un eje fijo* o bien *desplazando linealmente una superficie normal a un plano*. Una *región* en la *recta* es un *segmento*, en el *plano* es un *área* y en el *espacio* es un *volumen*.

Función escalar multivariable y vectorial. Una *función escalar multivariable* es aquella que consta de 2 o más *variables independientes*, por ejemplo: $z=f(x,y)$ o $u=f(x,y,z)$ y se denota por $f:U \subset \mathbb{R}^3 \rightarrow \mathbb{R}^1$. Una *función vectorial* $\mathbf{E}=\mathbf{E}(x,y,z)$ *asigna vectores a cada punto del espacio*, es decir es un vector cuyas componentes (E_1, E_2, E_3) son funciones escalares de x,y,z : $E_1=f(x,y,z)$, $E_2=g(x,y,z)$, $E_3=h(x,y,z)$, se denota por $f:U \subset \mathbb{R}^3 \rightarrow \mathbb{R}^3$. En general una función f de dominio U , se define por $f:U \subset \mathbb{R}^n \rightarrow \mathbb{R}^m$, y será *función de valores escalares* si $m=1$ y *función de valores vectoriales* si $m>1$.

Derivada parcial. La derivada parcial da la *razón de cambio de una función en una dirección*. Consideremos la función $f(x,y,z)$, *continua (en el intervalo I) y diferenciable*, se define la *derivada parcial* $\partial f / \partial x$ de la función respecto de x (considerando las otras variables y,z como constantes al derivar), la *derivada parcial sucesiva de orden k* ($k \in \mathbb{N}$) tiene la forma $\partial^k f / \partial x^k$, la *derivada parcial iterada (o mixta) de orden 3* es $\partial^3 f / \partial x \partial y \partial z$ (el orden de derivación x,y,z es irrelevante), la *diferencial total* de f es df , en símbolos:

$$\begin{aligned}
 f &= f(x, y, z) \\
 \frac{\partial f}{\partial x}(x, y, z) &= f_x = \lim_{\Delta x \rightarrow 0} \frac{f(x + \Delta x, y, z) - f(x, y, z)}{\Delta x} \longleftarrow \text{Derivada parcial} \\
 \frac{\partial^3 f}{\partial x^3} &= \frac{\partial}{\partial x} \left(\frac{\partial^2 f}{\partial x^2} \right) = \frac{\partial}{\partial x} \left(\frac{\partial}{\partial x} \left(\frac{\partial f}{\partial x} \right) \right) \longleftarrow \text{Derivada parcial sucesiva de orden 3} \\
 \frac{\partial^3 f}{\partial xyz} &= \frac{\partial^3 f}{\partial yzx} = \frac{\partial^3 f}{\partial zyx} \dots = \frac{\partial}{\partial x} \left(\frac{\partial}{\partial y} \left(\frac{\partial f}{\partial z} \right) \right) \longleftarrow \text{Derivada parcial iterada de orden 3} \\
 df &= \frac{\partial f}{\partial x} dx + \frac{\partial f}{\partial y} dy + \frac{\partial f}{\partial z} dz \longleftarrow \text{Diferencial total}
 \end{aligned} \tag{70}$$

Regla de la cadena. Se aplica al derivar funciones de funciones. Sea una función $g(x,y,z)$, se tienen 2 casos: (1) que (x,y,z) sean funciones de (t_1,t_2,t_3) o bien, (2) que (x,y,z) sean funciones del mismo parámetro (t) , entonces:

$$g = g(x,y,z), \quad x = x(t_1), \quad y = y(t_2), \quad z = z(t_3)$$

$$\frac{\partial g}{\partial t_1} = \frac{\partial g}{\partial x} \frac{\partial x}{\partial t_1} + \frac{\partial g}{\partial y} \frac{\partial y}{\partial t_1} + \frac{\partial g}{\partial z} \frac{\partial z}{\partial t_1}, \quad \frac{\partial g}{\partial t_2} = \frac{\partial g}{\partial x} \frac{\partial x}{\partial t_2} + \frac{\partial g}{\partial y} \frac{\partial y}{\partial t_2} + \frac{\partial g}{\partial z} \frac{\partial z}{\partial t_2}, \quad \frac{\partial g}{\partial t_3} = \frac{\partial g}{\partial x} \frac{\partial x}{\partial t_3} + \frac{\partial g}{\partial y} \frac{\partial y}{\partial t_3} + \frac{\partial g}{\partial z} \frac{\partial z}{\partial t_3} \quad (71)$$

$$g = g(x,y,z), \quad x = x(t), \quad y = y(t), \quad z = z(t)$$

$$\frac{dg}{dt} = \frac{\partial g}{\partial x} \frac{dx}{dt} + \frac{\partial g}{\partial y} \frac{dy}{dt} + \frac{\partial g}{\partial z} \frac{dz}{dt}$$

Derivada vectorial. Si $u=u(t)$ es un función escalar y $\mathbf{A}=\mathbf{A}(t)$, $\mathbf{B}=\mathbf{B}(t)$ son funciones vectoriales, cuyas *componentes son funciones del parámetro t* , entonces:

$$\vec{A} = (A_1(t), A_2(t), A_3(t)) \quad \vec{B} = (B_1(t), B_2(t), B_3(t)) \quad u = u(t)$$

$$D_t(\vec{A} + \vec{B}) = D_t(\vec{A}) + D_t(\vec{B})$$

$$D_t(u\vec{A}) = uD_t(\vec{A}) + D_t(u)\vec{A}$$

$$D_t(\vec{A} \cdot \vec{B}) = \vec{A} \cdot D_t(\vec{B}) + D_t(\vec{A}) \cdot \vec{B} \quad (72)$$

$$D_t(\vec{A} \times \vec{B}) = \vec{A} \times D_t(\vec{B}) + D_t(\vec{A}) \times \vec{B}$$

$$D_t(\vec{A}) = (D_t A_1, D_t A_2, D_t A_3) \rightarrow d\vec{A} = \hat{i}dA_1 + \hat{j}dA_2 + \hat{k}dA_3$$

Campo escalar y vectorial. Un *campo* es una región (porción o parte) del espacio que se puede describir por una función escalar o vectorial en cada uno de sus puntos. Sea $\mathbf{r}=\mathbf{r}(x,y,z)$ un vector de posición que define cada punto del espacio en el que actúa un campo.

Un *campo escalar* o *campo potencial* $\Phi = \Phi(\mathbf{r})$ asigna un escalar a cada punto del espacio (*campos escalares: temperatura, densidad*), las superficies con $\Phi = cte$, se denominan *equipotenciales*.

Un *campo vectorial* $\mathbf{E} = \mathbf{E}(\mathbf{r})$ asigna un vector a cada punto del espacio (*campos vectoriales: campo eléctrico, campo magnético, campo gravitacional*), gráficamente los vectores que salen de la fuente campo se representan como *líneas de flujo* con magnitud proporcional y dirección igual al vector en ese punto. Un campo vectorial $\mathbf{E}$ en el espacio $(\mathbb{R}^3)$, tiene tres componentes (E_1, E_2, E_3) , si cada componente es una función de k variables, se dice que el campo es de clase C^k . Por ejemplo $\mathbf{E}(x,y,z)$ es un *campo* cuyas componentes son funciones

de 3 variables x, y, z y es de *clase 3* (C^3); $\mathbf{E}(t)$ es un *campo* cuyas componentes son funciones de 1 *variable* t , y es de *clase 1* (C^1).

Gradiente, derivada direccional, divergencia, laplaciano y rotacional. Se define al *operador nabla* o del ∇ como un *vector de derivadas parciales* $\nabla = (\partial / \partial x, \partial / \partial y, \partial / \partial z)$. Sea $\mathbf{r} = \mathbf{r}(x, y, z)$ un *vector de posición*, y $\Phi = \Phi(\mathbf{r})$ un *campo escalar* y $\mathbf{E} = \mathbf{E}(\mathbf{r})$ un *campo vectorial*.

El **gradiente** $\text{grad}\Phi = \nabla\Phi$ es el *producto del operador nabla con la función escalar* y genera un vector cuyas componentes son las derivadas parciales de Φ , el *gradiente de un campo escalar*, es un vector que representa tanto la magnitud como la dirección de la máxima rapidez de incremento del campo.

La **derivada direccional** de Φ en dirección del vector unitario de $\mathbf{E}$, es $D_{\hat{\mathbf{E}}}\Phi = \nabla\Phi \cdot \hat{\mathbf{E}}$. La **divergencia** es el *producto punto del operador nabla con una función vectorial* $\text{div}\mathbf{E} = \nabla \cdot \mathbf{E}$, la *divergencia de un campo vectorial* en un punto dado puede considerarse como una medida del grado en que el campo diverge o emana de tal punto.

El **laplaciano** de un campo escalar Φ , es la *divergencia del gradiente* de Φ : $\nabla^2\Phi = \nabla \cdot \nabla\Phi$ y es un escalar. El **rotacional** es el *producto cruz del operador nabla con una función vectorial* $\text{rot}\mathbf{E} = \nabla \times \mathbf{E}$, el *rotacional de un campo vectorial* en un punto dado puede considerarse como el grado en que el campo gira alrededor de tal punto. En símbolos:

$$\text{del} = \nabla = (\partial / \partial x, \partial / \partial y, \partial / \partial z) = \hat{i} \frac{\partial}{\partial x} + \hat{j} \frac{\partial}{\partial y} + \hat{k} \frac{\partial}{\partial z}$$

$$\Phi = \Phi(\mathbf{r}) = \Phi(x, y, z) \quad \leftarrow \text{campo escalar}$$

$$\mathbf{E} = \mathbf{E}(\mathbf{r}) = E(x, y, z) \quad \leftarrow \text{campo vectorial}$$

$$\text{grad}\Phi = \nabla\Phi = \hat{i} \frac{\partial\Phi}{\partial x} + \hat{j} \frac{\partial\Phi}{\partial y} + \hat{k} \frac{\partial\Phi}{\partial z} \quad \leftarrow \text{gradiente}$$

$$D_{\hat{\mathbf{E}}}\Phi = \nabla\Phi \cdot \hat{\mathbf{E}} \quad \leftarrow \text{derivada direccional}$$

$$\text{div}\mathbf{E} = \nabla \cdot \mathbf{E} = (\partial / \partial x, \partial / \partial y, \partial / \partial z) \cdot (E_1, E_2, E_3) = \frac{\partial E_1}{\partial x} + \hat{j} \frac{\partial E_2}{\partial y} + \hat{k} \frac{\partial E_3}{\partial z} \quad \leftarrow \text{divergencia}$$

$$\nabla^2\Phi = \nabla \cdot \nabla\Phi = \frac{\partial^2\Phi}{\partial x^2} + \frac{\partial^2\Phi}{\partial y^2} + \frac{\partial^2\Phi}{\partial z^2} \quad \leftarrow \text{laplaciano}$$

$$\text{rot}\mathbf{E} = \nabla \times \mathbf{E} = \begin{vmatrix} \hat{i} & \hat{j} & \hat{k} \\ \frac{\partial}{\partial x} & \frac{\partial}{\partial y} & \frac{\partial}{\partial z} \\ E_1 & E_2 & E_3 \end{vmatrix} \quad \leftarrow \text{rotacional}$$

(73)

Algunas identidades básicas del análisis vectorial se listan a continuación:

$$\begin{aligned}
 \nabla(f+g) &= \nabla f + \nabla g \\
 \nabla cf &= c\nabla f \\
 \nabla(f * g) &= f\nabla g + g\nabla f \\
 \nabla(f / g) &= (g\nabla f - f\nabla g) / g^2 \\
 \operatorname{div}(\vec{F} + \vec{G}) &= \operatorname{div}(\vec{F}) + \operatorname{div}(\vec{G}) \\
 \operatorname{rot}(\vec{F} + \vec{G}) &= \operatorname{rot}(\vec{F}) + \operatorname{rot}(\vec{G}) \\
 \operatorname{div}(f * \vec{G}) &= f * \operatorname{div}(\vec{G}) + \vec{F} \cdot \nabla f \\
 \operatorname{div}(\vec{F} \times \vec{G}) &= \vec{G} \cdot \operatorname{rot} \vec{F} - \vec{F} \cdot \operatorname{rot} \vec{G} \\
 \operatorname{div}(\operatorname{rot} \vec{F}) &= 0 \\
 \operatorname{rot}(f * \vec{F}) &= f * \operatorname{rot} \vec{F} + \nabla f \times \vec{F} \\
 \operatorname{rot}(\nabla f) &= 0 \\
 \nabla^2 f(f * g) &= f * \nabla^2 g + g * \nabla^2 f + 2(\nabla f \cdot \nabla g) \\
 \operatorname{div}(\nabla f \times \nabla g) &= 0 \\
 \operatorname{div}(f * \nabla g - g * \nabla f) &= f * \nabla^2 g - g * \nabla^2 f
 \end{aligned}
 \tag{74}$$

Integrales dobles y triples. Sea $z=f(x,y)$ una *superficie* elevada sobre el plano XY hacia el eje Z^+ ; S una *región en el plano XY*, limitada por dos rectas paralelas al eje Y: $x=x_1$, $x=x_2$ tales que $x_2 > x_1 \forall x$, y limitada por dos curvas $\phi_1(x)$, $\phi_2(x)$ tales que $\phi_2(x) > \phi_1(x) \forall x$, entonces la *integral doble* de f sobre la *región S* representa el *volumen V* del *sólido*, cuya *altura* es z y su *base* es S , si $f(x,y)=1$ se obtiene el *área A* de la *región S*:

$$\begin{aligned}
 V &= \iint_R f(x,y) dA = \int_{x_1}^{x_2} \int_{\phi_1(x)}^{\phi_2(x)} f(x,y) dy dx = \int_{x_1}^{x_2} \left[\int_{\phi_1(x)}^{\phi_2(x)} f(x,y) dy \right] dx \\
 A &= \iint_R dA = \int_{x_1}^{x_2} \int_{\phi_1(x)}^{\phi_2(x)} dy dx = \int_{x_1}^{x_2} \left[\int_{\phi_1(x)}^{\phi_2(x)} dy \right] dx
 \end{aligned}
 \tag{75}$$

Consideremos una *región W* en el *espacio*, cuya *densidad* este dada por una función $f(x,y,z)$, entonces el *volumen V* y la *masa m* de tal *región* están dadas por la *integral triple de f* en la *región S*, en símbolos:

$$\begin{aligned}
 V &= \iiint_W dx dy dz \\
 m &= \iiint_W f(x,y,z) dx dy dz
 \end{aligned}
 \tag{76}$$

En general las *integrales múltiples* son *integrales iteradas* y cada integral es una *integral parcial*, no importando el orden de integración siempre que se tomen los límites apropiadamente.

Elementos diferenciales vectoriales de longitud, área y volumen. Un *desplazamiento diferencial* es el espacio es un vector cuyos componentes son desplazamientos diferenciales a lo largo de los ejes coordenados $d\vec{l} = (dx, dy, dz)$.

El *área normal diferencial* es un vector perpendicular al plano de dos elementos diferenciales de longitud $d\vec{S} = (dydz, 0, 0) = (0, dydz, 0) = (0, 0, dxdz) = dS\hat{n}$. El *volumen diferencial* es un escalar $dv = dxdydz$.

Integral de trayectoria, línea y superficie. Las integrales que involucran la suma de contribuciones que un campo escalar o vectorial provoca a lo largo de una curva, a través de una superficie o en el interior de un volumen tienen aplicaciones físicas importantes.

Una *línea* es el contorno de una *curva abierta* (alisada o alisada por partes) o *cerrada* (simple o no simple con nodos). Consideremos dos puntos P y Q a lo largo de una curva C y un *campo escalar* $f=f(x(t), y(t), z(t))$ y un *campo vectorial* $\mathbf{E}=\mathbf{E}(E_1, E_2, E_3)$ en el espacio donde se sitúa tal curva (la línea podría ser una *espira de alambre* inmersa en un *campo magnético* emitido por dos *magnetos permanentes*).

Se define la *integral de trayectoria* como la suma de las contribuciones del campo escalar f a lo largo de la curva C. Se define la *integral de línea* como la suma de los componentes del campo a lo largo de la curva $\sum \mathbf{E} \cdot d\vec{l}$ desde P hasta Q. Si la curva es cerrada se denota la integral por el símbolo $\oint$ y se denomina *circulación del campo*.

La *integral de superficie* de un campo vectorial que pasa a través de una superficie S, es la cantidad neta de fluido que fluye a través de la superficie por unidad de tiempo, es decir, la razón de flujo: $\iint_S \mathbf{F} \cdot d\vec{S}$.

Teoremas del cálculo vectorial. El *teorema de divergencia* o de *Gauss-Ostrogradsky*, establece que: *el flujo neto del campo vectorial $\mathbf{E}$ que sale de la superficie cerrada S es igual a la integral de volumen de la divergencia del campo $\mathbf{E}$:*

$$\oint_S \mathbf{E} \cdot d\vec{S} = \int_V \nabla \cdot \mathbf{E} dV \quad (77)$$

El *teorema de rotacional* o *teorema de Stokes* establece que: *la circulación del campo vectorial $\mathbf{E}$ alrededor de un contorno cerrado C es igual a la integral del rotacional de $\mathbf{E}$ sobre la superficie abierta S limitada por el contorno C :*

$$\oint_C \mathbf{E} \cdot d\vec{l} = \int_S \nabla \times \mathbf{E} \cdot d\vec{S} \quad (78)$$

Electromagnetismo y las ecuaciones de Maxwell

Gran parte de la formalización matemática del cálculo vectorial proviene de los estudios que *Maxwell* hizo en la electricidad y el magnetismo para formular la *teoría electromagnética*. Sus resultados se simplifican en 4 ecuaciones a continuación expuestas:

$$\begin{aligned} (a) \quad \nabla \times \vec{E} &= -\frac{\partial \vec{B}}{\partial t} \leftrightarrow \oint_C \vec{E} \cdot d\vec{l} = -\int_S \frac{\partial \vec{B}}{\partial t} \cdot d\vec{S} \leftarrow \text{Ley de Faraday} \\ (b) \quad \nabla \times \vec{H} &= \vec{J} + \frac{\partial \vec{D}}{\partial t} \leftrightarrow \oint_C \vec{H} \cdot d\vec{l} = \oint_S \vec{J} \cdot d\vec{S} + \int_S \frac{\partial \vec{D}}{\partial t} \cdot d\vec{S} \leftarrow \text{Ley de Ampere} \\ (c) \quad \nabla \cdot \vec{B} &= 0 \leftrightarrow \oint_S \vec{B} \cdot d\vec{S} = 0 \leftarrow \text{Ley de Gauss Campo Magnético} \\ (d) \quad \nabla \cdot \vec{D} &= \rho \leftrightarrow \oint_S \vec{D} \cdot d\vec{S} = \int_V \rho dV \leftarrow \text{Ley de Gauss Campo Eléctrico} \end{aligned} \quad (79)$$

Donde:

$\mathbf{E}$ es la intensidad decampo eléctrico (V/m)
 $\mathbf{H}$ es la intensidad de campo magnético (A/m)
 $\mathbf{B}$ es la densidad de flujo magnético (T=Wb/m²)
 $\mathbf{D}$ es la densidad de flujo eléctrico (C/m²)
 $\mathbf{J}$ es la densidad de corriente (A/m²)
 ρ es la densidad de carga (C/m³)

En (a) el signo negativo indica que la *fem inducida* se opone al cambio de flujo (*Ley de Lenz*).

En (a) y (b) S es la superficie abierta limitada por el contorno cerrado C .

En (c) y (d) V es el volumen limitado por la superficie cerrada S .

Las ecuaciones de Maxwell junto con la ecuación de conservación de la carga y la ecuación de fuerza de Lorentz brindan una descripción completa de todas las interacciones electromagnéticas clásicas:

$$\begin{aligned}\nabla \cdot \mathbf{J} &= -\frac{\partial \rho}{\partial t} && \leftarrow \text{conservación de la carga} \\ \mathbf{F} &= q[\mathbf{E} + \mathbf{v} \times \mathbf{B}] && \leftarrow \text{fuerza de Lorentz}\end{aligned}\quad (80)$$

Las densidades de flujo eléctrico $\mathbf{D}$ y magnético $\mathbf{B}$ se relacionan con las intensidades de campo eléctrico $\mathbf{E}$ y magnético $\mathbf{H}$ mediante la permitividad ϵ y la permeabilidad μ :

$$\begin{aligned}\bar{\mathbf{D}} &= \epsilon \bar{\mathbf{E}} = \epsilon_r \epsilon_0 \bar{\mathbf{E}} && \leftarrow \epsilon_0 = 8.854 \times 10^{-12} \approx \frac{10^{-9}}{36\pi} \text{ (F/m)} \\ \bar{\mathbf{B}} &= \mu \bar{\mathbf{H}} = \mu_r \mu_0 \bar{\mathbf{H}} && \leftarrow \mu_0 = 4\pi \times 10^{-7} \text{ (H/m)} \\ \text{donde : } \epsilon_r &\leftarrow \text{constante dieléctrica,} && \mu_r \leftarrow \text{permeabilidad relativa} \\ c &= 1/\sqrt{\epsilon_0 \times \mu_0} && \leftarrow \text{constante de la luz}\end{aligned}\quad (81)$$

La ecuación (a) es un caso especial de la **ley de inducción de Faraday**, y representa la *fuerza electromotriz inducida* (fem) en una *espira cerrada estacionaria* debida a una tasa de cambio de la densidad de flujo magnético respecto al tiempo y se conoce como *fem de transformación*. Si un conductor se mueve a una *velocidad v* en un *campo magnético B*, se induce una fem conocida como *fem movimiento*. La *fem total inducida de una espira que se mueve en un campo magnético*, es la suma de la *fem de transformación* y la *fem de movimiento*:

$$e = e_t + e_m = -\int_S \frac{\partial \bar{\mathbf{B}}}{\partial t} \cdot d\bar{\mathbf{S}} + \int_C (\bar{\mathbf{v}} \times \bar{\mathbf{B}}) \cdot d\bar{\mathbf{l}}\quad (82)$$

En términos del *flujo magnético total que pasa a través de la espira* ϕ , la ecuación anterior se escribe en forma compacta como:

$$e = -\frac{d\phi}{dt} \quad \leftrightarrow \quad \phi = \int_S \bar{\mathbf{B}} \cdot d\bar{\mathbf{S}}\quad (83)$$

Ejemplo 1.12. Una *espira cuadrada* (figura 1.1) con lados de 10 cm (0.1 m) de longitud está en un campo magnético con variación senoidal de intensidad 100 A/m y frecuencia de 50 MHz. El plano de la espira es perpendicular a la

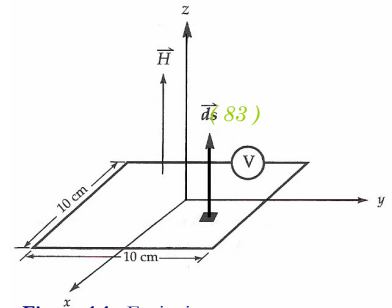


Figura 1.1 Espira inmersa en un campo magnético con variación senoidal $\mathbf{B}$

dirección del campo magnético. Si se conecta un *voltímetro* en serie con la espira, ¿cuál es su lectura?

Solución. Puesto que la espira es estacionaria, la *fem inducida* se debe a la *fem de transformación*, considerando la *permeabilidad del espacio libre* μ_0 y sabiendo que la lectura del voltímetro será un valor *rms* se tiene que:

$$\begin{aligned} H &= 100 \sin \omega t \hat{z} [\text{A/m}] \rightarrow \vec{B} = \mu_0 H = 100 \mu_0 \sin \omega t \hat{z} [\text{T}] \\ \omega &= 2\pi f, \quad f = 50 \times 10^6 \\ \frac{\partial \vec{B}}{\partial t} &= 100 \mu_0 \omega \cos \omega t \hat{z} = 100 \times 4\pi \times 10^{-7} \times 2\pi \times 50 \times 10^6 \cos \omega t \hat{z} \\ \frac{\partial \vec{B}}{\partial t} &= 39,478.4 \cos \omega t \hat{z} \\ d\vec{S} &= dx dy \hat{z} \\ e &= -39,478.4 \cos \omega t \int_{-0.05}^{0.05} dx \int_{-0.05}^{0.05} dy \\ e &= -39,478.4 \cos \omega t \times [0.05 + 0.05] \times [0.05 + 0.05] = -394.784 \cos \omega t \\ \rightarrow V_{rms} &= \frac{394.784}{\sqrt{2}} = 279.15 [\text{V}] \quad \clubsuit \end{aligned}$$

La ecuación (b) es la definición matemática de la **ley de Ampere**. Establece que la integral de línea de la intensidad de campo magnético alrededor de un ciclo cerrado es igual a la corriente total encerrada. La corriente total es la suma de la corriente de conducción I_c y la corriente de desplazamiento I_d . En un conductor por el que fluye corriente a bajas frecuencias, la corriente de desplazamiento es mínima.

$$I = I_c + I_d = \int_s \vec{J} \cdot d\vec{S} + \int_s \frac{\partial \vec{D}}{\partial t} \cdot d\vec{S} \quad (84)$$

La **ley de fuerza de Ampere** establece que un conductor por el que circula una corriente I inmerso en un campo magnético $\vec{B}$, recibe una fuerza dada por:

$$\vec{F} = \int_c I d\vec{l} \times \vec{B} \quad (85)$$

De esta manera para el trozo de conductor de longitud $2b$ mostrado en la **figura 1.2**, portador de una corriente I , situado en un campo magnético $\vec{B}$, siente una

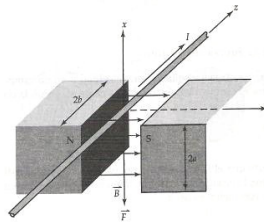


Figura 1.2 Fuerza en un conductor que conduce una corriente I en un campo magnético uniforme $\vec{B}$

fuerza en la dirección $-x$ y esta dada por:

$$\begin{aligned}\vec{F} &= \int_C I d\vec{l} \times \vec{B} = \int_{-b}^b IB dz (\hat{a}_z \times \hat{a}_y) \\ \vec{F} &= IB[b+b](-\hat{a}_x) = -2IBb\hat{a}_x [\text{N}]\end{aligned}\quad (86)$$

Consideremos la espira de alambre mostrada en la **figura 1.3**, con 2 lados de longitud L (paralelos al eje X) y dos lados de longitud W , y teniendo capacidad para girar libremente, inmersa en un *campo magnético constante* $\mathbf{B}$ en dirección Z ; sea $\mathbf{a}_n$ el vector unitario normal a su superficie; sea $\mathbf{I}$ la *corriente que circula en su interior*.

Cuando θ es el ángulo que forma $\mathbf{a}_n$ con el vector de campo $\mathbf{B}$, la fuerza que sienten los lados L_1 y L_2 de la espira esta dada por:

$$F_a = F_b = -BIL\hat{a}_y \quad (87)$$

De esta manera los pares $\mathbf{r} \times \mathbf{F}$ ejercidos sobre los conductores L_1 y L_2 generan un par total en la espira que está dada por:

$$\vec{T}_a = \vec{T}_b = BIL(W/2) \sin \theta \hat{a}_x, \quad \vec{T} = BIA \sin \theta \hat{a}_x \quad \leftrightarrow \quad A = LW \quad (88)$$

La anterior es la *ecuación fundamental que rige el desarrollo del par en todas las máquinas eléctricas*, para N espiras se tiene: $\vec{T} = BIAN \sin \theta \hat{a}_x$.

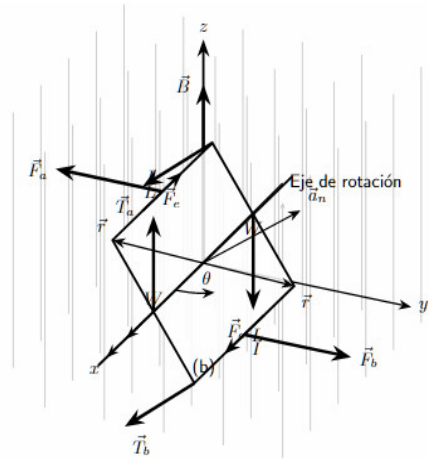


Figura 1.3 Espira inmersa en un campo magnético uniforme $\mathbf{B}$, formando un ángulo θ

Transformación de sistemas coordenados

Las coordenadas de un punto P, en los tres *sistemas ortogonales* de mayor uso en la ingeniería son: *rectangular* (R) $P(x,y,z)$, *cilíndrico* (C) $P(\rho,\theta,z)$ y *esférico* (E) $P(r,\phi,\theta)$. Los *vectores unitarios* $(\mathbf{a}_1, \mathbf{a}_2, \mathbf{a}_3)$, de los 3 sistemas son: *rectangular* $(\mathbf{a}_x, \mathbf{a}_y, \mathbf{a}_z)$, *cilíndrico* $(\mathbf{a}_\rho, \mathbf{a}_\theta, \mathbf{a}_z)$ y *esférico* $(\mathbf{a}_r, \mathbf{a}_\phi, \mathbf{a}_\theta)$, tienen magnitud unitaria, y su dirección es la de los ejes coordenados positivos del sistema; mediante estos el vector $\mathbf{A}$ con origen $O(0,0,0)$ de componentes (A_1, A_2, A_3) , se representa por: $\mathbf{A} = A_1\mathbf{a}_1 + A_2\mathbf{a}_2 + A_3\mathbf{a}_3$.

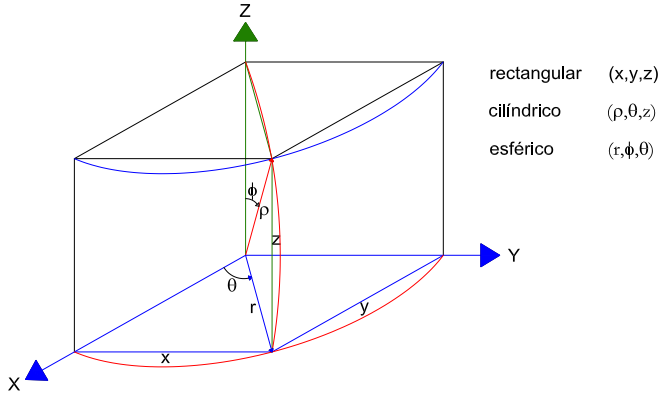


Figura 1.4 Sistemas coordenados ortogonales (rectangular, cilíndrico y esférico)

Las relaciones entre las coordenadas de los sistemas se obtienen directamente de la **figura 1.4** aplicando las *funciones trigonométricas* correspondientes y el *teorema de Pitágoras*. Por ejemplo:

$$r = \sqrt{x^2 + y^2} \leftrightarrow \tan \theta = \frac{y}{x} \rightarrow \theta = \tan^{-1} \frac{y}{x} \leftrightarrow \rho = \sqrt{x^2 + y^2 + z^2} \leftrightarrow \cos \phi = \frac{z}{\rho} \rightarrow \phi = \cos^{-1} \frac{z}{\rho} \quad (89)$$

Las componentes de $\mathbf{A}$ en el *sistema coordinado de vectores unitarios* $\{\mathbf{b}_1, \mathbf{b}_2, \mathbf{b}_3\}$ se pueden hallar mediante la siguiente *relación de transformación*:

$$\begin{bmatrix} B_1 \\ B_2 \\ B_3 \end{bmatrix} = \begin{bmatrix} b_1 \cdot a_1 & b_2 \cdot a_1 & b_3 \cdot a_1 \\ b_1 \cdot a_2 & b_2 \cdot a_2 & b_3 \cdot a_2 \\ b_1 \cdot a_3 & b_2 \cdot a_3 & b_3 \cdot a_3 \end{bmatrix} \begin{bmatrix} A_1 \\ A_2 \\ A_3 \end{bmatrix} \quad (90)$$

La ecuación de transformación anterior puede escribirse matricialmente como $\mathbf{B}=\mathbf{MA}$, siendo $\mathbf{A}$ y $\mathbf{B}$ *vectores columna de las componentes*, y $\mathbf{M}$ la *matriz de los productos escalares de los vectores unitarios*, obsérvese que $\mathbf{A}=\mathbf{M}^{-1}\mathbf{B}$. El listado siguiente proporciona las *ecuaciones matriciales de transformación coordenada* para la conversión entre los 3 sistemas expuestos.

$$\begin{aligned}
 C \leftarrow R \quad \begin{bmatrix} A_r \\ A_\phi \\ A_z \end{bmatrix} &= \begin{bmatrix} \cos \theta & \sin \theta & 0 \\ -\sin \theta & \cos \theta & 0 \\ 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} A_x \\ A_y \\ A_z \end{bmatrix} \\
 R \leftarrow C \quad \begin{bmatrix} A_x \\ A_y \\ A_z \end{bmatrix} &= \begin{bmatrix} \cos \theta & -\sin \theta & 0 \\ \sin \theta & \cos \theta & 0 \\ 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} A_r \\ A_\phi \\ A_z \end{bmatrix} \\
 E \leftarrow R \quad \begin{bmatrix} A_\rho \\ A_\phi \\ A_\theta \end{bmatrix} &= \begin{bmatrix} \sin \phi \cos \theta & \sin \phi \sin \theta & \cos \phi \\ -\cos \phi \cos \theta & \cos \phi \sin \theta & -\sin \phi \\ -\sin \phi & \cos \theta & 0 \end{bmatrix} \begin{bmatrix} A_x \\ A_y \\ A_z \end{bmatrix} \\
 R \leftarrow E \quad \begin{bmatrix} A_x \\ A_y \\ A_z \end{bmatrix} &= \begin{bmatrix} \sin \phi \cos \theta & \cos \phi \cos \theta & -\sin \theta \\ \sin \phi \sin \theta & \cos \phi \sin \theta & \cos \theta \\ \cos \theta & -\sin \phi & 0 \end{bmatrix} \begin{bmatrix} A_\rho \\ A_\phi \\ A_\theta \end{bmatrix} \\
 E \leftarrow C \quad \begin{bmatrix} A_\rho \\ A_\phi \\ A_\theta \end{bmatrix} &= \begin{bmatrix} \sin \phi & 0 & \cos \phi \\ \cos \phi & 0 & -\sin \phi \\ \cos \theta & -\sin \phi & 0 \end{bmatrix} \begin{bmatrix} A_r \\ A_\phi \\ A_z \end{bmatrix} \\
 C \leftarrow E \quad \begin{bmatrix} A_r \\ A_\theta \\ A_z \end{bmatrix} &= \begin{bmatrix} \sin \phi & \cos \phi & 0 \\ 0 & 0 & 1 \\ \cos \phi & -\sin \phi & 0 \end{bmatrix} \begin{bmatrix} A_\rho \\ A_\phi \\ A_\theta \end{bmatrix} \tag{91}
 \end{aligned}$$

Transformación fasorial

La *transformación fasorial* se aplica únicamente a ondas *sinusoidales* (*coseno* o *seno*) y da como resultado una *notación de variable compleja*, fácilmente manejable mediante algebra lineal. La *transformación fasorial se emplea ampliamente en sistemas eléctricos debido a que la energía eléctrica se genera y transmite en forma de onda sinusoidal*, la frecuencia típica en México es de $f=60$ Hz ($\omega=2\pi f=2\pi \cdot 60 \approx 377$ rad/s) .

La transformación fasorial se basa en la *identidad de Euler*, obtenida por observación de *series de Maclaurin* de las *funciones seno, coseno y exponencial*:

$$\begin{aligned}\cos(\theta) &= 1 - \frac{\theta^2}{2!} + \frac{\theta^4}{4!} - \frac{\theta^6}{6!} + \dots = \sum_{k=0}^{\infty} (-1)^k \frac{\theta^{2k}}{(2k)!} \\ \sin(\theta) &= \theta - \frac{\theta^3}{3!} + \frac{\theta^5}{5!} - \frac{\theta^7}{7!} + \dots = \sum_{k=0}^{\infty} (-1)^k \frac{\theta^{2k+1}}{(2k+1)!} \\ j \sin(\theta) &= j\theta - j \frac{\theta^3}{3!} + j \frac{\theta^5}{5!} - j \frac{\theta^7}{7!} + \dots = j \sum_{k=0}^{\infty} (-1)^k \frac{\theta^{2k+1}}{(2k+1)!} \\ e^{\theta} &= 1 + \theta + \frac{\theta^2}{2!} + \frac{\theta^3}{3!} + \frac{\theta^4}{4!} + \dots = \sum_{k=0}^{\infty} \frac{\theta^k}{k!} \\ e^{j\theta} &= 1 + j\theta - \frac{\theta^2}{2!} - j \frac{\theta^3}{3!} + \frac{\theta^4}{4!} + \dots\end{aligned}\tag{92}$$

La observación de que la función $e^{j\theta}$ equivale a la suma de las funciones $\cos(\theta)$ y $j \sin \theta$, condujo a Euler a la siguiente relación conocida como *identidad de Euler*:

$$e^{j\theta} = \cos \theta + j \sin \theta \leftrightarrow e^{-j\theta} = \cos \theta - j \sin \theta\tag{93}$$

De las relaciones anteriores se obtiene la forma exponencial compleja para las funciones seno y coseno:

$$\begin{aligned}e^{j\theta} &= \cos \theta + j \sin \theta & e^{j\theta} &= \cos \theta + j \sin \theta \\ + e^{-j\theta} &= \cos \theta - j \sin \theta & -e^{-j\theta} &= -\cos \theta + j \sin \theta \\ \frac{e^{j\theta} + e^{-j\theta}}{2} &= \cos \theta & \frac{e^{j\theta} - e^{-j\theta}}{2j} &= \sin \theta \\ \Rightarrow \cos \theta &= \frac{1}{2}(e^{j\theta} + e^{-j\theta}) & \Rightarrow \sin \theta &= \frac{1}{2j}(e^{j\theta} - e^{-j\theta})\end{aligned}$$

Valor medio cuadrático de la función senoidal

Una *función senoidal* (seno o coseno) de *frecuencia angular* $\omega=2\pi/T$ y de *periodo* $T=2\pi/\omega$, alterna su sentido cada $T/2$ segundos (con $f=60\text{Hz}$, $T=1/f \approx 16.67\text{ ms}$); puesto que la función es simétrica, el *valor promedio* de la función es cero. Una alternativa al valor promedio de la función senoidal, es el *valor medio cuadrático RMS (root mean square)* o *valor eficaz*, definido por:

$$\begin{aligned}i(t) &= I_{\max} \cos(\omega t + \phi) \\ I_{\text{rms}} &= \sqrt{\frac{1}{T} \int_0^T i^2 dt} = \sqrt{\frac{1}{T} \int_0^T [I_{\max} \cos(\omega t + \phi)]^2 dt}\end{aligned}\tag{94}$$

La evaluación de la integral resulta fácil empleando *identidades trigonométricas*:

$$\cos^2(\omega t + \phi) = \frac{1}{2}[1 + \cos(2\omega t + 2\phi)]$$

$$\text{sen}(n\pi + \phi) = \text{sen}(\phi) \quad n \in \{\mathbb{Z}^+\}$$

$$I_{rms} = \sqrt{\frac{1}{2\pi/\omega} \int_0^{2\pi/\omega} \frac{1}{2} [1 + \cos(2\omega t + 2\phi)] dt} = \sqrt{\frac{I_{max}^2}{2\pi/\omega} \frac{1}{2} \int_0^{2\pi/\omega} [1 + \cos(2\omega t + 2\phi)] dt}$$

$$I_{rms} = \frac{I_{max}}{2} \sqrt{\frac{\omega}{\pi} \int_0^{2\pi/\omega} [1 + \cos(2\omega t + 2\phi)] dt} = \frac{I_{max}}{2} \sqrt{\frac{\omega}{\pi} \left[t + \frac{1}{2\omega} \text{sen}(2\omega t + 2\phi) \right]_0^{2\pi/\omega}}$$

$$I_{rms} = \frac{I_{max}}{2} \sqrt{\frac{\omega}{\pi} \left[\frac{2\pi}{\omega} + \frac{1}{2\omega} (\text{sen}(4\pi + 2\phi) - \text{sen}(2\phi)) \right]} = \frac{I_{max}}{2} \sqrt{\frac{\omega}{\pi} \frac{2\pi}{\omega}} = \frac{I_{max}}{2} \sqrt{2} = \frac{\sqrt{2} I_{max}}{\sqrt{2} \sqrt{2}} = \frac{I_{max}}{\sqrt{2}} \approx 0.707 I_{max}$$

Transformación fasorial

La *transformación fasorial* consiste en representar una *función cosenoidal* en el dominio del tiempo $i(t)$, de *amplitud* I_{max} , *periodo* T , *ángulo de fase* ϕ , y *frecuencia angular* $\omega = 2\pi/T$, como un *número complejo* en su forma polar $I(\omega)$, el cual se dice que esta en *dominio de la frecuencia*. Es decir, la *transformación fasorial* transforma una *función senoidal* en una *función exponencial compleja* (cuya base es la identidad de Euler):

$$i(t) = I_{max} \cos(\omega t + \phi) \Rightarrow \tilde{I} = I_{rms} \angle \phi$$

$$I_{rms} = \frac{I_{max}}{\sqrt{2}} \approx 0.707 I_{max} \quad (95)$$

La *transformación fasorial inversa* consiste en dado un fasor de *frecuencia angular* ω , *valor efectivo* o *valor medio cuadrático* I_{rms} , *ángulo de fase* ϕ expresarlo en el dominio del tiempo $i(t)$, como una función cosenoidal:

$$\tilde{I} = I_{rms} \angle \phi \Rightarrow i(t) = I_{max} \cos(\omega t + \phi)$$

$$I_{max} = \sqrt{2} I_{rms} \approx 1.414 I_{rms} \quad (96)$$

Si la función esta expresada en términos de la función senoidal es fácil convertirla en función cosenoidal, recordando las siguientes identidades trigonométricas:

$$\begin{array}{llll} \sin(-\theta) = -\sin(\theta) & \sin(\pi/2 \pm \theta) = \cos(\theta) & \text{y} & \cos(-\theta) = \cos(\theta) \quad \cos(\pi/2 - \theta) = \sin(\theta) \\ \sin(\theta + \pi/2) = \cos(\theta) & \sin(\theta - \pi/2) = -\cos(\theta) & & \cos(\theta - \pi/2) = \sin(\theta) \quad \cos(\pi/2 + \theta) = -\sin(\theta) \\ \sin(\pi - \theta) = \sin(\theta) & \sin(\theta - \pi) = -\sin(\theta) & & \cos(\pi - \theta) = -\cos(\theta) \quad \cos(\theta - \pi) = -\cos(\theta) \end{array}$$

Una vez transformado en fasor, el número complejo resultante obedece a todas las leyes y principios de números complejos ordinarios. Por lo antes dicho es fácil obtener la representación fasorial para la derivada e integral de la función cosenoidal, como se muestra a continuación:

$$\begin{aligned}
 i &= I_{\max} \cos(\omega t + \phi) \rightarrow \tilde{I} = I_{rms} \angle \phi = I_{rms} e^{(j\omega t + \phi)} \\
 \frac{di}{dt} &= \frac{d(I_{rms} e^{(j\omega t + \phi)})}{dt} = j\omega I_{rms} e^{(j\omega t + \phi)} \rightarrow j\omega I_{rms} \angle \phi = \omega I_{rms} \angle \phi + 90^\circ \\
 \int i dt &= \int I_{rms} e^{(j\omega t + \phi)} dt = \frac{I_{rms} e^{(j\omega t + \phi)}}{j\omega} \rightarrow \frac{1}{j\omega} I_{rms} \angle \phi = -j \frac{1}{\omega} I_{rms} \angle \phi = \frac{1}{\omega} I_{rms} \angle \phi - 90^\circ
 \end{aligned} \tag{97}$$

Transformación fasorial en circuitos eléctricos de corriente alterna

La aplicación más extendida de los fasores es en circuitos eléctricos sujetos a funciones de excitación senoidales, conocidos como circuitos de corriente alterna **ca**. Considérese el *circuito RLC serie*, mostrado en la **figura 1.5**. Hallar el modelo del circuito y obtener su solución.

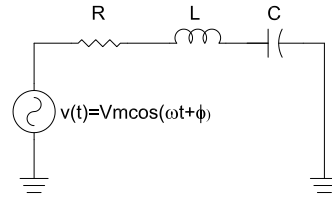


Figura 1.5 Circuito RLC serie

Aplicando la LVK a la malla y la convención pasiva de los signos se tiene la siguiente ecuación diferencial, cuya primitiva $i(t)$, es su solución:

$$-V_m \cos(\omega t + \alpha) + Ri + L \frac{di}{dt} + \frac{1}{C} \int i dt = 0$$

Considerando que la función de excitación es senoidal $v(t)$, y el circuito es lineal, entonces se espera una respuesta $i(t)$ senoidal, de esta forma aplicando *transformación fasorial* a la *ecuación diferencial* obtenida, resulta en:

$$L \frac{di}{dt} + \frac{1}{C} \int idt + Ri = V_m \cos(\omega t + \alpha) \rightarrow j\omega L I_{rms} \angle \beta + \frac{1}{j\omega C} I_{rms} \angle \beta + R I_{rms} \angle \beta = V_{rms} \angle \alpha$$

$$I_{rms} \angle \beta * \left(j(\omega L - 1/\omega C) + R \right) = V_{rms} \angle \alpha \rightarrow I_{rms} \angle \beta = \frac{1}{R + j(\omega L - 1/\omega C)} V_{rms} \angle \alpha$$

$$I_{rms} \angle \beta = \frac{1}{K \angle \theta} V_{rms} \angle \alpha = \frac{V_{rms}}{K} \angle (\alpha - \theta)$$

donde :

$$K = \sqrt{R^2 + (\omega L - 1/\omega C)^2}$$

$$\theta = \arctan[(\omega L - 1/\omega C) / R]$$

Aplicando la *transformación fasorial inversa*, resulta:

$$i(t) = \frac{\sqrt{2}}{K} V_{rms} \cos(\omega t + \alpha - \theta)$$

Al aplicar la transformación fasorial a los circuitos eléctricos con fuentes de excitación senoidal (*corriente alterna senoidal*), el análisis se simplifica del manejo de ecuaciones integrodiferenciales al manejo de expresiones con números complejos.

El trabajo se simplifica aún más al aplicar la *ley de ohm fasorial*, a los bloques funcionales eléctricos pasivos (*resistor, inductor y capacitor*) sometidos a voltajes y corrientes fasoriales:

$$\begin{aligned} \tilde{V} &= V_{rms} \angle \alpha, & \tilde{I} &= I_{rms} \angle \beta \\ v &= Ri \Rightarrow \tilde{V} = R\tilde{I} \rightarrow \tilde{I} = \frac{1}{R} \tilde{V} \\ v &= L \frac{di}{dt} \Rightarrow \tilde{V} = j\omega L \tilde{I} \rightarrow \tilde{I} = -j \frac{1}{\omega L} \tilde{V} = \frac{1}{\omega L} V_{rms} \angle (\alpha - \pi/2) \\ v &= \frac{1}{C} \int idt \Rightarrow \tilde{V} = \frac{1}{j\omega C} \tilde{I} = -j \frac{1}{\omega C} \tilde{I} \rightarrow \tilde{I} = j\omega C \tilde{V} = \omega C V_{rms} \angle (\alpha + \pi/2) \end{aligned} \quad (98)$$

De las ecuaciones anteriores se tiene para los *elementos pasivos* que: en el **resistor**, la corriente y voltaje están en fase; en el **inductor** la corriente está 90° ($\pi/2$ rad) atrasada respecto al voltaje; en el **capacitor** la corriente está 90° ($\pi/2$ rad) adelantada respecto al voltaje.

En términos fasoriales, se denomina impedancia ($\mathbf{Z}$) a la razón del voltaje fasorial ($\mathbf{V}$) a la corriente fasorial ($\mathbf{I}$), es un numero complejo cuya parte real es la *resistencia* (R) y cuya parte imaginaria es la *reactancia* (X), que puede ser *reactancia inductiva* (X_L) o *reactancia capacitiva* (X_C). El inverso de la impedancia se denomina *admitancia* ($\mathbf{Y}$), su parte real es la *conductancia* (G) y su parte compleja la *susceptancia* (B). Tanto la impedancia, como la admitancia no son fasores, son números complejos puros:

$$\begin{aligned}\bar{V} &= \hat{Z}\bar{I} \rightarrow \hat{Z} = \frac{\bar{V}}{\bar{I}} \\ \hat{Z} &= R + jX, \quad X_L = \omega L = 2\pi fL, \quad X_C = \frac{1}{\omega C} = \frac{1}{2\pi fC} \\ \bar{I} &= \hat{Y}\bar{V} \rightarrow \hat{Y} = \frac{\bar{I}}{\bar{V}}, \\ \hat{Y} &= G + jB = \frac{1}{\hat{Z}} = \frac{1}{R + jX} = \frac{1 \angle 0}{\sqrt{R^2 + X^2} \angle (\tan^{-1} b/a)} = \left(1/\sqrt{R^2 + X^2}\right) \angle (-\tan^{-1} b/a)\end{aligned}\tag{99}$$

Para efectos de cálculo, la *reactancia inductiva* y la *reactancia capacitiva* pueden transformarse en *impedancias complejas* de la siguiente manera:

$$\begin{aligned}\hat{Z}_L &= jX_L = j\omega L = j2\pi fL = \omega L \angle 90^\circ \\ \hat{Z}_C &= -jX_C = -j\frac{1}{\omega C} = -j\frac{1}{2\pi fC} = \frac{1}{\omega C} \angle -90^\circ\end{aligned}\tag{100}$$

El equivalente de n impedancias conectadas en serie está dado por:

$$\hat{Z}_{eq} = \sum_{k=1}^n \hat{Z}_k = \sum_{k=1}^n R_k + j \sum_{k=1}^n X_k\tag{101}$$

El equivalente de n admitancias conectadas en paralelo está dado por:

$$\hat{Y}_{eq} = \sum_{k=1}^n \hat{Y}_k = \sum_{k=1}^n G_k + j \sum_{k=1}^n B_k\tag{102}$$

Resonancia

Se dice que un circuito de ca RLC en serie o en paralelo, está en *resonancia* a una *frecuencia de resonancia* f_R , cuando la *reactancia inductiva* (X_L) es igual a la *reactancia capacitiva* (X_C). En resonancia la *amplitud* de la corriente es máxima y se define α como la *frecuencia neperiana* o *coeficiente de amortiguamiento exponencial*.

$$X_L = X_C \rightarrow \omega L = \frac{1}{\omega C} \rightarrow \omega^2 LC = 1 \rightarrow \omega = \frac{1}{\sqrt{LC}} \Rightarrow f_R = \frac{1}{2\pi\sqrt{LC}}, \quad \alpha = \frac{1}{2RC} \quad (103)$$

Potencia y factor de potencia

La *potencia compleja* (S) de un elemento en circuito de ca, es el producto del *voltaje fasorial de sus terminales* (V), por el conjugado complejo de la *corriente fasorial* (I^*) que circula por el elemento. Su parte real se llama *potencia promedio* o *potencia real* (P), se mide en *watts* (W), su parte imaginaria se llama *potencia reactiva* (Q), se mide en *volts amperes reactivos* (VAR), la magnitud de la potencia compleja se denomina *potencia aparente* (S), se mide en *volts amperes* (VA). El coseno del ángulo entre el fasor de voltaje y corriente se denomina *factor de potencia* $FP = \cos(\theta)$.

$$\begin{aligned} \text{Sea } \bar{V} &= V_{rms} \angle \alpha, \quad \bar{I} = I_{rms} \angle \beta \\ \hat{S} &= \bar{V} \bar{I}^* = V_{rms} \angle \alpha * I_{rms} \angle -\beta = V_{rms} I_{rms} \angle (\alpha - \beta) \\ \hat{S} &= P + jQ = |\hat{S}| \cos(\theta) + j|\hat{S}| \sin(\theta) \\ P &= |\hat{S}| \cos(\theta), \quad Q = |\hat{S}| \sin(\theta) \rightarrow |\hat{S}| = V_{rms} I_{rms} = \sqrt{P^2 + Q^2} \\ \text{Si } \theta &= \alpha - \beta \\ FP &= \cos(\theta) = \cos(\alpha - \beta) = \frac{P}{|\hat{S}|} \rightarrow |\hat{S}| = \frac{P}{FP} \leftrightarrow \theta = \arccos(FP) \\ \therefore \hat{S} &= |\hat{S}| \angle \theta = \frac{P}{FP} \angle \arccos(FP) \end{aligned} \quad (104)$$

Cuando el fasor corriente adelanta el voltaje ($\beta > \alpha$) se dice que el *factor de potencia esta adelantado*, en caso contrario ($\beta < \alpha$), se dice que el *factor de potencia esta atrasado*. Cualquier decremento en el FP aumenta la corriente para la misma potencia activa, esto provoca mayores *pérdidas por efecto Joule* en los conductores eléctricos.

Por el *principio de conservación de energía* se puede establecer el **principio de las potencias parciales**: “la potencia compleja entregada a varias cargas interconectadas es la suma de las potencias complejas entregadas a cada una de las cargas individuales, independientemente de cómo estén conectadas (serie o paralelo)”.

Corrección del factor de potencia

Para evitar pérdidas por *efecto Joule* (disipación de calor en los resistores) el factor de potencia debe ser cercano a la unidad. El *Reglamento para el Suministro de Energía Eléctrica*, establece que el consumidor está obligado a mantener un factor de potencia tan cercano a la unidad como sea posible. Existe una multa por factor de potencia menor a 0.85, los valores deseables son de 0.90 y 0.95. Como se puede prever el factor de potencia se corrige conectando capacitores en paralelo con la carga inductiva, con la intención de que el voltaje de la carga no varíe.

La corrección del FP se puede plantear como sigue: se tiene una carga inductivo-resistiva conectada a la línea de voltaje V , consumiendo una potencia activa P , trabajando a un factor de potencia FP_1 , se desea aumentar el factor de potencia hasta FP_2 , conectando una carga capacitiva en paralelo. Determinar la potencia reactiva del capacitor (Q_C) y su capacitancia (C).

La solución se obtiene del siguiente razonamiento: el incremento del factor de potencia desde FP_1 hasta FP_2 mediante la conexión en paralelo de una carga capacitiva de capacitancia C , provocará que la potencia reactiva disminuya desde Q_1 hasta Q_2 , reduciendo el ángulo de fase desde θ_1 hasta θ_2 , sin alterar el voltaje de línea V y la potencia activa P consumida por la parte resistiva de la carga:

$$\begin{aligned}
 FP &= \cos \theta, & S &= \frac{P}{\cos \theta} = \frac{Q}{\sin \theta} \rightarrow Q = P \frac{\sin \theta}{\cos \theta} = P \tan \theta \\
 \theta_1 &= \arccos(FP_1), & Q_1 &= P \tan(\theta_1), & S_1 &= \frac{P}{FP_1} \\
 \theta_2 &= \arccos(FP_2), & Q_2 &= P \tan(\theta_2), & S_2 &= \frac{P}{FP_2} \\
 Q_C &= Q_1 - Q_2 = P [\tan(\cos^{-1}(FP_1)) - \tan(\cos^{-1}(FP_2))] \\
 Q_C &= \frac{V^2}{X_C} \rightarrow X_C = \frac{V^2}{Q_C} = \frac{1}{\omega C} = \frac{1}{2\pi f C} \\
 C &= \frac{Q_C}{2\pi f V^2} \quad \text{para } f = 60\text{Hz} \rightarrow C = \frac{Q_C}{377V^2} = \frac{P [\tan(\cos^{-1}(FP_1)) - \tan(\cos^{-1}(FP_2))]}{377V^2}
 \end{aligned}
 \tag{105}$$

 **Ejemplo 1.13.** Cierta carga toma una corriente de 10 A con un factor de potencia de 0.5 en atraso desde una fuente de 120 V y 60 Hz. Calcule el tamaño del capacitor para corregir el factor de potencia hasta 0.8.

Solución. Primero se calcula la potencia activa consumida a partir de la corriente y el voltaje, después se sustituyen los valores en la fórmula antes obtenida:

$$\begin{aligned}
 P &= VI \cos \theta = VI \times FP = 120 \times 10 \times 0.5 = 600W \\
 C &= \frac{P[\tan(\cos^{-1}(FP_1)) - \tan(\cos^{-1}(FP_2))]}{377V^2} \\
 &= \frac{600[\tan(\cos^{-1}(0.5)) - \tan(\cos^{-1}(0.8))]}{377 \times 120^2} = 1.0853 \times 10^{-4} F = 108.53 \mu F \quad \clubsuit
 \end{aligned}$$

Transformación de fuentes de ca

Leyes de circuitos. Una *red eléctrica* de cualquier complejidad operando en *estado estacionario sinusoidal* (ca) a bajas frecuencias, se puede representar *fasorialmente* por un conjunto de *elementos activos* (*fuentes de voltaje y corriente*) y *elementos pasivos* (*impedancias*). Los puntos donde 2 o más elementos tienen una *conexión común* de denominan *nodos*, *barras* o *buses*.

Un *lazo* es una *trayectoria cerrada*, una *mall*a es un lazo que no aloja ningún otro lazo. El *análisis de circuitos eléctricos lineales* se sustenta en *dos leyes de conservación* (LVK, LCK), *el principio de superposición* (PS) y *dos teoremas de transformación* (TT, TN).

Ley de voltajes de Kirchhoff (LVK): “*en una malla la suma de voltajes es igual a cero*”. **Ley de corrientes de Kirchhoff** (LCK): “*en un nodo la suma de corrientes es igual a cero*”; estas leyes expresan la *conservación de la energía* y la *conservación de la carga* respectivamente.

La *polaridad de voltaje* y *circulación de corriente* en los elementos de circuito se basa en la *convención pasiva de los signos* (en las fuentes de voltaje la corriente circula de $-$ a $+$, en las impedancias de $+$ a $-$). Se dice que una *red está muerta o en estado pasivo* cuando no actúa ninguna fuente de voltaje o corriente; es decir, las *fuentes de corriente se reemplacen por circuito abierto*, y las *fuentes de voltaje por cortocircuito*.

$$\sum_{k=1}^n \hat{V}_k = 0 \leftarrow \text{Ley de voltajes de Kirchhoff}$$

$$\sum_{k=1}^n \hat{I}_k = 0 \leftarrow \text{Ley de corrientes de Kirchhoff} \quad (106)$$

Principio de superposición y transformación de fuentes. En estado estable es posible aplicar el **principio de superposición** (PS), donde las *funciones de excitación* son las *fuentes de corriente y voltaje*, actuando una a la vez; la *respuesta total*, es la *suma de las respuestas parciales*.

Una *impedancia Z*, se *transforma* en una *admitancia Y=1/Z*, e inversamente $Z=1/Y$. Una *fente de voltaje V*, en *serie con una impedancia Z*, se *transforma en una fuente de corriente I=V/Z*, en *paralelo con la misma impedancia*.

Una *fente de corriente I*, en *paralelo con una impedancia Z*, se *transforma en una fuente de voltaje V=ZI*, en *serie con la misma impedancia*.

Teoremas de circuitos. Cuando solamente interesa conocer la respuesta en un elemento del circuito, se reemplaza el circuito de poco interés por un equivalente, esto se sustenta en dos teoremas:

Teorema de Thévenin (TT): “*toda red activa lineal de dos terminales a-b, equivale a la conexión en serie de una fuente de voltaje y una impedancia; la fuente de voltaje suministra un voltaje igual al existente entre las terminales de la red original en circuito abierto y la impedancia, es la impedancia equivalente de la red muerta vista desde las terminales a-b*”.

Teorema de Norton (TN): “*toda red activa lineal de dos terminales a-b, equivale a la conexión en paralelo de una fuente de corriente y una impedancia; la fuente de corriente suministra la corriente igual a la corriente que suministraría la red original cuando las terminales a-b están en cortocircuito, la impedancia, es la impedancia equivalente de la red muerta vista desde las terminales a-b*”.

$$\sum_{k=1}^n \hat{V}_k = \hat{V}_T \leftarrow \text{Principio de superposición}$$

$$\sum_{k=1}^n \hat{I}_k = 0 \leftarrow \text{Ley de corrientes de Kirchhoff} \quad (107)$$

División de voltaje. Consideremos un conjunto de n impedancias conectadas en serie cuyo voltaje total en sus terminales es $\mathbf{V}_T$, y cuya impedancia equivalente es $\mathbf{Z}_T$, entonces el voltaje en la k -ésima impedancia $\mathbf{Z}_k$, es una fracción del voltaje total:

$$\mathcal{V}_k = (\mathbf{Z}_k / \mathbf{Z}_T) \cdot \mathcal{V}_T \quad (108)$$

División de corriente. Supongamos que una corriente $\mathbf{I}_T$ alimenta un nodo al que se conectan n admitancias en paralelo, cuya admitancia equivalente es $\mathbf{Y}_T$ entonces la corriente que fluye por la rama que contiene la k -ésima admitancia $\mathbf{Y}_k$ es una fracción de la corriente total:

$$\mathcal{I}_k = (\mathbf{Y}_k / \mathbf{Y}_T) \cdot \mathcal{I}_T \quad (109)$$

Análisis de nodos. Considere un circuito de ca con n nodos operando en estado estable, con fuentes de tensión expresadas fasorialmente $\mathbf{E}_i$ e impedancias $\mathbf{Z}_i$. El análisis sistemático de nodos consiste en hallar los voltajes en cada nodo respecto de un nodo de referencia (designado comúnmente como tierra) y se realiza como sigue:

- (1) Tomar un nodo de referencia y asignarle 0, numerar los nodos, entonces el voltaje fasorial del i -ésimo nodo respecto al nodo de referencia es $\mathbf{V}_{i0}$.
- (2) Transformar las fuentes de voltaje en serie con una impedancia, en fuente de corriente $\mathbf{I}_i$ en paralelo con una impedancia, transformar todas las impedancias en admitancias $\mathbf{Y}_i$.
- (3) Escribir las ecuaciones nodales matriciales de la red de la forma $(\mathbf{Y}_{ij})_{n \times n}(\mathbf{V}_i)_n = (\mathbf{I}_i)_n$; donde $(\mathbf{Y}_{ij})$ es la matriz de admitancias nodal, y sus términos de la diagonal principal $\mathbf{Y}_{kk}$ (autoadmitancia o admitancia de excitación) son la suma de las admitancias conectadas al nodo k , los demás términos $\mathbf{Y}_{ij}$ (admitancia mutua o de transferencia) son el negativo de las admitancias conectadas entre los nodos i y j ($i \neq j$). Puesto que $\mathbf{Y}_{ij} = \mathbf{Y}_{ji}$, la matriz de admitancias es simétrica.

- (3) Entonces el vector de voltajes nodales está dado por: $(\mathbf{V}_i) = (\mathbf{Y}_{ij})^{-1}(\mathbf{I}_i)$

Análisis de mallas. Es el dual del análisis de nodos y sistema matricial formado tiene la forma $(\mathbf{Z}_{ij})_{n \times n}(\mathbf{I}_i)_n = (\mathbf{V}_i)_n$, las impedancias $\mathbf{Z}_{kk}$ son la suma de las impedancias alrededor de la malla k , mientras que $\mathbf{Z}_{ij}$ es la impedancia entre las mallas i y j . El vector de corrientes de malla está dado por: $(\mathbf{I}_i) = (\mathbf{Z}_{ij})^{-1}(\mathbf{V}_i)$.

Transformación delta y estrella

Cuando los voltajes y corrientes se expresan fasorialmente se supone que implícitamente la magnitud que acompaña al fasor es un *valor medio cuadrático* o *valor efectivo*, como se ha expuesto; es decir:

$$V_{\text{ef}} = V \angle \theta \Leftrightarrow V = V_{\text{rms}} = V_{\text{max}} / \sqrt{2} \quad (110)$$

Circuitos trifásicos

Los *circuitos eléctricos* pueden ser *monofásicos* o *polifásicos* (*bifásicos*, *trifásicos*, *tetrafásicos* etc.) según tengan una o mas de una fase. Las *fuentes polifásicas* se forman conectando *fuentes simples* en configuraciones especiales, el *número de fases es igual al número de fuentes*. Son de especial interés por su amplio uso los *circuitos trifásicos* que constan de *fuentes trifásicas* y *cargas trifásicas* en configuración *estrella* (Y) o *delta* (Δ) existiendo las siguientes posibilidades *fuentes-carga*: Y-Y, Y- Δ , Δ -Y, Δ - Δ .

Fuentes y cargas trifásicas. Una *fuentes trifásica* es la conexión de 3 fuentes que cumplen dos condiciones: **a)** *el voltaje eficaz*, V ($V = V_{\text{rms}} = V_{\text{max}} / \sqrt{2}$) *de cada fuente es el mismo*; **b)** *el voltaje en cada fuente esta 120° fuera de fase respecto a las otras dos*.

La conexión trifásica puede realizarse en Y o Δ . Cada fuente se designa como una *fase* y se representa por una letra *a*, *b* o *c*; el voltaje medido entre la terminal positiva de cada fuente y el neutro, se denomina *voltaje de fase*; el voltaje medido entre dos fases se denomina *voltaje de línea*; la corriente que fluye a través de una fuente se denomina *corriente de fase* y la corriente que fluye por una línea se llama *corriente de línea*. Se tienen dos *secuencias de fase*: *positiva* o *abc* y *negativa* o *cba*.

La *secuencia positiva* (**figura 1.6a**) se distingue por que el voltaje de la *fase a* adelanta 120° el de la *fase b* y este a su vez adelanta 120° el de la *fase c*; es decir: $\angle v_a > \angle v_b > \angle v_c$.

La *secuencia negativa* (**figura 1.6b**) se distingue por que el voltaje de la *fase a* se atrasa en 120° respecto al de la *fase b* y este a su vez se atrasa en 120°

respecto al de la fase b ; es decir: $\angle v_c > \angle v_b > \angle v_a$; en ambas secuencias los ángulos medidos en sentido horario son positivos y en sentido antihorario negativos. En lo sucesivo se trabajara con la secuencia de fases positiva.

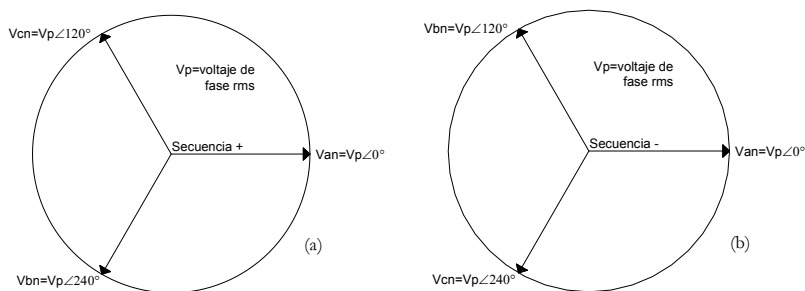


Figura 1.6 Fasores trifásicos de voltaje con secuencia de fases (a) positiva y (b) negativa

En la *configuración estrella* el *punto común* se designa por n (*neutro*) y generalmente se *aterrija* (*conexión a potencial 0*); los *voltajes de fase* o *voltajes al neutro* se designan por v_{an} , v_{bn} y v_{cn} o simplemente v_a , v_b y v_c ; en la *conexión Y* el *voltaje de línea adelanta en 30° el voltaje de fase y lo multiplica por $\sqrt{3}$* ; es decir:

$$V_l = \sqrt{3} \cdot V \quad \leftarrow \quad \text{Magnitud del voltaje de línea}$$

$$\vec{V}_{an} = V \angle 0 \quad \vec{V}_{bn} = V \angle -120 \quad \vec{V}_{cn} = V \angle 120$$

$$\begin{aligned} \vec{V}_{ab} &= \vec{V}_{an} - \vec{V}_{bn} = \sqrt{3}V \angle 30 = V_l \angle 30 \\ \vec{V}_{bc} &= \vec{V}_{bn} - \vec{V}_{cn} = \sqrt{3}V \angle -90 = V_l \angle -90 \\ \vec{V}_{ca} &= \vec{V}_{cn} - \vec{V}_{an} = \sqrt{3}V \angle 150 = V_l \angle 150 \end{aligned} \quad (111)$$

$$I_{na} = I_{aA} \quad I_{nb} = I_{bB} \quad I_{nc} = I_{cC} \quad \leftarrow \quad I_{na} + I_{nb} + I_{nc} = 0$$

En la *configuración delta* no existe punto común a las tres fases y tampoco hay presencia de neutro; en la *conexión Δ* la *corriente de línea se atrasa en 30° respecto a la de fase y la multiplica por $\sqrt{3}$* ; es decir:

$$\tilde{V}_{ab} = V \angle 0 \quad \tilde{V}_{bc} = V \angle -120 \quad \tilde{V}_{ca} = V \angle 120 \quad \leftarrow \quad \tilde{V}_{ab} + \tilde{V}_{bc} + \tilde{V}_{ca} = 0$$

$$I_l = \sqrt{3} \cdot I \quad \leftarrow \quad \text{Magnitud de la corriente de línea}$$

$$\tilde{I}_{ba} = I_l \angle \phi = \sqrt{3} \cdot I \angle \phi \leftarrow \text{Corriente de la fase 'ba'} \quad (112)$$

$$\tilde{I}_{aA} = \tilde{I}_{ba} - \tilde{I}_{ac} = \sqrt{3} \tilde{I}_{ba} \angle -30 = \sqrt{3} I \angle (\phi - 30) = I_l \angle (\phi - 30)$$

$$\tilde{I}_{bB} = I_l \angle (\phi - 150)$$

$$\tilde{I}_{cC} = I_l \angle (\phi + 90)$$

Las *cargas trifásicas* (*motores, transformadores...*) pueden conectarse en Δ o Y, resulta conveniente balancear las cargas trifásicas y mantener impedancias de igual magnitud por fase con el objeto de utilizar conductores de igual calibre para las líneas de transmisión, subtransmisión, distribución o utilización.

Análisis de circuitos trifásicos. El análisis de circuitos trifásicos con cargas balanceadas en configuración Y-Y, con neutros unidos se lleva a cabo por fase (figura 1.7), la unión de los neutros mediante un hilo conductor de impedancia cero no afecta el circuito. En el circuito Y-Y es suficiente el análisis de la *fase a*, de manera que las fases b y c se deducen de los resultados de esta. Sea Z_s el equivalente serie de las impedancias de línea y carga, entonces:

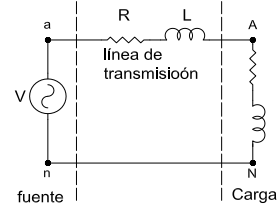


Figura 1.7 Circuito por fase

$$\hat{Z}_s = \hat{Z}_{\text{carga}} + \hat{Z}_{\text{línea}} \quad \tilde{I}_{aA} = \frac{\tilde{V}_a}{\hat{Z}_s} \quad \tilde{V}_{\text{línea}} = \hat{Z}_{\text{línea}} \tilde{I}_{aA} \quad \tilde{V}_{\text{carga}} = \hat{Z}_{\text{carga}} \tilde{I}_{aA} \quad (113)$$

Para las configuraciones que involucran Δ en el lado de la fuente o de la carga es necesario realizar transformaciones iniciales para el cálculo de las corrientes de línea; después se lleva el circuito a su forma original y se calculan los voltajes, las potencias y corrientes necesarias. Fuentes y cargas balanceadas pueden transformarse de Δ a Y y viceversa mediante:

$$\tilde{V}_Y = \frac{\tilde{V}_\Delta}{\sqrt{3}} \leftrightarrow \tilde{V}_\Delta = \sqrt{3} \tilde{V}_Y \quad (114)$$

$$\hat{Z}_Y = \frac{\hat{Z}_\Delta}{3} \leftrightarrow \hat{Z}_\Delta = 3 \hat{Z}_Y$$

Ejemplo 1.14. Una fuente trifásica balanceada (figura 1.8) de 707 V y 60Hz, conectada en Y, alimenta a una carga equilibrada conectada en Δ por medio de una línea de transmisión de 3 hilos de 100 km de longitud, la impedancia de cada hilo de la línea de transmisión es de $1+j2\Omega$. La impedancia de la carga por fase es $177-j246\Omega$. Si la secuencia de fases es positiva, determinar las corrientes de línea y de fase, la potencia absorbida por la carga y la potencia disipada por la línea de transmisión. V_{AB}

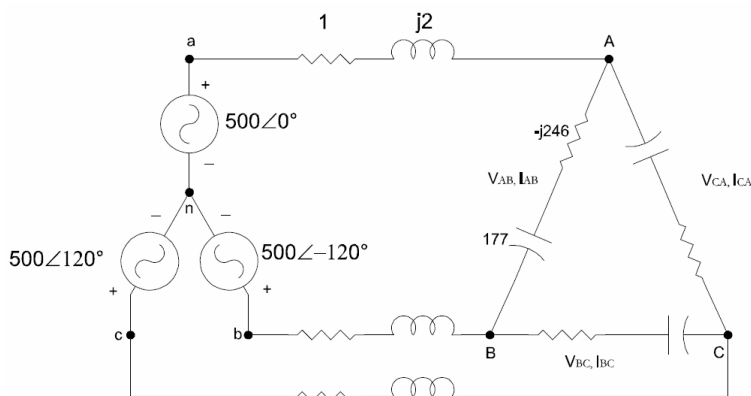


Figura 1.8 Circuito trifásico en conexión Y- Δ

Solución. Primero se expresa el voltaje de cada fuente como fasor, luego se transforma la carga Δ en carga Y, se procede a calcular la impedancia total, suma de la impedancia de la línea y la impedancia de carga, entonces se calcula la corriente de línea para el sistema monofásico resultante de la fase a y se determinan las restantes corrientes. Luego se determina la corriente de fase en Δ a partir de la corriente de línea mediante $I_{AB} = I_{na} \times (1/\sqrt{3})\angle 30^\circ$, así mismo los voltajes de fase. La potencia disipada por fase es la potencia real y la potencia total de la carga es 3 veces la de una fase; de igual forma se calcula la potencia disipada en la línea, en símbolos:

$$\tilde{V}_a = \frac{707}{\sqrt{2}} \angle 0 = 500 \angle 0 \quad \leftarrow \text{Voltaje de fase rms}$$

$$\tilde{Z}_Y = \frac{\tilde{Z}_\Delta}{3} = \frac{177 - j246}{3} = 59 - j82 \quad \leftarrow \text{Conversión carga } \Delta \text{ a } Y$$

$$\tilde{Z}_T = \tilde{Z}_L + \tilde{Z}_Y = (1 + j2) + (59 - j82) = 60 - j80 = 100 \angle -53.13$$

$$\text{fp} = \cos(\theta_Z) = \cos(-53.13) = 0.6 \text{ (adelantado)}$$

$$\tilde{I}_l = \frac{\tilde{V}_a}{\tilde{Z}_T} = \frac{500 \angle 0}{100 \angle -53.13} = 5 \angle 53.13 \quad \leftarrow \text{Corriente de línea en Y fase a}$$

$$\tilde{I}_{na} = 5 \angle 53.13$$

$$\tilde{I}_{nb} = 5 \angle (53.13 - 120) = 5 \angle -66.87$$

$$\tilde{I}_{nc} = 5 \angle (-66.87 - 120) = 5 \angle -186.87 = 5 \angle 173.13$$

$$\tilde{I}_{AB} = \frac{\tilde{I}_{na} \angle 30}{\sqrt{3}} = \frac{5 \angle 53.13 \angle 30}{\sqrt{3}} = 2.887 \angle 83.13 \quad \leftarrow \text{Corriente en la carga AB en } \Delta$$

$$\tilde{I}_{BC} = 2.887 \angle (83.13 - 120) = 2.887 \angle -36.87$$

$$\tilde{I}_{CA} = 2.887 \angle (-36.87 - 120) = 2.887 \angle -156.87$$

$$\tilde{V}_{AB} = \tilde{I}_{AB} \tilde{Z}_{AB} = 2.887 \angle 83.13 \times [177 - j246] = 874.93 \angle 28.87 \quad \leftarrow \text{Voltaje de carga en } \Delta$$

$$\tilde{V}_{BC} = 874.93 \angle (28.87 - 120) = 874.93 \angle -93.13$$

$$\tilde{V}_{CA} = 874.93 \angle (93.13 - 120) = 874.93 \angle 148.87$$

$$P_{AB} = I_{AB}^2 \times 177 = 1475.25 \quad \leftarrow \text{Potencia real de carga en } \Delta$$

$$P_{\text{carga}} = 3 \times P_{AB} = 3 \times 1475.25 = 4425.75$$

$$P_{\text{línea}} = 3 \times I_{aA}^2 \times 1 = 3 \times 5^2 \times 1 = 75$$

$$P_{\text{total}} = P_{\text{carga}} + P_{\text{línea}} = 4425.75 + 75 = 4500.75 \quad \clubsuit$$

Medición de la potencia. La potencia en un *circuito monofásico de cd* es el producto del voltaje por la corriente en la fuente o carga en cuestión, en un *circuito monofásico de ca* la potencia activa o real es el producto de los valores efectivos de voltaje y corriente; en el campo se emplea un *vatímetro* de dos bobinas: *bobina de corriente* y *bobina de potencial* para medir ambos, su lectura es la potencia activa del circuito y se calcula por:

$$P = V_{\text{rms}} I_{\text{rms}} \cos \theta = \text{Re}[\tilde{V} \tilde{I}^*]$$

(115)

La potencia de un circuito trifásico Y-Y de 4 hilos (con hilo neutro) **figura 1.9a**, se obtiene conectando 3 vatímetros entre cada fase y el neutro; entonces la potencia total es la suma de las tres lecturas y se calcula por:

$$\begin{aligned}
 P &= P_1 + P_2 + P_3 \\
 P_1 &= \text{Re}[\tilde{V}_{an} \tilde{I}_{aA}^*] \\
 P_2 &= \text{Re}[\tilde{V}_{bn} \tilde{I}_{bB}^*] \\
 P_3 &= \text{Re}[\tilde{V}_{cn} \tilde{I}_{cC}^*]
 \end{aligned} \quad (116)$$

Si la carga está balanceada la lectura será $P=3 \times P_1$.

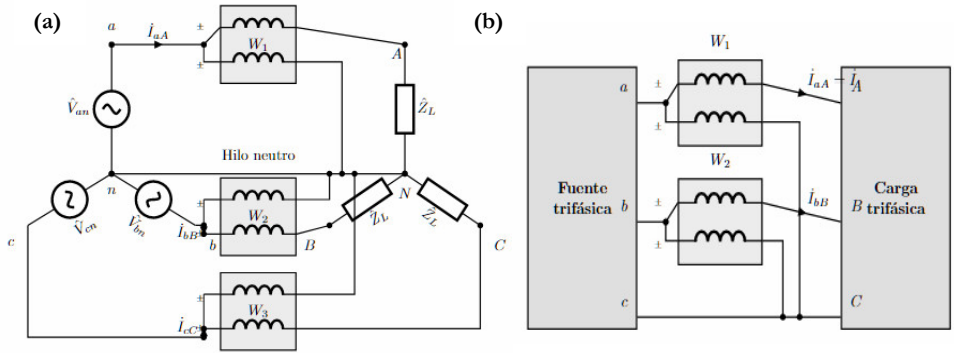


Figura 1.9 Medición de potencia en circuitos trifásicos (a) de 4 hilos; (b) de 3 hilos

Para un *circuito trifásico de 3 hilos* (que involucra una Δ en la fuente y/o en la carga) **figura 1.9b**, no se dispone de hilo neutro; en tal caso la potencia se mide con dos vatímetros y la potencia total es la suma algebraica de las lecturas en ambos.

$$\begin{aligned}
 P_1 &= \text{Re}[\tilde{V}_{ac} \tilde{I}_{aA}^*] = \text{Re}[(\sqrt{3}V\angle -30)(I\angle -\theta)] = \sqrt{3}VI \cos(30+\theta) \\
 P_2 &= \text{Re}[\tilde{V}_{bc} \tilde{I}_{bB}^*] = \text{Re}[(\sqrt{3}V\angle -90)(I\angle -\theta+120)] = \sqrt{3}VI \cos(30-\theta) \\
 P &= P_1 + P_2 = \sqrt{3}VI [\cos(30+\theta) + \cos(30-\theta)] = \sqrt{3}VI \cos\theta
 \end{aligned} \quad (117)$$

Transformación unidad

En los *sistemas eléctricos de potencia* donde se trabajan valores del orden de los kV y MVA resulta útil para efectos de cálculo y más informativo expresar las cantidades eléctricas como razones de sus cantidades bases eliminando así el uso de unidades, una cantidad por unidad se define por:

$$\text{Cantidad por unidad} = \frac{\text{Cantidad real}}{\text{Cantidad base}} \quad (118)$$

El voltaje V , la corriente I , la potencia S y la impedancia Z , se relacionan entre sí de manera tal que la selección de los valores base para dos de ellos determina los valores base de los dos restantes. El uso de valores por unidad simplifica el circuito equivalente del transformador, de tal forma que las tensiones, las corrientes, las impedancias y las admitancias externas expresadas por unidad, no cambian cuando se refieren de uno de los lados del transformador hacia el otro.

Transformación por unidad en circuitos ca

Para los *circuitos de ca monofásicos* (1Σ) o *trifásicos* (3Σ) las siguientes cantidades por unidad se definen de igual manera, considerando que en un circuito 1Σ el voltaje y la potencia son de fase; mientras que en un circuito 3Σ , el voltaje es línea a línea y la potencia es trifásica, $P_{3\Sigma} = 3P_{1\Sigma}$:

$$\begin{aligned} V_{pu} &= \frac{\text{voltaje real}}{\text{voltaje base, } V_b} & I_{pu} &= \frac{\text{corriente real}}{\text{corriente base, } I_b} \\ V_{A \text{ base}} &= \frac{\text{voltaje base, } V_b}{\text{corriente base, } I_b} & V_{Apu} &= \frac{V_{A \text{ real}}}{V_{A \text{ base}}}, & Z_{\text{real}} &= Z_{pu} \frac{V_b}{I_b} \end{aligned} \quad (119)$$

La impedancia por unidad para un circuito 1Σ y para un circuito 3Σ se definen por:

$$\begin{aligned} Z_{pu \ 1\phi} &= \frac{\text{impedancia real, } Z}{\text{impedancia base, } Z_b} = \frac{I_b \cdot Z}{I_b \cdot Z_b} = \frac{I_b \cdot Z}{V_b} \\ Z_{pu \ 3\phi} &= \frac{S_{b \ 3\phi} \cdot Z}{V_{b \ LL}^2} \end{aligned} \quad (120)$$

Para cambiar un valor de impedancia por unidad Z_{pu1} con valores base V_{b1} y S_{b1} a nuevos valores base de voltaje y potencia V_{b2} y S_{b2} se emplea la siguiente relación de transformación:

$$Z_{pu_2} = Z_{pu_1} \left(\frac{V_{b1}}{V_{b2}} \right)^2 \left(\frac{S_{b2}}{S_{b1}} \right) \quad (121)$$

Ejemplo 1.15. Para el circuito 1Σ de 3 zonas mostrado en la **figura 310**, dividido por dos transformadores T_1 y T_2 , cuyas *capacidades nominales* se muestran, y usando *valores base* de 30kVA y 240 V en la zona 1, dibujar en circuito por unidad y determinar la corriente de la carga. Para los transformadores, las *resistencias de los devanados* y las ramas de *admitancia en derivación* se desprecian.

Solución. (1) La potencia base para toda la red es la misma $S_b=30$ kVA; (2) se determinan los voltajes base en cada zona, sabiendo que $V_{b1}=240$ V y conociendo las relaciones de transformación; (3) se determinan las impedancias base; (4) se expresan las reactancias pu de los transformadores en los valores base del sistema; (5) se determinan las impedancias pu de las líneas de transmisión; (6) se determinan los valores del circuito necesarios (el circuito con valores por unidad se muestra en la **figura 1.11**):

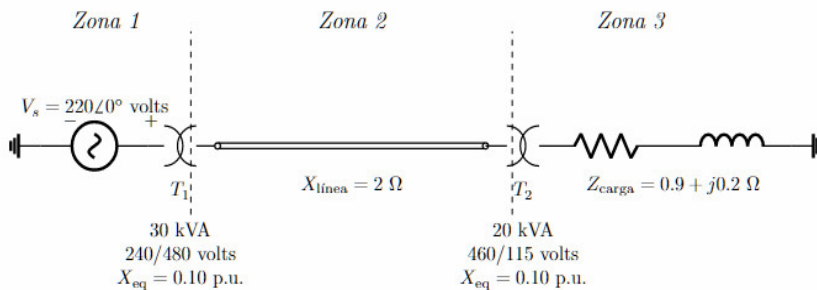
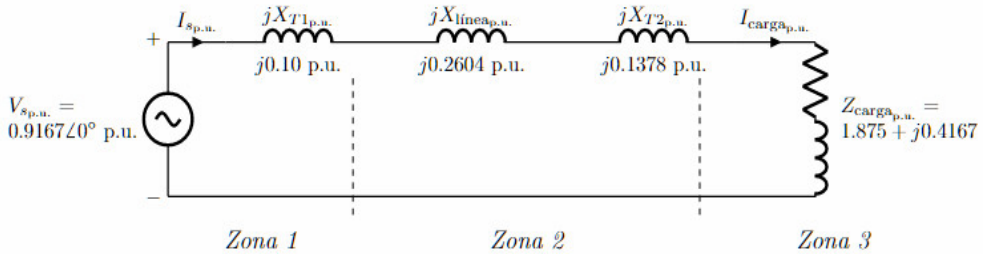


Figura 1.10 Circuito monofásico de 3 zonas

- (1) $S_b = 30,000 \text{ VA}$
 - (2) $V_{b_1} = 240 \text{ V}; \quad V_{b_2} = \frac{V_2}{V_1} V_{b_1} = \frac{480}{240} 240 = 480 \text{ V}; \quad V_{b_3} = \frac{V_3}{V_2} V_{b_2} = \frac{115}{460} 480 = 120 \text{ V}$
 - (3) $Z_{b_1} = \frac{V_{b_1}^2}{S_b} = \frac{240^2}{30,000} = 1.92 \Omega, \quad Z_{b_2} = \frac{V_{b_2}^2}{S_b} = \frac{480^2}{30,000} = 7.68 \Omega, \quad Z_{b_3} = \frac{V_{b_3}^2}{S_b} = \frac{120^2}{30,000} = 0.48 \Omega$
 - (4) $X_{T_1} \text{ pu} = 0.10, \quad X_{T_2} \text{ pu} = 0.10 \left(\frac{460}{480} \right)^2 \left(\frac{30,000}{20,000} \right) = 0.10 \left(\frac{115}{120} \right)^2 \left(\frac{30,000}{20,000} \right) = 0.1378$
 - (5) $X_{\text{linea}} \text{ pu} = \frac{X_{\text{linea}}}{Z_{b_2}} = \frac{2}{7.68} = 0.2604, \quad Z_{\text{carga}} \text{ pu} = \frac{Z_{\text{carga}}}{Z_{b_3}} = \frac{0.9 + j0.2}{0.48} = 1.875 + j0.4167$
 - (6) $V_s \text{ pu} = \frac{220 \angle 0}{240} = 0.917 \angle 0$
- $$I_{\text{carga}} \text{ pu} = I_s \text{ pu} = \frac{V_s \text{ pu}}{j(X_{T_1} \text{ pu} + X_{\text{linea}} \text{ pu} + X_{T_2} \text{ pu}) + Z_{\text{carga}} \text{ pu}} = \frac{0.9167 \angle 0}{j(0.10 + 0.2604 + 0.1378) + (1.875 + j0.4167)}$$
- $$I_{\text{carga}} \text{ pu} = \frac{0.9167 \angle 0}{2.086 \angle 26.01} = 0.4395 \angle -26.01$$
- $$I_{b_3} = \frac{S_b}{V_{b_3}} = \frac{30,000}{120} = 250 \text{ A}$$
- $$I_{\text{carga}} = I_{\text{carga}} \text{ pu} \times I_{b_3} = (0.4395 \angle -26.01) \times 250 \text{ A} = 109.9 \angle -26.01 \text{ A} \quad \clubsuit$$



$$\begin{array}{lll}
 V_{\text{base1}} = 240 \text{ volts} & V_{\text{base2}} = 480 \text{ volts} & V_{\text{base3}} = 120 \text{ volts} \\
 Z_{\text{base1}} = \frac{(240)^2}{30,000} = 1.92 \Omega & Z_{\text{base2}} = \frac{(480)^2}{30,000} = 7.68 \Omega & Z_{\text{base3}} = \frac{(120)^2}{30,000} = 0.48 \Omega \\
 S_{\text{base}} = 30 \text{ kVA} & & I_{\text{base3}} = \frac{30,000}{120} = 250 \text{ A}
 \end{array}$$

Figura 1.11 Circuito monofásico con valores por unidad (pu)

Transformación en componentes simétricas

En el estudio de componentes simétricas resulta útil introducir el *operador complejo* $a = 1\angle 120^\circ = \frac{-1}{2} + j\frac{\sqrt{3}}{2}$, y sus siguientes identidades:

$$\begin{aligned} a &= 1\angle 120^\circ, & a^2 &= 1\angle 240^\circ, & a^3 &= 1\angle 0^\circ, & a^4 &= a = 1\angle 120^\circ, & ja &= 1\angle 210^\circ \\ 1 + a + a^2 &= 0, & 1 - a &= \sqrt{3}\angle -30^\circ, & 1 - a^2 &= \sqrt{3}\angle 30^\circ, & a^2 - a &= \sqrt{3}\angle 270^\circ \\ 1 + a - a^2 &= 1\angle 60^\circ, & 1 + a^2 - a &= 1\angle -60^\circ, & a + a^2 &= -1 = 1\angle 180^\circ \end{aligned} \quad (122)$$

Componentes simétricas de los voltajes de fase

La *transformación en componentes simétricas* es una técnica desarrollada por C. L. Fortescue en 1918 para analizar sistemas trifásicos desbalanceados; consiste en descomponer cada *voltaje o corriente de fase* del sistema trifásico en *tres componentes de secuencia (cero, positiva y negativa)*, de manera que cada *fase de fase* es la suma de sus 3 componentes de secuencia (**figura 1.12**):

$$\begin{aligned} V_a &= V_{a0} + V_{a1} + V_{a2} \\ V_b &= V_{b0} + V_{b1} + V_{b2} \\ V_c &= V_{c0} + V_{c1} + V_{c2} \end{aligned} \quad (123)$$

(0) Componentes de secuencia cero. Tres fasores de magnitud igual y desplazamiento de fase cero (V_{a0} , V_{b0} , V_{c0}).

$$V_{a0} = V_{b0} = V_{c0} \quad (124)$$

(1) Componentes de secuencia positiva. Tres fasores de magnitud igual, desplazamiento de fase ± 120 y secuencia positiva (V_{a1} , V_{b1} , V_{c1}).

$$\begin{aligned} V_{b1} &= V_{a1} \times 1\angle 240 = a^2 V_{a1} \\ V_{c1} &= V_{a1} \times 1\angle 120 = a V_{a1} \end{aligned} \quad (125)$$

(2) Componentes de secuencia negativa. Tres fasores de magnitud igual, desplazamiento de fase ± 120 y secuencia negativa (V_{a2} , V_{b2} , V_{c2}).

$$\begin{aligned} V_{b2} &= V_{a2} \times 1\angle 120 = a V_{a2} \\ V_{c2} &= V_{a2} \times 1\angle 240 = a^2 V_{a2} \end{aligned} \quad (126)$$

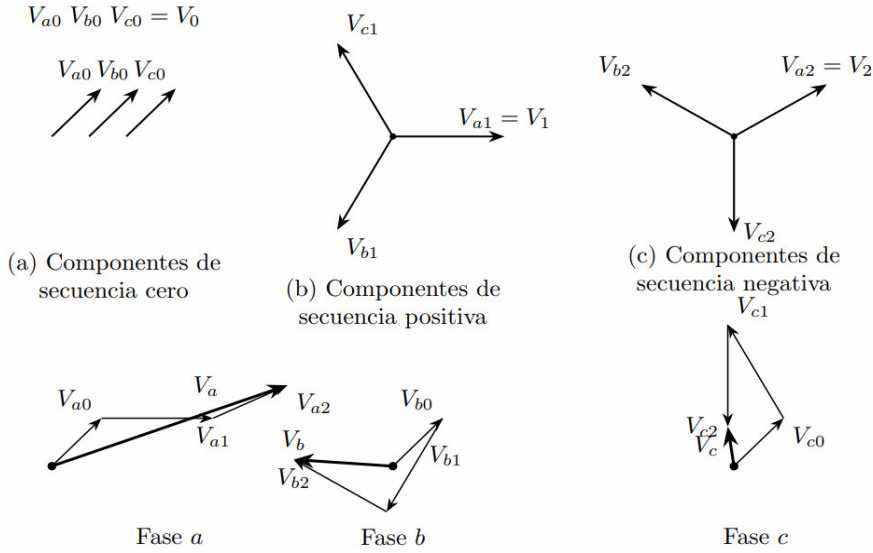


Figura 1.12 Voltajes de fase en términos de sus componentes de secuencia

Escribiendo *matricialmente* los *fasores de fase* en términos de sus *componentes de la fase a*, se tiene lo siguiente:

$$\begin{bmatrix} V_a = V_{a0} + V_{a1} + V_{a2} \\ V_b = V_{a0} + a^2 V_{a1} + a V_{a2} \\ V_c = V_{a0} + a V_{a1} + a^2 V_{a2} \end{bmatrix} \rightarrow \begin{bmatrix} V_a \\ V_b \\ V_c \end{bmatrix} = \begin{bmatrix} 1 & 1 & 1 \\ 1 & a^2 & a \\ 1 & a & a^2 \end{bmatrix} \begin{bmatrix} V_{a0} \\ V_{a1} \\ V_{a2} \end{bmatrix} \quad (127)$$

Definiendo un *vector de fase* $\mathbf{V}_p$, un *vector de componentes para la fase a* $\mathbf{V}_{sa}$ y una *matriz 3-cuadrada de transformación A*:

$$\mathbf{V}_p = \begin{bmatrix} V_a \\ V_b \\ V_c \end{bmatrix}, \quad \mathbf{V}_{sa} = \begin{bmatrix} V_{a0} \\ V_{a1} \\ V_{a2} \end{bmatrix}, \quad \mathbf{A} = \begin{bmatrix} 1 & 1 & 1 \\ 1 & a^2 & a \\ 1 & a & a^2 \end{bmatrix} \rightarrow \mathbf{A}^{-1} = \frac{1}{3} \begin{bmatrix} 1 & 1 & 1 \\ 1 & a & a^2 \\ 1 & a^2 & a \end{bmatrix} \quad (128)$$

De esta manera se tienen las siguientes ecuaciones de conversión entre fasores de fase y de secuencia para la *fase a*:

$$\mathbf{V}_p = \mathbf{A} \mathbf{V}_{sa} \quad \Leftrightarrow \quad \mathbf{V}_{sa} = \mathbf{A}^{-1} \mathbf{V}_p \quad (129)$$

En un *sistema 3φ balanceado* la componente de secuencia cero es cero ya que $\mathbf{v}_a + \mathbf{v}_b + \mathbf{v}_c = 0$. En un *sistema 3φ desbalanceado*, los *voltajes línea a neutro* pueden

tener componente de secuencia cero; pero los *voltajes línea a línea* no tienen componente de secuencia cero.

Ejemplo 1.16. Para un sistema 3ϕ balanceado en secuencia positiva con $v_{an}=227$ volts, calcular las componentes de secuencia para la *fase a*.

Solución. Empleando las ecuaciones antes deducidas se tiene:

$$V_p = \begin{bmatrix} \tilde{V}_a \\ \tilde{V}_b \\ \tilde{V}_c \end{bmatrix} = \begin{bmatrix} 227\angle 0 \\ 227\angle -120 \\ 227\angle 120 \end{bmatrix}$$

$$V_{sa} = \begin{bmatrix} \tilde{V}_{a0} \\ \tilde{V}_{a1} \\ \tilde{V}_{a2} \end{bmatrix} = \begin{bmatrix} \frac{1}{3}(\tilde{V}_a + \tilde{V}_b + \tilde{V}_c) \\ \frac{1}{3}(\tilde{V}_a + a\tilde{V}_b + a^2\tilde{V}_c) \\ \frac{1}{3}(\tilde{V}_a + a^2\tilde{V}_b + a\tilde{V}_c) \end{bmatrix} = \begin{bmatrix} \frac{1}{3}(227\angle 0 + 227\angle -120 + 227\angle 120) \\ \frac{1}{3}(227 + 227 + 227) \\ \frac{1}{3}(227 + 227\angle 120 + 227\angle 240) \end{bmatrix} = \begin{bmatrix} 0 \\ 227 \\ 0 \end{bmatrix} \quad \clubsuit$$

Componentes simétricas de las corrientes de fase

La transformación en componentes simétricas también es aplicable a corrientes fasoriales de fase, como se muestra a continuación:

$$I_p = AI_{sa} \Leftrightarrow I_{sa} = A^{-1}I_p \quad I_p = \begin{bmatrix} \tilde{I}_a \\ \tilde{I}_b \\ \tilde{I}_c \end{bmatrix}, \quad I_{sa} = \begin{bmatrix} \tilde{I}_{a0} \\ \tilde{I}_{a1} \\ \tilde{I}_{a2} \end{bmatrix} \quad (130)$$

En un sistema 3ϕ conectado en Y de 4 hilos, la corriente neutra esta dada por: $\tilde{I}_n = \tilde{I}_a + \tilde{I}_b + \tilde{I}_c = 3\tilde{I}_{a0}$, de manera que en un sistema balanceado conectado en Y, las corrientes de línea no tienen componente de secuencia cero, puesto que la corriente de neutro es cero. De lo anterior se deduce que en un sistema 3ϕ de 3 hilos Δ o Y no aterrizado las corrientes de línea no tienen componente de secuencia cero.

Ejemplo 1.17. Una línea 3ϕ que alimenta una carga equilibrada conectada en Y con neutro aterrizado, tiene abierta la fase b. Dadas las *corrientes de línea desbalanceadas*, calcular las *corrientes de secuencia para la fase a* y la *corriente al neutro*:

$$I_p = \begin{bmatrix} \tilde{I}_a \\ \tilde{I}_b \\ \tilde{I}_c \end{bmatrix} = \begin{bmatrix} 10\angle 0 \\ 0 \\ 10\angle 120 \end{bmatrix}$$

Solución. Empleando las ecuaciones definitorias de corrientes de secuencia se tiene que:

$$I_{sa} = \begin{bmatrix} \mathcal{I}_{a0} \\ \mathcal{I}_{a1} \\ \mathcal{I}_{a2} \end{bmatrix} = \begin{bmatrix} \frac{1}{3}(\mathcal{I}_a + \mathcal{I}_b + \mathcal{I}_c) \\ \frac{1}{3}(\mathcal{I}_a + a\mathcal{I}_b + a^2\mathcal{I}_c) \\ \frac{1}{3}(\mathcal{I}_a + a^2\mathcal{I}_b + a\mathcal{I}_c) \end{bmatrix} = \begin{bmatrix} \frac{1}{3}(10 + 0 + 10 \angle 120) \\ \frac{1}{3}(10 + 0 + 10) \\ \frac{1}{3}(10 + 0 + 10 \angle 240) \end{bmatrix} = \begin{bmatrix} 3.333 \angle 60 \\ 6.667 \\ 3.333 \angle -60 \end{bmatrix}$$

$$\mathcal{I}_n = \mathcal{I}_a + \mathcal{I}_b + \mathcal{I}_c = 10 + 0 + 10 \angle 120 = 10 \angle 60 = 3 \times \mathcal{I}_{a0} \quad \clubsuit$$

Transformación de Laplace

La *transformada de Laplace* es una *operación lineal* que transforma una función en del dominio del tiempo $f(t)$ en una función $F(s)$ en el dominio s de los números complejos y *existe si la integral converge para $t > 0$* ; se define por:

$$F(s) = L[f(t)] = \int_0^\infty f(t)e^{-st} dt \quad L[\] \leftarrow \text{Operador de Laplace}$$

(131)

$$\left[\begin{array}{l} L[af(t)] = aL[f(t)] \quad a \in \Re \vee C \\ L[a_1f(t) \pm a_2g(t)] = a_1L[f(t)] \pm a_2L[g(t)] \end{array} \right] \leftarrow \text{Linealidad}$$

La integral converge si $f(t)$ es *seccionalmente continua* en cada *intervalo finito* en el rango $t > 0$ y si es de *orden exponencial* conforme $t \rightarrow \infty$, esto es existe una constante $k \in R$, tal que $\lim_{t \rightarrow \infty} e^{-kt} |f(t)| = 0$.

Teoremas de la transformada de Laplace

Los teoremas de la transformada de Laplace: (1) *traslación o retardo en el tiempo*; (2) *traslación exponencial*; (3) *cambio de escala en el tiempo*; (4) *diferenciación real*; (5) *valor final*; (6) *valor inicial*; (7) *integración real*; (8) *diferenciación compleja*; ayudan a simplificar el cálculo de las *transformadas de funciones, derivadas e integrales*.

- (1) $L[f(t-t_0) \cdot u(t-t_0)] = e^{-st_0} F(s)$
 donde : $u(t-t_0) \leftarrow$ función · impulso · unitario
 $u(t-t_0) = 1 \cdot \text{para } t \geq t_0 \quad y \quad u(t-t_0) = 0 \cdot \text{para } t < t_0$
- (2) $L[e^{at} f(t)] = \int_0^\infty e^{at} e^{-st} f(t) dt = \int_0^\infty e^{-(s-a)t} f(t) dt = F(s-a), \quad a \in R \cdot \forall \cdot a \in C$
- (3) $L[f(\frac{t}{k})] = kF(ks), \quad k \in R^+ \wedge k > 0$
- (4a) $L\left[\frac{df(t)}{dt}\right] = sF(s) - f(0), \quad L\left[\frac{d^2 f(t)}{dt^2}\right] = s^2 F(s) - sf(0) - f'(0),$
- (4b) $L\left[\frac{d^n f(t)}{dt^n}\right] = s^n F(s) - s^{n-1} f(0) - s^{n-2} f'(0) \dots - sf^{(n-2)}(0) - f^{(n-1)}(0)$

(continua)

- (5) $\lim_{t \rightarrow 0} f(t) = \lim_{t \rightarrow 0} sF(s)$
- (6) $f(0^+) = \lim_{s \rightarrow \infty} sF(s)$
- (7) $L\left[\int f(t) dt\right] = \frac{F(s)}{s} + \frac{g(0^+)}{s} \quad \leftarrow g(0^+) = \int f(t) \cdot dt \Big|_{t=0} \quad (132)$
 si las condiciones iniciales son nulas $g(0^+) = 0 \therefore L\left[\int f(t) dt\right] = \frac{F(s)}{s}$

- (8) $L[t^n f(t)] = (-1)^n \frac{d^n F(s)}{ds^n} \quad n = 1, 2, 3, \dots$

así mismo la función delta se define como una integral compleja :

$$\delta(t) = \int_{-\infty}^{\infty} e^{j2\pi ft} df = \int_{-\infty}^{\infty} e^{-j2\pi ft} df \quad \leftarrow f = \text{frecuencia dominio del tiempo}$$

$$\delta(t) = \frac{1}{2\pi} \int_{-\infty}^{\infty} e^{j\omega t} d\omega = \frac{1}{2\pi} \int_{-\infty}^{\infty} e^{-j\omega t} d\omega \quad \leftarrow \omega = 2\pi f = \text{frecuencia angular}$$

Función impulso unitario

Antes de continuar nuestra exposición se estudiará una función de gran interés y aplicación. La función impulso unitario o función delta es una función del tiempo que vale cero cuando su argumento ($t-t_0$) es menor que cero; que vale cero cuando su argumento es mayor que cero; que es infinito cuando su argumento vale cero , y que tiene un área unitaria, matemáticamente se expresa por:

$$\delta(t - t_0) = 0 \quad \text{para } t \neq t_0$$

$$\int_{-\infty}^{\infty} \delta(t - t_0) dt = \int_{t_0^-}^{t_0^+} \delta(t - t_0) dt = 1 \quad (133)$$

La *amplitud del impulso es infinita*. Cuando se multiplica por una función $f(t)$, su *amplitud o peso* cambia de 1 a $f(t_0)$, es decir:

$$\text{Intensidad}[f(t)\delta(t - t_0)] = \int_{-\infty}^{\infty} f(t)\delta(t - t_0) dt = f(t_0) \quad (134)$$

En términos de la *frecuencia f* , y para efectos de *análisis espectrales* y su uso con *transformadas de Fourier*, se tienen una definición alternativa altamente útil para la función delta:

$$\begin{aligned} \delta(t) &= \int_{-\infty}^{\infty} e^{j2\pi ft} df = \int_{-\infty}^{\infty} e^{-j2\pi ft} df \\ \delta(t) &= \frac{du(t)}{dt} \leftrightarrow u(t) \leftarrow \text{escalón unitario} \end{aligned} \quad (135)$$

Transformación de funciones

A continuación se presentan y transforman funciones de amplio uso en la ingeniería de control:

Función **escalón**:

$$\begin{aligned} f(t) &= \begin{cases} 0 & \text{para } t < 0 \\ A & \text{para } t > 0 \end{cases} \\ L[A] &= \int_0^{\infty} A e^{-st} dt = \frac{A}{-s} [e^{-\infty} - e^0] = \frac{A}{s} \end{aligned} \quad (136)$$

Función **rampa**:

$$\begin{aligned} f(t) &= \begin{cases} 0 & \text{para } t < 0 \\ At & \text{para } t \geq 0 \end{cases} \\ L[At] &= \int_0^{\infty} Ate^{-st} dt = (-1) \frac{d}{ds} \left(\frac{A}{s} \right) = -A \frac{d}{ds} (s^{-1}) = \frac{A}{s^2} \end{aligned} \quad (137)$$

Función **pulso**: es una función escalón de altura A/T y longitud T

$$f(t) = \begin{cases} 0 & \text{para } t < 0, T < t \\ \frac{A}{T} & \text{para } 0 < t < T \end{cases} \rightarrow f(t) = \frac{A}{T} 1(t) - \frac{A}{T} 1(t-T)$$

$$L[f(t)] = L\left[\frac{A}{T} 1(t)\right] - L\left[\frac{A}{T} 1(t-T)\right] = \frac{A}{Ts} - \frac{A}{Ts} e^{-sT} = \frac{A}{Ts} (1 - e^{-sT})$$
(138)

Función **impulso**: es un caso limitado de la función pulso

$$f(t) = \begin{cases} 0 & \text{para } t < 0, T < t \\ \lim_{T \rightarrow 0} \frac{A}{T} & \text{para } 0 < t < T \end{cases}$$

$$L[f(t)] = \lim_{T \rightarrow 0} \left[\frac{A}{Ts} (1 - e^{-sT}) \right] = \lim_{T \rightarrow 0} \left[\frac{\frac{d}{dT} [A(1 - e^{-sT})]}{\frac{d}{dT} [Ts]} \right] = \frac{As}{s} = A \leftarrow \text{Area bajo el impulso}$$
(139)

Función **impulso unitario** o **delta de Dirac**: tiene *magnitud infinita* y *duración de cero*

$$f(t) = \begin{cases} \delta(t-T) = 0 & \text{para } t \neq T \\ \delta(t-T) = \infty & \text{para } t = T \end{cases} \rightarrow \int_{-\infty}^{\infty} \delta(t-T) dt = 1 \quad \delta(t-T) = \frac{d}{dt} 1(t-T)$$

$$L[\delta(t-T)] = 1$$
(140)

Función **exponencial**:

$$f(t) = \begin{cases} 0 & \text{para } t < 0 \\ Ae^{-at} & \text{para } t \geq 0 \end{cases} \quad a, A = \text{ctes}$$

$$L[Ae^{-at}] = \int_0^{\infty} Ae^{-at} e^{-st} dt = \int_0^{\infty} Ae^{-(a+s)t} dt = \frac{A}{-(s+a)} [e^{-\infty} - e^0] = \frac{A}{s+a}$$
(141)

Función **senoidal y cosenoidal**:

$$f(t) = \begin{cases} 0 & \text{para } t < 0 \\ A \sin \omega t & \text{para } t \geq 0 \end{cases}$$

$$L[A \sin \omega t] = \int_0^{\infty} (A \sin \omega t) e^{-st} dt = \frac{A}{2j} \int_0^{\infty} (e^{j\omega t} - e^{-j\omega t}) e^{-st} dt = \frac{A}{2j} \left(\frac{1}{s-j\omega} - \frac{1}{s+j\omega} \right) = \frac{A\omega}{s^2 + \omega^2}$$

$$L[A \cos \omega t] = \int_0^{\infty} (A \cos \omega t) e^{-st} dt = \frac{A}{2} \int_0^{\infty} (e^{j\omega t} + e^{-j\omega t}) e^{-st} dt = \frac{As}{s^2 + \omega^2}$$
(142)

Transformación inversa de Laplace

El *proceso inverso* de encontrar la *función del tiempo* $f(t)$ a partir de la *transformada de Laplace* $F(s)$ se denomina *transformada inversa de Laplace* y se define por:

$$f(t) = L^{-1}[F(s)] = \frac{1}{2\pi j} \int_{c-j\infty}^{c+j\infty} F(s)e^{st} dt \quad \text{para } t > 0 \quad (143)$$

Donde $c \in \Re$ es la *abscisa de convergencia*, y cumple que $c > \Re[s_i]$, siendo s_i cualquier punto singular de $F(s)$. La integral es del campo de los números complejos; y su evaluación es complicada; en su lugar se emplea una tabla de funciones $f(t)$ y sus transformadas $F(s)$ (ver **tabla 1.5**). En algunos casos es necesario expresar $F(s)$ como una *suma de fracciones parciales* para transformar cada término separadamente; este es el método preferido.

 **Ejemplo 1.18.** Hallar la *transformada inversa de Laplace* de $F(s) = 5[s(s+1)(s+2)]^{-1}$

Solución. Expresando la *función racional* como una *suma de fracciones parciales* e invirtiendo término a término se tiene:

$$\begin{aligned} F(s) &= \frac{5}{s(s+1)(s+2)} = \frac{a_0}{s} + \frac{b_0}{s+1} + \frac{c_0}{s+2} \\ a_0 &= \frac{1}{0!} \frac{d^0}{ds^0} \left(\frac{P(s)}{R(s)} \right) \Bigg|_{s=0} = \left(\frac{5}{(s+1)(s+2)} \right) \Bigg|_{s=0} = \left(\frac{5}{(1)(2)} \right) = 2.5 \\ b_0 &= \frac{1}{0!} \frac{d^0}{ds^0} \left(\frac{P(s)}{R(s)} \right) \Bigg|_{s=-1} = \left(\frac{5}{s(s+2)} \right) \Bigg|_{s=-1} = \left(\frac{5}{-1(-1+2)} \right) = -5 \\ c_0 &= \frac{1}{0!} \frac{d^0}{ds^0} \left(\frac{P(s)}{R(s)} \right) \Bigg|_{s=-2} = \left(\frac{5}{s(s+1)} \right) \Bigg|_{s=-2} = \left(\frac{5}{-2(-2+1)} \right) = 2.5 \\ F(s) &= \frac{2.5}{s} - \frac{5}{s+1} + \frac{2.5}{s+2} \\ f(t) &= L^{-1}[F(s)] = L^{-1} \left[\frac{2.5}{s} \right] - L^{-1} \left[\frac{5}{s+1} \right] + L^{-1} \left[\frac{2.5}{s+2} \right] = 2.5u(t) - 5e^{-t} + 2.5e^{-2t} \quad \clubsuit \end{aligned}$$

Con frecuencia resultan transformadas inversas de Laplace que involucran *polos complejos conjugados*, en tal caso no es necesario la expansión en fracciones parciales; sino como una suma de *funciones seno y coseno amortiguadas*, tomando en cuenta que:

$$L[e^{-at} \sin \omega t] = \frac{\omega}{(s+a)^2 + \omega^2}, \quad L[e^{-at} \cos \omega t] = \frac{s+a}{(s+a)^2 + \omega^2} \quad (144)$$

 **Ejemplo 1.19.** Hallar la *transformada inversa de Laplace* de $F(s) = \frac{2s+12}{s^2+2s+5}$

Solución. La factorización del denominador es $s^2+2s+5 = (s+1+j2)(s+1-j2)$, lo cual también se factoriza como $s^2+2s+5 = (s+1)^2+2^2$, de aquí que:

$$F(s) = \frac{2s+12}{s^2+2s+5} = \frac{2s+12}{(s+1)^2+2^2} = \frac{10+2(s+1)}{(s+1)^2+2^2} = 5 \frac{2}{(s+1)^2+2^2} + 2 \frac{s+1}{(s+1)^2+2^2}$$

$$L^{-1}[F(s)] = 5L^{-1}\left[\frac{2}{(s+1)^2+2^2}\right] + 2L^{-1}\left[\frac{s+1}{(s+1)^2+2^2}\right] = 5e^{-t} \sin 2t + 2e^{-t} \cos 2t \quad \clubsuit$$

Convolución

La *convolución* es un tipo especial de producto conmutativo de dos funciones y se define como una integral, es decir, la convolución de $f(t)$ y $g(t)$ se denota y está dada por:

$$(f * g)(t) = f(t) * g(t) = \int_0^t f(\tau)g(t-\tau)d\tau \quad \rightarrow \quad f * g = g * f \quad (145)$$

La transformada de la convolución de f y g es producto de las transformadas individuales de f y g . La transformada inversa del producto de las transformadas F y G , es la convolución de f y g .

$$\begin{aligned} L[f * g](t) &= L[f(t)] \cdot L[g(t)] \\ L^{-1}[F(s) \cdot G(s)] &= (f * g)(t) \end{aligned} \quad (146)$$

Este último hecho permite hallar la transformada inversa de una función $H(s)$ siempre que: (1) $H(s)$ pueda descomponerse en dos funciones $F(s)$ y $G(s)$; (2) las transformadas inversas $f(t)$ y $g(t)$ existan; (3) sea posible evaluar las integrales de convolución de f y g .

 **Ejemplo 1.20.** Hallar la transformada inversa de $H(s) = [(s+a)(s+b)]^{-1}$, usando la técnica de convolución.

Solución. Se aplican los tres pasos enunciados y se tiene que:

$$\begin{aligned} (1) \quad H(s) &= F(s) \cdot G(s) \rightarrow F(s) = \frac{1}{(s+a)}, \quad G(s) = \frac{1}{(s+b)} \\ (2) \quad f(t) &= L^{-1}\left[\frac{1}{s+a}\right] = e^{-at}, \quad g(t) = L^{-1}\left[\frac{1}{s+b}\right] = e^{-bt} \\ (3) \quad (f * g)(t) &= \int_0^t f(\tau)g(t-\tau)d\tau = \int_0^t e^{-a\tau} e^{-b(t-\tau)} d\tau = \int_0^t e^{-a\tau-b(t-\tau)} d\tau \\ &= \int_0^t e^{(b-a)\tau-bt} d\tau = \frac{1}{b-a} e^{(b-a)\tau-bt} \Big|_0^t = \frac{1}{b-a} [e^{(b-a)t-bt} - e^{-bt}] \\ &= \frac{1}{b-a} [e^{(b-a-b)t} - e^{-bt}] = \frac{1}{b-a} [e^{-at} - e^{-bt}] \\ \therefore L^{-1}[H(s)] &= (f * g)(t) = \frac{1}{b-a} (e^{-at} - e^{-bt}) \quad \clubsuit \end{aligned}$$

Transformación de ecuaciones diferenciales lineales invariantes en el tiempo

Las ecuaciones diferenciales lineales invariantes en el tiempo con o sin condiciones iniciales se resuelven de forma sencilla mediante la transformación de Laplace. Dada una ED se realizan los siguientes pasos: (1) se aplica la transformación de Laplace a cada término de la ED; (2) se resuelve para la variable independiente $F(s)$; (3) de ser necesario se aplica transformación inversa de Laplace para expresar la función en el dominio del tiempo $f(t)$.

 **Ejemplo 1.21.** Hallar la función $y(t)$ que satisface la ED $D^2y + 3Dy + 2y = 0$, sujeta a las condiciones iniciales: $y(0) = a$, $y'(0) = b$.

Solución. Se aplican los tres pasos antes numerados y se tiene que:

$$(1) \quad L[y + 3y' + 2y'' = 0] = [s^2Y(s) - sy'(0) - y(0)] + 3[sY(s) - y(0)] + 2[Y(s)] = 0$$

$$s^2Y - as - b + 3sY - 3a + 2Y = 0 \rightarrow Y(s^2 + 3s + 2) - as - b - 3a = 0$$

$$(2) \quad Y = \frac{a(s+3) + b}{(s^2 + 3s + 2)} = \frac{a_0}{(s+1)} + \frac{b_0}{(s+2)}$$

$$a_0 = \frac{1}{0!} \frac{d^0}{ds^0} \left(\frac{P(s)}{R(s)} \right) \Bigg|_{s=-1} = \left(\frac{a(s+3) + b}{s+2} \right) \Bigg|_{s=-1} = \frac{a(-1+3) + b}{-1+2} = 2a + b$$

$$b_0 = \frac{1}{0!} \frac{d^0}{ds^0} \left(\frac{P(s)}{R(s)} \right) \Bigg|_{s=-2} = \left(\frac{a(s+3) + b}{s+1} \right) \Bigg|_{s=-2} = \frac{a(-2+3) + b}{(-2+1)} = -(a+b)$$

$$Y = \frac{2a+b}{(s+1)} - \frac{a+b}{(s+2)}$$

$$(3) \quad y(t) = L^{-1} \left[\frac{2a+b}{(s+1)} \right] - L^{-1} \left[\frac{a+b}{(s+2)} \right] = (2a+b)e^{-t} + (a+b)e^{-2t} \quad \clubsuit$$

Tabla 1.5 Transformadas directas e inversas de Laplace

$F(s)=L[f(t)]$	$f(t)=L^{-1}[F(s)]$	$F(s)=L[f(t)]$	$f(t)=L^{-1}[F(s)]$
1	$\delta(t)$	$\frac{\omega}{s^2 + \omega^2}$	$\sin \omega t$
$\frac{1}{s}$	1	$\frac{s}{s^2 + \omega^2}$	$\cos \omega t$
$\frac{A}{Ts}(1 - e^{-sT})$	$\frac{A}{T}$	$\frac{\omega}{s^2 - \omega^2}$	$\sinh \omega t$
$\frac{1}{s^2}$	t	$\frac{s}{s^2 - \omega^2}$	$\cosh \omega t$
$\frac{1}{s^n}$	$\frac{t^{n-1}}{(n-1)!}$	$\frac{\omega}{(s+a)^2 + \omega^2}$	$e^{-at} \sin \omega t$
$\frac{n!}{s^{n+1}}$	t^n	$\frac{s+a}{(s+a)^2 + \omega^2}$	$e^{-at} \cos \omega t$
$\frac{1}{s+a}$	e^{-at}	$\frac{\omega^2}{s(s^2 + \omega^2)}$	$1 - \cos \omega t$
$\frac{1}{(s+a)^2}$	te^{-at}	$\frac{\omega_n^2}{s^2 + 2\xi\omega_n s + \omega_n^2}$	$\frac{\omega_n}{K} e^{-\xi\omega_n t} \sin(\omega_n Kt), K = \sqrt{1 - \xi^2}$
$\frac{n!}{(s+a)^{n+1}}$	$t^n e^{-at}$	$\frac{s}{s^2 + 2\xi\omega_n s + \omega_n^2}$	$-\frac{1}{K} e^{-\xi\omega_n t} \sin(\omega_n Kt - \phi)$ $K = \sqrt{1 - \xi^2}, \phi = \tan^{-1} K/\xi$
$\frac{1}{s(s+a)}$	$\frac{1}{a}(1 - e^{-at})$	$\frac{\omega_n^2}{s(s^2 + 2\xi\omega_n s + \omega_n^2)}$	$1 - \frac{1}{K} e^{-\xi\omega_n t} \sin(\omega_n Kt + \phi)$ $K = \sqrt{1 - \xi^2}, \phi = \tan^{-1} K/\xi$
$\frac{1}{(s+a)(s+b)}$	$\frac{1}{b-a}(e^{-at} - e^{-bt})$	$\frac{\omega^3}{s^2(s^2 + \omega^2)}$	$\omega t - \sin \omega t$
$\frac{s}{(s+a)(s+b)}$	$\frac{1}{b-a}(be^{-bt} - ae^{-at})$	$\frac{2\omega^3}{(s^2 + \omega^2)^2}$	$\sin \omega t - \omega t \cos \omega t$
$\frac{1}{s(s+a)(s+b)}$	$\frac{1}{ab}\left[1 + \frac{1}{a-b}(be^{-at} - ae^{-bt})\right]$	$\frac{s}{(s^2 + \omega^2)^2}$	$\frac{1}{2\omega} t \sin \omega t$
$\frac{1}{s(s+a)^2}$	$\frac{1}{a^2}(1 - e^{-at} - ate^{-at})$	$\frac{s^2 - \omega^2}{(s^2 + \omega^2)^2}$	$t \cos \omega t$
$\frac{1}{s(s+a)^2}$	$\frac{1}{a^2}(1 - e^{-at} - ate^{-at})$	$\frac{s^2}{(s^2 + \omega_1^2)(s^2 + \omega_2^2)}$	$\frac{1}{\omega_2 - \omega_1}(\cos \omega_1 t - \cos \omega_2 t)$
$\frac{1}{s^2(s+a)}$	$\frac{1}{a^2}(at - 1 + e^{-at})$	$\frac{s^2}{(s^2 + \omega^2)^2}$	$\frac{1}{2\omega}(\sin \omega t - \omega t \cos \omega t)$

Transformación en series de potencias

Sucesión. Es una *función* $S(k)$ cuyo *dominio* es el conjunto de los *enteros no negativos*, $k \in \{Z^{0+}\}$, y cuyos *términos* son los *elementos de la sucesión (contradominio)*, pudiendo ser *finita o de N términos*: $k = 0, 1, 2, \dots, N-1$; o bien *infinita*: $k = 0, 1, 2, \dots, \infty$; siendo estas últimas de mayor interés:

$$\begin{aligned} S(0), S(1), S(2), \dots, S(k), \dots = a_0, a_1, a_2, \dots, a_k, \dots = \{a_k\}_0^\infty = \{a_k\} &\leftarrow \text{Sucesión infinita} \\ S(0), S(1), S(2), \dots, S(N-1) = a_0, a_1, a_2, \dots, a_{N-1}, \dots = \{a_k\}_0^{N-1} &\leftarrow \text{Sucesión finita} \end{aligned} \quad (147)$$

Una sucesión puede ser *convergente* o *divergente* según tienda o no hacia un valor L cuando $k \rightarrow \infty$:

$$\lim_{k \rightarrow \infty} a_k = L \quad (148)$$

El criterio de *convergencia de Cauchy* establece que: una sucesión $\{a_k\}$ converge si y sólo si para un $\varepsilon > 0$ existe un N tal que $|a_i - a_j| < \varepsilon$ para toda $i, j > N$. Una sucesión es *monótona* si es: *creciente* $\{a_k < a_{k+1}\}$, *decreciente* $\{a_k > a_{k+1}\}$, *no creciente* $\{a_k \leq a_{k+1}\}$, *no decreciente* $\{a_k \geq a_{k+1}\}$.

Serie. Es la *suma de los términos de una sucesión*, pueden ser representadas por la notación sigma y es de interés primordial saber si son convergentes o divergentes:

$$\begin{aligned} S &= S(0) + S(1) + S(2) + \dots + S(k) + \dots = a_0 + a_1 + a_2 + \dots + a_k + \dots = \sum_{k=0}^{\infty} a_k \\ S_N &= S(0) + S(1) + S(2) + \dots + S(N-1) = a_0 + a_1 + a_2 + \dots + a_{N-1} = \sum_{k=0}^{N-1} a_k \end{aligned} \quad (149)$$

La serie $\sum_{k=0}^{\infty} a_k$ converge si la sucesión de *sumas parciales* $S_k = \sum_{k=0}^{N-1} a_k$, converge; es decir:

$$\lim_{N \rightarrow \infty} S_N = \lim_{N \rightarrow \infty} \sum_{k=0}^{N-1} a_k = L \quad (150)$$

Si el límite no existe entonces la serie diverge. Una condición necesaria para la convergencia de la serie es:

$$\lim_{k \rightarrow \infty} a_k = 0 \quad (151)$$

Una serie tiene *convergencia absoluta* si $\sum_{k=0}^{\infty} |a_k|$ converge; de otra forma tiene *convergencia condicional*. El orden de los sumandos de una serie de convergencia absoluta no altera la suma; pero si la altera en una serie de convergencia condicional.

Series especiales

La *serie geométrica* se define por y tiene las siguientes propiedades:

$$\begin{aligned}
 S &= \sum_{k=0}^{\infty} z^k = 1 + z + z^2 + \cdots + z^k + \cdots \\
 S_N &= \sum_{k=0}^{N-1} z^k = 1 + z + z^2 + \cdots + z^{N-1} = \frac{1 - z^N}{1 - z} \\
 \lim_{N \rightarrow \infty} S_N &= \lim_{N \rightarrow \infty} \frac{1 - z^N}{1 - z} = \frac{1}{1 - z} \text{ siempre que } |z| < 1 \\
 S &= 1 - 1 + 1 - 1 + 1 - 1 + \cdots - \cdots + \text{ la serie es oscilante para } z = -1
 \end{aligned}
 \tag{152}$$

La *serie harmónica* es *absolutamente convergente* para $\text{Re}[\alpha] > 1$, y *absolutamente divergente* para $\text{Re}[\alpha] \leq 1$; se define por:

$$\begin{aligned}
 S &= \sum_{k=1}^{\infty} \frac{1}{k^\alpha} = 1 + \frac{1}{2^\alpha} + \frac{1}{3^\alpha} + \cdots + \frac{1}{k^\alpha} + \cdots \\
 S_N &= \sum_{k=1}^N \frac{1}{k^\alpha} = 1 + \frac{1}{2^\alpha} + \frac{1}{3^\alpha} + \cdots + \frac{1}{N^\alpha}
 \end{aligned}
 \tag{153}$$

La *serie harmónica alternante* es *absolutamente convergente* para $\text{Re}[\alpha] > 1$, y *absolutamente divergente* para $\text{Re}[\alpha] \leq 1$; se define por:

$$\begin{aligned}
 S &= \sum_{k=1}^{\infty} \frac{(-1)^{k+1}}{k^\alpha} = 1 - \frac{1}{2^\alpha} + \frac{1}{3^\alpha} - \cdots + \frac{1}{k^\alpha} - \cdots \\
 S_N &= \sum_{k=1}^N \frac{(-1)^{k+1}}{k^\alpha} = 1 - \frac{1}{2^\alpha} + \frac{1}{3^\alpha} - \cdots + \frac{1}{N^\alpha}
 \end{aligned}
 \tag{154}$$

La *función zeta de Riemann* se define como la *suma de la serie harmónica infinita*.

$$\zeta(\alpha) = \sum_{k=1}^{\infty} \frac{1}{k^\alpha}
 \tag{155}$$

Pruebas de convergencia

Prueba de comparación. La serie de términos positivos $\sum a_k$ converge, si existe una serie convergente $\sum b_k$ tal que $a_k \leq b_k, \forall k$. Similarmente, la serie $\sum a_k$ diverge, si existe una serie divergente $\sum b_k$ tal que $a_k \geq b_k, \forall k$.

Prueba de integral. Si los coeficientes a_k de una serie $\sum_{k=0}^{\infty} a_k$, son decrecientes y pueden ser extendidos a una *función decreciente de variable continua* x : $a(x) = a_k$ para $x \in \mathbb{Z}^+$, entonces *la serie converge o diverge con la integral*: $\int_0^{\infty} a(x)dx$.

Prueba de límites. Estas pruebas son de *cociente* de *raíz* y de *Raabe* se aplican a la serie $\sum a_k$. El caso mostrado indica que la serie converge absolutamente, de lo contrario diverge, si el límite es la unidad, la prueba falla:

$$\begin{aligned} \lim_{k \rightarrow \infty} \left| \frac{a_{k+1}}{a_k} \right| &< 1 && \leftarrow \text{Prueba de cociente} \\ \lim_{k \rightarrow \infty} |a_k|^{1/k} &< 1 && \leftarrow \text{Prueba de raíz} \\ \lim_{k \rightarrow \infty} n \left(1 - \left| \frac{a_{k+1}}{a_k} \right| \right) &> 1 && \leftarrow \text{Prueba de Raabe} \end{aligned} \quad (156)$$

Ejemplo 1.22: Aplicar la prueba del cociente a la serie $\sum_{k=1}^{\infty} \frac{e^k}{k!}$

$$\lim_{k \rightarrow \infty} \left| \frac{a_{k+1}}{a_k} \right| = \lim_{k \rightarrow \infty} \left| \frac{e^{k+1} k!}{e^k (k+1)!} \right| = \lim_{k \rightarrow \infty} \left| \frac{e}{k+1} \right| = 0 \therefore \text{ la serie es absolutamente convergente } \clubsuit$$

Prueba de Gauss. Considere la serie $\sum a_k$, si $\frac{a_{k+1}}{a_k} = 1 - \frac{L}{k} + \frac{b_k}{k^2}$, donde b_k esta definido, entonces la serie converge absolutamente si $L > 1$; de otra forma la serie diverge o converge condicionalmente.

Convergencia uniforme

Una *función* $f(z)$ es *continua* en un dominio cerrado si $\lim_{\zeta \rightarrow z} f(\zeta) = f(z)$, para toda z del dominio. Considere una serie en la cual los términos son funciones de z , $\sum_{k=0}^{\infty} a_k(z)$. La serie converge en un dominio, si la serie converge por cada punto z en el dominio. Si definimos la función continua $f(z) = \sum_{k=0}^{\infty} a_k(z)$, podemos establecer un *criterio de convergencia*; para un ε dado existe una función $N(z)$ tal que:

$$\left| f(z) - S_{N(z)}(z) \right| = \left| f(z) - \sum_{k=0}^{N(z)-1} a_k(z) \right| < \varepsilon \quad \text{para todo } z \text{ en el dominio} \quad (157)$$

Una serie de la forma $\sum_{k=0}^{\infty} a_k(z)$ es *uniformemente convergente* en el dominio si para un $\varepsilon > 0$ dado, existe un N independiente de z tal que:

$$|f(z) - S_N(z)| = \left| f(z) - \sum_{k=1}^N a_k(z) \right| < \varepsilon \quad \text{para todo } z \text{ en el dominio} \quad (158)$$

Prueba de Weierstrass. Una serie $\sum_{k=0}^{\infty} a_k(z)$ es uniforme y absolutamente convergente, si existe una serie convergente de términos positivos $\sum_{k=0}^{\infty} M_k$ tal que: $|a_k(z)| \leq M_k$ para toda z del dominio.

Prueba de Dirichlet. Considere una secuencia decreciente, de constantes positivas c_k , con límite cero. Si todas las sumas parciales de $a_k(z)$, están definidas en un dominio cerrado, es decir: $\left| \sum_{k=1}^N a_k(z) \right| < \text{constante}$, para toda N ; entonces $\sum_{k=1}^{\infty} c_k a_k(z)$, es uniformemente convergente en ese dominio cerrado. Esta prueba no implica convergencia absoluta.

En general: *una serie uniformemente convergente de términos continuos representa una función continua.*

Series de potencias de convergencia uniforme

Las *series de potencia* tienen la forma:

$$\sum_{k=0}^{\infty} a_k(z - z_0)^k = a_0 + a_1(z - z_0) + a_2(z - z_0)^2 + \dots + a_k(z - z_0)^k + \dots \quad (159)$$

El *dominio de convergencia* de una serie de potencia es un círculo en el plano complejo, y su *radio de convergencia* se obtiene al aplicar la *prueba de convergencia del cociente*:

$$\sum_{k=0}^{\infty} a_k z^k \quad \rightarrow \quad R = \lim_{k \rightarrow \infty} \frac{|a_k|}{|a_{k+1}|} \quad (160)$$

Si la serie de potencias $\sum_{k=0}^{\infty} a_k z^k$ converge para $z = z_0$, entonces la serie converge absolutamente para $|z| = |z_0|$. **Ejemplo 1.23**, hallar el radio de convergencia para la serie de potencias $\sum_{k=1}^{\infty} k! z^{k!}$:

$$R = \lim_{k \rightarrow \infty} \frac{|a_k|}{|a_{k+1}|} = \lim_{k \rightarrow \infty} \frac{|(k+1)! z^{(k+1)!}|}{|k! z^{k!}|} = \lim_{k \rightarrow \infty} (k+1) |z|^{(k+1)! - k!} = \lim_{k \rightarrow \infty} (k+1) |z|^{(k)!} < 1$$

$$R = \lim_{k \rightarrow \infty} (\ln(k+1) + (k)k! \ln|z|) < 0, \rightarrow \ln|z| < \lim_{k \rightarrow \infty} \frac{-\ln(k+1)}{(k)k!}, \rightarrow \ln|z| < 0, \rightarrow |z| < 1$$

$\therefore$ la serie converge absolutamente para $|z| < 1$: $-1 < z < 1$ ♣

Función analítica. Una *función* $f(z)$ es *analítica* en un punto $z=a$ si es diferenciable en a y en todos los puntos de una vecindad de a . De forma equivalente $f(z)$ es analítica en a si puede representarse mediante una serie de potencias, con radio de convergencia (R):

$$f(z) = \sum_{k=0}^{\infty} c_k (z-a)^k, \quad R > 0 \quad (160a)$$

Una propiedad fundamental de las funciones analíticas es que, si son analíticas en un dominio simplemente conexo (*región sin huecos*), la integral sobre cualquier *contorno cerrado simple* C es nula (cero):

$$\oint_C f(z) dz = 0 \quad (160b)$$

Esto significa que, al recorrer una trayectoria cerrada y regresar al punto de partida, el cambio neto acumulado es cero. Por ejemplo, en un circuito eléctrico ideal, la suma algebraica de todas las caídas y elevaciones de voltaje alrededor de una malla es cero, principio utilizado en análisis de circuitos, campos electromagnéticos, etc.

Si una función es analítica en $z=a$, entonces a es un (1) *punto ordinario*, de otra forma es un (2) *punto singular*, el cual puede ser (2a) *singular regular*, (2b) *singular irregular*.

Derivación e integración de series de potencias

La *derivación de una serie de potencias* se realiza en cada término de la serie como se muestra:

$$\begin{aligned}
 \frac{d}{dz} \sum_{k=0}^{\infty} a_k z^k &= \frac{d}{dz} a_0 + \frac{d}{dz} a_1 z + \frac{d}{dz} a_2 z^2 + \cdots + \frac{d}{dz} a_k z^k + \cdots \\
 &= \sum_{k=1}^{\infty} k a_k z^{k-1} = \sum_{k=0}^{\infty} (k+1) a_{k+1} z^k \\
 \text{Si } y &= \sum_{k=0}^{\infty} a_k z^k \\
 \rightarrow y' &= y^{(1)} = \sum_{k=1}^{\infty} k a_k z^{k-1} = \sum_{k=0}^{\infty} (k+1) a_{k+1} z^k \\
 \rightarrow y'' &= y^{(2)} = \sum_{k=2}^{\infty} k(k-1) a_k z^{k-2} \\
 \rightarrow y''' &= y^{(3)} = \sum_{k=3}^{\infty} k(k-1)(k-2) a_k z^{k-3} \\
 \Rightarrow y^{(n)} &= \sum_{k=n}^{\infty} k(k-1)(k-2) \cdots (k-n+1) a_k z^{k-n} \\
 \Rightarrow &= \sum_{k=0}^{\infty} (k+n)(k+n-1)(k+n-2) \cdots (k+1) a_{k+n} z^k \\
 \therefore y^{(n)} &= \sum_{k=0}^{\infty} \left(\prod_{j=0}^{n-1} (k+n-j) \right) a_{k+n} z^k
 \end{aligned} \tag{161}$$

La *integración de una serie de potencias* se realiza en cada término de la serie como se muestra:

$$\begin{aligned}
 \int \sum_{k=0}^{\infty} a_k z^k dz &= \int a_0 dz + \int a_1 z dz + \int a_2 z^2 dz + \cdots + \int a_k z^k dz + \cdots = \sum_{k=0}^{\infty} a_k \int z^k dz \\
 &= \sum_{k=0}^{\infty} \frac{a_k}{k+1} z^{k+1} + C
 \end{aligned} \tag{162}$$

Desarrollo de funciones en series de potencias

Brook Taylor demostró que toda función $f(z)$ continua y diferenciable puede ser aproximada en la vecindad de $z=a$ por una serie de potencias conocida como *serie de Taylor*:

$$f(z) = \sum_{k=0}^{\infty} c_k (z-a)^k = c_0 + c_1 (z-a) + c_2 (z-a)^2 + c_3 (z-a)^3 + \cdots + c_k (z-a)^k + \dots \tag{163}$$

Taylor descubrió que $f^{(k)}(a) = k! c_k$ es decir: $c_k = f^{(k)}(a) / k!$, de esta manera la serie de Taylor que representa la función $f(z)$ en la vecindad de $z=a$ esta dada por:

$$f(z) = \sum_{k=0}^{\infty} c_k (z-a)^k = \sum_{k=0}^{\infty} \frac{f^{(k)}(a)}{k!} (z-a)^k \tag{164}$$

La serie de Taylor se emplea para representar la función con un número finito de términos, que recibe el nombre de *polinomio de Taylor de grado n*. El resto de la función se representa por un *residuo* $R(z)$:

$$P_n(z) = f(a) + \frac{f^{(1)}(a)}{1!}(z-a) + \frac{f^{(2)}(a)}{2!}(z-a)^2 + \dots + \frac{f^{(n)}(a)}{n!}(z-a)^n$$

$$R_n(z) = \frac{f^{(n+1)}(a)}{(n+1)!}(z-a)^{n+1} \quad (165)$$

 **Ejemplo 1.24.** Hallar el polinomio de Taylor para $\cos x$, en $z = \pi/4$, con 4 términos.

Solución. Aplicando la serie de Taylor en las cercanías de $z = \pi/4$ se tiene:

$$f(z) = \cos z, \quad a = \pi/4$$

$$P_n(z) = f(a) + f^{(1)}(a)(z-a) + \frac{f^{(2)}(a)}{2!}(z-a)^2 + \frac{f^{(3)}(a)}{3!}(z-a)^3 + \dots$$

$$f(a) = \cos(\pi/4) = 1/\sqrt{2}, \quad f^{(1)}(a) = -\sin(\pi/4) = -1/\sqrt{2}$$

$$f^{(2)}(a) = -\cos(\pi/4) = -1/\sqrt{2}, \quad f^{(3)}(a) = \sin(\pi/4) = 1/\sqrt{2}$$

$$\cos z = \frac{1}{\sqrt{2}} - \frac{1}{\sqrt{2}}(z - \pi/4) - \frac{1}{\sqrt{2}} \frac{(z - \pi/4)^2}{2!} + \frac{1}{\sqrt{2}} \frac{(z - \pi/4)^3}{3!} + \dots$$

$$= 0.70711 \left[1 - (z - \pi/4) - \frac{1}{2}(z - \pi/4)^2 + \frac{1}{6}(z - \pi/4)^3 + \dots \right] \quad \clubsuit$$

Para $z=0$, la serie de Taylor se conoce como *serie de Maclaurin*:

$$f(z) = \sum_{k=0}^{\infty} c_k z^k = \sum_{k=0}^{\infty} \frac{f^{(k)}(0)}{k!} z^k \quad (166)$$

La **tabla 1.6** muestra las *series de Maclaurin para varias funciones elementales*, siendo $z \in \Re \vee C$:

Tabla 1.6 Series de Maclaurin para funciones elementales

$(1-z)^{-1} = \sum_{k=0}^{\infty} z^k, z < 1$	$\log\left(\frac{1+z}{1-z}\right) = 2 \sum_{k=1}^{\infty} \frac{z^{2k-1}}{2k-1}, z < 1$	$\cos^{-1} z = \frac{\pi}{2} - \sin^{-1} z, z < 1$
$(1-z)^{-2} = \sum_{k=0}^{\infty} (k+1)z^k, z < 1$	$\sin z = \sum_{k=0}^{\infty} \frac{(-1)^k z^{2k+1}}{(2k+1)!}, z < \infty$	$\tan^{-1} z = \sum_{k=1}^{\infty} \frac{(-1)^{k+1} z^{2k-1}}{2k-1}, z < 1$
$(1+z)^{\alpha} = \sum_{k=0}^{\infty} \binom{\alpha}{k} z^k, z < 1$	$\cos z = \sum_{k=0}^{\infty} \frac{(-1)^k z^{2k}}{(2k)!}, z < \infty$	$\sinh z = \sum_{k=0}^{\infty} \frac{z^{2k+1}}{(2k+1)!}, z < \infty$
$e^z = \sum_{k=0}^{\infty} \frac{z^k}{k!}, z < \infty$	$\tan z = z + \frac{z^3}{3} + \frac{2z^5}{15} + \frac{17z^7}{315} + \dots, z < \frac{\pi}{2}$	$\cosh z = \sum_{k=0}^{\infty} \frac{z^{2k}}{(2k)!}, z < \infty$
$\log(1-z) = -\sum_{k=1}^{\infty} \frac{z^k}{k}, z < 1$	$\sin^{-1} z = z + \frac{z^3}{2 \cdot 3} + \frac{1 \cdot 3 z^5}{2 \cdot 4 \cdot 5} + \frac{1 \cdot 3 \cdot 5 z^7}{2 \cdot 4 \cdot 6 \cdot 7} + \dots, z < 1$	$\tanh z = z - \frac{z^3}{3} + \frac{2z^5}{15} - \frac{17z^7}{315} + \dots, z < \frac{\pi}{2}$

Solución de ecuaciones diferenciales por series de potencias

Como hemos visto cualquier función puede ser aproximada por una serie de potencias; este hecho se aplica a la EDO de cualquier tipo que no es sino una *igualdad de funciones*:

 **Ejemplo 1.25.** Resolver la ED $(x-3)y'+2y=0$, alrededor del punto $x=0$, empleando series de potencias.

Solución. La ED es analítica en $x=0$, y su solución por hipótesis es una serie de la forma $y = \sum_{k=0}^{\infty} c_k x^k$. El *método de solución* es el siguiente: (1) se sustituye cada función y sus derivadas por series de potencias; (2) se agrupan términos convenientemente de manera tal que se pueda aplicar el *principio de identidad para sumas*; (3) se obtiene una ecuación de recurrencia para los coeficientes; (4) se encuentra un patrón para los coeficientes en términos de las n primeras constantes arbitrarias; (5) se expresa la solución como una serie de potencias; (6) se determina el radio de convergencia para la solución.

SOLUCIÓN DE UNA ED POR SERIES DE POTENCIAS : $(x-3)y' + 2y = 0$

$$y = \sum_{k=0}^{\infty} c_k x^k, \quad y' = \sum_{k=1}^{\infty} k c_k x^{k-1} = \sum_{k=0}^{\infty} (k+1) c_{k+1} x^k$$

$$(1) \quad (x-3) \sum_{k=0}^{\infty} (k+1) c_{k+1} x^k + 2 \sum_{k=0}^{\infty} c_k x^k = 0$$

$$\rightarrow x \sum_{k=0}^{\infty} (k+1) c_{k+1} x^k - 3 \sum_{k=0}^{\infty} (k+1) c_{k+1} x^k + 2 \sum_{k=0}^{\infty} c_k x^k = 0$$

$$(2) \quad \sum_{k=0}^{\infty} (k+1) c_{k+1} x^{k+1} - \sum_{k=0}^{\infty} 3(k+1) c_{k+1} x^k + \sum_{k=0}^{\infty} 2c_k x^k = 0$$

$$\rightarrow \sum_{k=0}^{\infty} (k c_k x^k - 3(k+1) c_{k+1} x^k + 2c_k x^k) = \sum_{k=0}^{\infty} (k c_k - 3(k+1) c_{k+1} + 2c_k) x^k = 0$$

$$(3) \quad (k c_k - 3(k+1) c_{k+1} + 2c_k) = 0 \Rightarrow c_{k+1} = \frac{k c_k + 2c_k}{3(k+1)} = \frac{k+2}{3(k+1)} c_k$$

$$(4) \quad c_1 = \frac{0+2}{3(0+1)} c_0 = \frac{2}{3} c_0$$

$$c_2 = \frac{1+2}{3(1+1)} c_0 = \frac{3}{6} c_1 = \frac{3}{9} c_0$$

$$c_3 = \frac{2+2}{3(2+1)} c_0 = \frac{4}{9} c_2 = \frac{4}{27} c_0$$

$$\therefore c_k = \frac{k+1}{3^k} c_0$$

$$(5) \quad y = c_0 \sum_{k=0}^{\infty} \frac{k+1}{3^k} x^k$$

$$(6) \quad R = \lim_{k \rightarrow \infty} \left| \frac{c_k}{c_{k+1}} \right| = \lim_{k \rightarrow \infty} \left| \frac{c_k}{c_{k+1}} \right|, \quad c_{k+1} = \frac{k+2}{3(k+1)} c_k = \frac{k+2}{3(k+1)} \times \frac{k+1}{3^k} c_0$$

$$R = \lim_{k \rightarrow \infty} \left| \frac{\frac{(k+1)c_0}{3^k}}{\frac{(k+2) \times (k+1)c_0}{3(k+1) \times 3^k}} \right| = \lim_{k \rightarrow \infty} \left| \frac{3(k+1)}{k+2} \right| = \lim_{k \rightarrow \infty} \left| \frac{\frac{d}{dk}(3k+3)}{\frac{d}{dk}(k+2)} \right| = \lim_{k \rightarrow \infty} \left| \frac{3}{1} \right| = 3 \leftarrow \text{Regla de L'Hospital}$$

La solución converge para : $-3 < x < 3$

♣

Solución de ED lineales de segundo orden con coeficientes variables por series de potencias

Una aplicación práctica de las series de potencia, es la solución de ED lineales de segundo orden con coeficientes variables:

$$(1) \quad A(x)y'' + B(x)y' + C(x)y = 0$$

$$(2) \quad y'' + P(x)y' + Q(x)y = 0 \quad \rightarrow \quad P(x) = B(x)/A(x), \quad Q(x) = C(x)/A(x) \quad (167)$$

Si tanto P como Q son analíticas en $x=a$, entonces a es un punto ordinario, y la ecuación (1) tiene dos soluciones linealmente independientes de la forma:

$$y(x) = \sum_{k=0}^{\infty} c_k (x-a)^k.$$

SOLUCIÓN GENERAL DE UNA ED POR SERIES DE POTENCIAS : $(x^2 + 1)y'' + xy' - y = 0$, en $x = 0$

$$y = \sum_{k=0}^{\infty} c_k x^k, \quad y' = \sum_{k=1}^{\infty} k c_k x^{k-1}, \quad y'' = \sum_{k=2}^{\infty} k(k-1) c_k x^{k-2}$$

$$(1) \Rightarrow (x^2 + 1) \sum_{k=2}^{\infty} k(k-1) c_k x^{k-2} + x \sum_{k=1}^{\infty} k c_k x^{k-1} - \sum_{k=0}^{\infty} c_k x^k = 0$$

$$\rightarrow \sum_{k=2}^{\infty} k(k-1) c_k x^k + \sum_{k=2}^{\infty} k(k-1) c_k x^{k-2} + \sum_{k=1}^{\infty} k c_k x^k - \sum_{k=0}^{\infty} c_k x^k = 0$$

$$(2) \Rightarrow \sum_{k=2}^{\infty} k(k-1) c_k x^k + [2c_2 x^0 + 6c_3 x + \sum_{k=4}^{\infty} k(k-1) c_k x^{k-2}] + [c_1 x + \sum_{k=2}^{\infty} k c_k x^k] - [c_0 x^0 + c_1 x \sum_{k=2}^{\infty} c_k x^k] = 0$$

$$\sum_{k=2}^{\infty} k(k-1) c_k x^k + [2c_2 + 6c_3 x + \sum_{k=4}^{\infty} k(k-1) c_k x^{k-2}] + [c_1 x + \sum_{k=2}^{\infty} k c_k x^k] - [c_0 + c_1 x \sum_{k=2}^{\infty} c_k x^k] = 0$$

$$[2c_2 + 6c_3 x + c_1 x - c_0 - c_1 x] + \sum_{k=2}^{\infty} [k(k-1) c_k + ((k+2)(k+1) c_{k+2}) + k c_k + k c_k] x^k = 0$$

De donde se observa que : $2c_2 - c_0 = 0 \rightarrow c_2 = \frac{1}{2} c_0 \leftrightarrow c_0 = 2c_2, c_3 = 0$ entonces :

$$\sum_{k=2}^{\infty} [k(k-1) c_k + ((k+2)(k+1) c_{k+2}) + k c_k - c_k] x^k = 0$$

$$(3) \Rightarrow k(k-1) c_k + (k+2)(k+1) c_{k+2} + k c_k - c_k = k(k-1) c_k + (k+2)(k+1) c_{k+2} + (k-1) c_k = 0$$

$$(k+1)(k-1) c_k + (k+2)(k+1) c_{k+2} = 0 \rightarrow c_{k+2} = \frac{-(k+1)(k-1) c_k}{(k+2)(k+1)} = \frac{(1-k)}{(k+2)} c_k$$

$$(4) \Rightarrow c_4 = -\frac{1}{4} c_2 = -\frac{1}{2 \cdot 4} c_0 = -\frac{1}{2^2 \cdot 2!} c_0$$

$$c_5 = -\frac{2}{5} c_3 = 0$$

$$c_6 = -\frac{3}{6} c_4 = \frac{3}{2 \cdot 4 \cdot 6} c_0 = \frac{1 \cdot 3}{2^3 \cdot 3!} c_0$$

$$c_7 = -\frac{4}{7} c_5 = 0$$

$$c_8 = -\frac{5}{8} c_6 = -\frac{3 \cdot 5}{2 \cdot 4 \cdot 6 \cdot 8} c_0 = -\frac{1 \cdot 3 \cdot 5}{2^4 \cdot 4!} c_0$$

$$(5) \Rightarrow y = c_0 + c_1 x + c_2 x^2 + c_3 x^3 + c_4 x^4 + c_5 x^5 + c_6 x^6 + c_7 x^7 + c_8 x^8 \dots$$

$$y = c_1 x + c_0 \left[1 + \frac{1}{2} x^2 - \frac{1}{2^2 \cdot 2!} x^4 + \frac{1 \cdot 3}{2^3 \cdot 3!} x^6 - \frac{1 \cdot 3 \cdot 5}{2^4 \cdot 4!} x^8 + \dots \right]$$

$$\therefore y = c_1 x + c_0 \left[1 + \frac{1}{2} x^2 + \sum_{k=2}^{\infty} (-1)^{k-1} \frac{1 \cdot 3 \cdot 5 \dots (2n-3)}{2^n \cdot n!} x^{2n} \right]$$

(6) Se tiene un punto singular en $x^2 + 1 = 0 \rightarrow s = \pm i$,

por lo tanto el radio de convergencia es de al menos $|x| < 1$ ♣

Función Gamma

Una generalización de la *función factorial* que se extiende para todo $z | z \in \mathbb{C} \wedge \text{Re}[z] > 0$ es la *función Gamma* definida por:

$$\Gamma(z) = \int_0^{\infty} e^{-t} t^{z-1} dt \quad \leftrightarrow \quad \text{Re}[z] > 0 \quad (168)$$

Usando integración por partes se demuestra la siguiente propiedad más importante de la función Gamma:

$$\begin{aligned} \Gamma(z+1) &= \int_0^{\infty} e^{-t} t^z dt = \left[-e^{-t} t^z \right]_0^{\infty} - \int_0^{\infty} -e^{-t} z t^{z-1} dt \\ \text{El término } \left[-e^{-t} t^z \right]_0^{\infty} &\text{ se descarta ya que la función se define para } \text{Re}[z] > 0, \text{ entonces :} \\ \Gamma(z+1) &= z \int_0^{\infty} e^{-t} t^{z-1} dt = z \Gamma(z) \end{aligned} \quad (169)$$

La *fórmula de Gauss* para la función Gamma esta dada por:

$$\Gamma(z) = \frac{1}{z} \prod_{k=1}^{\infty} \left[\left(1 + \frac{1}{k} \right)^z \left(1 + \frac{z}{k} \right)^{-1} \right] \quad (170)$$

La *fórmula de Weierstrass* para la función Gamma esta dada por:

$$\begin{aligned} \frac{1}{\Gamma(z)} &= z e^{\gamma z} \prod_{k=1}^{\infty} \left[\left(1 + \frac{z}{k} \right)^{-1} e^{-z/k} \right] \\ \gamma &= \lim_{k \rightarrow \infty} \left[\left(1 + \frac{1}{2} + \frac{1}{3} + \dots + \frac{1}{k} + \dots \right) - \log k \right] = 0.5772 \dots \leftarrow \text{Constante Euler - Mascheroni} \end{aligned} \quad (171)$$

Transformación de Fourier

Si $f(t)$ es una *función periódica*, de *periodo* T_0 (frecuencia angular fundamental $\omega_0 = 2\pi / T_0$), $f(t) = f(t + T_0)$; entonces la integración de f a lo largo de un intervalo de longitud igual a su periodo, es independiente de los límites de integración siempre que: $L_{\text{sup}} - L_{\text{inf}} = T_0$:

$$\int_0^{T_0} f(t) dt = \int_{-T_0/2}^{T_0/2} f(t) dt = \int_{-T_0}^0 f(t) dt \quad (172)$$

La siguiente es una lista de *integrales que simplifican el análisis de Fourier*:

$$\begin{aligned} \int_0^{T_0} \sin(k\omega_0 t) dt &= \int_0^{T_0} \cos(k\omega_0 t) dt = 0, & k, n \in \mathbb{Z}^+ \wedge k \neq n \\ \int_0^{T_0} \sin(k\omega_0 t) \cos(k\omega_0 t) dt &= \int_0^{T_0} \sin(k\omega_0 t) \sin(n\omega_0 t) dt = \int_0^{T_0} \cos(k\omega_0 t) \cos(n\omega_0 t) dt = 0 \\ \int_0^{T_0} \sin^2(k\omega_0 t) dt &= \int_0^{T_0} \cos^2(k\omega_0 t) dt = \frac{T_0}{2} \end{aligned} \quad (173)$$

Serie de Fourier

Una *función de tiempo periódica* $f(t) = f(t + T_0)$, de *periodo fundamental* T_0 , y *frecuencia fundamental* $f_0 = \omega_0 / 2\pi \rightarrow \omega_0 = 2\pi f_0 = 2\pi / T_0 \leftrightarrow T_0 = 1 / f_0$, puede ser representada por una *serie trigonométrica* conocida como *serie de Fourier* (demostrada por **Joseph Fourier**, en su tratado *The Analytical Theory of Heat*, 1822); se define por:

$$\begin{aligned} f(t) = f(t + T_0) &= A_0 + \sum_{k=1}^{\infty} \left(A_k \cos \frac{2\pi k}{T_0} t + B_k \sin \frac{2\pi k}{T_0} t \right), & A_0 &= \frac{1}{T_0} \int_{-T_0/2}^{T_0/2} f(t) dt \\ A_k &= \frac{2}{T_0} \int_{-T_0/2}^{T_0/2} f(t) \cos \left(\frac{2\pi k}{T_0} t \right) dt, & B_k &= \frac{2}{T_0} \int_{-T_0/2}^{T_0/2} f(t) \sin \left(\frac{2\pi k}{T_0} t \right) dt \end{aligned} \quad (174)$$

Una *función con simetría par* $f(-t) = f(t)$, conserva sólo los *términos coseno* de la serie; una *función con simetría impar* $f(-t) = -f(t)$, conserva sólo los *términos senos* de la serie. De manera alternativa, la serie de Fourier se representa por:

$$\begin{aligned} f(t) &= M_0 + \sum_{k=1}^{\infty} \left(M_k \cos \frac{2\pi k}{T_0} t - \phi_k \right) \\ M_0 &= A_0, & M_k &= \sqrt{A_k^2 + B_k^2}, & \phi_k &= \tan^{-1} \frac{B_k}{A_k} \end{aligned} \quad (175)$$

De la expresión anterior se observa que la serie de fourier es una sumatoria de las *armónicas* ($k f_0$) de la *frecuencia fundamental* (f_0). Los coeficientes M_k , se denominan *amplitudes espectrales*.

Convergencia. Si $f(t)$ es *suave por partes*, entonces su *serie de Fourier converge*: (a) al valor $f(t)$ para todo t donde $f(t)$ es *continua*; (b) al valor $\frac{1}{2}[f(t^+) + f(t^-)]$ para toda t donde $f(t)$ es *discontinua*. Una *función* $f(t)$ es *continua por partes* para toda t cuando es *continua por partes en cada intervalo acotado* y en los puntos de discontinuidad existen *limites laterales finitos*:

$$f(t^+) = \lim_{u \rightarrow t^+} f(u), \quad f(t^-) = \lim_{u \rightarrow t^-} f(u) \quad (176)$$

Forma exponencial de la serie de fourier. La serie de Fourier se puede expresar en términos de la función exponencial compleja considerando la equivalencia exponencial de las funciones seno y coseno:

$$f(t) = \sum_{k=-\infty}^{\infty} C_k e^{j \frac{2\pi k}{T_0} t}, \quad C_k = \frac{1}{T_0} \int_{-T_0/2}^{T_0/2} f(t) e^{-j \frac{2\pi k}{T_0} t} dt \quad (177)$$

$$C_0 = M_0, \quad C_k = \frac{M_k}{2} e^{-j\phi_k}$$

Siendo C_k la *amplitud espectral* de los *componentes espectrales* $e^{j2\pi k f_0 t}$.

Ejemplo 1.26. Determinar la *serie de fourier* para el *tren de impulsos* de longitud I , mostrado en la **figura 1.13**. El impulso en $t=0$ se define por $I\delta(t)$, siendo $\delta(t)$ la *función delta* definida por: $\delta(t - T_0) = 1$ en $t = T_0$ y $\delta(t) = 0$ para $t \neq T_0$, además $\int_{-\infty}^{\infty} \delta(t) dt = 1$, su *longitud es el área bajo el impulso* 1. El *tren de impulsos periódicos*, de *periodo* T_0 se define por: $f(t) = I \sum_{k=-\infty}^{\infty} \delta(t - kT_0)$:

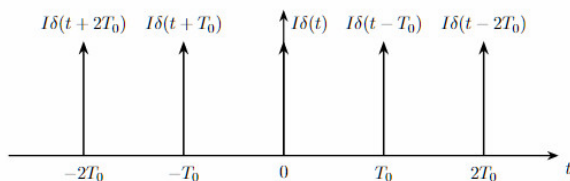


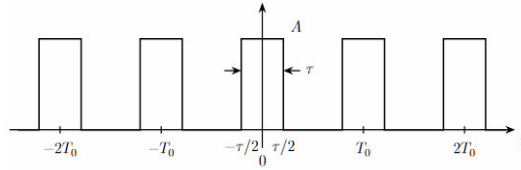
Figura 1.13 Tren de impulsos de longitud I y periodo T_0 .

Solución. Se determina la *serie de Fourier* en sus formas: (1) *trigonométrica*; (2) *trigonométrica compacta* y (3) *exponencial compleja*.

$$\begin{aligned}
 (1) \quad A_0 &= \frac{I}{T_0} \int_{-T_0/2}^{T_0/2} \delta(t) dt = \frac{I}{T_0} \\
 A_k &= \frac{2I}{T_0} \int_{-T_0/2}^{T_0/2} \delta(t) \cos\left(\frac{2\pi k}{T_0} t\right) dt = \frac{2I}{T_0} \\
 B_k &= \frac{2I}{T_0} \int_{-T_0/2}^{T_0/2} \delta(t) \sin\left(\frac{2\pi k}{T_0} t\right) dt = 0 \\
 f(t) &= A_0 + \sum_{k=1}^{\infty} \left(A_k \cos \frac{2\pi k}{T_0} t + B_k \sin \frac{2\pi k}{T_0} t \right) = \frac{I}{T_0} + \frac{2I}{T_0} \sum_{k=1}^{\infty} \cos\left(\frac{2\pi k}{T_0} t\right) \\
 (2) \quad M_0 &= \frac{I}{T_0}, \quad M_k = \frac{2I}{T_0}, \quad \phi_k = 0, \\
 f(t) &= M_0 + \sum_{k=1}^{\infty} \left(M_k \cos \frac{2\pi k}{T_0} t - \phi_n \right) = \frac{I}{T_0} + \sum_{k=1}^{\infty} \left(\frac{2I}{T_0} \cos \frac{2\pi k}{T_0} t \right) \\
 (3) \quad C_0 &= C_k = \frac{I}{T_0} \\
 f(t) &= \sum_{k=-\infty}^{\infty} C_k e^{\left(\frac{j2\pi k}{T_0} t\right)} = \frac{I}{T_0} \sum_{k=-\infty}^{\infty} e^{\left(\frac{j2\pi k}{T_0} t\right)} \quad \clubsuit
 \end{aligned}$$

Ejemplo 1.27. La función $f(t)$ es un *tren de pulsos rectangulares* de amplitud A y duración τ , que se repite periódicamente cada T_0 segundos, como se muestra en la **figura 1.14**. Determinar la *serie exponencial de Fourier* para $f(t)$:

Figura 1.14 Tren de pulsos rectangulares de amplitud A , duración τ y periodo T_0 .



Solución. Se encuentran los parámetros A_0 , A_k , B_k y se sustituyen en la serie de Fourier trigonométrica:

$$\begin{aligned}
 A_0 &= M_0 = C_0 = \frac{1}{T_0} \int_{-T_0/2}^{T_0/2} f(t) dt = \frac{A\tau}{T_0} \\
 A_k &= M_k = 2C_k = \frac{2}{T_0} \int_{-T_0/2}^{T_0/2} f(t) \cos\left(\frac{2\pi k t}{T_0}\right) dt = \frac{2A\tau}{T_0} \frac{\sin(k\pi\tau/T_0)}{k\pi\tau/T_0}, \quad B_k = 0, \quad \phi_k = 0 \\
 \rightarrow f(t) &= \frac{A\tau}{T_0} + \frac{2A\tau}{T_0} \sum_{k=1}^{\infty} \frac{\sin(k\pi\tau/T_0)}{k\pi\tau/T_0} \cos\left(\frac{2\pi k t}{T_0}\right) = \frac{A\tau}{T_0} \sum_{k=-\infty}^{\infty} \frac{\sin(k\pi\tau/T_0)}{k\pi\tau/T_0} e^{\left(\frac{j2\pi k t}{T_0}\right)} \quad \clubsuit
 \end{aligned}$$

La función pulso se convierte en la función impulso, cuando su anchura tiende a cero, esto es fácil mostrar mediante límites, ya que:

$$\lim_{\tau \rightarrow 0} \frac{\sin(k\pi\tau / T_0)}{k\pi\tau / T_0} = 1 \quad (178)$$

La función muestreo

La *función muestreo* escrita como $Sa(x)$, se dibuja en la **figura 1.15a**, tiene simetría par, se emplea en el *análisis espectral*, se define y tiene las propiedades mostradas a continuación:

$$\begin{aligned} Sa(x) &= \frac{\sin x}{x}, & \text{Max } Sa(x) &= Sa(0) = 1, & Sa(x = \pm k\pi) &= 0 \\ Sa[\pm(k + \frac{1}{2})\pi] &= \frac{2(-1)^n}{(2n+1)\pi} \end{aligned} \quad (179)$$

La **figura 1.15b** muestra las amplitudes espectrales y la *envolvente* de la representación en series de Fourier para el tren de pulsos rectangulares de la **figura 1.14**.

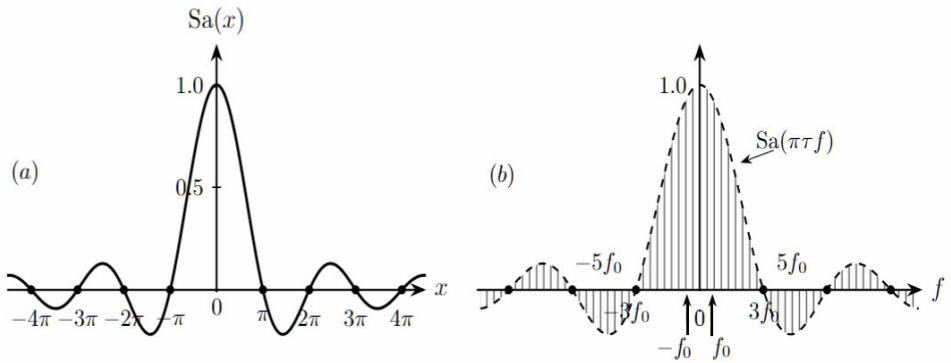


Figura 1.15 (a) función muestreo $Sa(x)$; (b) amplitudes espectrales C_k para la representación en series de Fourier del tren de pulsos rectangulares.

La transformación de Fourier

La *transformada de Fourier* es una operación que transforma una función del dominio del tiempo al dominio de la frecuencia; la operación inversa es la *transformada inversa de Fourier* y se definen por:

$$\begin{aligned} (1) \quad F(j\omega) &= \int_{-\infty}^{\infty} e^{-j\omega t} f(t) dt & (2) \quad F(f) &= \int_{-\infty}^{\infty} e^{-j2\pi \cdot f t} f(t) dt \Leftrightarrow f = \frac{\omega}{2\pi} & (180) \\ f(t) &= \frac{1}{2\pi} \int_{-\infty}^{\infty} e^{-j\omega t} F(j\omega) d\omega & f(t) &= \int_{-\infty}^{\infty} e^{-j2\pi \cdot f t} F(f) df \Leftrightarrow f = \frac{\omega}{2\pi} \end{aligned}$$

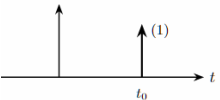
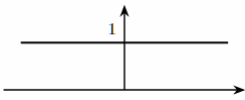
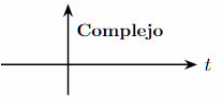
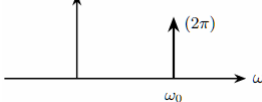
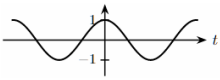
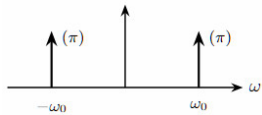
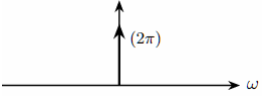
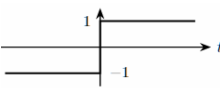
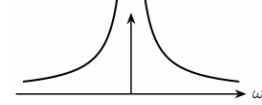
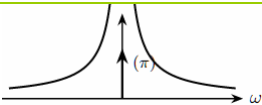
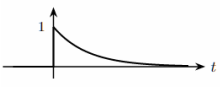
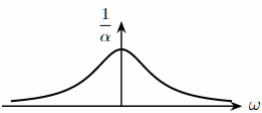
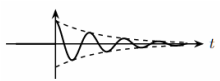
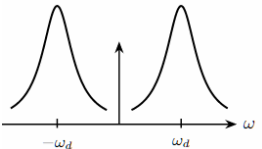
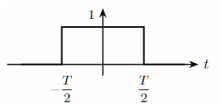
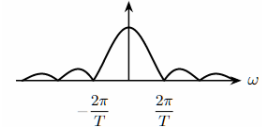
 **Ejemplo 1.28.** Determinar la transformada de Fourier para la función $f(t) = \cos \omega_0 t$.

Solución. Aplicando la ecuación definitoria (2) e identificando la función impulso unitario, se tiene que:

$$\begin{aligned} f(t) = \cos \omega_0 t &= \frac{1}{2} e^{j\omega_0 t} + \frac{1}{2} e^{-j\omega_0 t}, \quad \omega_0 = \frac{2\pi}{T_0} \\ F(f) &= \int_{-\infty}^{\infty} \cos(\omega_0 t) e^{-j2\pi f t} dt = \frac{1}{2} \int_{-\infty}^{\infty} e^{-j2\pi(f-f_0)t} dt + \frac{1}{2} \int_{-\infty}^{\infty} e^{-j2\pi(f+f_0)t} dt \\ &= \frac{1}{2} \delta(f-f_0) + \frac{1}{2} \delta(f+f_0) \quad \clubsuit \end{aligned}$$

Basado en las integrales que definen la transformada directa e inversa de Fourier, se han desarrollado fórmulas para funciones comunes, como la mostrada en **tabla 1.7**:

Tabla 1.7 Transformada de Fourier directa e inversa para funciones comunes

Gráfico $f(t)$	$f(t)$	$F[f(t)] = F(j\omega)$	Gráfico $ F(j\omega) $
	$\delta(t - t_0)$	$e^{-j\omega t_0}$	
	$e^{j\omega_0 t}$	$2\pi\delta(\omega - \omega_0)$	
	$\cos \omega_0 t$	$\pi\delta(\omega + \omega_0) + \pi\delta(\omega - \omega_0)$	
	1	$2\pi\delta(\omega)$	
	$\text{sgn}(t) = \begin{cases} -1 & t < 0 \\ 1 & t > 0 \end{cases}$	$\frac{2}{j\omega}$	
	$u(t)$	$\pi\delta(\omega) + \frac{1}{j\omega}$	
	$e^{-\alpha t} u(t)$	$\frac{1}{\alpha + j\omega}$	
	$e^{-\alpha t} \cos \omega_d t \cdot u(t)$	$\frac{\alpha + j\omega}{(\alpha + j\omega)^2 + \omega_d^2}$	
	$u(t + \frac{1}{2}T) - u(t - \frac{1}{2}T)$	$T \frac{\sin(\omega T / 2)}{\omega T / 2}$	

 **Ejemplo 1.29.** Emplear la tabla para determinar la transformada de Fourier de $f(t) = 3e^t \cos 4t \cdot u(t)$.

Solución. Identificando la función en la tabla se tiene que:

$$e^{-\alpha t} \cos \omega_d t \cdot u(t) \Leftrightarrow \frac{\alpha + j\omega}{(\alpha + j\omega)^2 + \omega_d^2}, \quad \alpha = 1 \quad \omega_d = 4 \rightarrow F(j\omega) = 3 \frac{1 + j\omega}{(\alpha + j\omega)^2 + 16} \quad \clubsuit$$

El *teorema de Parseval* nos permite hallar la *densidad de energía o energía por unidad de ancho de banda* (J/Hz) para una función $f(t)$:

$$\int_{-\infty}^{\infty} f^2(t) dt = \frac{1}{2\pi} \int_{-\infty}^{\infty} |F(j\omega)|^2 d\omega \quad (181)$$

Para un circuito, se define la *función del sistema* $h(t)$ como la *transformada de Fourier de la respuesta al impulso unitario en $t=0$* del circuito. En el dominio de la frecuencia compleja, la función del sistema $H(j\omega)$ es idéntica a la *función de transferencia del sistema* $G(s)$ y es una relación entre la transformada de Fourier y la transformada de Laplace para una función lineal, es decir:

$$H(j\omega) = G(s) = \frac{V_o}{V_i} = \frac{L[y]}{L[x]} \quad (182)$$

Sean dos funciones $y(t)$ y $x(t)$ y $Y(j\omega)$, $X(j\omega)$ sus respectivas transformadas de Fourier; entonces se define la *convolución* de y y x como la integral cuya transformada de Fourier es el producto de las transformadas de Y y X , es decir:

$$\begin{aligned} Y(j\omega) &= F[y(t)], & X(j\omega) &= F[x(t)], \\ F[v(t)] &= Y(j\omega) * X(j\omega) \rightarrow v(t) = F^{-1}[Y(j\omega) * X(j\omega)] \quad \text{donde :} \\ v(t) &= \int_{-\infty}^{\infty} y(\tau)x(t-\tau)d\tau = \int_{-\infty}^{\infty} x(\tau)y(t-\tau)d\tau \leftarrow \text{Convolución de } x \text{ y } y \end{aligned} \quad (183)$$

Si en una red lineal, se conoce la entrada al circuito $x(t)$ y la *función del sistema o respuesta al impulso* $h(t)$, entonces la salida $y(t)$ está dada por la convolución de x y h , y se define por:

$$y(t) = \int_{-\infty}^{\infty} x(\tau)h(t-\tau)d\tau \quad (184)$$

🔗 **Ejemplo 1.30.** Para el circuito RL mostrado en la **figura 1.16**, determinar el voltaje en el inductor cuando el voltaje de entrada es un pulso exponencial decreciente. Emplear las propiedades de convolución.

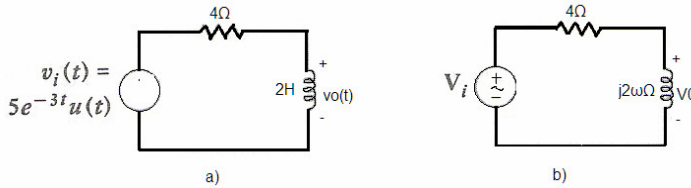


Figura 1.16 Circuito RL: (a) dominio del tiempo; (b) dominio de la frecuencia $j\omega$

Solución. (1) Se determina la función del sistema mediante un análisis senoidal permanente; (2) se determina la transformada de Fourier de la función de excitación; (3) se aplica el principio de convolución para determinar la respuesta del sistema:

$$(1) \quad H(j\omega) = \frac{V_o}{V_i} = \frac{j2\omega}{4 + j2\omega}$$

$$(2) \quad F[5e^{-3t} \cdot u(t)] = \frac{5}{3 + j\omega}$$

$$(3) \quad F[v_o(t)] = H(j\omega)F[v_i(t)] = \frac{j2\omega}{4 + j2\omega} \cdot \frac{5}{3 + j\omega} = \frac{15}{3 + j\omega} - \frac{10}{2 + j\omega} = 15e^{-3t} \cdot u(t) - 10e^{-2t} \cdot u(t) \quad \clubsuit$$

Transformación zeta

Las *computadoras digitales*, los *controladores lógicos programables* (PLC) y los *microcontroladores* poseen una *interfaz de entrada* que convierte *señales continuas* en *señales discretas*; el *microprocesador*, su núcleo cerebral, solo trabaja señales discretas. Una forma de considerar estas señales es como *señales continuas del tiempo que se han muestreado a intervalos regulares*; el resultado es una *secuencia en tiempo discreto*. La transformada zeta es un método matemático para analizar señales discretas.

Sistemas de datos discretos

Considérese el sistema de control mostrado en la **figura 1.17**, donde se usa un microprocesador programado para implementar la acción de control. (1) la entrada al sistema es una señal analógica, que se convierte en una señal discreta mediante un *convertidor analógico-digital* (ADC); (2) el *microprocesador* aplica la estrategia de control, de acuerdo a un programa almacenado en una memoria ROM con la cual se comunica; (3) la *señal de salida digital* del microprocesador es convertida en una señal analógica mediante un *convertidor digital-analógico* (DAC); (4) la *señal analógica* resultante se usa para manejar la unidad de corrección y así controlar la planta variable.

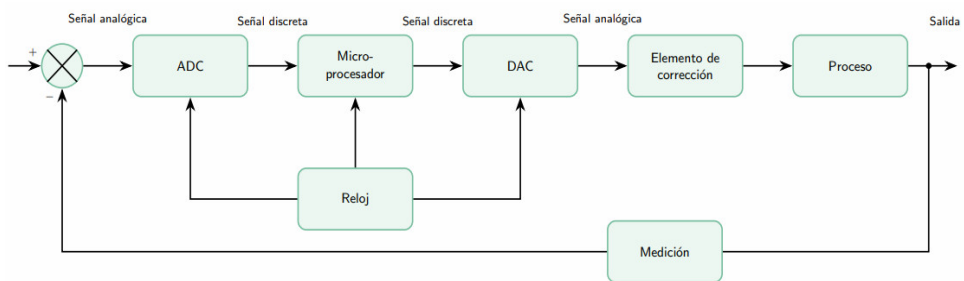


Figura 1.17 Sistema de control con datos muestreados

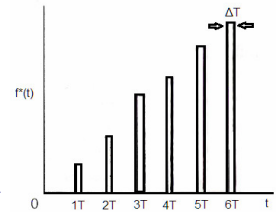
Un *reloj* alimenta un pulso cada T segundos (*periodo de muestreo*), cada vez que el ADC recibe un pulso, este muestrea la señal de error. La entrada al microprocesador es entonces una serie de pulsos.

Consideremos una señal de tiempo continuo $f(t)$ que se muestrea en intervalos regulares de tiempo con periodo de muestreo T y se representa por $f^*(T)$, para cada T , el valor de $f(kT)$ se toma durante un intervalo corto $\Delta t \ll T$, lo suficientemente pequeño para que su magnitud se considere constante durante el tiempo que se toma la muestra. La *secuencia de impulsos* se denota por:

$$f(t=0), f(t=1T), f(t=2T), \dots, f(t=kT), \dots = f(0), f(T), f(2T), \dots, f(kT), \dots = f[kT] \quad (185)$$

La *función delta* se emplea para representar señales discretas, esta función se define por $\delta(t - T_0) = 1$ en $t = T_0$ y $\delta(t) = 0$ para $t \neq T_0$. De esta manera la señal de pulsos de la *función rampa* mostrada en la **figura 1.18** se representa por:

Figura 1.18
Función rampa



$$0\delta(t), 1\delta(t=1T), 2\delta(t=2T), 3\delta(t=3T), \dots, k\delta(t=kT), \dots \quad (186)$$

En general la *señal de datos muestreados*, empleando la notación $f[kT] \rightarrow k \in \mathbb{Z}^{0+}$, se puede describir mediante la siguiente *función discreta*:

$$f^*(t) = f[0]\delta(t) + f[1]\delta(t-1T) + \dots + f[k]\delta(t-kT) + \dots + f[n]\delta(t-nT) = \sum_{k=0}^n f[k]\delta(t-kT) \quad (187)$$

Ecuaciones en diferencias

En un *sistema de procesamiento de señales en tiempo discreto* existe la entrada de una secuencia de pulsos y una salida de pulsos; las expresiones matemáticas que involucran secuencias discretas se denominan *ecuaciones en diferencias*. Así por ejemplo la *secuencia de entrada* $x[k]$ al controlador, se puede relacionar con su *secuencia de salida* $y[k]$ mediante la siguiente ecuación en diferencias:

$$y[k] = a \cdot y[k-1] + x[k], \quad a \equiv \text{Factor de amplificación} \rightarrow a \in \mathbb{R}^+ \quad (188)$$

Las ecuaciones en diferencias se pueden representar mediante los *bloques funcionales*: punto, suma, retardo unitario y *amplificación a*. La **figura 1.19** muestra el bloque funcional de la ecuación anterior.

Cuando se desea *diferenciar la secuencia de entrada* se emplea la *aproximación tangencial* entre dos muestras sucesivas de entrada $x[k-1]$, $x[k]$. Cuando se desea *integrar la secuencia de entrada*, se emplea la *aproximación trapecoidal* para hallar el área bajo la curva de dos muestras sucesivas de entrada $x[k-1]$, $x[k]$. De esta forma la salida $y[k]$ se relaciona con la derivada e integral de su entrada $x[k]$, mediante las siguientes ecuaciones en diferencias:

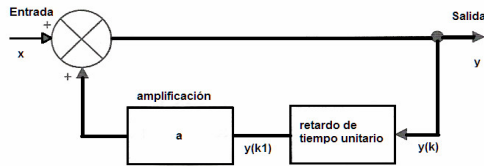


Figura 1.19 Diagrama de bloques con amplificación y retardo unitario

$$\begin{aligned} y[k] &= (x[k] - x[k-1]) / T \quad \Rightarrow \quad x[k] = x[k-1] + T \times y[k] & \leftarrow \text{Diferenciación} \\ y[k] &= y[k-1] + \frac{1}{2} T \times (x[k-1] + x[k]) & \leftarrow \text{Integración} \end{aligned} \quad (189)$$

Las ecuaciones en diferencias para funciones discretas, son el análogo a las ecuaciones diferenciales para funciones continuas.

Transformada zeta

Sea $f^*(t)$ la función discreta que describe una secuencia de impulsos. Aplicando la transformada de Laplace y al hacer $z = e^{Ts}$, se obtiene una *función transformada en términos de la variable zeta* $F(z)$:

$$\begin{aligned} F(s) &= \mathcal{L}[f^*(t)] = \mathcal{L}\left[\sum_{k=0}^n f[k] \delta(t - kT)\right] = f[0] + f[1]e^{-Ts} + \dots + f[k]e^{-kTs} + \dots + f[n]e^{-nTs} \\ &= \sum_{k=0}^n f[k] e^{-kTs} \\ \text{si } z &= e^{Ts} \rightarrow s = \frac{1}{T} \ln z \text{ entonces:} \\ F(z) &= Z[f(k)] = f[0] + f[1]z^{-1} + \dots + f[k]z^{-k} + \dots + f[n]z^{-n} = \sum_{k=0}^n f[k] z^{-k} \end{aligned} \quad (190)$$

🔗 **Ejemplo 1.31.** Hallar la transformada zeta para (a) la *función escalón unitario*, $f(t) = 1 \leftrightarrow t > 0 \wedge f(t) = 0 \leftrightarrow t \leq 0$; (b) la *función rampa*, $f(t) = t \leftrightarrow t \wedge f(t) = 0 \leftrightarrow t < 0$.

Solución. Aplicando la transformación zeta se tiene que:

$$(a) \quad F(z) = \sum_{k=0}^{\infty} f[k]z^{-k} = 1z^0 + 1z^{-1} + 1z^{-2} + 1z^{-3} + \dots = 1 + z^{-1} + z^{-2} + z^{-3} + \dots = 1 + (1/z) + (1/z)^2 + (1/z)^3 + \dots$$

Por comparación con la serie geométrica: $1 + x + x^2 + \dots = \frac{1}{1-x}$, para $|x| < 1$; se tiene que:

$$F(z) = \frac{1}{1-1/z} = \frac{z}{z-1} \text{ para } |z| > 1$$

$$(b) \quad F(z) = \sum_{k=0}^{\infty} f[k]z^{-k} = 0z^0 + 1z^{-1} + 2z^{-2} + 3z^{-3} + \dots = T(z^{-1} + 2z^{-2} + 3z^{-3} + \dots) \\ = Tz^{-1}(1 + 2z^{-1} + 3z^{-2} + \dots) \rightarrow \frac{zF(z)}{T} = 1 + 2z^{-1} + 3z^{-2} + \dots$$

Por comparación con la serie binomial: $(1-x)^{-2} = 1 + 2x + 3x^2 + \dots$ se tiene que:

$$\frac{zF(z)}{T} = \frac{1}{(1-1/z)^2} \rightarrow F(z) = \frac{T}{z(1-1/z)^2} = \frac{T}{z \left[\frac{z}{z} (1-1/z) \right]^2} = \frac{T}{z \left[\frac{1}{z} (z-1) \right]^2} = \frac{T}{z \frac{1}{z^2} (z-1)^2} = \frac{Tz}{(z-1)^2} \text{ para } |z| > 1$$

♣

La transformación zeta es una (1) *operación lineal*, y se aplican los siguientes teoremas: (2) *corrimiento*; (3) *translación compleja*; (4) *valor inicial*; (5) *valor final*:

$$(1) \quad Z[f(k) + g(k)] = Z[f(k)] + Z[g(k)], \quad Z[af(k)] = aZ[f(k)]$$

$$(2) \quad Z[f(k+n)] = z^n F(z) - (z^n f[0] + z^{n-1} f[1] + z^{n-2} f[2] + \dots + z f[n-1]) \\ = z^n F(z) - \sum_{k=0}^{n-1} z^{n-k} f[k]$$

$$Z[f(k+1)] = z^1 F(z) - \sum_{k=0}^{1-1} z^{1-k} f[k] = zF(z) - zf[0]$$

$$Z[f(k+2)] = z^2 F(z) - \sum_{k=0}^{2-1} z^{2-k} f[k] = zF(z) - z^2 f[0] - zf[1]$$

(191)

$$(3) \quad Z[e^{-akT} f(k)] = F[e^{aT} z] \quad \leftrightarrow \quad F(z) = Z[f(k)]$$

$$(4) \quad f[0] = \lim_{t \rightarrow 0} f[k] = \lim_{z \rightarrow \infty} F[z]$$

$$(5) \quad f[\infty] = \lim_{t \rightarrow \infty} f[k] = \lim_{z \rightarrow 1} (1 - z^{-1}) F(z) = \lim_{z \rightarrow 1} \frac{z-1}{z} F(z)$$

Transformada zeta inversa

Dada una función $F(z)$, se obtiene la función discreta $f^*(t)$ mediante la *transformada zeta inversa*, definida por

$$f^*(t) = Z^{-1}[F(z)] = \sum_{k=0}^{\infty} f[k] \delta(t - kT) \quad (192)$$

La **tabla 1.8** muestra la transformada zeta de: **(a)** funciones muestreadas con periodo T ; **(b)** secuencias.

Tabla 1.8 Transformadas zeta de: **(a)** funciones muestreadas con periodo T ; **(b)** secuencias

a) $f(t)$	$F(z)$	$f(t)$	$F(z)$	b) $f[k]$	$f[0], f[1], f[2], f[3], \dots$	$F(z)$
$\delta(t)$	1	$1 - e^{-at}$	$\frac{z(1 - e^{-aT})}{(z-1)(z - e^{-aT})}$	$lu[k]$	1,1,1,1,...	$\frac{z}{z-1}$
$\delta(t - kT)$	z^{-k}	te^{-at}	$\frac{Tze^{-aT}}{(z - e^{-aT})^2}$	a^k	$a^0, a^1, a^2, a^3, \dots$	$\frac{z}{z-a}$
$l(t)$	$\frac{z}{z-1}$	$e^{-at} - e^{-bt}$	$\frac{(e^{-aT} - e^{-bT})z}{(z - e^{-aT})(z - e^{-bT})}$	k	0,1,2,3,...	$\frac{z}{(z-1)^2}$
$l(t - kT)$	$\frac{z}{z^k(z-1)}$	$\sin \omega t$	$\frac{z \sin \omega T}{z^2 - 2z \cos \omega T + 1}$	ka^k	$0, a^1, 2a^2, 3a^3, \dots$	$\frac{az}{(z-a)^2}$
t	$\frac{Tz}{(z-1)^2}$	$\cos \omega t$	$\frac{z(z - \cos \omega T)}{z^2 - 2z \cos \omega T + 1}$	ka^{k-1}	$0, a^0, 2a^2, 3a^3, \dots$	$\frac{z^2}{(z-a)^2}$
t^2	$\frac{T^2 z(z+1)}{(z-1)^3}$	$e^{-at} \sin \omega t$	$\frac{ze^{-aT} \sin \omega T}{z^2 - 2ze^{-aT} \cos \omega T + e^{-2aT}}$	e^{-ak}	$e^0, e^{-a}, e^{-2a}, e^{-3a}, \dots$	$\frac{z}{z - e^{-a}}$
e^{-at}	$\frac{z}{z - e^{-aT}}$	$e^{-at} \cos \omega t$	$\frac{z(z - \cos \omega T)}{z^2 - 2z \cos \omega T + 1}$			

La *determinación de la transformada zeta inversa* se puede hallar mediante tres formas: **(1)** usando la ecuación definitoria; **(2)** usar un listado como en la tabla 1.7, simplificando $F(z)$ mediante fracciones parciales cuando así convenga; **(3)** realizar una *división larga* $F(z) = N(z)/D(z) \cdot$

🔗 **Ejemplo 1.32.** Determinar la *transformada zeta inversa* de $F(z) = \frac{z}{(z-1)(z-0.5)}$

mediante *expansión por fracciones parciales*.

Solución. En la transformada zeta, se prefiere obtener las fracciones parciales para $F(z)/z$ en vez de $F(z)$ directamente, ya que esto conduce a funciones comunes, es decir:

$$\begin{aligned}\frac{F(z)}{z} &= \frac{1}{(z-1)(z-0.5)} = \frac{a_0}{(z-1)} + \frac{b_0}{(z-0.5)} \\ a_0 &= \frac{1}{0!} \frac{d^0}{dz^0} \left(\frac{1}{z-0.5} \right) \Big|_{z=1} = \frac{1}{1-0.5} = 2 \\ b_0 &= \frac{1}{0!} \frac{d^0}{dz^0} \left(\frac{1}{z-1} \right) \Big|_{z=0.5} = \frac{1}{0.5-1} = -2 \\ \rightarrow F(z) &= \frac{2z}{(z-1)} - \frac{2z}{(z-0.5)} \Rightarrow f[k] = 2u[k] - 2 \times 0.5^k \quad \clubsuit\end{aligned}$$

🔗 **Ejemplo 1.33.** Determinar la *transformada zeta inversa* de $F(z) = \frac{2z}{z^2+z+1}$ por *división larga*.

Solución. Se realiza la división larga y se tiene que:

$$\begin{array}{r} 2z^{-1} - 2z^{-2} + 2z^{-4} - 2z^{-5} + \dots \\ z^2 + z + 1 \overline{) 2z} \\ \underline{2z + 2z^0 + 2z^{-1}} \\ -2z^0 - 2z^{-1} \\ \underline{-2z^0 - 2z^{-1} - 2z^{-2}} \\ 2z^{-2} \\ \underline{2z^{-2} + 2z^{-3} + 2z^{-4}} \\ -2z^{-3} - 2z^{-4} \\ \underline{-2z^{-3} - 2z^{-4} - 2z^{-5}} \\ 2z^{-5} \end{array}$$

$$\therefore F(z) = 0z^0 + 2z^{-1} - 2z^{-2} + 0z^{-3} + 2z^{-4} - 2z^{-5} + \dots \Rightarrow y[k] = 0, 2, -2, 0, 2, -2 \quad \clubsuit$$

Capítulo 2

Bloques Funcionales de Sistemas Físicos

Una vez sentadas las bases matemáticas estamos en disposición de explorar los componentes físicos elementales cuyo comportamiento se rige por leyes físicas bien definidas y que en unión dan lugar a la formación de sistemas físicos complejos. En este capítulo se persiguen los siguientes objetivos:

- Presentar los fundamentos teóricos que rigen el comportamiento de los sistemas físicos básicos: *mecánicos, eléctricos, fluidicos, térmicos.*
- Observar las equivalencias que existe entre los bloques funcionales de sistemas físicos.
- Enunciar las leyes y principios de la física en un lenguaje matemático.

Introducción

Un **bloque funcional** es un *elemento físico* simple, que se caracteriza por tener una *propiedad física* que se puede describir mediante las *leyes y principios de la física*. La descripción se logra formando la ecuación descriptiva que lo representa. Un *sistema físico* que consta de *elementos lineales* se estudia y representa por un *sistema de ecuaciones diferenciales lineales*; en ingeniería cuando un sistema no es lineal se buscan métodos para *aproximarlo a la linealidad* bajo ciertas *condiciones impuestas*. Un sistema se considera lineal si es posible aplicar el **principio de superposición**, que establece: “*la respuesta producida por la aplicación simultánea de n funciones de excitación diferentes es la suma de las n respuestas producidas por las n funciones actuando individualmente*”.

La manera en que se presentarán los bloques funcionales físicos es la siguiente: (1) se enunciarán los *principios y leyes fundamentales*; (2) se presentará su equivalente en términos de ecuaciones diferenciales con cantidades vectoriales y/o escalares según sea su naturaleza haciendo notar su *comportamiento lineal*.

Sistema General de Unidades de Medida

Se pretende presentar las ecuaciones físicas en su forma más simple y por lo tanto se presentaran las unidades de las cantidades en conjunto con sus ecuaciones, en su lugar se presentan a continuación las tablas de las *cantidades físicas fundamentales, cantidades derivadas* más empleadas y *constantes físicas* ampliamente usadas. La **Norma Oficial Mexicana NOM-008-SCFI-1993, Sistema General de Unidades de Medida** (publicada el 8 de octubre de 1993 en el *Diario Oficial de la Federación*), tiene por objeto *establecer un lenguaje común que responda a las exigencias actuales de las actividades científicas, tecnológicas, educativas, industriales y comerciales*; y está basada en el *Sistema Internacional de Unidades*.

Cantidades fundamentales

Con el objeto de simplificar el intercambio de información, disminuir la incertidumbre y ser consistentes en la representación de cantidades se ha desarrollado un sistema de unidades *universal, unificado y coherente* basado en el *Sistema Métrico Decimal, m-kg-s* desarrollado en 1790 por la *Academia de Ciencias de París*.

La *Conferencia General de Pesas y Medidas (CGPM)* es el organismo internacional encargado de la definición de los nuevos patrones de medida. Se han sostenido 20 reuniones desde 1889 hasta 1995, a partir de 1960 se conoce como *Sistema Internacional de Unidades, SI*, y consta de las siete *cantidades fundamentales* mostradas en la **tabla 2.1**. Con el objeto de expresar las cantidades en un amplio rango numérico la CGPM adoptó los prefijos mostrados en la **tabla 2.2**.

Tabla 2.1 Cantidades fundamentales del SI

Cantidad	Tiempo	Longitud	Masa	Cantidad de sustancia	Temperatura	Corriente eléctrica	Intensidad lumínica
Dimensión	[T]	[L]	[M]	[N]	[Θ]	[I]	[J]
Unidad	segundo	metro	kilogramo	mol	kelvin	ampere	candela
Símbolo	s	m	kg	mol	K	A	cd

Tabla 2.2 Prefijos adoptados por el SI

10 ⁻¹⁸	<i>atto</i>	a	10 ⁻⁶	<i>micro</i>	μ	10 ¹	<i>deca</i>	da	10 ⁹	<i>giga</i>	G
10 ⁻¹⁵	<i>femto</i>	f	10 ⁻³	<i>mili</i>	m	10 ²	<i>hecto</i>	h	10 ¹²	<i>tera</i>	T
10 ⁻¹²	<i>pico</i>	p	10 ⁻²	<i>centi</i>	c	10 ³	<i>kilo</i>	k	10 ¹⁵	<i>peta</i>	P
10 ⁻⁹	<i>nano</i>	n	10 ⁻¹	<i>deci</i>	d	10 ⁶	<i>mega</i>	M	10 ¹⁸	<i>exa</i>	E

Cantidades derivadas y análisis dimensional

Las **magnitudes o cantidades derivadas** son aquellas que surgen de la interacción natural de los sistemas físicos. *Estas magnitudes describen los diferentes fenómenos y comportamientos de la naturaleza*, y sus unidades se obtienen combinando las unidades fundamentales del **SI**. La **tabla 2.3** muestra las cantidades derivadas de mayor uso.

El **análisis dimensional** es una técnica que permite determinar las dimensiones y las unidades fundamentales de una magnitud física derivada.

El procedimiento consiste en: (1) *identificar las fórmulas físicas correspondientes*, (2) *descomponer las variables en magnitudes básicas*, (3) *sustituir cada magnitud por sus símbolos dimensionales y simplificar algebraicamente hasta obtener la ecuación dimensional*, y (4) *sustituir cada dimensión de la ecuación dimensional obtenida por su unidad fundamental del SI*.

Ejemplo, utilizar análisis dimensional para encontrar las dimensiones y unidades fundamentales de la *fuerza* (F) y el *voltaje* (V):

Fuerza:

$$(1) \rightarrow F = ma$$

$$(2) \rightarrow m \cdot \left(\frac{v}{t}\right) = m \cdot \left(\frac{d}{t \cdot t}\right) = m \cdot \left(\frac{d}{t^2}\right)$$

$$(3) \rightarrow [F] = [M] \cdot \left[\frac{L}{T^2}\right] = [M][L][T^{-2}]$$

$$(4) \rightarrow [M][L][T^{-2}] = \text{kg} \cdot \text{m} \cdot \text{s}^{-2} = \text{kg} \frac{\text{m}}{\text{s}^2} \leftarrow \text{Newton (N)}$$

Voltaje:

$$(1) \rightarrow V = \frac{U}{q}$$

$$(2) \rightarrow \frac{F \cdot d}{i \cdot t} = \frac{(m \cdot a) \cdot d}{i \cdot t} = \frac{m \left(\frac{d}{t^2}\right) \cdot d}{i \cdot t} = \frac{m \cdot d^2}{i \cdot t^3}$$

$$(3) \rightarrow [V] = \frac{[M] \cdot [L]^2}{[I] \cdot [T]^3} = [M][L]^2[T]^{-3}[I]^{-1}$$

$$(4) \rightarrow [M][L]^2[T]^{-3}[I]^{-1} = \text{kg} \cdot \text{m}^2 \cdot \text{s}^{-3} \cdot \text{A}^{-1} = \frac{\text{kg} \cdot \text{m}^2}{\text{A} \cdot \text{s}^3} \leftarrow \text{Voltio (V)}$$

Tabla 2.3 Cantidades derivadas del SI

Cantidad	Unidad	Símbolo	Dimensiones
Área	—	—	m ²
Volumen	—	—	m ³
Frecuencia	Hertz	Hz	s ⁻¹
Densidad de masa	—	—	kg/ m ³
Velocidad	—	—	m/s
Velocidad angular	—	—	rad/s
Aceleración	—	—	m/ s ²
Aceleración angular	—	—	rad/ s ²
Fuerza	Newton	N	kg·m/s ²
Presión mecánica	Pascal	Pa	N/m ²
Viscosidad cinemática	—	—	·m ² /s
Viscosidad dinámica	—	—	N·s/m ²
Trabajo, energía, calor	Joule	J	N·m
Potencia	Watt	W	J/s
Potencia <i>reactiva</i>	VAR	VAR	J/s
Cantidad de electricidad	Coulomb	C	A·s
Fuerza electromotriz, diferencia de potencial	Volt	V	W/A, J/C
Campo eléctrico	—	—	V/m
Resistencia eléctrica, <i>reactancia, impedancia</i>	Ohm	Ω	V/A
Conductancia eléctrica, <i>susceptancia, admitancia</i>	Siemen	S	A/V
Capacitancia	Faradio	F	A·s/V
Flujo magnético	Weber	Wb	V·s
Inductancia	Henry	H	V·s/A
Densidad de flujo magnético	Tesla	T	Wb/m ²
Campo magnético	—	—	A/m
Fuerza magnetomotriz	Ampere	A	A
Flujo luminoso	Lumen	lm	cd·sr
Luminancia	—	—	cd/m ²
Iluminancia	Lux	lx	lm/m ²
Entropía	—	—	J/K
Capacidad específica del calor	—	—	J/(kg·K)
Conductividad térmica	—	—	W/(m·K)
Intensidad radiactiva	—	—	W/sr
Actividad radiactiva	Becquerel	Bq	s ⁻¹
Angulo plano	Radian	rad	unidad suplementaria
Ángulo sólido	Steradian	sr	unidad suplementaria

Unidades no oficiales comúnmente empleadas

Algunas unidades que no forman parte del SI, son ampliamente usadas en sectores científicos, industriales y comerciales. La **tabla 2.4** muestra las unidades no oficiales comúnmente empleadas en ingeniería.

Tabla 2.4 Unidades del Sistema Británico

Cantidad	Unidad	Equivalencia	Cantidad	Unidad	Equivalencia
Longitud	Ángstrom	$\text{\AA}=10^{-10} \text{ m}$	Fuerza	dina	10^{-5} N
	año luz	$9.461 \times 10^{15} \text{ m}$	Presión	bar	10^5 Pa
Volumen	litro	10^{-3} m^3		mm Hg	torr= 133.3 Pa
Tiempo	minuto	min= 60 s	Energía	ergio	erg= 10^{-7} J
	hora	h= 3600 s		caloría	cal= 4.186 J
	día	d= 86400 s		kWh	kWh= $3.6 \times 10^6 \text{ J}$
	año	a= $3.156 \times 10^7 \text{ s}$	Temperatura	°C	K-273.15
Ángulo	revolución	rev= $360^\circ=2\pi \text{ rad}$	Vel. angular	rpm	rpm=rev/min

Otros sistemas de unidades

Existen otros sistemas locales o específicos de alguna ciencia entre los cuales se encuentran: *británico*, *gaussiano* y *cgs*. La mayoría de los países han adoptado el SI oficialmente; sin embargo y debido a que gran cantidad de instrumentos de control están basados en el *Sistema Británico*, en la [tabla 2.5](#) se presentan las equivalencias para las unidades más comunes.

Tabla 2.5 Unidades del Sistema Británico

Cantidad	Unidad	Símbolo	Equivalencia
Distancia	pulgada	in	2.54 cm
	pie	ft	30.48 cm
	yarda	–	91.44 cm
	milla	–	1609 m
	milla náutica	–	1852 m
Masa	slug	–	14.59 kg
	onza	–	28.35 g
Fuerza	libra	lb	4.448 N
Energía	btu	btu	1055 J
Potencia	hp	hp	745.7 W
Temperatura	grado Fahrenheit	°F	1.8°C+32
Volumen	galón	–	3.788 litros

Constantes físicas

En la definición de ecuaciones aparecen ciertos valores numéricos cuyo valor ha sido aceptado como constante, las **tablas 2.6 y 2.7** muestra las constantes más comúnmente empleadas en la ingeniería.

Tabla 2.6 Constantes físicas empleadas en ingeniería

Descripción	Símbolo	Valor	Unidades
Velocidad de la luz	c	2.99792458×10^8	m/s
Carga del electrón	e	1.602177×10^{-19}	C
Constante Gravitacional	G	6.67259×10^{-11}	N·m ² ·kg ²
Constante de Planck	h	$6.6260755 \times 10^{-34}$	J·s
Constante de Boltzmann	k	1.38066×10^{-23}	J/K
Número de Avogadro	N _A	6.022×10^{23}	moléculas/mol
Constante de Gas	R	8.314510	J/(mol·K)
Masa del electrón	m _e	9.10939×10^{-31}	kg
Masa del neutrón	m _n	1.67262×10^{-27}	kg
Masa de protón	m _p	1.67492×10^{-27}	kg
Permitividad en el vacío	ε ₀	8.854×10^{-12}	C ² /(N·m ²)
Permeabilidad en el vacío	μ ₀	$4\pi \times 10^{-7}$	Wb/(A·m)
Equivalente mecánico del calor		4.186	J/cal
Presión atmosférica estándar	atm	1.013×10^5	Pa
Cero absoluto	0 K	-273.15	°C
Electronvolt	1 eV	1.602×10^{-19}	J
Unidad de masa atómica, uma	1 u	1.66054×10^{-27}	kg
Energía del electrón en reposo	m·c ²	0.511	MeV
Equivalente energético de uma	M·c ²	931.494	MeV
Volumen de molar de gas ideal	V	22.4	litros/mol
Aceleración debida a la gravedad	g	9.78049	m/s ²

Tabla 2.7 Constantes matemáticas empleadas en ingeniería

Razón de la circunferencia	π	3.1416	--
Número de Euler	e	2.71828	--
Unidad imaginaria	i, j	$i=j=\sqrt{-1}$	--

Sistemas mecánicos

El estudio del *movimiento de los cuerpos* se agrupa convenientemente en dos categorías: *movimiento traslacional* y *rotación con respecto a un eje fijo*. Las *propiedades mecánicas* más simples que manifiestan la mayoría de objetos materiales son las siguientes: *rigidez*, *amortiguamiento* y *fricción* e *inercia*. Estas propiedades están representadas en elementos que exhiben fuertemente cada una de ellas: *resorte*, *amortiguador* y *masa*. Estos elementos se consideran *bloques funcionales* y se encuentran regulados por las leyes de la mecánica clásica.

Fundamentos de cinemática y dinámica traslacional

El *desplazamiento traslacional* es el cambio de posición de un objeto entre dos puntos a lo largo de una *trayectoria rectilínea* o *curvilínea* de tal manera que todas las partículas que forman el objeto se desplazan a lo largo de *trayectorias paralelas*, y poseen la misma velocidad y aceleración en todo momento. La posición del objeto como función del tiempo se denota por $\mathbf{x}(t)$, la *velocidad* $\mathbf{v}(t)$, es el *cambio de posición con respecto al tiempo*, la *aceleración* $\mathbf{a}(t)$ es el *cambio de velocidad respecto al tiempo* y se definen por:

$$\bar{\mathbf{v}} = \frac{d\bar{\mathbf{x}}}{dt} \quad (1)$$

$$\bar{\mathbf{a}} = \frac{d^2\bar{\mathbf{x}}}{dt^2} \quad (2)$$

En función de la velocidad ($\mathbf{v}$) es fácil conocer la posición ($\mathbf{x}$) del objeto en el tiempo t , usando integración:

$$\bar{\mathbf{x}} = \int \bar{\mathbf{v}} dt \quad (3)$$

Si el movimiento traslacional es *uniformemente acelerado*, las siguientes relaciones describen el movimiento del objeto:

$$\bar{\mathbf{x}} = \bar{\mathbf{x}}_0 + \frac{1}{2}(\bar{\mathbf{v}}_0 + \bar{\mathbf{v}})t \quad \bar{\mathbf{x}} = \bar{\mathbf{x}}_0 + \bar{\mathbf{v}}_0 t + \frac{1}{2}\bar{\mathbf{a}}t^2 \quad \bar{\mathbf{v}} = \bar{\mathbf{v}}_0 + \bar{\mathbf{a}}t \quad \bar{\mathbf{v}}^2 = \bar{\mathbf{v}}_0^2 + 2\bar{\mathbf{a}}(\bar{\mathbf{x}} - \bar{\mathbf{x}}_0) \quad (4)$$

El estudio de las causas del movimiento de objetos es materia de la *dinámica*, las *leyes de Newton* son su base teórica fundamental.

Primera Ley de Newton o ley de la inercia. Si la fuerza resultante que actúa sobre un objeto es cero, entonces el objeto permanecerá en su estado de movimiento. Un cuerpo se acelera sólo si la fuerza resultante que actúa sobre el es diferente de cero.

Segunda ley de Newton. Si la fuerza resultante que actúa sobre un objeto es diferente de cero, entonces este se acelera; la aceleración (**a**) es proporcional a la fuerza (**f**) e inversamente proporcional a la masa (*m*) del objeto.

$$\vec{f} = m\vec{a} = m \frac{d^2\vec{x}}{dt^2} \quad (5)$$

Tercera ley de Newton o ley de acción-reacción. Por cada fuerza que actúa sobre un cuerpo, existe otra de igual magnitud, pero en sentido opuesto actuando sobre algún otro cuerpo.

Ley de la gravitación universal de Newton. Establece que dos partículas de masa *m* y *M*, separadas entre sí por una distancia *r*, se atraen mutuamente con fuerzas de magnitudes iguales y opuestas en sentido:

$$\vec{f} = G \frac{Mm}{r^2} \hat{r} \quad (6)$$

Siendo *G* la constante de gravitación universal.

El *peso* de un objeto de masa *m*, es **P=mg** y describe la fuerza gravitacional que ejerce el planeta tierra sobre el objeto. En este caso se define la *aceleración debida a la gravedad (g)* dirigida al centro de la tierra, como:

$$\vec{g} = G \frac{M}{R^2} \hat{R} \quad (7)$$

Siendo *R=6.38x10⁶* el radio promedio de la tierra y *M=5.97x10²⁴* su masa.

La *energía (E)* es la capacidad que tiene un cuerpo para realizar un trabajo. El trabajo físico (*W*), es un escalar que se determina por el desplazamiento (**x**) que provee una fuerza o un campo de fuerzas (**f**) para desplazar un objeto.

$$dW = \vec{f} \cdot d\vec{x} \rightarrow W = \int \vec{f} \cdot d\vec{x} \quad (8)$$

La potencia (p) es la razón a la que se efectúa el trabajo (W) o se consume energía (E) y queda definida por:

$$p = \frac{dW}{dt} = \frac{dE}{dt} \rightarrow E = \int p dt \quad (9)$$

La conservación de la energía mecánica enuncia: “En cualquier sistema aislado de objetos que interactúan sólo por fuerzas conservativas, tales como el bloque y resorte, la energía puede ser transferida una y otra vez de cinética a potencial; pero el cambio total es de cero; la suma de las energía cinética y potencial permanece constante”.

La *energía potencial* (U), es la energía almacenada en un sistema a causa de la posición relativa u orientación de las partes de un sistema y está definida sólo para fuerzas conservativas:

$$f = -\frac{dU}{dx} \rightarrow dU = -\vec{f} \cdot d\vec{x} \quad (10)$$

El *teorema trabajo-energía* establece que: “El trabajo neto efectuado por las fuerzas que actúan sobre una partícula es igual al cambio en la energía cinética de las partículas”. Para una partícula de masa m que mueve con una *velocidad* $\mathbf{v}$, la *energía cinética* K , y el *momento lineal* o *cantidad de movimiento* $\mathbf{p}$ de la partícula se definen por:

$$K = \frac{1}{2} m |\mathbf{v}|^2 = \frac{1}{2} m \left| \frac{d\vec{x}}{dt} \right|^2 \quad (11)$$

$$\mathbf{p} = m\mathbf{v} \quad (12)$$

El *principio de conservación de momento lineal*, establece que: “si la fuerza externa resultante que actúa sobre un sistema de objetos es cero, entonces la suma vectorial de la cantidad de movimiento de los objetos permanece constante”. Si dos objetos de masa m_1 y m_2 chocan, el momento lineal antes y después del impacto se conserva.

$$m_1 \vec{v}_{i1} + m_2 \vec{v}_{i2} = m_1 \vec{v}_{f1} + m_2 \vec{v}_{f2} \quad (13)$$

El *coeficiente de restitución* se define por $e = (\bar{v}_{f2} - \bar{v}_{f1}) / (\bar{v}_{i1} - \bar{v}_{i2})$. Para una *colisión perfectamente elástica* $e=1$, *inelástica* $e<1$, si los cuerpos permanecen *unidos después de la colisión* $e=0$. La segunda ley de Newton en términos de la cantidad de movimiento es:

$$\sum \vec{F} = \frac{d\vec{p}}{dt} \quad (14)$$

Bloques funcionales de sistemas mecánicos con movimientos traslacionales

Resorte

La **rigidez** de un resorte se describe mediante la relación entre la fuerza F empleada para estirar o comprimir un resorte y la deformación resultante x , ya sea de estiramiento o compresión. *Roberto Hooke* en 1678 enunció la siguiente ley “*la deformación de un resorte es directamente proporcional a la fuerza aplicada*”.

$$\vec{f} = k\vec{x} = k \int \vec{v} dt \quad (15)$$

El objeto que aplica la fuerza se encuentra sujeto a una fuerza de igual magnitud y de sentido contrario, como lo establece la *tercera ley de Newton*.

La *energía potencial del resorte* de acuerdo a lo visto, queda definida por:

$$U = \int \vec{f} \cdot d\vec{x} = \int (k\vec{x}) \cdot d\vec{x} = \frac{1}{2} kx^2 = \frac{1}{2} \frac{f^2}{k} \quad (16)$$

Amortiguador

El **amortiguamiento** y **fricción** son fuerzas de oposición no conservativas, que actúan sobre la superficie de un cuerpo con el objeto de impedir el desplazamiento. Un *amortiguador* es un elemento físico que presenta la propiedad de amortiguamiento y se puede describir como sigue: “*un émbolo que comprime un fluido alojado dentro de recipiente cerrado, recibe una fuerza de oposición al desplazamiento que es directamente proporcional a la velocidad del émbolo*”.

$$\vec{f} = c \frac{d\vec{x}}{dt} \quad (17)$$

La fuerza del amortiguador es no-conservativa, por lo tanto el émbolo no regresa a su posición original cuando ha cesado la fuerza. Sin embargo la energía se transforma y disipa en forma de calor, la *potencia disipada* esta dada por:

$$P = c \left| \frac{d\vec{x}}{dt} \right|^2 \quad (18)$$

Masa

La **inercia** es la capacidad que tiene un cuerpo debido a su *masa* a cambiar su estado de movimiento o reposo. La fuerza inercial esta determinada por la *segunda ley de Newton* o *ley de la inercia*: “la aceleración producida por una fuerza dada es inversamente proporcional a la masa que es acelerada”

$$\vec{f} = m \frac{d^2 \vec{x}}{dt^2} \quad (19)$$

Como se enunció anteriormente la *energía cinética* que presenta una masa acelerada esta dada por:

$$K = \frac{1}{2} m \left[\frac{d\vec{x}}{dt} \right]^2 \quad (20)$$

Fundamentos de cinemática y dinámica rotacional

Las consideraciones expuestas son válidas para sistemas cuyos movimientos sean traslacionales, si el movimiento es rotacional los bloques funcionales se convierten en: *resorte torsional*, *amortiguador rotatorio* y *momento de inercia*.

El *movimiento de rotación* es aquel que las partículas que constituyen el objeto se desplazan en planos paralelos, a lo largo círculos centrados en el mismo eje fijo llamado *eje de rotación*.

Consideremos el cuerpo rígido de densidad uniforme mostrado en la **figura 2.1**, que gira en torno al eje Z, en sentido antihorario. Sea P un punto inmerso en este cuerpo, descrito por el vector $\mathbf{r}$, en el sistema coordenado con origen O. El punto P describe una trayectoria circular de radio $AP = r \cdot \sin(\phi)$ en el plano XY, y su desplazamiento angular en radianes se denota por θ . La longitud de arco s , el ángulo θ y el *radio de giro* AP , se relacionan por:

$$\theta = \frac{s}{AP} \rightarrow s = r \overline{AP} \quad (21)$$

La *velocidad angular* (ω), y *aceleración angular* (α) del punto P como funciones del tiempo (t), están dadas por:

$$\omega = \frac{d\theta}{dt} \rightarrow \theta = \int \omega dt \quad (22)$$

$$\alpha = \frac{d\omega}{dt} = \frac{d^2\theta}{dt^2} = \omega \frac{d\omega}{d\theta} \quad (23)$$

En términos del *radio vector* $\mathbf{r}$, la *velocidad* $\mathbf{v}$ y la *aceleración* $\mathbf{a}$ del punto P, quedan determinados por:

$$\mathbf{v} = \frac{d\mathbf{r}}{dt} = \omega \times \mathbf{r} = \omega \cdot r \sin(\phi) \quad (24)$$

$$\mathbf{a} = \frac{d\mathbf{v}}{dt} = \frac{d(\omega \times \mathbf{r})}{dt} = \frac{d\omega}{dt} \times \mathbf{r} + \omega \times \frac{d\mathbf{r}}{dt} = \alpha \times \mathbf{r} + \omega \times \mathbf{v} = \alpha \times \mathbf{r} + \omega \times (\omega \times \mathbf{r}) \quad (25)$$

Es común descomponer la aceleración en dos vectores, *aceleración tangencial* $\mathbf{a}_T$ y *aceleración radial* $\mathbf{a}_R$, como se muestra a continuación.

$$\mathbf{a} = \mathbf{a}_T + \mathbf{a}_R \quad \mathbf{a}_T = \alpha \times \mathbf{r} \quad \mathbf{a}_R = \omega \times \mathbf{v} \quad (26)$$

$$\mathbf{a}_T = \alpha \times \mathbf{r} \rightarrow |\mathbf{a}_T| = \alpha \cdot r \sin(\phi) \quad (27)$$

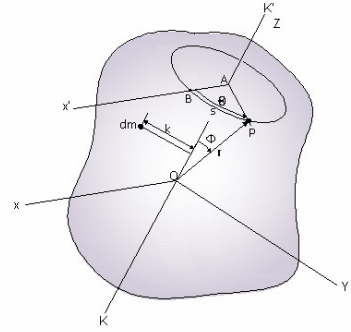


Figura 2.1 Movimiento rotacional

$$\vec{a}_R = \omega \times \vec{r} = \omega \times (\omega \times \vec{r}) \rightarrow |\vec{a}_R| = \omega^2 \cdot \sin(\phi) \quad (28)$$

Si la rotación es *uniformemente acelerada*, las siguientes relaciones describen el movimiento del punto P:

$$\vec{\theta} = \vec{\theta}_0 + \vec{\omega}t \quad \vec{\theta} = \vec{\theta}_0 + \vec{\omega}_0 t + \frac{1}{2} \vec{\alpha} t^2 \quad \vec{\omega} = \vec{\omega}_0 + \vec{\alpha}t \quad \vec{\omega}^2 = \vec{\omega}_0^2 + 2\vec{\alpha}(\vec{\theta} + \vec{\theta}_0) \quad (29)$$

El *momento de inercia* (I) de un cuerpo es una medida de su inercia rotacional. Si un objeto se considera constituido por masas infinitesimales $m1, m2, m3 \dots$ a las distancias respectivas $r1, r2, r3 \dots$ a partir de un eje, su momento de inercia con respecto a este eje es:

$$I = \sum m_i r_i^2 = \int r^2 dm \quad (30)$$

Consideremos nuevamente la **figura 2.1**. Definimos al *radio de giro* del cuerpo como la distancia k paralela al *eje de giro*, donde se debe concentrar toda la masa del cuerpo si su *momento de inercia* I, respecto al eje KK' debe permanecer inalterado.

$$I = k^2 m \rightarrow k = \sqrt{\frac{I}{m}} \quad (31)$$

Los *momentos de inercia* respecto a los ejes X, Y, Z están definidos por:

$$I_x = \int (y^2 + z^2) dm \quad I_y = \int (x^2 + z^2) dm \quad I_z = \int (x^2 + y^2) dm \quad (32)$$

El *teorema de los ejes paralelos* establece que: “la inercia de rotación (I), de cualquier cuerpo en torno a un eje arbitrario es igual a la inercia de rotación alrededor de un eje paralelo que pase por el centro de masa (I_{cm}), más la masa total (m), por la distancia (d) entre los dos ejes elevada al cuadrado”:

$$I = I_{cm} + md^2 \quad (33)$$

En términos del momento de inercia, la *energía cinética* K, del cuerpo rígido en rotación, esta dada por:

$$K = \frac{1}{2} I \omega^2 \quad (34)$$

La *torca* o *momento de fuerza* τ , que actúa sobre un punto P de un cuerpo, está representada por el producto vectorial del *brazo de palanca* $\mathbf{r}$ y el *vector fuerza* $\mathbf{f}$ que ocasiona la rotación del cuerpo a través de un eje que pasa por P y es perpendicular al plano formado por $\mathbf{r}$ y $\mathbf{F}$.

$$\tau = \mathbf{r} \times \mathbf{f} \quad (35)$$

La *aceleración angular* α , y el momento de inercia I , y el momento de fuerza τ , quedan relacionados por la *segunda ley de Newton para rotación*:

$$\sum \tau = I\alpha \quad (36)$$

El *trabajo* (W) efectuado para desplazar el cuerpo desde θ_1 hasta θ_2 , y la *potencia* consumida (P), están dados por:

$$W = \int_{\theta_1}^{\theta_2} \tau \cdot d\theta \quad P = |\tau| |\dot{\theta}| \quad (37)$$

El *momento angular* $\mathbf{L}$, con respecto al punto O , de una partícula de masa m y *velocidad* $\mathbf{v}$, es:

$$\vec{L} = \mathbf{r} \times m\mathbf{v} \rightarrow L = I\omega \quad (38)$$

Bloques funcionales de sistemas mecánicos con movimientos rotacionales

Resorte torsional

En un *resorte torsional* el desplazamiento angular θ es proporcional al par T :

$$T = k\theta \quad (39)$$

Amortiguador rotatorio

En un *amortiguador rotatorio*, un disco que gira en un fluido, el par rotatorio es proporcional a la velocidad angular:

$$T = c\omega = c \frac{d\theta}{dt} \quad (40)$$

Momento de inercia

La segunda ley de Newton aplicada a sistemas con movimiento rotacional se describe por:

$$T = I\alpha = I \frac{d^2\theta}{dt^2} \quad (41)$$

La **tabla 2.8** lista las ecuaciones que describen las propiedades de bloques funcionales mecánicos para movimientos traslacionales y rotacionales.

Tabla 2.8 Propiedades de los bloques funcionales mecánicos

Bloque funcional	Ecuación descriptiva	
	Movimiento traslacional	Movimiento rotacional
Resorte	$F = kx$	$T = k\theta$
<i>Posee energía en virtud de su posición. Energía potencial</i>	$U = \frac{1}{2} kx^2 = \frac{1}{2} \frac{F^2}{k}$	$U = \frac{1}{2} k\theta^2 = \frac{1}{2} \frac{T^2}{k}$
Amortiguador	$F = c \frac{dx}{dt}$	$T = c \frac{d\theta}{dt}$
<i>Disipa energía</i>	$P = c \left(\frac{dx}{dt} \right)^2$	$P = c \left(\frac{d\theta}{dt} \right)^2$
Masa/Momento de Inercia	$F = m \frac{d^2x}{dt^2}$	$T = I \frac{d^2\theta}{dt^2}$
<i>Posee energía en virtud de su movimiento. Energía cinética</i>	$K = \frac{1}{2} m \left(\frac{dx}{dt} \right)^2$	$K = \frac{1}{2} I \left(\frac{d\theta}{dt} \right)^2$

Sistemas eléctricos

Los *elementos pasivos* básicos que constituyen la base de los circuitos eléctricos son: *inductor*, *resistor* y *capacitor*. Estos elementos exhiben fuertemente las siguientes propiedades: *inductancia*, *resistencia*, *capacitancia*.

Elementos activos y elementos pasivos

En circuitos eléctricos consideramos como *elementos activos* aquellos que son capaces de proporcionar potencia promedio mayor que cero, por ejemplo las fuentes de corriente y voltaje y el amplificador operacional. Un elemento pasivo se define como aquel que no suministra una potencia promedio mayor que cero, en un intervalo de duración infinita.

Convección de corriente eléctrica

Para efectos de análisis de circuitos se ha adoptado por convención que la dirección de la corriente sea la misma que la de *portadores de carga positivos*. De esta manera: *en un elemento activo la corriente circula de la terminal de menor potencial a la de mayor potencial* ($- \rightarrow +$); *en un elemento pasivo la corriente circula de la terminal de mayor potencial a la de menor potencial* ($+ \rightarrow -$). Lo cual se conoce como *convención pasiva de los signos*.

Fundamentos de electrostática y electrodinámica

La *potencia eléctrica instantánea* para un elemento eléctrico por el cual circula una *corriente* $i(t)$ con una *diferencia de potencial* $v(t)$ a través de sus terminales, se define por:

$$p = vi$$

(42)

Una *fuerza electromotriz* práctica es un *generador* que bajo el principio de inducción electromagnética y la conversión de energía mecánica a eléctrica es capaz de producir una diferencia de potencial $v(t)$ en sus terminales. La *ley de Faraday* establece que: *La fem ($\mathcal{E}$) inducida en un circuito es*

igual al negativo de la velocidad con que cambia con el tiempo el flujo magnético (ϕ_B) a través del circuito:

$$e = -\frac{d\phi_B}{dt} \quad (43)$$

Esto es se produce una diferencia de potencial en las terminales de un conductor que se haya sometido en un campo magnético, cuando existe un movimiento relativo entre ambos.

La *corriente eléctrica*, se define como el *la tasa de flujo de cargas eléctricas a través de un conductor eléctrico* que se encuentra sometido al campo potencial de una fuente cuya fuerza electromotriz produce una diferencia de potencial $v(t)$ en las terminales del conductor. *Por naturaleza los portadores de carga negativa (electrones) fluyen de la terminal con menor potencial hacia la terminal con mayor potencial, es decir del cátodo (-) al ánodo (+).*

$$i = \frac{dq}{dt} \quad (44)$$

Bloques funcionales de sistemas eléctricos

Inductor

Una *bobina* o *inductor* es un alambre conductor en forma de espira, tal configuración le propicia una propiedad eléctrica conocida como *inductancia*. La inductancia (L) es la propiedad por la cual un material eléctrico se opone a cambio súbitos de corriente a través de él. Cuando por las terminales de un inductor fluye una corriente $i(t)$, entonces en las terminales de este se induce una *fuerza contraelectromotriz* dada por:

$$v = L \frac{di(t)}{dt} \rightarrow i = \frac{1}{L} \int v dt \quad (45)$$

Por ser un elemento pasivo, el inductor disipa energía cuando sobre él circula una corriente. La energía disipada se obtiene fácilmente al considerar las ecuaciones definitorias de potencia-energía y voltaje en las terminales del inductor.

$$U = \int p dt = \int v i dt = \int \left(L \frac{di}{dt} \right) i dt = L \int i di = \frac{1}{2} L i^2 \quad (46)$$

Resistor

Un *resistor* es un material que opone *resistencia* al paso de la corriente eléctrica a través de él. La resistencia de un conductor eléctrico es función de la configuración del conductor, para un alambre metálico se tiene que:

$$R = \rho \frac{L}{S} \quad (47)$$

Siendo ρ la resistividad del material, L la longitud del alambre y S el área de la sección transversal. La resistividad es una propiedad de los materiales y es función de la temperatura de acuerdo con:

$$\rho = \rho_0 [1 + \alpha(T - T_0)] \quad (48)$$

Siendo ρ_0 la resistividad del material a la temperatura T_0 y α el *coeficiente de temperatura* propio del material.

Un *material ohmico* es aquel que obedece la *ley de Ohm*:

$$V = Ri(t) \quad (49)$$

Siendo $i(t)$ corriente que fluye por el material con resistencia R expuesto a una diferencia de potencial V en sus terminales.

Un resistor disipa energía en forma de calor, tal fenómeno es conocido como *efecto Joule*, la energía disipada en forma de potencia, está dada por:

$$p = vi = (Ri)i = Ri^2 = \frac{1}{R} v^2 \quad (50)$$

Capacitor

Un *capacitor* es elemento pasivo cuya propiedad característica es la capacitancia. La *capacitancia* (C) es la propiedad de un material por la cual se almacena energía eléctrica en forma de campo eléctrico.

Un capacitor sencillo es un elemento de dos placas conductoras metálicas paralelas distanciadas una de la otra. Cuando las terminales de las placas se conectan a una fuente de voltaje, las placas adquieren cargas iguales en magnitud y opuestas en signo. La *capacitancia* se define como la razón entre la magnitud de la carga en cualquiera de las placas y la diferencia de potencial entre ellas:

$$C = \frac{q}{v} \rightarrow v = \frac{1}{C} q \quad (51)$$

De la definición de capacitancia se obtiene la relación entre voltaje y corriente del capacitor:

$$\frac{dv}{dt} = \frac{1}{C} \frac{dq}{dt} = \frac{1}{C} i \rightarrow i = C \frac{dv}{dt} \quad (52)$$

La energía almacenada entre las placas del capacitor se obtiene fácilmente al considerar la definición de potencia eléctrica:

$$U = \int p dt = \int v i dt = \int C v \frac{dv}{dt} dt = C \int v dv = \frac{1}{2} C v^2 \quad (53)$$

La **tabla 2.9** muestra las ecuaciones descriptivas de los bloques funcionales eléctricos.

Tabla 2.9 Propiedades de los bloques funcionales eléctricos

Bloque funcional	Ecuación descriptiva	
Inductor	$v = L \frac{di}{dt} \rightarrow i = \frac{1}{L} \int v dt$	$U = \frac{1}{2} L i^2$
Resistor	$V = R i$	$p = R i^2 = \frac{1}{R} v^2$
Capacitor	$i = C \frac{dv}{dt} \rightarrow v = \frac{1}{C} \int i dt$	$U = \frac{1}{2} C v^2$

Sistemas fluídicos

La *mecánica de fluidos* nos provee la base teórica para el estudio de fluidos. Los fluidos pueden clasificarse en dos categorías atendiendo su *compresibilidad*. Los *fluidos hidráulicos* y *fluidos neumáticos*; no compresibles y compresibles respectivamente.

Fundamentos de hidrostática e hidrodinámica

Los *fluidos hidráulicos* son fluidos en estado líquido. El estado líquido se distingue porque sus moléculas tienen cierta libertad para desplazarse unas respecto a otras, pero están sometidas a la acción limitante de las fuerzas de cohesión que les impide extenderse indefinidamente; debido a ello su volumen es constante, a temperatura constante; su forma se adapta a la del recipiente que los contiene.

La densidad ρ se define como la cantidad de materia (m) por unidad de volumen (V), y en los fluidos líquidos es constante:

$$\rho = \frac{m}{V} \quad (54)$$

La presión (p) es la fuerza (f) por unidad de área (A):

$$p = \frac{f}{A} \rightarrow f = pA \quad (55)$$

La *presión hidrostática* (P_h) debida a una columna de un fluido de altura h y densidad de masa ρ , esta dada por:

$$p_h = \rho gh \quad (56)$$

El *principio de Pascal* establece que: “la presión aplicada a un fluido encerrado es transmitida a cada porción del fluido así como a las paredes del recipiente que lo contiene”:

$$p = \frac{f_1}{A_1} = \frac{f_2}{A_2} \quad (57)$$

El *principio de Arquímedes* afirma que: “Cuando un cuerpo es sumergido en un fluido, el fluido ejerce una fuerza hacia arriba (*fuerza boyante*) igual al peso del fluido que el cuerpo desplaza, el volumen de líquido desplazado es igual al volumen de la parte sumergida del objeto”. La fuerza boyante f_b que recibe un objeto que se sumerge un volumen V_O en un fluido con densidad ρ_f está dado por:

$$F_b = g\rho_f V_O \quad (58)$$

Cuando un fluido que llena un tubo de sección transversal A corre a lo largo de este con una velocidad promedio v , entonces el *gasto*, *flujo*, *descarga* o *razón de flujo volumétrico*, esta dado por $Q=Av$, además si el fluido es incompresible, se cumple la siguiente relación conocida como *ecuación de continuidad*:

$$Q = A_1 v_1 = A_2 v_2 \quad (59)$$

La *viscosidad* es la resistencia natural que ofrece un fluido para desplazarse sobre una superficie.

La *ley de Poiseuille* establece que el flujo volumétrico de un fluido que corre a través de un tubo cilíndrico de longitud L y sección transversal de radio r y con presiones p_1 y p_2 en los extremos del tubo (*siendo K la constante de conductancia hidráulica*), esta dado por:

$$Q = \frac{\pi \cdot r^4 (p_1 - p_2)}{8\eta L} = K(p_1 - p_2) \longleftarrow K = \frac{\pi \cdot r^4}{8\eta L} \quad (60)$$

Para un flujo a lo largo de una tubería de una corriente continua de un fluido de viscosidad despreciable es válida la siguiente relación conocida como *ecuación de Bernoulli*:

$$p_1 + \frac{1}{2}\rho v_1^2 + \rho g h_1 = p_2 + \frac{1}{2}\rho v_2^2 + \rho g h_2 \quad (61)$$

Donde los subíndices 1 y 2 denotan dos puntos a lo largo de la tubería, ρ , g y h , son la densidad, la aceleración debida a la gravedad y la altura respectivamente.

Bloques funcionales de sistemas fluidicos hidráulicos

Las propiedades de los fluidos hidráulicos son: *resistencia hidráulica*, *capacitancia hidráulica* e *inertancia hidráulica*.

Resistencia hidráulica

Es la resistencia a fluir que se presenta como resultado de un flujo de líquido a través de válvulas o cambios de diámetro de las tuberías. *La diferencia de presiones entre dos puntos dentro de la tubería que transporta un fluido hidráulico es directamente proporcional al flujo volumétrico*, constante de proporcionalidad se conoce como resistencia hidráulica y es la constante que acompaña la ley de Poiseuille:

$$p_1 - p_2 = Rq \quad (62)$$

Capacitancia hidráulica

Considérese la **figura 2.2**, se tiene un contenedor cúbico de altura h y área de base A , alimentado por un flujo q_1 , con una presión p_1 en la altura máxima de líquido y un flujo q_2 a la salida, con una presión p_2 al fondo del contenedor. La razón de cambio de volumen (V) es igual a la diferencia de los flujos volumétricos (q) que entran y salen del contenedor:

$$q_1 - q_2 = \frac{dV}{dt} = \frac{d(Ah)}{dt} = A \frac{dh}{dt} \quad (63)$$

La diferencia de presiones según lo visto, está dada por:

$$p = p_1 - p_2 = h\rho g \rightarrow h = \frac{p}{\rho g} \quad (64)$$

Sustituyendo h de la ecuación anterior se tiene que:

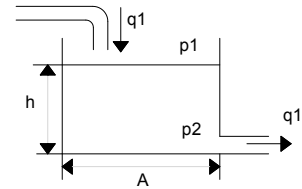


Figura 2.2 Capacitancia hidráulica

$$q_1 - q_2 = A \frac{dh}{dt} = A \frac{d\left(\frac{p}{\rho g}\right)}{dt} = \frac{A}{\rho g} \frac{dp}{dt} = C \frac{dp}{dt} \xrightarrow{\text{siendo}} C = \frac{A}{\rho g} \quad (65)$$

Donde la constante C se conoce como *capacitancia hidráulica*. Reescribiendo la ecuación en su forma integral, resulta:

$$p = \frac{1}{C} \int (q_1 - q_2) dt \quad (66)$$

Inertancia hidráulica

Considérese el bloque mostrado en la **figura 2.3**, representa un volumen de líquido cuya densidad es ρ , con longitud L y área de sección transversal A , del lado izquierdo se haya sometido a una presión p_1 y del lado derecho a una presión p_2 .

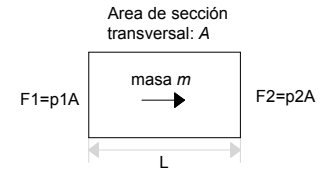


Figura 2.3 Inertancia hidráulica

La fuerza neta que actúa sobre el bloque de líquido está dado por:

$$F_1 - F_2 = p_1 A - p_2 A = (p_1 - p_2) A = m \frac{dv}{dt} = AL\rho \frac{dv}{dt} \rightarrow p_1 - p_2 = L\rho \frac{dv}{dt} \quad (67)$$

Considerando la definición de flujo volumétrico $q = Av \rightarrow v = q/A$ se tiene:

$$p_1 - p_2 = \frac{L\rho}{A} \frac{dq}{dt} = I \frac{dq}{dt} \quad (68)$$

Donde la constante I se denomina *inertancia hidráulica* y se define por (es el equivalente del resorte en sistemas mecánicos y el inductor en sistemas eléctricos):

$$I = \frac{L\rho}{A} \quad (69)$$

Las propiedades de los fluidos neumáticos son: *resistencia neumática*, *capacitancia neumática* e *inertancia neumática*. El estudio de los fluidos neumáticos es menester de la *Termodinámica*.

Fundamentos de fisicoquímica y termodinámica

Un *mol* es una unidad de cantidad que equivale a 6.0225×10^{23} partículas. Consideremos una *masa* (m) de un gas en su máximo estado de pureza, de *masa molecular* M , encerrado en un recipiente de *volumen* V , bajo una *presión* p y con una *temperatura* T , la *ley de los gases ideales*, establece que:

$$pV = nRT \rightarrow \frac{p_1 V_1}{T_1} = \frac{p_2 V_2}{T_2} \quad (70)$$

Siendo $n = m/M$ el número de moles y R la *constante universal de gases ideales*. El *principio de Avogadro* establece que: “*volúmenes iguales de todos los gases, a la misma presión y temperatura contienen igual número de moléculas*”.

La *ley de Dalton de las presiones parciales*, afirma que “*la presión total (p_t) de una mezcla de n gases es igual a la suma de las presiones parciales (p_k) de los componentes individuales de la mezcla*”.

$$p_t = \sum_{i=1}^n p_i = p_1 + p_2 + p_3 + \dots \quad (71)$$

La *teoría cinética de los gases ideales* establece que un gas de n' *moléculas* cada una de *masa* m' , moviéndose a una *velocidad promedio* v_{rms} , sometidos a una *presión* p en un *volumen* V , entonces se cumple que:

$$pV = \frac{1}{3} m' n' v_{rms}^2 \quad (72)$$

Así mismo la *energía cinética promedio* de cada molécula está dada por $E_k = \frac{3}{2} kT$, así mismo la *velocidad por molécula* es $v_{rms} = \sqrt{\frac{3RT}{M}} = \sqrt{\frac{3P}{\rho}}$.

Para *gases reales* sin embargo, se han desarrollado ecuaciones de estado que semejen el caso ideal, esto se ha logrado parcialmente mediante factores de corrección, una de las cuales es la *ecuación de Van Der Waals*, con a y b como constantes:

$$\left(p + \frac{n^2 a}{V^2}\right)(V - nb) = nRT \quad (73)$$

Las *leyes de la termodinámica* establecen los fundamentos de la dinámica de fluidos en estado gasificado. Las *propiedades termodinámicas* de las sustancias son: *temperatura* (T), *presión* (P), *densidad* (ρ), *entropía* (S). Un *proceso es reversible* cuando su dirección puede invertirse en cualquier punto por un cambio infinitesimal en las condiciones externas. En la práctica todos los procesos son irreversibles.

La *energía térmica* (Q), es la energía que fluye de un cuerpo a otro debido a la diferencia de temperaturas. La *energía interna* (U), es la energía total contenida en un sistema, suma de las energías cinética, potencial, eléctrica, nuclear que poseen átomos y moléculas del sistema.

La *conversión de trabajo* en formas de *energía potencial, cinética, eléctrica, o mecánica* se puede llevar a cabo con un *eficiencia* cercana al 100% eliminando la *fricción* que disipa energía en forma de calor. Sin embargo lo contrario, la *conversión de calor en trabajo* u otra forma aprovechable de energía tiene una *eficiencia* ($\text{eficiencia} = \text{trabajo realizado} / \text{calor introducido}$) que no supera el 40%. La segunda ley de la termodinámica trata de explicar este hecho.

El *trabajo* (W) necesario para expandir un gas, a presión constante p , de un volumen de V_1 a un volumen V_2 , es:

$$dW = p dV \rightarrow W = \int_{V_1}^{V_2} p dV \quad (74)$$

La *entalpía* (H) es una relación simple entre la *energía interna* (U) y el *trabajo* ($W=PV$) de un sistema, que se ha definido para efectos prácticos como:

$$dH = dU + d(PV) \rightarrow \Delta H = \Delta U + \Delta(PV) \quad (75)$$

Ley cero de la termodinámica. Si dos cuerpos están en equilibrio con un tercero, entonces estos están en equilibrio térmico entre si.

Primera Ley de la termodinámica. Si una cantidad de energía térmica dQ fluye dentro de un sistema, entonces esta debe aparecer como un incremento de la energía interna dU del sistema y/o como un trabajo dW por el sistema sobre sus alrededores.

$$dQ = dU + dW = dU + pdV \quad (76)$$

Segunda Ley de la termodinámica. “Es imposible para cualquier sistema sufrir un proceso en el cual se absorba calor de un depósito a una temperatura única y convierta el calor completamente en trabajo mecánico, con el sistema terminando en mismo estado en el cual comenzó”. Es decir, si un sistema experimenta cambios espontáneos, éste cambiará en tal forma que su *entropía* (S) se incrementa, o bien permanezca constante, pero nunca disminuye. *La entropía es una medida del desorden espacial y térmico de un sistema.* El calor siempre fluye espontáneamente de los cuerpos de mayor a los de menor temperatura. Cualquier *proceso irreversible* produce cambios del sistema hacia un *estado de mayor entropía*. Para un *proceso ideal reversible e isotérmico* se tiene:

$$dS = \frac{dQ}{T} \rightarrow \Delta S = \int_1^2 \frac{dQ}{T} \quad (77)$$

Tercera ley de la termodinámica. El cero absoluto no puede ser alcanzado por ningún procedimiento que conste de un número finito de pasos.

El *calor específico* o *capacidad calorífica*, $c = \Delta Q / (m \Delta T)$, de una sustancia es la cantidad de calor requerida para elevar la temperatura de una unidad de masa m , en un grado. Para el agua $c = 4180 \text{ J/kg}$. En caso de gases se tienen dos capacidades caloríficas, *calor específico a volumen constante* (C_V) y *calor específico a presión constante* (C_P). Para un gas ideal de *masa molecular* M , se cumple que:

$$C_P - C_V = \frac{R}{M} \quad (78)$$

La *razón de calor específico* para gases es $\gamma = C_P / C_V$, cuyo valor para *gases monoatómicos* es $\gamma = 1.67$, y para gases *diatómicos* $\gamma = 1.40$.

En los cálculos de ingeniería los gases que están sujetos a presiones de unos cuantos *bar* pueden considerarse como ideales. Para un *mol de gas ideal* en un proceso donde *no hay flujo* y es *mecánicamente reversible*, se aplican las siguientes ecuaciones:

$$dU = dQ + dW \rightarrow \Delta U = Q + W$$

$$dW = p dV \rightarrow W = \int p dV$$

$$dU = C_v dT \rightarrow \Delta U = \int C_v dT \quad (79)$$

$$dH = C_p dT \rightarrow \Delta H = \int C_p dT$$

Se dice que el proceso es *politrópico* si la relación presión-volumen, está dada por: $PV^\delta = K$, siendo K una constante. En tal proceso, el *trabajo* W , y el *calor* Q , se obtienen con las siguientes ecuaciones:

$$W = \frac{RT_1}{\delta - 1} \left[\left(\frac{P_2}{P_1} \right)^{(\delta-1)/\delta} - 1 \right]$$

$$Q = \frac{(\delta - \gamma)RT_1}{(\delta - 1)(\gamma - 1)} \left[\left(\frac{P_2}{P_1} \right)^{(\delta-1)/\delta} - 1 \right] \quad (80)$$

Para valores particulares de la constante δ , las ecuaciones anteriores se reducen a los siguientes *procesos termodinámicos*, para un gas ideal: *isobárico* (presión constante, $\delta=0$), *isotérmico* (temperatura constante, $dU=0$, $\delta=1$), *adiabático* (sin transferencia de calor, $dQ=0$, $\delta=\gamma$), *isocórico* o *isovolumétrico* (volumen constante, $dW=0$, $\delta=\infty$),

Los problemas cuyas soluciones dependen exclusivamente de la *conservación de la masa* y de las *leyes de la termodinámica* suelen dejarse fuera del estudio de la *mecánica de fluidos* (que involucra el principio de *momento lineal*) y son estudiados por la *termodinámica*.

Un *volumen de control* (V_C), es un volumen aislado de un sistema físico en estudio, con una entrada y una salida para comunicarse con sus alrededores. Consideremos un volumen de control que aloja una cantidad contable X de fluido. Una *ecuación de balance de masa* expresa la relación entre la *rapidez de transporte* ($\dot{X}_T$) hacia (+) o desde el fluido (-), con la *rapidez con la que se genera* ($\dot{X}_G$) la sustancia dentro del volumen y con la *rapidez de cambio de la sustancia* dentro del volumen de control ($d\dot{X}_{vc}/dt$):

$$\dot{X}_T + \dot{X}_G = \frac{d\dot{X}_{vc}}{dt} \quad (81)$$

La *rapidez de flujo másico* m con la que fluye un fluido en fase gaseosa, de *densidad* ρ , a través de una tubería de *área de sección transversal* A , con *velocidad promedio* v , esta dada por:

$$m = \rho v A \quad (82)$$

El proceso de *flujo en estado estable* es aquel en el que el flujo másico de entrada y salida es el mismo. Dicho estado obedece la *ecuación de continuidad*, la cual se puede escribir en términos del *volumen específico* ($V_e = 1/\rho$).

$$m = \rho_1 v_1 A_1 = \rho_2 v_2 A_2 = \frac{v_1 A_1}{V_{e1}} = \frac{v_2 A_2}{V_{e2}} \quad (83)$$

Bloques funcionales de sistemas fluidicos neumáticos

Resistencia neumática

Considere un sistema en fase gaseosa en flujo estable, fluyendo a través de una tubería. La *resistencia neumática* se define en términos de la razón de *flujo másico* m y la diferencia de presiones en dos puntos 1 y 2, a lo largo de la tubería, como:

$$p_1 - p_2 = Rm \quad (84)$$

Capacitancia neumática

Consideremos un volumen de control V_c , que aloja un gas, con las siguientes propiedades, *masa* m , *densidad* ρ , *temperatura* T , al cual entra un flujo másico m_1 y sale un flujo másico m_2 . Supongamos que las condiciones del gas estén dadas por la ecuación del gas ideal.

$$pV = mRT \rightarrow p = (m/V)RT = \rho RT \rightarrow \rho = (1/RT)p$$

El *cambio de masa* en el volumen de control, sabiendo que $m = \rho V$, y considerando la ecuación anterior es:

$$m_1 - m_2 = \frac{d(m)}{dt} = \frac{d(\rho V)}{dt} = \rho \frac{dV}{dt} + V \frac{d\rho}{dt} = \rho \frac{dV}{dt} + V \frac{d(\frac{1}{RT} p)}{dt} = \rho \frac{dV}{dt} + \frac{V}{RT} \frac{dp}{dt}$$

Lo cual se puede reescribir como:

$$m_1 - m_2 = \rho \frac{dV}{dp} \frac{dp}{dt} + \frac{V}{RT} \frac{dp}{dt} = \left(\rho \frac{dV}{dp} + \frac{V}{RT} \right) \frac{dp}{dt} = (C_1 + C_2) \frac{dp}{dt} \quad (85)$$

Donde las constantes C_1 y C_2 son la *capacitancia neumática debida al cambio de volumen* y la *capacitancia debida a la compresibilidad del gas*, respectivamente.

$$C_1 = \rho \frac{dV}{dp}, \quad C_2 = \frac{V}{RT} \quad (86)$$

En términos de la diferencia de presiones, la ecuación puede escribirse como:

$$p_1 - p_2 = \frac{1}{C_1 + C_2} \int (m_1 - m_2) dt \quad (87)$$

Inertancia neumática

Consideremos un bloque de gas de *área de la sección transversal* A , *longitud* L . Sea su *densidad uniforme* ρ , y su *masa* $m = \rho AL$. Apliquemos la *segunda ley de Newton* para estudiar las relaciones entre la fuerza que se provee al bloque de gas para acelerarlo y sus propiedades. Consideremos que la fuerza actúa en las caras laterales del bloque, ejerciendo presiones p_1 y p_2 .

$$f = (p_1 - p_2)A = ma = \frac{d(mv)}{dt} = \frac{d(\rho LA v)}{dt} = L \frac{d(\rho q)}{dt} \quad (88)$$

Donde $q = Av$, es el *flujo volumétrico*, así mismo el *flujo másico* esta dado por $m = \rho q$. En estas condiciones se tiene que:

$$p_1 - p_2 = \frac{L}{A} \frac{dm}{dt} = I \frac{dm}{dt} \Leftrightarrow I = \frac{L}{A} \quad (89)$$

La constante $I = L/A$ se denomina *inertancia neumática*.

Sistemas térmicos

Para sistemas térmicos solo se consideran dos bloques funcionales: *resistencia térmica* y *capacitancia térmica*.

Fundamentos de transferencia de calor y calorimetría

Las *escalas de temperatura* mas comunes (con el *punto fusión* y *ebullición* del agua indicado en paréntesis) son: *Celsius* (0-100 °C), *Fahrenheit* (32-212 °F), *Rankine* (492-672 °R) y la *escala absoluta* o *Kelvin* (273.15-373.15 K). Se relacionan entre si, como se indica a continuación:

$$\begin{aligned} K &= ^\circ C + 273.15 \\ ^\circ F &= 1.8^\circ C + 32 \rightarrow ^\circ C = (^\circ F - 32)/1.8 \\ ^\circ R &= ^\circ F + 460 \end{aligned} \quad (90)$$

El *punto triple del agua* se define a 0.01°C y equivale a 273.16K. Si un material de *coeficiente lineal de expansión* α , y *coeficiente volumétrico de expansión* β ($\beta=3\alpha$ para sólidos), con *longitud* L_0 y *volumen* V_0 , se somete a un *cambio de temperatura*, ΔT , este resiente un cambio en longitud y volumen dados por:

$$\Delta L = \alpha L_0 \Delta T \quad \Delta V = \beta V_0 \Delta T \quad (91)$$

El *calor transferido de un cuerpo a otro* es el resultado de la *diferencia de temperaturas*. El calor necesario para elevar la temperatura de un material de *capacidad calorífica* c y *masa* m , una cantidad ΔT es: $Q = mc\Delta T$.

Los tres mecanismos para la *transferencia de calor* son: *conducción*, *convección* y *radiación*. La *conducción* ocurre cuando la energía calorífica pasa a través de un material como resultado de las colisiones entre la moléculas del mismo. La *convección* de energía calorífica involucra movimientos de masa de una región a otra, el material de menor temperatura es desplazado. La *radiación* es transferencia de energía a través de ondas electromagnéticas.

Conducción. Consideremos dos cuerpos de temperatura T_A y T_B , siendo $T_A > T_B$, y sea un *material de conducción* aislado de *longitud* L y *área de sección*

transversal A y sea k su *conductividad*, entonces la *corriente de calor por conducción* (H) es:

$$H = \frac{dQ}{dT} = -kA \frac{dT}{dx} = \frac{Ak(T_A - T_B)}{L} = \frac{A(T_A - T_B)}{R} \Leftrightarrow R = \frac{L}{k} \quad (92)$$

Convección. El proceso de transferencia de calor por convección es complejo y no existe una ecuación general, por el contrario se desarrollan ecuaciones particulares dependiendo de su aplicación. Por ejemplo la *convección natural de aire* desde superficies verticales y planos horizontales esta dado por:

$$\frac{hL}{k_f} = b \left[\frac{L^2 \rho_f^2 g \beta_f \Delta T}{\mu_f^2} \left(\frac{c_p \mu}{k} \right)_f \right]^n \quad (93)$$

Siendo: h =coeficiente medio de transferencia de calor, k_f =conductividad calorífica del fluido, c_p =calor específico a presión constante del fluido, ρ_f =densidad del fluido, β_f =coeficiente de expansión térmica del fluido, g =aceleración debida a la gravedad, ΔT =diferencia de temperaturas entre la lamina y el fluido, μ_f =viscosidad del fluido.

Radiación. Un *cuerpo negro* es un cuerpo que absorbe toda la energía radiante que incide sobre el. En *equilibrio térmico* un cuerpo emite la misma cantidad de energía radiante que la que absorbe. A partir de la *ley de Stefan-Boltzmann*, se tiene la *corriente de calor por radiación* (H) para una *superficie* A , de un cuerpo a *temperatura absoluta* T , esta dada por:

$$H = \frac{dQ}{dT} = Ae\sigma T^4 \quad (94)$$

Donde σ es la constante de *Stefan-Boltzmann*, y e es la *emisividad de la superficie*. Para un cuerpo negro $e=1$ y $e<1$ para una superficie real. Cuando un cuerpo con temperatura absoluta T , está en una región donde la temperatura absoluta es T_0 , la energía radiada por el cuerpo es:

$$H = \frac{dQ}{dT} = Ae\sigma(T^4 - T_0^4) \quad (95)$$

Bloques funcionales de sistemas térmicos

Resistencia térmica

Como sabemos de la termodinámica, sólo hay flujo de calor neto entre dos puntos si hay una diferencia de temperaturas entre ellos. El *flujo de calor* q , debido a una *superficie* A , entre dos puntos de *temperatura* T_1 y T_2 se relacionan por:

$$q = \frac{T_2 - T_1}{R} \quad (96)$$

Donde R es la *resistencia térmica*. El valor de la resistencia depende del modo en el que se transfiere el calor. En la conducción a través de un sólido $R = L/k$. En la convección $R = 1/(Ah)$.

Capacitancia térmica

La *capacitancia térmica* (C) es una medida del almacenamiento de energía interna de un sistema. Sea q_1 la razón de flujo en el interior de un sistema y q_2 la razón de flujo que sale de este, entonces:

$$\Delta q = q_1 - q_2 = mc \frac{dT}{dt} = C \frac{dT}{dt} \Leftrightarrow C = mc \quad (97)$$

Equivalencias entre bloques funcionales de sistemas físicos

Hasta ahora nuestro esfuerzo se ha concentrado en la determinación de los bloques funcionales de los siguientes sistemas físicos: *mecánicos traslacionales*, *mecánicos rotacionales*, *eléctricos*, *fluidicos hidráulicos*, *fluidicos neumáticos* y *térmicos*. La intención que se ha tenido en mente es presentarlos de manera tal que se observen las equivalencias en los modelos matemáticos que los presentan, como se resume en la **tabla 2.10**.

Tabla 2.10 Equivalencias de los bloques funcionales de sistemas físicos

Sistema físico	Bloques funcionales		
	Almacenamiento de energía potencial	Disipador de energía	Almacenamiento de energía cinética
<i>Mecánico trasnacional</i>	Resorte traslacional $F = k \int \bar{v} dt$	Amortiguador traslacional $F = c \bar{v}$	Masa $F = m \frac{d\bar{v}}{dt}$
<i>Mecánico rotacional</i>	Resorte torsional $T = k \int \bar{\omega} dt$	Amortiguador rotacional $T = c \bar{\omega}$	Momento de inercia $T = I \frac{d\bar{\omega}}{dt}$
<i>Eléctrico</i>	Inductor $i = \frac{1}{L} \int v dt$	Resistor $i = \frac{1}{R} v$	Capacitor $i = C \frac{dv}{dt}$
<i>Fluídico hidráulico</i>	Inertancia hidráulica $q = \frac{1}{I} \int (p_1 - p_2) dt$	Resistencia hidráulica $q = \frac{1}{R} (p_1 - p_2)$	Capacitancia hidráulica $q = C \frac{d(p_1 - p_2)}{dt}$
<i>Fluídico neumático</i>	Inertancia neumática $m = \frac{1}{I} \int (p_1 - p_2) dt$	Resistencia neumática $m = \frac{1}{R} (p_1 - p_2)$	Capacitancia neumática $m = C \frac{d(p_1 - p_2)}{dt}$
<i>Térmico</i>	Sin equivalente	Resistencia térmica $q = \frac{1}{R} (T_1 - T_2)$	Capacitancia térmica $q_1 - q_2 = C \frac{dT}{dt}$

Capítulo 3

Modelado de Sistemas Dinámicos

Los sistemas físicos se presentan en la naturaleza como sistemas dinámicos. El modelo matemático que representa un sistema físico se obtiene a partir de un análisis del estado que guardan los bloques funcionales que componen el sistema, aplicando los principios de *conservación de energía, masa y momento (angular y lineal)*. El objetivo de este capítulo es presentar un conjunto de técnicas que permitan la obtención de *modelos matemáticos* para los siguientes sistemas físicos:

- Sistemas mecánicos.
- Sistemas eléctricos.
- Sistemas fluidicos.
- Sistemas térmicos.
- Sistemas electromecánicos.

El *modelado* se presenta mediante la función de transferencia del sistema en el dominio de s , empleando la transformada de Laplace; se presenta una introducción al análisis en espacio de estados.

Modelado de sistemas dinámicos

Una *ecuación diferencial de un sistema dinámico*, es un modelo matemático que representa la relación entre las variables de entrada del sistema y las variables de salida para una característica del sistema en particular, como una función del tiempo. Un *modelo de estado* es una representación matricial que relaciona las variables de: entrada, salida y de estado; de manera que determina el comportamiento del sistema en cualquier tiempo. En general las *leyes de conservación de la masa* (M), *energía* (E) y *momento* (P) se aplican a los componentes del sistema para formar una ecuación de balance que se puede escribir como:

$$\text{flujode[ME]de entrada} - \text{flujode[ME]de salida} = \text{tasa de acumulación de[ME]} \quad (1)$$

En ingeniería de control se emplean varias técnicas matemáticas bien fundamentadas: (1) modelado mediante *función de transferencia*, aplicable a *sistemas lineales invariantes con el tiempo* (ED lineales con coeficientes constantes); (2) representación del sistema mediante *diagramas de bloques* y (3) modelado de sistemas complejos mediante *espacio de estados*.

Modelado mediante función de transferencia

La *teoría de control convencional* emplea la función de transferencia para el análisis de sistemas lineales invariantes con el tiempo sujetos a una sola entrada y una sola salida; y el análisis se realiza en el dominio de la frecuencia ($s, j\omega$). Una *función de transferencia* o *ganancia* $G(s)$, de un sistema se obtiene al aplicar la *transformada de Laplace* a la *ED invariante con el tiempo*, que representa el sistema (transformando la ED del dominio del tiempo al dominio de los números complejos) y se expresa como una razón de la *salida* o *respuesta* $Y(s)$, entre la *entrada* o *excitación* $X(s)$; bajo la suposición de que todas las condiciones iniciales son cero:

$$Y(s) = G(s) \cdot X(s) \quad \rightarrow \quad G(s) = \frac{Y(s)}{X(s)} = \frac{L[y(t)]}{L[x(t)]} \quad (2)$$

El *orden del sistema* queda determinado por el orden de $X(s)$ y en general el orden de $X(s)$ será siempre menor que el orden de $Y(s)$.

Ejemplo 3.1. Obtener la *función de transferencia* para el sistema mecánico mostrado en la figura cuya entrada es la *fuerza externa* $u(t)$ y su salida el *desplazamiento* $y(t)$.

Solución. Cuando se aplica la *fuerza externa* $u(t)$ a la masa, en la dirección indicada, las *fuerzas de resorte* y *amortiguamiento* se oponen al desplazamiento. De esta manera: (1) aplicando la *segunda ley de Newton*, $F=ma$; (2) *transformando* la ED lineal e invariante resultante; (3) hallando la *función de transferencia* suponiendo condiciones iniciales iguales a cero, se tiene:

$$\begin{aligned} u(t) - ky(t) - b \frac{dy(t)}{dt} &= m \frac{d^2 y(t)}{dt^2} \quad \rightarrow \quad m \frac{d^2 y}{dt^2} + b \frac{dy}{dt} + ky = u \\ ms^2 Y(s) + bsY(s) + kY(s) &= U(s) \quad \rightarrow \quad Y(s) \times [ms^2 + bs + k] = U(s) \\ \Rightarrow G(s) &= \frac{Y(s)}{U(s)} = \frac{1}{ms^2 + bs + k} \quad \clubsuit \end{aligned}$$

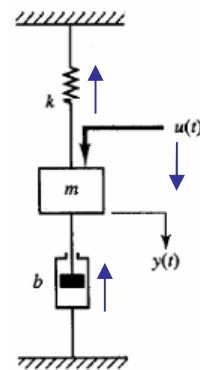


Figura 3.1 Sistema bloque-resorte-amortiguador

Representación del sistema mediante diagramas de bloques

Los *diagramas de bloques* son representaciones gráficas de la dinámica de un sistema de control, en su forma más sencilla representan una función de transferencia; y constan de los siguientes elementos (**figura 3.2**): (1) **flechas**: se emplean para representar la dirección del flujo de la señal; (2) **punto suma**: es el lugar donde dos o más señales se suman o restan dependiendo del signo indicado; (3) **punto de separación**: es el lugar donde la señal se separa en dos o más direcciones de flujo; (4) **bloques**: son recuadros que alojan una operación matemática de transformación para la señal de entrada.

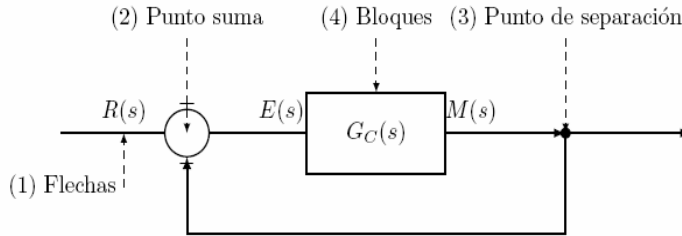


Figura 3.2 Elementos de un diagrama de bloques

Para un **sistema de lazo abierto**, que consta de elementos conectados en serie, la función de transferencia global es el producto de las funciones de transferencia de los elementos individuales:

$$G = \prod_{k=1}^n G_k = G_1 G_2 \cdots G_n \quad (3)$$

Para un **sistema de lazo cerrado** consideremos el *diagrama de bloques* mostrado en la **figura 3.3**. Se tienen las siguientes relaciones:

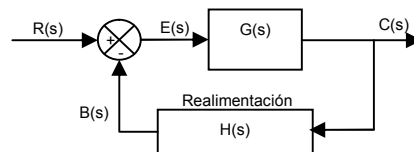


Figura 3.3 Función de transferencia de un sistema de lazo cerrado

$$\begin{aligned}
 C(s) &= G(s)E(s) \\
 E(s) &= R(s) - B(s) = R(s) - H(s)C(s) \\
 \rightarrow \frac{C(s)}{R(s)} &= \frac{G(s)}{1 + G(s)H(s)}
 \end{aligned}
 \tag{4}$$

En la ecuación anterior G se conoce como *función de transferencia de la trayectoria directa*, GH se conoce como *función de transferencia de lazo*.

El lenguaje de los diagramas de bloques

El álgebra de bloques (**figura 3.4**) es el conjunto de reglas que permite representar, analizar y simplificar gráficamente las interconexiones entre bloques de un sistema sin alterar las ecuaciones que describen su comportamiento y se fundamenta en las siguientes propiedades generales:

- (1) **Conmutativa:** El orden de los bloques o señales no altera el resultado (aplica en sumadores y en conexiones en serie).
- (2) **Asociativa:** La agrupación de operaciones no modifica el efecto global sobre la señal.
- (3) **Distributiva:** Un bloque que multiplica a una suma o resta de señales puede aplicarse a cada señal por separado.
- (4) **Superposición:** La respuesta total a múltiples entradas es la suma de las respuestas individuales.

Las reglas derivadas de estas propiedades son las siguientes.

Regla 1: Combinación y desplazamiento de sumadores. Cuando a una señal principal se le restan varias entradas secundarias, el resultado no cambia ya sea que las restas se realicen en pasos consecutivos o bien se agrupen en un único nodo con múltiples entradas, ya que las operaciones de suma cumplen las propiedades *conmutativa* y *asociativa*.

Regla 2: Bloques en cascada o serie. Cuando una señal pasa por varios bloques conectados en serie, el resultado no cambia aunque se intercambie el orden de los bloques o se sustituyan por un único bloque equivalente cuya función de transferencia sea el producto de todos ellos ($G_1 \dots G_2$), ya que la conexión en serie cumple las propiedades *conmutativa* y *asociativa*.

Regla 3: Distributividad respecto a un bloque. Cuando un bloque multiplica una suma o resta de señales, el resultado no cambia si el bloque se aplica por separado a cada una de ellas y posteriormente se combinan nuevamente mediante un sumador, ya que la multiplicación cumple la *propiedad distributiva*.

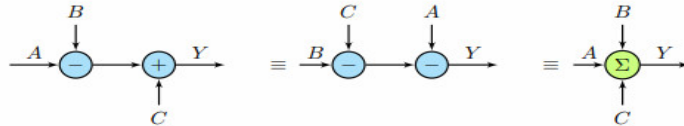
Esto permite desplazar los sumadores antes o después de los bloques, siempre que se conserven las relaciones entre las señales.

Regla 4: Bloques en paralelo. Cuando una misma señal se divide para atravesar varios bloques en paralelo y posteriormente sus salidas se combinan en un sumador, el resultado equivale a reemplazar todos los bloques por un único bloque cuya función de transferencia sea la suma algebraica de ellos, ya que las contribuciones de cada camino se acumulan.

Regla 5: Superposición de entradas múltiples. Cuando un sistema lineal recibe varias entradas independientes de manera simultánea, la salida total es igual a la suma de las respuestas que produciría cada entrada por separado, ya que los sistemas lineales cumplen el *principio de superposición*.

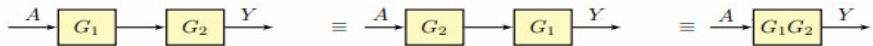
$$1. Y = A - B - C$$

R1: combinación y desplazamiento de sumadores



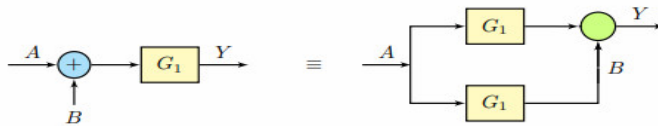
$$2. Y = G_1 G_2 A$$

R2: bloques en cascada o serie



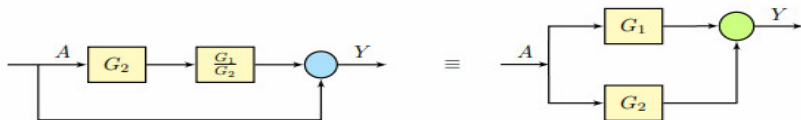
$$3. Y = G_1(A - B)$$

R3: distributividad respecto a un bloque



$$4. Y = (G_1 + G_2)A$$

R4: bloques en paralelo



$$5. Y = G_1 A + G_2 B$$

R5: superposición de entradas múltiples

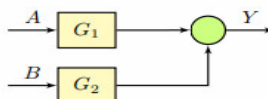


Figura 3.4 Reglas algebraicas de los diagramas de bloques

Efectos de las perturbaciones

Una **perturbación** es una señal no deseada la cual afecta la salida del sistema. Las perturbaciones pueden venir de fuentes exógenas (medio ambiente) o del interior del sistema, por ejemplo *ruido eléctrico*.

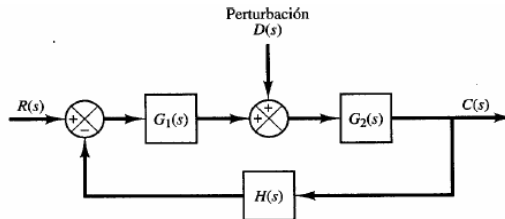


Figura 3.5 Sistema de control en lazo cerrado con perturbación

Consideremos el sistema de lazo cerrado sujeto a una perturbación $D(s)$, mostrado en la **figura 3.5**. La respuesta del sistema $C(s)$ es la superposición de las respuestas de la señal de entrada $C_R(s)$ y la perturbación $C_D(s)$, esto es:

$$C(s) = C_R(s) + C_D(s) = \frac{G_2(s)}{1 + G_1(s)G_2(s)H(s)} [G_1(s)R(s) + D(s)] \quad (5)$$

En general los sistemas de lazo cerrado son poco sensibles a los efectos de las perturbaciones.

Representación de un sistema dinámico en diagrama de bloques

Ejemplo 3.2. Para el circuito RC mostrado en la **figura 3.6**; cuya entrada es e_i y cuya salida es e_o . Hallar su representación en diagrama de bloques.

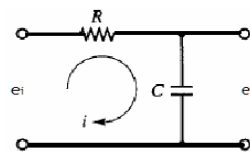
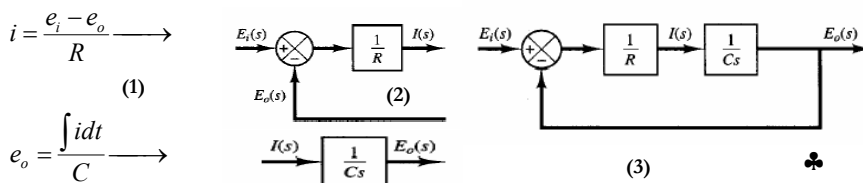


Figura 3.6 Circuito RC

Solución. (1) Se escriben ecuaciones que relacionen las variables de entrada y salida; (2) se transforman en diagramas de bloques; (3) se forma un diagrama de bloques general:



Modelado mediante espacio de estados

En la *teoría de control moderna* se emplea el concepto de espacio de estados para el análisis de sistemas dinámicos que pueden tener entradas y salidas múltiples y que pueden variar con el tiempo (ED no lineales); y el análisis se realiza en el dominio del tiempo (t).

El *estado* de un sistema queda determinado al conocer el *n-vector de estado*, cuyos componentes son *n-variables de estado* y que gráficamente se representa como un punto en el *n-espacio de estados*. Una *ecuación de estado* es una relación entre tres conjuntos de variables que determinan el comportamiento del sistema: (1) *variables de entrada*, (2) *variables de salida* y (3) *variables de estado*. Un sistema dinámico debe incorporar *elementos que memoricen los valores de la entrada*, estos son los *integradores* y *sus salidas son las variables de estado* $x_i = \int x'_i dx$.

Mediante el espacio de estados un *sistema dinámico*, se puede representar mediante el siguiente conjunto de ecuaciones matriciales:

$$\mathbf{x}'(t) = \mathbf{A}\mathbf{x}(t) + \mathbf{B}\mathbf{u}(t)$$

$$\mathbf{y}(t) = \mathbf{C}\mathbf{x}(t) + \mathbf{D}\mathbf{u}(t)$$

donde :

A = matriz de estado (6)

B = matriz de entrada

C = matriz de salida

D = matriz de transmisión directa

Las matrices **A**, **B**, **C** y **D** son funciones del tiempo si el sistema es variante con el tiempo, de lo contrario son matrices constantes. El *vector de entrada* (**u**), *vector de salida* (**y**) y *vector de estado* (**x**), se definen por:

$$\mathbf{u}(t) = \begin{bmatrix} u_1(t) \\ u_2(t) \\ \vdots \\ u_r(t) \end{bmatrix}, \quad \mathbf{y}(t) = \begin{bmatrix} y_1(t) \\ y_2(t) \\ \vdots \\ y_m(t) \end{bmatrix}, \quad \mathbf{x}(t) = \begin{bmatrix} x_1(t) \\ x_2(t) \\ \vdots \\ x_n(t) \end{bmatrix} \quad (7)$$

 **Ejemplo 3.3.** Representar el sistema masa-resorte-amortiguador, mediante ecuaciones de estado.

Solución. El punto clave para la representación del sistema mediante ecuaciones de estado es la selección adecuada de las variables de estado, y se lleva a cabo como sigue: (1) se determina la ED que representa el sistema dinámico y el grado de esta (n); (2) se eligen n -variables de estado, que pueden ser la función de salida y sus ($n-1$) derivadas; (3) se escriben relaciones de las variables de entrada y las variables de salida con las variables de estado; (4) se expresan matricialmente las ecuaciones de estado; (5) se representa el sistema en su forma matricial estándar identificando las matrices de coeficientes (**A**, **B**, **C**, **D**):

$$\begin{aligned}
 (1) \quad & m \frac{d^2 y}{dt^2} + b \frac{dy}{dt} + ky = u \rightarrow my'' + by' + ky = u \\
 (2) \quad & x_1(t) = y(t), \quad x_2(t) = y'(t) \quad \leftarrow \quad \text{variables de estado} \\
 (3) \quad & x_1' = x_2, \\
 & x_2' = \frac{1}{m}(-ky - by') + \frac{1}{m}u \rightarrow x_2' = -\frac{k}{m}x_1 - \frac{b}{m}x_2 + \frac{1}{m}u \\
 & Y = x_1 \\
 (4) \quad & \begin{bmatrix} x_1' \\ x_2' \end{bmatrix} = \begin{bmatrix} 0 & 1 \\ -\frac{k}{m} & -\frac{b}{m} \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} + \begin{bmatrix} 0 \\ \frac{1}{m} \end{bmatrix} u, \quad Y = \begin{bmatrix} 1 & 0 \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} \\
 (5) \quad & \mathbf{x}'(t) = \mathbf{Ax}(t) + \mathbf{Bu}(t) \\
 & \mathbf{y}(t) = \mathbf{Cx}(t) + \mathbf{Du}(t) \\
 & A = \begin{bmatrix} 0 & 1 \\ -\frac{k}{m} & -\frac{b}{m} \end{bmatrix}, \quad B = \begin{bmatrix} 0 \\ \frac{1}{m} \end{bmatrix}, \quad C = \begin{bmatrix} 1 & 0 \end{bmatrix}, \quad D = 0 \quad \clubsuit
 \end{aligned}$$

Representación de EDL con coeficientes constantes en espacio de estados

El método antes enunciado se generalizará para representar un sistema dinámico que es representado por una ecuación diferencial lineal de n -ésimo orden, en su *forma normal*:

$$y^{(n)} + a_1 y^{(n-1)} + \dots + a_{n-1} y^{(1)} + a_n y = b_1 u \quad (8)$$

Consideremos como variables de estado la función de salida y sus derivadas y definimos:

$$x_1 = y, \quad x_2 = y^{(1)}, \quad \dots, \quad x_n = y^{(n-1)} \quad (9)$$

Reescribiendo la ED en términos de las variables definidas:

$$x'_1 = x_2, \quad x'_2 = x_3, \quad \dots, \quad x'_{n-1} = x_n \quad (10)$$

Por lo que matricialmente en espacio de estados la ED lineal de n-ésimo orden se representa por dos ecuaciones lineales de primer orden:

$$\mathbf{x}' = \mathbf{Ax} + \mathbf{Bu}$$

$$\mathbf{y} = \mathbf{Cx} + \mathbf{D}$$

$$\mathbf{x} = \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_{n-1} \\ x_n \end{bmatrix}, \quad \mathbf{A} = \begin{bmatrix} 0 & 1 & 0 & \dots & 0 \\ 0 & 0 & 1 & \dots & 0 \\ 0 & 0 & 0 & \dots & 0 \\ 0 & 0 & 0 & \dots & 1 \\ -a_n & -a_{n-1} & -a_{n-2} & \dots & -a_1 \end{bmatrix}, \quad \mathbf{B} = \begin{bmatrix} 0 \\ 0 \\ \vdots \\ 0 \\ b_1 \end{bmatrix} \quad (11)$$

$$\mathbf{C} = [1 \ 0 \ 0 \ \dots \ 0], \quad \mathbf{D} = 0$$

Asimismo la función de transferencia del sistema está dada por:

$$G(s) = \frac{Y(s)}{U(s)} = \frac{b_1}{s^n + a_1 s^{n-1} + \dots + a_{n-1} s + a_n} \quad (12)$$

Relación entre funciones de transferencia y espacio de estados

Cuando se tiene un modelo en el espacio de estados de un sistema con una sola entrada y una sola salida, es posible obtener la función de transferencia a partir de los parámetros **A**, **B**, **C** y **D** como se muestra a continuación.

$$\begin{aligned} L[\mathbf{x}'(t) = \mathbf{Ax}(t) + \mathbf{Bu}(t)] &= sX(s) = \mathbf{AX}(s) + \mathbf{BU}(s) \rightarrow sX(s) - \mathbf{AX}(s) = \mathbf{BU}(s) \\ &= (sI - \mathbf{A})X(s) = \mathbf{BU}(s) \rightarrow X(s) = (sI - \mathbf{A})^{-1} \mathbf{BU}(s) \\ L[\mathbf{y}(t) = \mathbf{Cx}(t) + \mathbf{Du}(t)] &= Y(s) = \mathbf{CX}(s) + \mathbf{DU}(s) \\ &= \mathbf{C}(sI - \mathbf{A})^{-1} \mathbf{BU}(s) + \mathbf{DU}(s) = U(s) \times [\mathbf{C}(sI - \mathbf{A})^{-1} \mathbf{B} + \mathbf{D}] \\ G(s) &= \frac{Y(s)}{U(s)} = \mathbf{C}(sI - \mathbf{A})^{-1} \mathbf{B} + \mathbf{D} \end{aligned} \quad (13)$$

 **Ejemplo 3.4.** Obtener la función de transferencia para el sistema masa-resorte-amortiguador representado en espacio de estados.

Solución. Aplicando la ecuación definitoria se tiene que:

$$\begin{aligned}
 \mathbf{A} &= \begin{bmatrix} 0 & 1 \\ -\frac{k}{m} & -\frac{b}{m} \end{bmatrix}, & \mathbf{B} &= \begin{bmatrix} 0 \\ 1 \end{bmatrix}, & \mathbf{C} &= [1 \quad 0], & \mathbf{D} &= 0 \\
 (sI - \mathbf{A})^{-1} &= \left(\begin{bmatrix} s & 0 \\ 0 & s \end{bmatrix} - \begin{bmatrix} 0 & 1 \\ -\frac{k}{m} & -\frac{b}{m} \end{bmatrix} \right)^{-1} = \begin{bmatrix} s & -1 \\ \frac{k}{m} & s + \frac{b}{m} \end{bmatrix}^{-1} \\
 \det(sI - \mathbf{A}) &= s(s + b/m) + k/m, & \text{adj } A &= \begin{bmatrix} a_{22} & -a_{12} \\ -a_{21} & a_{11} \end{bmatrix} = \begin{bmatrix} s + \frac{b}{m} & 1 \\ -\frac{k}{m} & s \end{bmatrix} \\
 (sI - \mathbf{A})^{-1} &= \frac{\text{adj } A}{|A|} = \frac{1}{s^2 + \frac{b}{m}s + \frac{k}{m}} \begin{bmatrix} s + \frac{b}{m} & 1 \\ -\frac{k}{m} & s \end{bmatrix} = R \begin{bmatrix} s + \frac{b}{m} & 1 \\ -\frac{k}{m} & s \end{bmatrix} \leftrightarrow R = \frac{1}{s^2 + \frac{b}{m}s + \frac{k}{m}} \\
 G(s) = \frac{Y(s)}{U(s)} &= \mathbf{C}(sI - \mathbf{A})^{-1}\mathbf{B} + \mathbf{D} = [1 \quad 0] \begin{bmatrix} Rs + R\frac{b}{m} & R \\ -R\frac{k}{m} & Rs \end{bmatrix} \begin{bmatrix} 0 \\ 1 \end{bmatrix} + 0 = R \frac{1}{m} \\
 &= \frac{1}{s^2 + \frac{b}{m}s + \frac{k}{m}} \times \frac{1}{m} = \frac{1}{ms^2 + bs + k} \quad \clubsuit
 \end{aligned}$$

Sistemas mecánicos

Para los *sistemas mecánicos traslacionales* existen tres bloques funcionales básicos: *masa-resorte-amortiguador*; para *sistemas mecánicos rotacionales* los bloques funcionales básicos son: *momento de inercia-resorte rotacional-amortiguador rotacional*. En *análisis de sistemas mecánicos* se lleva a cabo aplicando la segunda ley de Newton a cada masa o momento de inercia según corresponda, con frecuencia es útil dibujar un *diagrama de cuerpo libre*, en el que se muestran las fuerzas o torques actuando sobre cada masa:

$$\sum F = m \frac{d^2 x}{dt^2}, \quad \sum T = I \frac{d^2 \theta}{dt^2}$$

(14)

Ejemplo 3.5. Obtener el modelo por función de transferencia del sistema mostrado en la **figura 3.7**, donde la fuerza externa es la variable de entrada y el desplazamiento la variable de salida.

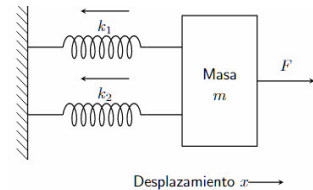


Figura 3.7 Sistema masa-resorte

Solución. Cuando la fuerza $f(t)$ se aplica a la masa en la dirección indicada, los resortes se oponen al movimiento con fuerzas indicadas. Aplicando la segunda ley de Newton, se tiene que:

$$f(t) - k_1 x - k_2 x = m \frac{d^2 x}{dt^2} \quad \rightarrow \quad m \frac{d^2 x}{dt^2} + (k_1 + k_2)x = f(t)$$

$$\mathcal{L} \left[m \frac{d^2 x}{dt^2} + (k_1 + k_2)x = f(t) \right] \quad \rightarrow \quad ms^2 X(s) + (k_1 + k_2)X(s) = F(s)$$

$$X(s) \times [ms^2 + (k_1 + k_2)] = F(s) \quad \rightarrow \quad G(s) = \frac{X(s)}{F(s)} = \frac{1}{ms^2 + (k_1 + k_2)} \quad \clubsuit$$

Ejemplo 3.6. Obtener el modelo por función de transferencia de un motor que hace girar una carga, donde el par externo T , es la variable de entrada y el desplazamiento angular la variable de salida.

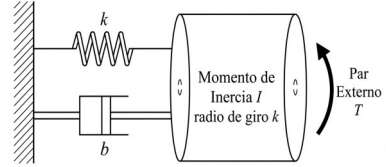


Figura 3.8 Masa que gira sujeta al extremo de un eje

Solución. En la **figura 3.8** se muestra el modelo de bloques. La masa es representada por el momento de inercia I , con radio de giro k y es hecha girar por un par externo T (internamente es un par aplicado a la armadura, debido a fuerzas electromagnéticas por interacción con el campo magnético, $T = B I A N \sin \theta_x$), el resorte y el amortiguador rotacionales generan pares que se oponen al movimiento de la masa; se aplica la *segunda ley de Newton en forma rotacional* $T = I \frac{d^2 \theta}{dt^2}$, y se tiene que:

$$T(t) - k\theta - c \frac{d\theta}{dt} = I \frac{d^2 \theta}{dt^2} \quad \rightarrow \quad I \frac{d^2 \theta}{dt^2} + c \frac{d\theta}{dt} + k\theta = T(t)$$

$$\mathcal{L} \left[I \frac{d^2 \theta}{dt^2} + c \frac{d\theta}{dt} + k\theta = T(t) \right] \quad \rightarrow \quad Is^2 \Theta(s) + cs \Theta(s) + k\Theta(s) = T(s)$$

$$\Theta(s) \times [Is^2 + cs + k] = T(s) \quad \rightarrow \quad G(s) = \frac{\Theta(s)}{T(s)} = \frac{1}{I \cdot s^2 + cs + k} \quad \clubsuit$$

🔗 **Ejemplo 3.7.** Obtener el modelo por función de transferencia del sistema de suspensión de un automóvil, **figura 3.9a**.

Solución. El sistema de suspensión es complejo, un acercamiento es como el mostrado en la **figura 3.9b**, una simplificación mayor conduce al sistema mostrado en la **figura 3.9c**. En el cual el desplazamiento x_i del punto P es la entrada del sistema y el desplazamiento x_o de la masa es la salida del sistema.

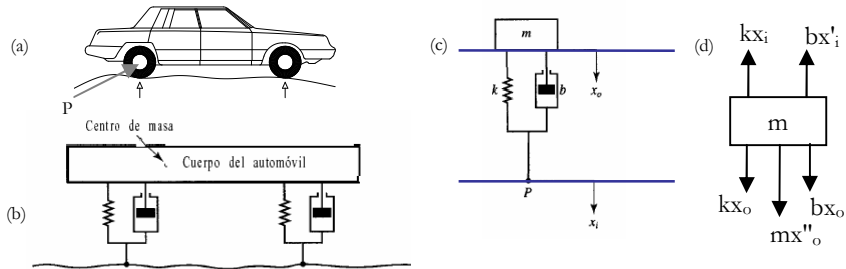


Figura 3.9 (a) Suspensión de un automóvil; (b) sistema de suspensión; (c) simplificación; (d) cuerpo libre

La suspensión de un automóvil tiene por objeto la estabilidad de este, de manera que se minimicen las vibraciones como resultado del relieve del terreno. Se compone de dos elementos principales *resortes* (o muelles) y *amortiguadores* además de brazos y ejes rígidos. Refiriéndose a la figura 3.9c el sistema se modela como un sistema *masa-resorte-amortiguador*. En términos de desplazamientos, cuando un desplazamiento x_i , se aplica al punto P, la masa m reacciona con desplazamiento x_o a partir de su posición de equilibrio; pero se amortigua mediante el amortiguador de constante b y el resorte de constante k . El diagrama de cuerpo libre se muestra en la **figura 3.9d**, en símbolos se tiene que:

$$\begin{aligned}
 m \frac{d^2 x_o}{dt^2} + kx_o + b \frac{dx_o}{dt} - kx_i - b \frac{dx_i}{dt} &= m \frac{d^2 x_o}{dt^2} + b \left(\frac{dx_o}{dt} - \frac{dx_i}{dt} \right) + k(x_o - x_i) = 0 \\
 L \left[m \frac{d^2 x_o}{dt^2} + b \left(\frac{dx_o}{dt} - \frac{dx_i}{dt} \right) + k(x_o - x_i) \right] &\rightarrow ms^2 X_o(s) + bs(X_o(s) - X_i(s)) + k(X_o(s) - X_i(s)) = 0 \\
 X_o(s) \times (ms^2 + bs + k) - X_i(s) \times (bs + k) &= 0 \rightarrow X_o(s) \times (ms^2 + bs + k) = X_i(s) \times (bs + k) \\
 G(s) = \frac{X_o(s)}{X_i(s)} &= \frac{bs + k}{ms^2 + bs + k} \quad \clubsuit
 \end{aligned}$$

Sistemas eléctricos

Los bloques funcionales básicos de sistemas eléctricos son: *inductor-resistor-capacitor*. El análisis de sistemas eléctricos se lleva a cabo aplicando las *leyes de voltajes y corrientes de Kirchhoff* (LVK, LCK), mediante dos métodos conocidos como análisis de mallas y análisis de nodos que expresan la conservación de la energía y la carga respectivamente:

$$\sum_{k=1}^n \hat{V}_k = 0 \leftarrow \text{LVK}$$

$$\sum_{k=1}^n \hat{I}_k = 0 \leftarrow \text{LCK}$$

(15)

Ejemplo 3.8. Obtener el modelo por función de transferencia del sistema RLC mostrado en la **figura 3.10**, la variable de entrada es el voltaje aplicado v , la variable de salida es el voltaje en el capacitor v_C .

Solución. La **figura 3.10a**, muestra el análisis de nodos, la **figura 3.10b** muestra el análisis de mallas, ambos conllevan al mismo resultado, es decir:

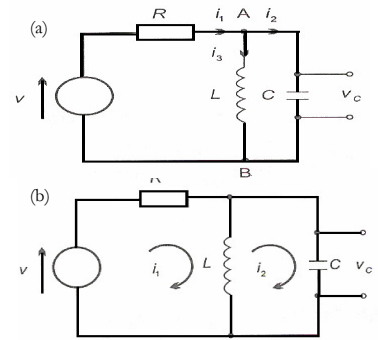


Figura 3.10 Circuito R-LτC, (a) análisis de nodos; (b) análisis de mallas

Análisis de nodos :

$$\begin{aligned} \sum i = 0: \quad i_1 - i_2 - i_3 &= 0, \quad i_1 = \frac{v - v_A}{R}, \quad i_2 = \frac{1}{L} \int V_A dt, \quad i_3 = C \frac{dv_A}{dt} \\ \frac{v - v_A}{R} - \frac{1}{L} \int v_A dt - C \frac{dv_A}{dt} &= 0 \quad \text{pero } v_C = v_A, \text{ entonces: } RC \frac{dv_C}{dt} + \frac{R}{L} \int v_C dt + v_C = v \\ L \left[RC \frac{dv_C}{dt} + \frac{R}{L} \int v_C dt + v_C = v \right] &\rightarrow RCsV_C(s) + \frac{R}{L} \frac{V_C(s)}{s} + V_C(s) = V(s) \\ V_C(s) \left(RCs + \frac{R}{sL} + 1 \right) = V(s) &\rightarrow G(s) = \frac{V_C(s)}{V(s)} = \frac{1}{RCs + \frac{R}{sL} + 1} \\ \therefore G(s) &= \frac{1}{LCs^2 + RCs + 1} \end{aligned}$$

Análisis de mallas :

$$\text{Malla 1 : } v = Ri_1 + L \frac{d(i_1 - i_2)}{dt}$$

$$\text{Malla 2 : } 0 = v_c + L \frac{d(i_2 - i_1)}{dt} = v_c - L \frac{d(i_1 - i_2)}{dt} \rightarrow v_c = L \frac{d(i_1 - i_2)}{dt} \rightarrow \int v_c dt = L(i_1 - i_2)$$

$$\rightarrow \int v_c dt = Li_1 - LC \frac{dv_c}{dt} \rightarrow i_1 = \frac{1}{L} \int v_c dt + C \frac{dv_c}{dt}$$

$$\Rightarrow v = R \left[\frac{1}{L} \int v_c dt + C \frac{dv_c}{dt} \right] + v_c \quad \therefore \quad v = RC \frac{dv_c}{dt} + \frac{R}{L} \int v_c dt + v_c \quad \clubsuit$$

Ejemplo 3.9. Escribir el circuito eléctrico análogo al sistema mecánico mostrado en la **figura 3.11**.

Solución. La equivalencia entre elementos mecánicos y eléctricos es como sigue: *resorte* $\rightarrow$ *inductor*, *amortiguador* $\rightarrow$ *resistor*, *masa* $\rightarrow$ *capacitor*, la fuerza F , es análoga a la intensidad de corriente eléctrica i . Los elementos de la rama 1 están en paralelo con el elemento de la rama 2, lo mismo ocurre en su circuito dual. El resultado se muestra en la **figura 3.12**.

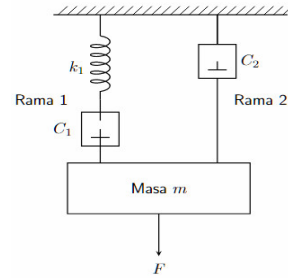


Figura 3.12 Sistema mecánico

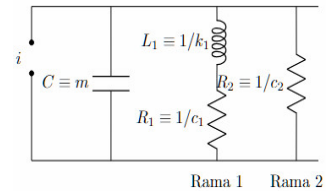


Figura 3.12 Sistema eléctrico

Transformada de Laplace de impedancias complejas

El análisis de circuitos eléctricos se facilita en gran medida cuando las impedancias del circuito se transforman por la técnica de Laplace. La impedancia se define como la razón $Z(s) = V(s)/I(s)$; para los elementos *resistor* R , *inductor* L y *capacitor* C , se tienen las impedancias complejas R , Ls y $1/Cs$ respectivamente, el recíproco de la impedancia es la *admitancia* $Y(s)$. El equivalente de impedancias en serie es su suma algebraica, el equivalente de admitancias en paralelo es su suma algebraica, es decir:

$$Z_{eq}(s) = \sum_{k=1}^n Z_k(s), \quad Y_{eq}(s) = \sum_{k=1}^n Y_k(s) \quad (16)$$

🔗 **Ejemplo 3.10.** Obtener la transformación de Laplace del circuito RLC mostrado en la **figura 3.13**.

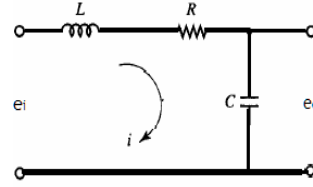


Figura 3.13 Circuito RLC

Solución. Se escribe la LVK para la malla y se transforma:

$$\begin{aligned}
 e_i - v_L - v_R - v_C &= 0, & v_L &= L \frac{di}{dt}, & v_R &= Ri, & v_C &= \frac{1}{C} \int idt \\
 e_i - L \frac{di}{dt} - Ri - \frac{1}{C} \int idt &= 0, & e_o &= \frac{1}{C} \int idt \\
 \mathcal{L} \left[e_i - L \frac{di}{dt} - Ri - \frac{1}{C} \int idt = 0 \right] &= E_i(s) - LsI(s) - RI(s) - \frac{I(s)}{sC} = E_i(s) - I(s) \left[Ls + R + \frac{1}{sC} \right] = 0 \\
 \mathcal{L} \left[e_o = \frac{1}{C} \int idt \right] &= E_o(s) = \frac{I(s)}{sC} \\
 \Rightarrow G(s) = \frac{E_o(s)}{E_i(s)} &= \frac{I(s)/sC}{I(s) \left[Ls + R + \frac{1}{sC} \right]} = \frac{1/sC}{Ls + R + \frac{1}{sC}} = \frac{1}{LCs^2 + RCs + 1} \quad \clubsuit
 \end{aligned}$$

🔗 **Ejemplo 3.11.** Obtener el modelo por función de transferencia del circuito mostrado en la **figura 3.14**, empleando el concepto de impedancia compleja. Siendo $Z_1(s) = Ls + R$, $Z_2(s) = 1/Cs$.

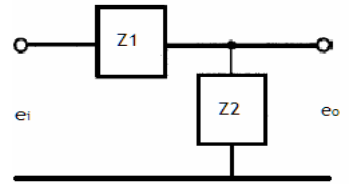


Figura 3.14 Circuito de impedancias complejas

Solución. Puesto que es un circuito de una sola malla, se aplica la LVK y se tiene que:

$$\begin{aligned}
 E_i(s) &= Z_{eq}(s)I(s) = [Z_1(s) + Z_2(s)] \times \frac{E_o(s)}{Z_2(s)} \\
 \Rightarrow G(s) = \frac{E_o(s)}{E_i(s)} &= \frac{Z_2(s)}{Z_1(s) + Z_2(s)} = \frac{1/Cs}{Ls + R + 1/Cs} = \frac{1}{LCs^2 + RCs + 1} \quad \clubsuit
 \end{aligned}$$

🔗 **Ejemplo 3.12.** Obtener el modelo en espacio de estados del circuito RLC mostrado en la **figura 3.13**.

Solución. Se aplica el método para EDL con coeficientes constantes, antes presentado, que consta de los siguientes pasos: (1) se obtiene la EDL y se escribe en su forma normal y se obtiene el orden n ; (2) se escriben las n variables de estado $x_1 \dots x_n$, que consta de la variable de salida y sus $n-1$ derivadas; (3) se escribe la variable de entrada u , y la variable de salida $y=x_n$; (4) se escribe las ecuaciones matriciales de primer orden que representan el estado del sistema de acuerdo al formato presentado; (5) se obtiene la función de transferencia empleando la ecuación antes presentada:

$$\begin{aligned}
 (1) \quad & e_i - L \frac{di}{dt} - Ri - \frac{1}{C} \int i dt = 0, \quad e_o = \frac{1}{C} \int i dt \rightarrow i = C \frac{de_o}{dt} \\
 & e_i - L \frac{d\left(C \frac{de_o}{dt}\right)}{dt} - R \left(C \frac{de_o}{dt}\right) - \frac{1}{C} \int \left(C \frac{de_o}{dt}\right) dt = e_i - LC \frac{d^2 e_o}{dt^2} - RC \frac{de_o}{dt} - e_o = 0 \\
 & \rightarrow e_o'' + \frac{R}{L} e_o' + \frac{1}{LC} e_o = \frac{1}{LC} e_i \quad \leftarrow \quad \text{EDL en su forma normal de orden } n = 2 \\
 & a_1 = \frac{R}{L}, \quad a_2 = \frac{1}{LC} \\
 (2) \quad & x_1 = e_o, \quad x_2 = e_o' \quad \leftarrow \quad \text{variables de estado} \\
 (3) \quad & u = e_i, \quad y = e_o = x_1 \quad \leftarrow \quad \text{variables de entrada y salida} \\
 (4) \quad & \begin{bmatrix} x_1' \\ x_2' \end{bmatrix} = \begin{bmatrix} 0 & 1 \\ -\frac{1}{LC} & -\frac{R}{L} \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} + \begin{bmatrix} 0 \\ \frac{1}{LC} \end{bmatrix} u, \\
 & y = \begin{bmatrix} 1 & 0 \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} \\
 (5) \quad & G(s) = \frac{E_o(s)}{E_i(s)} = \frac{1/LC}{s^2 + R/Ls + 1/LC} = \frac{1}{LCs^2 + RCs + 1} \quad \clubsuit
 \end{aligned}$$

Sistemas fluidicos

Un líquido que fluye dentro de una tubería se dice que tiene *flujo laminar* si su *número de Reynolds* es menor que 2000, si es mayor se dice que es *flujo turbulento*. El número de Reynolds es una cantidad adimensional que se define por:

$$NR = \frac{DG}{\mu}$$

donde : μ = viscosidad [poise = dina - seg/cm² = g/seg - cm]

G = velocidad másica del flujo

D = diámetro de la sección transversal circular

(17)

Los sistemas con flujo laminar se modelan con EDL mientras que los de flujo turbulento requieren EDP; los primeros son menester de nuestro estudio.

 **Ejemplo 3.13.** Para el sistema de nivel de líquidos con iteración mostrado en la **figura 3.15**, obtener las EDL que describen el cambio de las alturas de los líquidos en los contenedores respecto al tiempo, si la inercancia es despreciable.

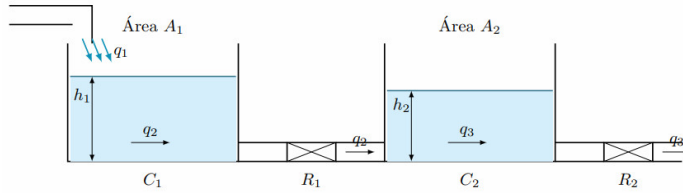


Figura 3.15 Sistema de nivel de líquidos comunicantes

Solución. Los elementos que componen el sistema mostrado son básicamente dos tipos: los *contenedores o capacitores* y las *válvulas*; los primeros tienen la propiedad de *capacidad* y los segundos de *resistencia*. Empleando las ecuaciones definitorias se tiene que:

Para el contenedor 1 :

$$q_1 - q_2 = C_1 \frac{dp}{dt}, \quad p = h_1 \rho g, \quad C_1 = A_1 / \rho g \rightarrow q_1 - q_2 = A_1 \frac{dh_1}{dt}$$

Para la válvula 1 :

$$\Delta p = R_1 q_2, \quad \Delta p = h_1 \rho g - h_2 \rho g \rightarrow (h_1 - h_2) \rho g = R_1 q_2 \rightarrow q_2 = \frac{\rho g}{R_1} (h_1 - h_2)$$

$$\Rightarrow q_1 - \frac{\rho g}{R_1} (h_1 - h_2) = A_1 \frac{dh_1}{dt} \rightarrow \frac{dh_1}{dt} + \frac{\rho g}{A_1 R_1} (h_1 - h_2) = \frac{q_1}{A_1}$$

Para el contenedor 2 :

$$q_2 - q_3 = C_2 \frac{dp}{dt}, \quad p = h_2 \rho g, \quad C_2 = A_2 / \rho g \rightarrow q_2 - q_3 = A_2 \frac{dh_2}{dt}$$

Para la válvula 2 : (el líquido sale a presión atmosférica a $p = 0$)

$$\Delta p = R_2 q_3, \quad \Delta p = h_2 \rho g - 0 \rightarrow h_2 \rho g = R_2 q_3 \rightarrow q_3 = \frac{\rho g}{R_2} h_2$$

$$\rightarrow q_2 - \frac{\rho g}{R_2} h_2 = A_2 \frac{dh_2}{dt} \rightarrow q_2 = A_2 \frac{dh_2}{dt} + \frac{\rho g}{R_2} h_2, \text{ sustituyen do } q_2 \text{ en la ec de la válvula 1 :}$$

$$\Rightarrow \frac{\rho g}{R_1} (h_1 - h_2) = A_2 \frac{dh_2}{dt} + \frac{\rho g}{R_2} h_2 \rightarrow \frac{dh_2}{dt} - \frac{\rho g}{R_1 A_2} (h_1 - h_2) + \frac{\rho g}{R_2 A_2} h_2 = 0 \quad \clubsuit$$

🔗 **Ejemplo 3.14.** Para el tubo en U, mostrado en la **figura 3.16**, que contiene un líquido, obtener la EDL que describe el cambio de altura con el tiempo, cuando se incrementa la presión por el extremo del tubo indicado.

Solución. Cuando se ejerce presión por el extremo izquierdo del tubo, desplazándose el nivel 1 una altura h , la diferencia de alturas entre niveles es $2h$, y el volumen desplazado es $V=Ah$, siendo A el área de sección transversal. *En cualquier instante, la diferencia de presiones entre los dos extremos debe ser igual a la caída de presión total a través del sistema*, si consideramos que el sistema tiene *inertancia, resistencia y capacitancia* entonces se tiene que:

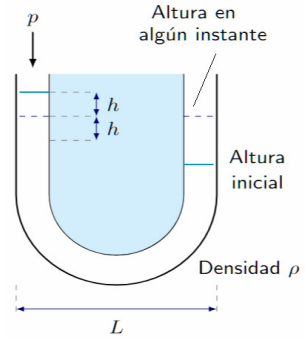


Figura 3.16 Tubo en U conteniendo líquido

$$p = I \frac{dq}{dt} + Rq + \frac{1}{C} \int q dt, \quad q = \frac{dV}{dt} = \frac{d(Ah)}{dt} = A \frac{dh}{dt} \rightarrow p = IA \frac{d^2 h}{dt^2} + RA \frac{dh}{dt} + \frac{A}{C} \int dh$$

$$\int dh = 2h \leftarrow \text{Diferencia entre los niveles}, \quad I = \frac{\rho L}{A}, \quad C = \frac{A}{\rho g}$$

$$\Rightarrow \rho L \frac{d^2 h}{dt^2} + RA \frac{dh}{dt} + 2\rho gh = p \quad \clubsuit$$

Sistemas térmicos

🔗 **Ejemplo 3.15.** Considérese un termómetro de temperatura T que se sumerge en un líquido de temperatura T_L , mostrado en la **figura 3.17**, determinar la EDL que describe la variación de temperatura del termómetro.

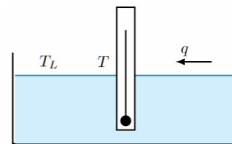


Figura 3.17 Termómetro

Solución. El termómetro tiene una *capacitancia térmica* C , y existe un flujo de calor del líquido al termómetro q , a través de una *resistencia térmica* R , es decir:

$$T_L - T = Rq \leftarrow \text{resistencia a térmica} \quad q = \frac{T_L - T}{R}$$

$$q = C \frac{dT}{dt} \leftarrow \text{flujo de calor del líquido al termómetro}$$

$$\frac{T_L - T}{R} = C \frac{dT}{dt} \rightarrow RC \frac{dT}{dt} + T = T_L \quad \clubsuit$$

🔗 **Ejemplo 3.16.** La **figura 3.18** muestra un sistema térmico que consta de un *calefactor eléctrico* CE en una habitación H. El CE emite calor a la razón q_1 y H pierde calor a la razón q_2 . Si el aire en H está a una *temperatura uniforme* T y que no se almacena calor en las paredes, obtener la EDL que describe como cambia T .

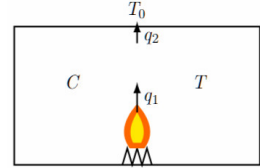


Figura 3.18 Calefacción

Solución. El aire en H tiene una capacitancia térmica C , la temperatura dentro de H es T y T_0 fuera de ella, empleando las ecuaciones definitorias, se tiene que:

$$q_1 - q_2 = C \frac{dT}{dt}, \quad Rq_2 = T - T_0 \rightarrow q_2 = \frac{1}{R}(T - T_0)$$

$$\Rightarrow q_1 - \frac{1}{R}(T - T_0) = C \frac{dT}{dt} \rightarrow RC \frac{dT}{dt} + T = Rq_1 - T_0 \quad \clubsuit$$

Sistemas electromecánicos

Gran parte de los dispositivos de aplicación práctica se componen de la agrupación de sistemas simples, el análisis de estos requiere la aplicación de leyes de conservación de las energías que intervienen en su operación.

🔗 **Ejemplo 3.17.** Determinar el modelo matemático del *motor de CD* (a) controlado por armadura; (b) controlado por campo. Dibujar el circuito, realizar un diagrama de representativo del modelo.

Solución. Un motor de CD es dispositivo que transforma energía eléctrica en energía mecánica. Se compone de partes principales un *embobinado de campo* y un *embobinado de armadura*; el primero genera un *campo magnético de densidad de flujo* B , por circulación de una *corriente* i_f ; cuando por el embobinado de la armadura de N *espiras*, circula una *corriente* i_a , sobre cada

espira de *longitud* L y *anchura* b , actúan *fuerzas* F , que la hacen girar a *velocidad angular* ω ; generando un *par* T . En las terminales de la armadura se induce un voltaje denominado *fuerza contraelectromotriz* v_b y se opone al *voltaje aplicado* v_a y es proporcional a ω .

(a) Para el *motor controlado por armadura*, la corriente de campo i_f es constante y el motor se controla ajustando el *voltaje de la armadura* v_a . El circuito de la armadura se puede considerar como una *resistencia* R_a en serie con una *inductancia* L_a y una *fuerza contraelectromotriz* v_b opuesta al voltaje aplicado a la armadura v_a , como se muestra en la **figura 3.19**. Al *par de la armadura* T , se opone el *par de amortiguamiento* T_m , proporcional a la velocidad de la armadura.

(b) Para el *motor controlado por campo*, la corriente de la armadura i_a es constante y el motor se controla ajustando el *voltaje del campo* v_f . El circuito del campo se puede considerar como una *resistencia* R_f en serie con una *inductancia* L_f , como se muestra en la **figura 3.19**. Al *par de la armadura*, se opone el *par de amortiguamiento*, T_m proporcional a la velocidad de la armadura. La **figura 3.20** muestra el diagrama en bloques en función del tiempo, la deducción matemática se muestra a continuación:

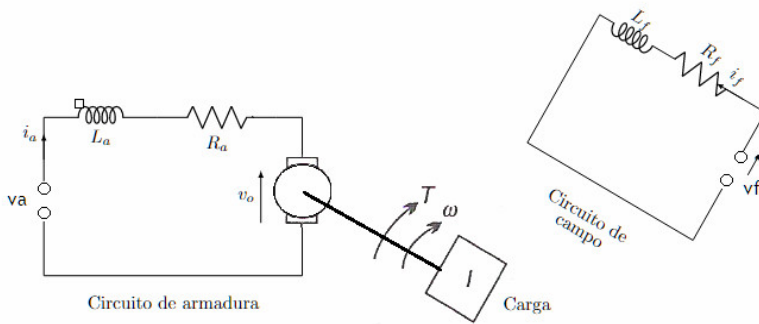


Figura 3.19 Circuitos de la armadura y campo para un motor de CD.

Comportamiento del motor de CD :

$$F = NBi_a L \quad \leftarrow \text{fuerza magnética sobre } N \text{ espiras}$$

$$T = Fb = NBi_a Lb = k_1 Bi_a \quad \leftarrow \text{par sobre } N \text{ espiras de la armadura}$$

$$v_b = k_2 B\omega \quad \leftarrow \text{fuerza contraelectromotriz inducida}$$

Motor CD controlado por armadura : $i_f = \text{constante} \therefore B = \text{constante}$

$$v_b = k_3 \omega$$

$$\sum v = 0 : \quad L_a \frac{di_a}{dt} + R_a i_a = v_a - v_b \rightarrow L_a \frac{di_a}{dt} + R_a i_a = v_a - k_3 \omega$$

$$\sum T = I \frac{d\omega}{dt} : \quad T = k_4 i_a, \quad T_m = c\omega \rightarrow I \frac{d\omega}{dt} = k_4 i_a - c\omega$$

Motor CD controlado por campo : $i_a = \text{constante}$

$$\sum v = 0 : \quad L_f \frac{di_f}{dt} + R_f i_f = v_f$$

$$\sum T = I \frac{d\omega}{dt} : \quad T = k_5 i_f, \quad T_m = c\omega \rightarrow I \frac{d\omega}{dt} = k_5 i_f - c\omega \quad \clubsuit$$

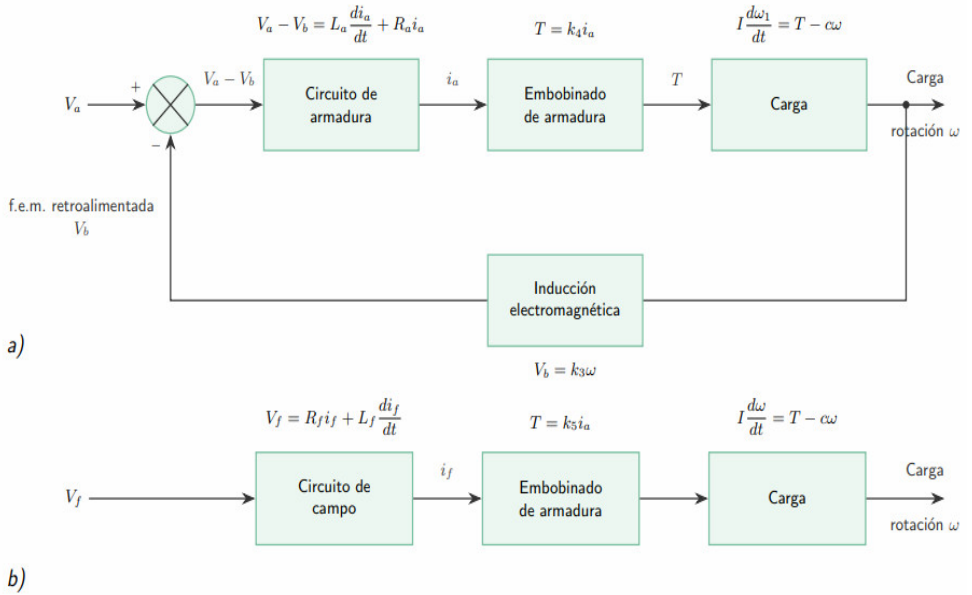


Figura 3.20 Diagrama en bloques en dominio del tiempo, para el motor de CD controlado por: (a) armadura; (b) campo.

Ejemplo 3.18. Para el motor de CD del **ejemplo 3.17**, obtener la función de transferencia por elemento, determinar la función de transferencia global y dibujar el diagrama de bloques en el dominio de s .

Solución. (1) Se obtiene la transformada de Laplace de cada elemento de la **figura 3.20** suponiendo condiciones iniciales iguales a cero; (2) se ordenan los términos convenientemente y se obtiene la función de transferencia global. El diagrama de bloques en el dominio de s , se muestra en la **figura 3.21**; los cálculos se muestran a continuación:

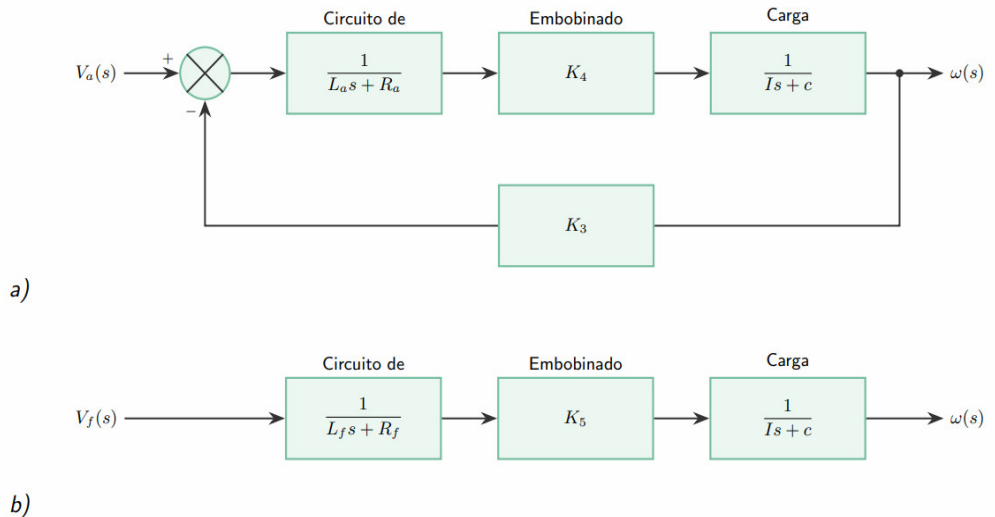


Figura 3.21 Diagrama en bloques en dominio de s , para el motor de CD controlado por: (a) armadura; (b) campo.

(1)

Circuitode la armadura:

$$V_a(s) - V_b(s) = L_a s I_a(s) + R_a I_a(s) = I_a(s) \times (L_a s + R_a) \rightarrow G(s) = \frac{I_a(s)}{V_a(s) - V_b(s)} = \frac{1}{L_a s + R_a}$$

Embobinado de la armadura:

$$G(s) = \frac{T(s)}{i_a(s)} = k_4$$

Carga:

$$I_s \omega(s) = T(s) - c \omega(s) \rightarrow G(s) = \frac{\omega(s)}{T(s)} = \frac{1}{I_s + c}$$

Realimentación:

$$H(s) = \frac{V_b(s)}{\omega(s)} = k_3$$

(2)

Motor CD controlado por armadura:

$$\text{Armadura: } \frac{1}{L_a s + R_a} = \frac{1/R_a}{(L_a/R_a)s + 1} = \frac{1/R_a}{\tau_1 s + 1} \leftrightarrow \tau_1 = \frac{L_a}{R_a}$$

$$\text{Carga: } \frac{1}{I_s + c} = \frac{1/c}{(I/c)s + 1} = \frac{1/c}{\tau_2 s + 1} \leftrightarrow \tau_2 = \frac{I}{c}$$

$$G_{eq}(s) = \frac{(1/R_a)k_4(1/c)/(\tau_1 s + 1)(\tau_2 s + 1)}{1 + k_3(1/R_a)k_4(1/c)/(\tau_1 s + 1)(\tau_2 s + 1)} = \frac{[(1/R_a)k_4(1/c)]/(\tau_1 \tau_2)}{s^2 + [(\tau_1 + \tau_2)/(\tau_1 \tau_2)]s + [k_3(1/R_a)k_4(1/c) + 1]/(\tau_1 \tau_2)}$$

$$G_{eq}(s) = \frac{K}{s^2 + 2\xi\omega_n s + \omega_n^2}$$

$$\xrightarrow{\text{donde}} K = [(1/R_a)k_4(1/c)]/(\tau_1 \tau_2) \leftarrow \text{ganancia de armadura en estado estacionario}$$

$$\omega_n = \sqrt{[k_3(1/R_a)k_4(1/c) + 1]/(\tau_1 \tau_2)} \leftarrow \text{relación de amortiguamiento}$$

$$\xi = (\tau_1 + \tau_2)/[2 \cdot \sqrt{\tau_1 \tau_2 [1 + k_3(1/R_a)k_4(1/c)]}] \leftarrow \text{frecuencia natural no amortiguada}$$

Motor CD controlado por campo:

$$\text{Campo: } \frac{1}{L_f s + R_f} = \frac{1/R_f}{(L_f/R_f)s + 1} = \frac{1/R_f}{\tau_1 s + 1} \leftrightarrow \tau_1 = \frac{L_f}{R_f}$$

$$\text{Carga: } \frac{1}{I_s + c} = \frac{1/c}{(I/c)s + 1} = \frac{1/c}{\tau_2 s + 1} \leftrightarrow \tau_2 = \frac{I}{c}$$

$$G(s) = \frac{\omega(s)}{V_f(s)} = \frac{1/R_f}{\tau_1 s + 1} \times k_5 \times \frac{1/c}{\tau_2 s + 1} = \frac{(1/R_f)k_5(1/c)}{(\tau_1 s + 1)(\tau_2 s + 1)}$$

$$G(s) = \frac{K_f}{(\tau_1 s + 1)(\tau_2 s + 1)}$$

$$\xrightarrow{\text{donde}} K_f = \frac{k_5}{R_f c} \leftarrow \text{ganancia de campo en estado estacionario}$$

♣

Capítulo 4

Sistemas de Control

En este capítulo se presentan los fundamentos de los sistemas de control de procesos industriales.

- ¿Qué es un sistema de control automático?
- Sistema de control lazo abierto y lazo cerrado.
- Tipos de señales de transmisión.
- Componentes de un sistema de control
- Dimensionamiento de válvulas
- Tipos de controladores por realimentación
- Controladores lógicos programables PLCs

Control automático

El **control de sistemas industriales** tiene como finalidad regular el comportamiento de ciertas **variables de proceso** (*temperatura, presión, nivel, flujo*, etc. ver **figura 4.1**), manteniéndolas en un valor deseado o haciéndolas seguir una referencia preestablecida, a pesar de la presencia de *perturbaciones externas* o de variaciones internas del propio sistema. Dicho valor de referencia se denomina punto de consigna o **punto de control**.

Un **sistema de control automático** puede definirse como aquel que compara continuamente el *valor real de la variable controlada* con el valor de referencia y genera las acciones correctivas necesarias para reducir su *diferencia numérica*, sin requerir la intervención directa del operador.

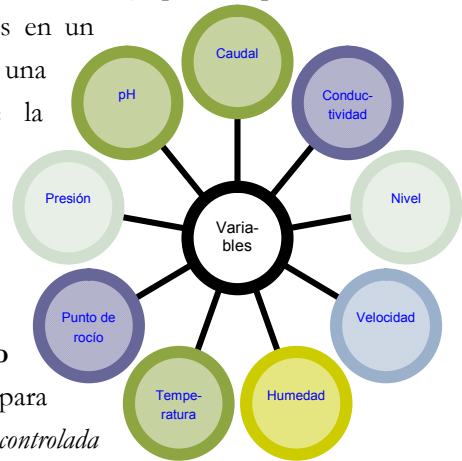


Figura 4.1 Cantidades variables comúnmente controladas

Los sistemas de control se clasifican en dos grupos: en *laço abierto* y en *laço cerrado*. El *sistema de laço abierto* carece de realimentación, el *sistema de laço cerrado* posee realimentación, la señal de realimentación es aquella que sale del controlador con la finalidad de realizar ajustes en la señal de entrada de manera que se establezca según convenga.

Sistema de lazo abierto

Los **sistemas de control en lazo abierto** operan generalmente con base en secuencias de acciones predefinidas, temporizadores o eventos programados, sin que la salida del sistema influya en la entrada. Los

componentes básicos de un sistema de control de lazo abierto son los siguientes (véase **figura 4.2**):

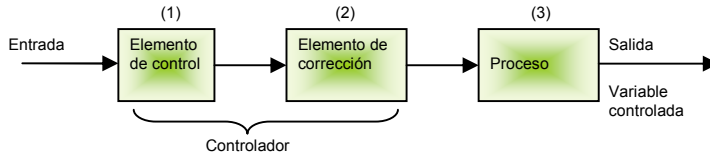


Figura 4.2 Elementos de un sistema de control de lazo abierto

1.- *Elemento de control*. Es el elemento que determina que acción se va a tomar dada la entrada al sistema.

2.- *Elemento de corrección*. Este elemento responde a la entrada que viene del elemento de control e inicia la acción para producir el cambio en la variable controlada al valor requerido.

3.- *Proceso*. Es el lugar donde se controla la variable.

Con los sistemas de control de lazo abierto los tipos de control más conocidos son de dos posiciones (ON/OFF) o secuencias de acciones conmutadas por tiempo.

Sistema de lazo cerrado

En un **sistema de control en lazo cerrado**, una *señal de realimentación* desde la salida se compara continuamente con el *valor de referencia*. La diferencia o variación entre ambos valores se utiliza para ajustar la entrada del sistema, con el propósito de mantener la salida dentro de un *rango de tolerancia* previamente establecido alrededor del valor deseado, aun cuando existan perturbaciones o variaciones en las condiciones de operación.

Los componentes de un sistema de control de lazo cerrado son los mostrados en la **figura 4.3**. La entrada del sistema es el valor requerido de la variable, y la salida es el valor real de la variable.

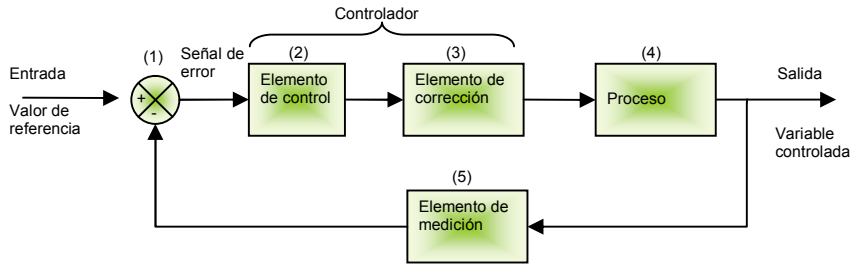


Figura 4.3 Elementos de un sistema de control de lazo cerrado

1.- *Elemento de comparación.* Compara el valor de referencia (*punto de control*) con el valor medido de la variable controlada, y genera una señal de error proporcional a la diferencia entre ambos.

$$\text{error} = (\text{valor de referencia}) - (\text{valor medido})$$

(1)

2.- *Elemento de control.* Decide que acción tomar cuando se recibe una señal de error. A menudo se utiliza el término *controlador* para un elemento que incorpora el elemento de control y la unidad de corrección.

3.- *Elemento de corrección.* Se utiliza para producir un cambio en el proceso al eliminar el error, y con frecuencia se denomina *actuador*.

4.- *Elemento proceso.* El proceso, o planta, es el sistema donde se va a controlar la variable.

5.- *Elemento de medición.* Produce una señal relacionada con la condición de la variable controlada, y proporciona la señal de realimentación al elemento de comparación para determinar si hay o no error.

Todo sistema de control de lazo cerrado posee *realimentación*; que es el medio a través del cual una señal relacionada con la variable real obtenida se realimenta para compararse con la señal de referencia. Se tiene una *realimentación negativa* cuando la señal realimentada se sustrae del valor de referencia:

$$\text{error} = (\text{valor de referencia}) - (\text{valor de realimentación}) \quad (2)$$

La *realimentación positiva* se presenta cuando la señal realimentada se adiciona al valor de referencia:

$$\text{error} = (\text{valor de referencia}) + (\text{valor de realimentación}) \quad (3)$$

Control de un proceso industrial

El bucle o lazo para un proceso industrial está formado por los siguientes elementos (**figura 4.4**): **proceso**, **transmisor**, **controlador** y **válvula de control**.

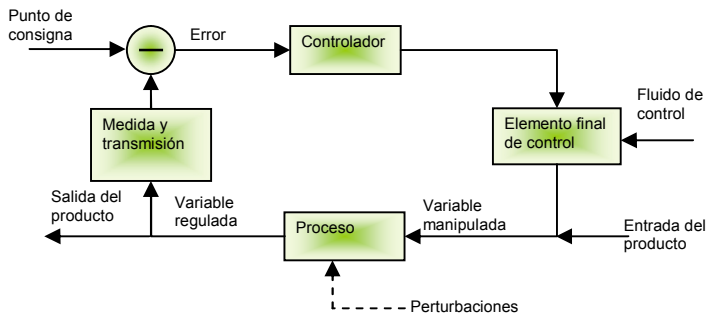


Figura 4.4 Elementos de un sistema de control industrial

El **proceso** consiste en un sistema que ha sido desarrollado para llevar a cabo un objetivo de terminado: tratamiento del material mediante una serie de operaciones específicas destinadas a llevar a cabo una transformación.

El **transmisor** (figura 4.5) es un instrumento que capta la *variable del proceso* (presión, temperatura, nivel, etc.) y la convierte en una *señal estandarizada* (neumática, electrónica o digital) para su transmisión a distancia a un instrumento receptor: *indicador, registrador, controlador* o una combinación de estos.

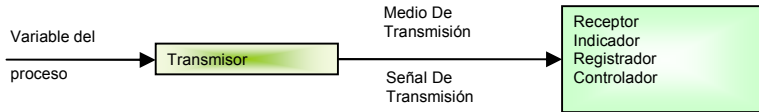


Figura 4.5 Elementos que intervienen en la transmisión

El **controlador** permite al proceso cumplir su objetivo de transformación del material y realiza dos funciones esenciales:

- ♦ Compara la variable medida con la de referencia o deseada (*punto de consigna*) para determinar el error.
- ♦ Estabiliza el funcionamiento dinámico del bucle de control mediante circuitos especiales para reducir o eliminar el error.

Los procesos presentan dos características principales que deben considerarse al automatizarlos:

- ♦ Los cambios en la variable controlada debido a alteraciones en las condiciones del proceso y llamados generalmente cambios de carga.
- ♦ El tiempo necesario para que la variable del proceso alcance un nuevo valor al ocurrir un cambio de carga. Este retardo se debe a una o varias propiedades del proceso: capacitancia, resistencia y tiempo de transporte.

La **válvula** o **elemento final** de control esta en contacto directo con la variable de campo y es la que ejecuta la acción de control de acuerdo a lo indicado por el controlador.

Señales de transmisión

En un sistema de control, la información se transmite entre los distintos elementos que lo componen mediante señales estandarizadas. Estas señales permiten comunicar el valor de las variables medidas, las órdenes del controlador y las acciones dirigidas al elemento final de control. Dependiendo de la aplicación, las señales pueden ser neumáticas, electrónicas, digitales, hidráulicas o telemétricas. La **tabla 4.1** presenta los tipos de señales de transmisión más utilizados en sistemas de control.

Tabla 4.1 Señales de transmisión

No.	Señal	Características
1	Neumática	3–15 PSI, 0.2–1 Bar, 0.21–1 kg/cm ²
2	Electrónica	4–20 mA CD, 0–20mA CD, 1–5 VCD
3	Digital	0 y 1 en grupos de 8, 16 y 32 bits
4	Hidráulica	Para transmisión de grandes potencias
5	Telemétrica	Para transmisión a largas distancias

Unidades de uso común en sistemas de control

En los sistemas de control es frecuente utilizar unidades que no pertenecen al *Sistema Internacional (SI)*, debido a que numerosos equipos, instrumentos y documentos técnicos, especialmente de origen estadounidense, emplean dichas unidades para representar variables de proceso como presión, temperatura y flujo. La **tabla 4.2** lista algunas de las unidades más comunes, junto con su descripción.

Tabla 4.2 Unidades de ingeniería de uso frecuente en sistemas de control

Símbolo	Unidad	Descripción
psig	libra por pulgada cuadrada manométrica	Presión relativa a la presión atmosférica.
psia	libra por pulgada cuadrada absoluta	Presión relativa al vacío absoluto.
Scfh	pie cúbico estándar por hora	Flujo volumétrico de gas en condiciones estándar (60°F, 14.7 psia).
gpm	galón por minuto	Flujo volumétrico de líquido (galón estadounidense).
lbm/hr	libra masa por hora	Flujo másico de gas o vapor.
°F	grado Fahrenheit	Escala de temperatura (uso común en EUA). $^{\circ}\text{F} = 1.8 \times ^{\circ}\text{C} + 32$
°C	grado Celsius	Escala de temperatura del SI. $^{\circ}\text{C} = (^{\circ}\text{F} - 32) / 1.8$
°R	grado Rankine	Escala de temperatura absoluta del sistema inglés. $^{\circ}\text{R} = ^{\circ}\text{F} + 459.67$

Nota: El cero absoluto corresponde a $0\text{ K} = -273.15\text{ }^{\circ}\text{C} = -459.67\text{ }^{\circ}\text{F} = 0\text{ }^{\circ}\text{R}$.

Componentes de un sistema de control

Los cuatro *componentes básicos de los sistemas de control* son los **sensores**, los **transmisores**, los **controladores** y los **elementos finales de control**; tales componentes desempeñan las *tres operaciones básicas de todo sistema de control*: **medición (M)**, **decisión (D)** y **acción (A)**.

Sensores y transmisores

Con los sensores y transmisores se realizan las operaciones de medición en el sistema de control. En el sensor se produce un fenómeno mecánico, eléctrico o similar, el cual se relaciona con la variable de proceso que se mide; el transmisor, a su vez, convierte este fenómeno en una señal que se puede transmitir y, por lo tanto, ésta tiene relación con la variable del proceso.

Existen tres términos importantes que se relacionan con la combinación sensor/transmisor S/T:

- (1) La **escala del instrumento** está definida por los valores inferior y superior de la variable de proceso que se desea medir. Por ejemplo, si un transmisor se calibra para medir una presión entre **20 y 50 psi**, se dice que su escala es de **20 a 50 psi**.
- (2) El **rango del instrumento** es la diferencia entre los valores superior e inferior de la escala. Para el ejemplo anterior, el rango es de **30 psi**. En consecuencia, la escala se define mediante dos valores (inferior y superior), mientras que el rango corresponde a la diferencia entre los dos valores.
- (3) El **valor inferior de la escala** se conoce como **cero del instrumento**. Este término no implica que dicho valor sea necesariamente cero; simplemente corresponde al límite inferior de la escala de medición. En el ejemplo anterior, el cero del instrumento es de **20 psi**.

Para un instrumento la **ganancia** es la relación del rango de salida respecto al rango de entrada:

$$G = \frac{\text{rango de salida}}{\text{rango de entrada}}, \quad \text{Ejemplo: } G = \frac{20\text{mA} - 4\text{mA}}{200\text{psi} - 0\text{psi}} = \frac{16\text{mA}}{200\text{psi}} = 0.08 \frac{\text{mA}}{\text{psi}} \quad (4)$$

Válvulas de control

La selección de una válvula de control debe considerar su comportamiento ante una falla en el suministro de energía, con el fin de garantizar una condición segura del proceso. Con base en esta característica, las válvulas se clasifican en:

- (1) **Abiertas en falla (AF):** *se abre al perderse la energía de accionamiento.*
- (2) **Cerradas en falla (CF):** *se cierra al perderse la energía de accionamiento.*

En la mayoría de las válvulas de control, la energía de accionamiento es aire comprimido (*actuador neumático*), por lo tanto:

- (1) Una válvula **AF** requiere aire para cerrarse $\rightarrow$ se denomina *aire para cerrar (AC)*.
- (2) Una válvula **CF** requiere aire para abrirse $\rightarrow$ se denomina *aire para abrir (AA)*.

Dimensionamiento de la válvula de control

El dimensionamiento de la válvula de control es el procedimiento mediante el cual se calcula el coeficiente de flujo de la válvula (C_v), definido como "*el caudal de agua, en galones estadounidenses por minuto (gpm), que atraviesa una válvula completamente abierta cuando la diferencia de presión entre su entrada y su salida es de 1 psi*".

La ecuación básica para el dimensionamiento de *válvulas de control en aplicaciones con líquidos* es común a la mayoría de los fabricantes:

$$C_v = q \sqrt{\frac{G_f}{\Delta P}} \quad (5)$$

donde:

q = flujo de líquido en gpm U.S.

ΔP = caída de presión, $P_1 - P_2$ en psi en la sección de la válvula

P_1 = presión de entrada a la válvula (corriente arriba) en psi

P_2 = presión de salida de la válvula (corriente abajo) en psi

G_f = gravedad específica del líquido a la temperatura en que fluye (para agua l a 60°F)

Las ecuaciones para el dimensionamiento de *válvulas de control utilizadas con gases y vapor de agua* varían entre fabricantes debido a la forma en que cada uno modela la condición de flujo crítico en fluidos compresibles. *El flujo crítico corresponde a la condición en la que el caudal deja de ser proporcional a la raíz cuadrada de la caída de presión a través de la válvula y pasa a depender únicamente de la presión de entrada.* Este fenómeno ocurre cuando la velocidad del fluido alcanza la velocidad del sonido en la vena contracta; bajo esta condición, las variaciones en la presión de salida dejan de influir en el caudal, el cual queda determinado exclusivamente por la presión de entrada. En consecuencia, se emplean las siguientes ecuaciones:

$$C_v = \frac{Q\sqrt{GT}}{836C_f P_1 \cdot (y - 0.148y^3)} \quad \leftarrow \text{flujo} \cdot \text{volumétrico} \cdot \text{de} \cdot \text{gas}$$

$$C_v = \frac{W}{2.8C_f P_1 \sqrt{G_f} \cdot (y - 0.148y^3)} \quad \leftarrow \text{flujo} \cdot \text{de} \cdot \text{gas} \cdot \text{por} \cdot \text{peso} \quad (6)$$

$$C_v = \frac{W(1 + 0.0007T_{SH})}{1.83C_f P_1 \cdot (y - 0.148y^3)} \quad \leftarrow \text{vapor} \cdot (\text{de} \cdot \text{agua})$$

donde :

Q = tasa de flujo de gas, en scfh bajo condiciones estándar (1.47 psia y 60 °F)

G = gravedad específica del gas bajo condiciones estándar

para gases ideales es $\rightarrow \frac{\text{peso} \cdot \text{molecular} \cdot \text{del} \cdot \text{gas}}{\text{peso} \cdot \text{molecular} \cdot \text{del} \cdot \text{aire}}$.

G_f = gravedad específica del gas a la temperatura del flujo $\rightarrow G_f = G \left(\frac{520}{T} \right)$

T = temperatura en °R

C_f = factor de flujo crítico, (0.6 a 0.95)

P_1 = presión de entrada a la válvula en psia

P_2 = presión de salida a la válvula en psia

$\Delta P = P_1 - P_2$

W = tasa de flujo, en lb / hr

T_{SH} = grados de sobrecalentamiento, en °F

El término y se utiliza para expresar la *condición crítica o subcrítica del flujo* y se define por:

$$y = \frac{1.63}{C_f} \sqrt{\frac{\Delta P}{P_1}} \quad (7)$$

El *dimensionamiento de la válvula*, mediante el *cálculo del coeficiente de flujo* (C_v), debe realizarse de forma que, con la válvula completamente abierta, sea capaz de conducir un caudal mayor que el requerido en condiciones normales de operación. *En consecuencia, es habitual considerar un factor de sobredimensionamiento que permita atender incrementos futuros en la demanda; en la mayoría de las aplicaciones, un factor de 2 constituye un criterio de diseño adecuado.*

$$q_{\text{diseño}} = 2.0 \cdot q_{\text{requerido}} \quad (8)$$

Consideraciones de caída de presión en la válvula

*Como criterio de diseño, la caída de presión en la válvula suele especificarse como el **25 %** de la caída dinámica total de presión del sistema de tuberías, o **10 psi**, el valor que resulte mayor. No obstante, el valor óptimo depende de las condiciones de operación y de los criterios de diseño establecidos por la organización. La*

caída de presión seleccionada tiene efecto en el desempeño de la válvula, como se analizará en la siguiente sección.

 **Ejemplo 4.1.** Dimensionar una válvula que será usada con gas; flujo nominal 25,000 lbm/hr; presión de entrada 250 psi; caída de presión de diseño 100 psi. Gravedad específica del gas 0.4, temperatura de flujo 150°F, peso molecular 12. Usar una válvula de acoplamiento ($C_f=0.92$).

Solución. Sustituyendo los valores en las ecuaciones se tiene que:

$$y = \frac{1.63}{C_f} \sqrt{\frac{\Delta P}{P_1}} = \frac{1.63}{0.92} \sqrt{\frac{100}{250}} = 1.12$$

$$W_{\text{diseño}} = 2W_{\text{nominal}} = 50,000 \frac{\text{lbm}}{\text{hr}}$$

$$C_v = \frac{W}{2.8C_f P_1 \sqrt{G_f \cdot (y - 0.148y^3)}}$$

$$C_v = \frac{5000}{2.8(0.92)(250)\sqrt{0.4 \cdot (1.12 - 0.148(1.12)^3)}} = 134.6$$

Controladores

El **controlador** es el elemento que toma las decisiones (D) en el sistema de control, sus funciones principales son:

- (1). Comparar la señal de la *variable controlada* (proveniente del *transmisor*) con el *punto de control*, generando la *señal de error*.
- (2). Calcular y enviar la *señal de corrección* adecuada al *elemento final de control* (*válvula, actuador, etc.*) para regular la *variable controlada* de manera que permanezca dentro de los límites especificados.

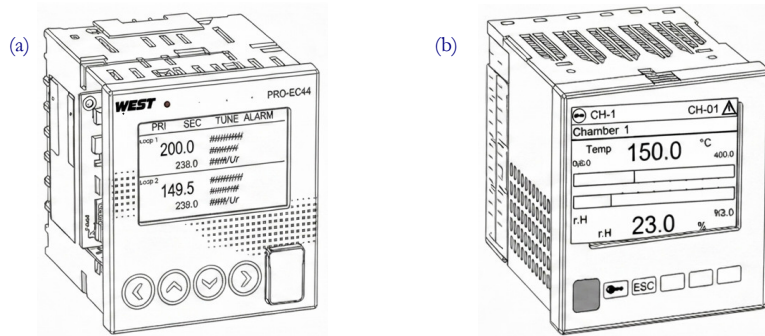


Figura 4.6 Controladores PID: (a) Temperatura/West Pro-EC44, (b) Multipropósito/West KS 98-2

Considérese ahora el circuito de control de nivel que se muestra en la **figura 4.7**. Si el nivel del líquido supera el *punto de ajuste*, el controlador debe abrir la válvula para restablecer el nivel al *punto de control*. Dado que la válvula es de *aire para abrir (AA)*, el controlador debe incrementar su señal de salida (ver las flechas en la figura); por lo tanto, debe configurarse en *acción directa*. Algunos fabricantes denominan esta configuración *acción de incremento*, ya que un aumento en la señal de entrada produce un aumento en la señal de salida del controlador. Para determinar la acción adecuada, el ingeniero de control debe conocer: (1) los requerimientos de control del proceso; (2) la acción de la válvula de control u otro elemento final de control.

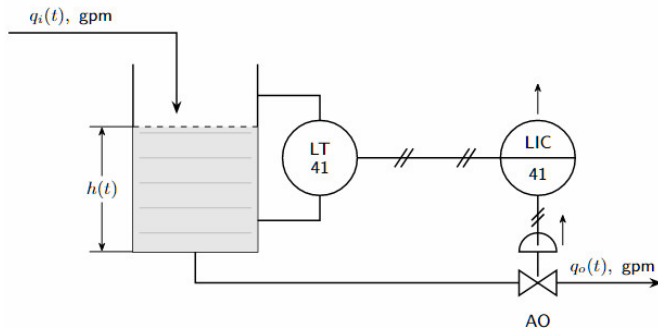


Figura 4.7 Circuito de control para nivel de líquido

Tipos de controladores por realimentación

Control proporcional. El controlador proporcional es el tipo más simple de controlador, con excepción del controlador de dos estados; la ecuación con que se describe su funcionamiento es la siguiente:

$$m(t) = \bar{m} + K_c(r(t)) - c(t), \quad m(t) = \bar{m} + K_c e(t) \quad (9)$$

donde:

$m(t)$ = salida del controlador en psig o mA

$r(t)$ = punto de control psig o mA

$c(t)$ = variable controlada (señal del transmisor), en psig o mA

$e(t)$ = señal de error, psig o mA; diferencia entre el punto de control y la variable controlada

K_c = ganancia del controlador, $\frac{psi}{psi} \cdot o \cdot \frac{mA}{mA}$

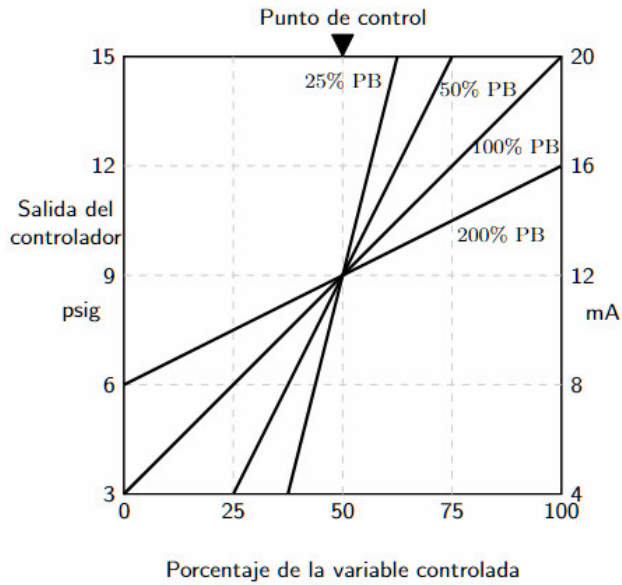
$\bar{m}$ = valor base psig o mA,

$\bar{m}$ es el valor de la salida del controlador cuando el error es cero

La sensibilidad de un controlador es la medida de qué tan intensamente responde su señal de salida ante una variación de la variable controlada respecto al valor de referencia. En un controlador proporcional esta sensibilidad se expresa comúnmente mediante la **ganancia proporcional** (K_c); sin embargo, muchos fabricantes utilizan el concepto de **Banda Proporcional (PB)**, definida como:

$$PB = \frac{100}{K_c} \quad (10)$$

En la **figura 4.8 se ilustra gráficamente la definición de la banda proporcional (PB)**. Una PB del 100 % significa que la salida del controlador recorre todo su rango cuando la variable controlada recorre todo su rango de medición. Una PB menor (por ejemplo, 50 %) implica mayor ganancia y mayor sensibilidad, mientras que una PB mayor (200 %) implica menor ganancia y menor sensibilidad.



	Salida del controlador		
	3 psig 4 mA	9 psig 12 mA	15 psig 20 mA
PB = 100%	0%	50%	100%
PB = 50%	25%	50%	75%
PB = 25%	37.5%	50%	62.5%
PB = 200%	—	50%	—

Figura 4.8 Definición de la banda proporcional BP.

La *función de transferencia* para este controlador es:

$$\frac{M(s)}{E(s)} = K_c \quad (11)$$

En los procesos en los que se permite una desviación dentro de una banda alrededor del punto de control, los controladores proporcionales son suficientes para mantener la variable de proceso dentro del rango especificado. Sin embargo, cuando el proceso requiere mantener la variable exactamente en el punto de control, los controladores proporcionales no proporcionan un control satisfactorio.

Control proporcional integral (PI). En procesos donde se requiere que la variable controlada coincida exactamente con el punto de consigna, el controlador proporcional resulta insuficiente debido al error en estado estacionario que presenta. Para eliminar esta desviación, se añade una acción integral al controlador, dando lugar al **controlador proporcional-integral (PI)**. La acción integral acumula el error a lo largo del tiempo y aplica una corrección hasta que el error se anula.

$$m(t) = \bar{m} + K_c[r(t) - c(t)] + \frac{K_c}{\tau_I} \int [r(t) - c(t)] dt, \quad m(t) = \bar{m} + K_c e(t) + \frac{K_c}{\tau_I} \int e(t) dt \quad (12)$$

Donde τ_I = tiempo de integración o reajuste minutos/repetición. Por lo tanto, el controlador PI tiene dos parámetros, K_C , y τ_I , que se deben ajustar para obtener un control satisfactorio. A continuación se muestran las ecuaciones con que algunos fabricantes describen la operación de sus controladores:

Foxboro·Co.

$$m(t) = \bar{m} + \frac{100}{PB} e(t) + \frac{100}{PB \cdot \tau_I} \int e(t) dt$$

Fisher·Controls

$$m(t) = \bar{m} + \frac{100}{PB} e(t) + \frac{100 \cdot c_R}{PB} \int e(t) dt \quad (13)$$

Taylor·Co., Honeywell·Inc.

$$m(t) = \bar{m} + K_C e(t) + K_C \cdot c_R \int e(t) dt$$

donde: $c_R = \frac{1}{\tau_I} \leftarrow$ repeticiones por minuto,

$\tau_I \leftarrow$ tiempo de integracion en minutos

La *función de transferencia* de este controlador está dada por:

$$\frac{M(s)}{E(s)} = K_c \left(1 + \frac{1}{\tau_i \cdot s} \right) \quad (14)$$

En general, los *controladores proporcional-integral (PI)* tienen dos parámetros de ajuste: la *ganancia* o *banda proporcional*, y el *tiempo de reajuste* o *rapidez de reajuste*. La principal ventaja de este controlador es que la acción integral elimina la desviación permanente.

Controlador proporcional-integral-derivativo (PID). En algunas aplicaciones, se incorpora una acción derivativa al controlador PI, dando lugar al controlador *proporcional-integral-derivativo (PID)*. La acción derivativa, también conocida como *acción de preactuación* o *acción de anticipación*, actúa sobre la tasa de cambio del error (derivada). Esto permite anticipar la tendencia del error y corregirla antes de que se desvíe excesivamente, mejorando así la estabilidad y velocidad de respuesta del sistema. La ecuación que describe este controlador es la siguiente::

$$m(t) = \bar{m} + K_c e(t) + \frac{K_c}{\tau_i} \int e(t) dt + K_c \cdot \tau_D \frac{de(t)}{dt} \quad (15)$$

Donde τ_D es el *tiempo de derivación* en minutos. Por lo tanto, el controlador PID tiene tres parámetros, K , o PB, τ_I o τ_I^R y τ_D , los cuales se deben ajustar para obtener un control satisfactorio. *Nótese que τ_D es el único parámetro asociado a la acción derivativa y se expresa siempre en minutos, independientemente del fabricante.*

Los controladores PID se utilizan en procesos con *constantes de tiempo largas*. Ejemplos típicos son los *procesos de control de temperatura* y *control de concentración*.

La *función de transferencia* de un controlador PID “ideal” es:

$$\frac{M(s)}{E(s)} = K_c \left(1 + \frac{1}{\tau_i \cdot s} \right) \left(\frac{\tau_D \cdot s + 1}{\alpha \cdot \tau_D \cdot s + 1} \right) \quad (16)$$

Los valores típicos de α están entre 0.05 y 0.1. En general, los controladores PID tienen tres parámetros de ajuste: la *ganancia* o *banda proporcional*, el *tiempo de reajuste* o *rapidez de reajuste*, y el *tiempo de derivación*. Este último se expresa siempre en minutos. Los controladores PID se recomiendan para procesos con *constantes de tiempo largas* y *bajo nivel de ruido*. La principal ventaja de la acción derivativa es que anticipa la respuesta del proceso a partir de la tasa de cambio del error.

Controlador proporcional/ derivativo (PD). Este controlador se utiliza en procesos en los que un controlador proporcional es suficiente, pero se requiere incorporar la *acción derivativa* para *anticipar la respuesta del proceso*. La ecuación que lo describe y su función de transferencia se presentan a continuación:

$$m(t) = \bar{m} + K_c e(t) + K_c \tau_D \frac{de(t)}{dt} \quad \frac{M(s)}{E(s)} = K_c (1 + \tau_D \cdot s) \quad (17)$$

Una desventaja del controlador PD es que presenta una *desviación permanente en la variable controlada*. Esta desviación sólo puede eliminarse mediante la *acción integral*. No obstante, la *acción derivativa* permite utilizar una mayor *ganancia proporcional*, reduciendo así la *desviación permanente* respecto a la obtenida con un *controlador proporcional*.

Controladores lógicos programables

Un **Controlador Lógico Programable** es una computadora industrial diseñada para controlar máquinas y procesos. Un PLC interactúa con el entorno por medio de módulos de *entradas y salidas*, el programa de controlador reside en su memoria, la información de sus entradas, es utilizada para actualizar los estados de sus salidas, es decir, un PLC implementa un control

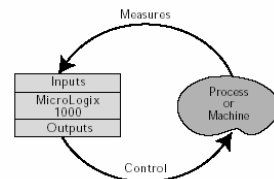


Figura 4.9 Operación del PLC.

de lazo cerrado para operar máquinas y procesos. El proceso conjunto de un PLC mostrado en la **figura 4.9** es muy simple, el PLC mide o sensa señales provenientes de máquinas o procesos y mediante un programa interno procesa esta información, generando acciones de salida que permitan modificar el comportamiento del sistema controlado. Los PLCs proporcionan muchos beneficios respecto a los sistemas de control electromecánicos. Uno de los principales beneficios es la facilidad para modificar la lógica de control mediante cambios en el programa, reduciendo el tiempo y costo asociados a modificaciones del sistema.

Componentes del PLC

Un PLC está conformado por dos componentes principales (**figura 4.10**): (1) el sistema de Entradas Salidas (E/S); (2) la Unidad Central de Procesamiento (CPU).

El sistema de entradas/salidas es la parte del PLC que interactúa físicamente con el mundo externo. La Unidad Central de Procesamiento, es donde el PLC almacena datos y realiza el proceso de cómputo.

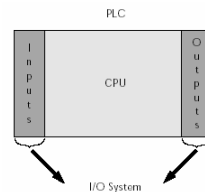


Figura 4.10 Componentes del PLC.

El **Sistema de entradas/salidas** se conforma de dos componentes, la interfaz de entrada y la interfaz de salida (**figura 4.11**):

Interfaz de entrada: La interfaz de entrada es el conjunto de terminales mediante los cuales se conectan dispositivos como botones, sensores e interruptores de límite al PLC. Su función es recibir las señales provenientes de estos dispositivos y

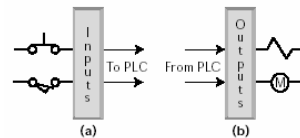


Figura 4.11 Interfaz I/O de PLC.

acondicionarlas para que puedan ser interpretadas por la CPU.

La interfaz de salida: La interfaz de salida es el conjunto de terminales mediante los cuales se conectan dispositivos como arrancadores, solenoides y válvulas al PLC. Su función es convertir las señales generadas por la CPU en señales adecuadas para controlar los dispositivos externos.

La **Unidad Central de Procesamiento** se conforma de los siguientes componentes (figura 4.12): (1) sistema de memoria; (2) procesador; (3) fuente de alimentación.

El **sistema de memoria** almacena el programa de control del PLC, así como los datos recibidos y enviados mediante el sistema de Entradas/Salidas. También conserva la información relacionada con la configuración de los dispositivos conectados a las interfaces de E/SS.

El **procesador** es el componente de la CPU encargado de ejecutar el programa de control. Procesa los datos almacenados en memoria y determina el estado de las salidas en función de la información obtenida desde las entradas.

La **fuerza de alimentación** proporciona la energía necesaria para el funcionamiento de la CPU y del sistema de Entradas/Salidas, permitiendo el funcionamiento del PLC.

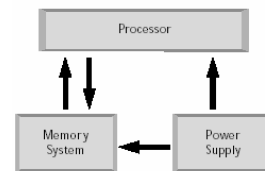


Figura 4.12 CPU del PLC.

Operación del PLC

El PLC ejecuta de forma continua una secuencia de tres pasos conocida como **ciclo de escaneo** (**figura 4.13a**), que consiste en lo siguiente:

- (1) **Lectura de entradas:** el PLC lee el estado de los dispositivos de entrada (sensores, interruptores, etc.) y almacena esta información en su memoria.
- (2) **Ejecución del programa:** el PLC ejecuta el programa de control almacenado en su memoria, procesando las entradas y determinando el estado de las salidas.
- (3) **Actualización de salidas:** el PLC actualiza el estado de los dispositivos de salida (válvulas, motores, etc.) según los resultados obtenidos.

El escaneo puede ser dividido en dos partes, el **escaneo de entradas/salidas** y el **escaneo del programa** (**figura 4.13b**). Durante el escaneo E/S, el PLC lee las entradas y actualiza las salidas. Durante el escaneo de programa, el PLC ejecuta el programa de control.

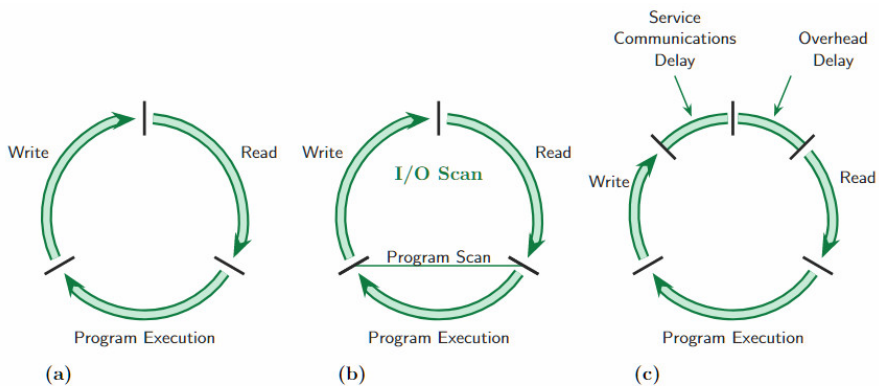


Figura 4.13 Escaneo del PLC.

El **tiempo de escaneo** es el intervalo requerido por el PLC para completar un ciclo completo de operación. Este tiempo incluye la lectura de entradas, la ejecución del programa y la actualización de salidas. En condiciones normales, el tiempo de escaneo es del orden de milisegundos. No obstante, cuando el PLC se encuentra en línea con un dispositivo de programación o monitoreo, pueden presentarse retardos adicionales (**figura 4.13c**), que son los siguientes:

Retardo por servicio de comunicación: tiempo requerido para transferir información entre el PLC y el dispositivo externo.

Retardo por seguridad de transmisión de datos: tiempo destinado a tareas internas del PLC, como la administración de memoria y la actualización de temporizadores.

La norma IEC 1131

La norma **IEC 1131** establece una estandarización para los lenguajes de programación utilizados en PLC, con el objetivo de proporcionar un entorno integrado de programación con múltiples lenguajes. La norma considera como lenguajes válidos los siguientes:

- (1) *Instruction List*, IL
- (2) *Structured Text*, ST
- (3) *Function Block Diagram*, FBD
- (4) *Sequential Function Chart*, SFC
- (5) *Ladder Diagram*, LD

Capítulo 5

Instrumentos de Control

En el presente capítulo se muestran los instrumentos básicos usados en sistemas de control así como las funciones de transferencia que los representan. Se tienen los siguientes objetivos:

- Presentar los instrumentos básicos de control de medición y corrección.
- Presentar la función de transferencia que representa el instrumento de control.
- Presentar los equipos básicos de medición altamente utilizados en centrales eléctricas.

Instrumentos de control industrial

Un *instrumento* es una herramienta o mecanismo para servicio científico y profesional. Los instrumentos básicos de un sistema de control pueden clasificarse en *instrumentos de medición* y *elementos de corrección*. Los **instrumentos de medición** son aquellos que sensan, detectan o miden las variables presentes en los procesos industriales. Los **instrumentos de corrección** son aquellos que ejecutan la acción correctiva dentro del sistema de control.

Instrumentos de medición

Potenciómetro

Un *potenciómetro* con una pista de resistencia uniforme produce una señal de voltaje proporcional al desplazamiento del contacto deslizante desde un extremo de la pista de resistencia del potenciómetro.

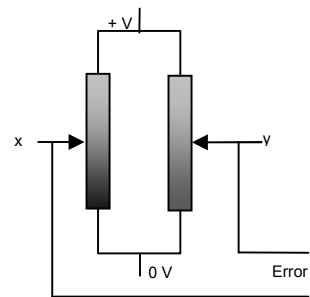


Figura 5.1 Sensor de error de dos potenciómetros

$$\text{Salida [Volts]} = K * (\text{Posición de entrada}) \quad (1)$$

Donde K es la *sensibilidad del potenciómetro*, expresada en *volts por unidad de desplazamiento* (radianes, grados o milímetros, según el caso).

Un *detector de error de dos potenciómetros* se puede construir con el arreglo mostrado en la **figura 5.1**. Un potenciómetro convierte el desplazamiento de entrada en un voltaje y el otro convierte el desplazamiento real en un voltaje. Si x e y son los desplazamientos angulares de cada potenciómetro, la señal de error en voltios está dada por

$$V_e = K * (x - y) \quad (2)$$

Transformador diferencial de variación lineal

Es un detector de posición LVDT (*Linear Variable Differential Transformer*). Proporciona una salida de voltaje alterno cuya amplitud está relacionada con la posición de un núcleo ferroso. La **figura 5.2** muestra el principio básico de funcionamiento. Cuando al devanado primario se aplica una corriente alterna, se inducen voltajes de *CA* en cada uno de los devanados secundarios de acuerdo con la *Ley de Faraday*. Las salidas de los dos devanados están conectadas de tal manera que la salida combinada es la diferencia entre ellas, como se muestra en la **figura 5.3**.

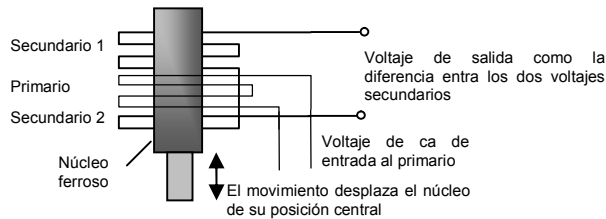


Figura 5.2 Transformador diferencial de variación lineal

A fin de que no se produzca una misma salida para desplazamientos diferentes, se usa un *demodulador con filtro pasa-bajas sensible a la fase* que proporcione un valor único para cada desplazamiento. La función de transferencia se obtiene aplicando las *Leyes de Kirchhoff* a los circuitos:

$$G(s) = \frac{\left[\frac{(M_1 - M_2)}{R_p} \right] \cdot s}{\tau \cdot s + 1} \quad (3)$$

Donde M_1 y M_2 son las inductancias mutuas de los devanados; $M_1 - M_2$ es la cantidad que varía en forma razonablemente lineal con el desplazamiento del núcleo R_p , la resistencia del circuito primario, y τ , es el cociente de la

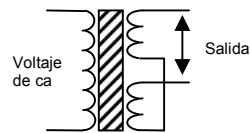


Figura 5.3 Conexión de los devanados

inductancia del primario entre sus resistencia.

Se emplean principalmente en sistemas de control donde se desean monitorear desplazamientos que están en el rango de ± 0.25 mm hasta ± 250 mm.

Extensómetro (galga extensométrica) de resistencia eléctrica

Cuando se estira un alambre metálico o una tira de semiconductor su resistencia cambia. El cambio fraccional en la resistencia es proporcional al cambio fraccional en la longitud, es decir, a la deformación unitaria (ε).

$$\frac{\Delta R}{R} = G \cdot \varepsilon \leftarrow \varepsilon = \frac{\Delta L}{L} \quad (4)$$

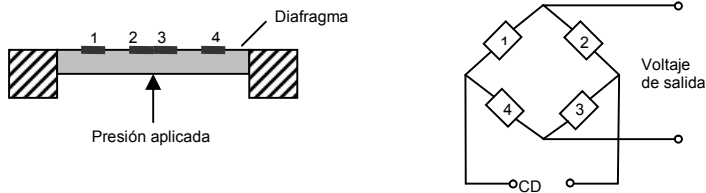


Figura 5.4 Celda de presión de galga extensométrica

Donde G es una constante denominada *factor de galga*. Este factor indica la sensibilidad del material: para metales suele ser ≈ 2 , mientras que para semiconductores puede alcanzar valores cercanos a 100, lo que los hace más sensibles pero también más no lineales. Las *galgas extensiométricas* se basan en este principio y se usan en puentes de *Wheatstone*, donde el voltaje de desbalance es una medida del cambio de resistencia.

La **figura 5.4** muestra como emplear el arreglo de puente de Wheatstone para monitorear la deformación de un diafragma sujeto a cambios de presión.

Tacómetro

Los tacómetros son sensores que producen una señal eléctrica proporcional a la velocidad de rotación de un eje. Existen dos tipos principales de CD y CA.

Tacómetro de CD. Es, en esencia, un generador de corriente continua con un imán permanente que produce un campo magnético dentro del cual gira un devanado (figura 5.5a). Cuando el devanado gira dentro del campo magnético, se induce en él una *fuerza electromotriz* (fem) alterna; cuanto mayor es la velocidad de giro, mayor es la magnitud de la *fem inducida*. Un conmutador, junto con las escobillas, rectifica la tensión alterna inducida para obtener una salida de corriente continua proporcional a la velocidad angular.

Tacómetro de CA. Consta de un rotor conductor y dos devanados dispuestos en ángulo recto (figura 5.5b). Al aplicar una corriente alterna a uno de los devanados, el otro genera una tensión cuya magnitud es proporcional a la velocidad de rotación del rotor.

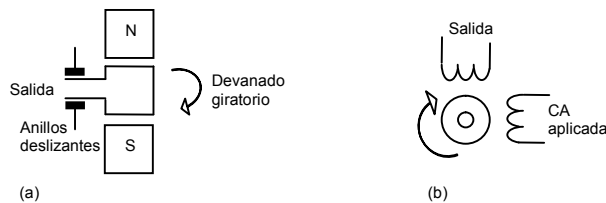


Figura 5.5 (a) Tacómetro de CD. (b) Tacómetro de CA

El voltaje de salida de un tacómetro es proporcional a la razón de cambio en la posición angular θ del eje del medidor, es decir:

$$V_s = K \frac{d\theta}{dt} = k\omega \quad (5)$$

Donde K es una constante y ω , la velocidad angular. De esta forma, la función de transferencia esta dada por:

$$G(s) = \frac{V(s)}{\theta(s)} = ks$$

(6)

Codificador

El término codificador se usa para un dispositivo que proporciona una salida digital como resultado de un desplazamiento angular o lineal. Un **codificador incremental** detecta

cambios en la posición (es decir, cuánto se ha movido desde un punto de referencia), mientras que un **codificador absoluto** proporciona la posición real en todo momento, sin necesidad de un punto de referencia inicial.

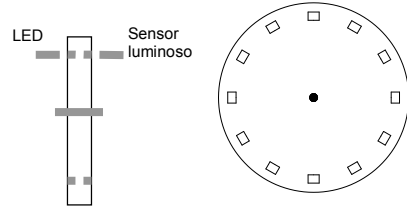


Figura 5.6 Codificador incremental

La **figura 5.6** muestra la forma básica de un decodificador incremental que podría usarse para medir el desplazamiento angular. En un disco se hace incidir un haz luminoso proveniente de un diodo emisor de luz *LED* (*Light Emitting Diode*), que es capaz de pasar a través de las ranuras para detectarlo mediante un sensor luminoso, por ejemplo un *fotodiodo* o un *fototransistor*. Cuando el disco gira, el haz de luz se transmite en forma intermitente proporcionando una salida en forma de pulsos.

El número de pulsos es proporcional al ángulo por el que gira el disco. La resolución es proporcional al número de ranuras del disco. Por ejemplo, con 60 ranuras alrededor del disco el movimiento de una ranura a otra es un giro de 6° . Con el uso de ranuras por compensación, es posible tener más de 1000 ranuras para una sola revolución.

El codificador absoluto difiere del incremental en que solo tiene un patrón de ranuras que sólo define cada posición angular mediante un patrón de señales encendido/apagado producidas.

Ejemplo. La **figura 5.7** describe un decodificador absoluto con 3 conjuntos de ranuras; los decodificadores por lo común tienen hasta 10 o 12 pistas. El

número de bits en la salida binaria resultante es igual al número de pistas, en este caso 3. De manera que el número de posiciones que se pueden detectar será de $2^3=8$, es decir, una resolución de $360^\circ/8=45^\circ$. Con 10 pistas se tendrían 10 bits, esto es $2^{10}=1,024$ posiciones y una resolución angular de $360^\circ/1024=0.35^\circ$.

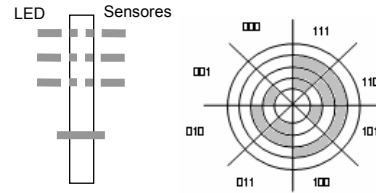


Figura 5.7 Codificador absoluto

Sincros

El sincro es un pequeño transformador giratorio de una sola fase que convierte el desplazamiento angular en voltaje de ca, o voltaje de ca en desplazamiento angular.

La **figura 5.8** muestra el principio básico de un elemento sincro. Consta de 3 devanados de estator espaciados 120°

alrededor de un estator cilíndrico, y de un devanado de rotor concéntrico interior que gira en forma libre entre estos devanados del estator. Una corriente alterna de entrada al devanado del rotor induce salidas en cada devanado secundario. Las salidas de cada uno de los 3 devanados del estator dependerán de la posición angular del rotor, por lo tanto éstos son una medida de la posición angular.

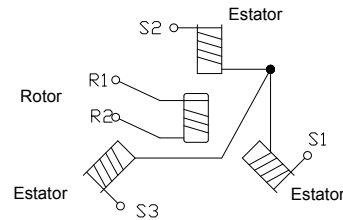


Figura 5.8 Elemento sincro

Con frecuencia los sincros se usan en pares (figura 5.9) como un medio para medir la diferencia entre el desplazamiento angular de dos ejes y así proporcionar una señal de error para un sistema de control.

Cuando los rotores del sincro transmisor y del sincro transformador están alineados en ángulo recto (90°), el voltaje en las terminales del rotor del transformador es cero. Si los ángulos difieren, el voltaje de salida es aproximadamente senoidal respecto a la diferencia angular. Para pequeñas desviaciones ($\leq \pm 5^\circ$), esta relación puede considerarse lineal, lo que permite usar el sincro como un detector de error proporcional.

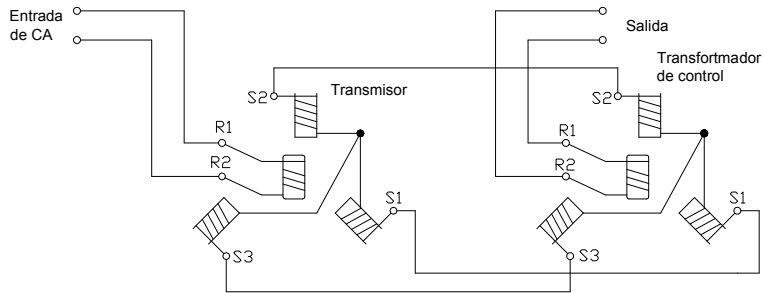


Figura 5.9 Detector de error con sincros

Detector de temperatura por resistencia

El detector de temperatura por resistencia, *RTD*, es un mecanismo muy común para sensar temperatura. *La resistencia eléctrica de los metales o semiconductores cambia con la temperatura.* La resistencia de los metales varía en forma lineal con la temperatura sobre un rango amplio de temperaturas, aunque el cambio real en la resistencia por grado es ligeramente pequeño. Por lo común los RTD se hacen de platino, níquel o aleaciones de níquel.

Los semiconductores, como los termistores, presentan cambios muy grandes en su resistencia con la temperatura, pero su respuesta es fuertemente no lineal, lo que limita su uso en aplicaciones de alta precisión sin compensación.

Los RTD se pueden usar en una rama de un puente de Wheatstone y la salida del puente se toma como medida de la temperatura (figura 5.10a), o se puede usar una fuente de corriente constante que pase a través del RTD y el voltaje se monitorea mediante este (figura 5.10b).

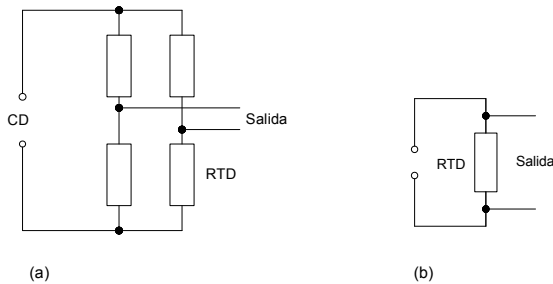


Figura 5.10 Detector de temperatura por resistencia. (a) Con puente; (b) Corriente constante

Sensores de razón de flujo de fluidos

Existen varias maneras de medir flujos, pero la más común se basa en la caída de presión que produce el flujo en una restricción (disminución del área de la sección transversal) ubicada en la trayectoria del fluido. Cuando un fluido que llena un tubo de sección transversal A corre a lo largo de este con una velocidad promedio v , entonces el *gasto, flujo, descarga o razón de flujo volumétrico*, esta dado por: $Q = Av$. Para fluidos incompresibles (densidad constante) se cumple la siguiente relación conocida como *ecuación de continuidad*:

$$Q = A_1 v_1 = A_2 v_2 \quad (7)$$

Siendo el fluido incompresible, y la corriente continua y estacionaria es posible aplicar la *ecuación de Bernoulli*:

$$p_1 + \frac{1}{2} \rho v_1^2 + h_1 \cdot \rho g = p_2 + \frac{1}{2} \rho v_2^2 + h_2 \rho g \quad (8)$$

Donde los subíndices 1 y 2 denotan dos puntos a lo largo de la tubería, ρ , g y h , son la densidad, la aceleración debida a la gravedad y la altura respectivamente. Si consideramos una tubería horizontal $h_1 = h_2$, simplificando y ordenando términos resulta:

$$\frac{v_2^2 - v_1^2}{2} = \frac{p_1 - p_2}{\rho} \quad (9)$$

Empleando la ecuación de continuidad, se tiene que:

$$Q = \frac{A_2}{\sqrt{1 - \left(\frac{A_2}{A_1}\right)^2}} \sqrt{\frac{2(p_2 - p_1)}{\rho}} \quad (10)$$

Puesto que las secciones transversales del instrumento (A_1 , A_2), la densidad del líquido (ρ), son conocidas, la ecuación permite conocer el flujo volumétrico en términos de la presión diferencial ($p_2 - p_1$). Sin embargo esta relación es no lineal, y para efectos prácticos el segundo miembro de la igualdad se ve afectado por una constante C conocida como *coeficiente de descarga*, que compensa las pérdidas de energía como resultado de la fricción.

Tubo Venturi

El tubo Venturi mostrado en la **figura 5.11** se basa en el principio de presión diferencial de la ecuación (13).

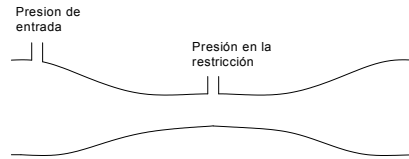


Figura 5.11 Tubo Venturi de presión diferencial

Consta de un tubo con un punto de restricción central. La diferencia de presión entre el flujo antes de la restricción y en esta se puede medir con un *celda de presión de diafragma*; las presiones se aplican en los lados del diafragma, de modo que la deflexión del diafragma es una medida de la diferencia de presión. El instrumento es de operación sencilla, tiene una exactitud de $\pm 0.5\%$ y una confiabilidad a largo plazo.

Medidor de flujo de boquilla

Una forma económica de un tubo Venturi es el *medidor de flujo de boquilla*. Los dos tipos de boquillas que se usan son la boquilla Venturi (**figura 5.12a**) y la boquilla de flujo (**figura 5.12b**). Las boquillas son más baratas que los tubos Venturi, proporcionan diferencias de presión similares y tienen una exactitud de $\pm 0.5\%$.

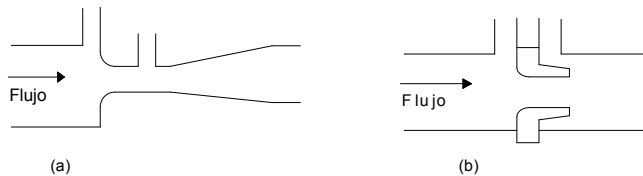


Figura 5.12 Medidor de flujo de boquilla. (a) Boquilla de Venturi; (b) Boquilla de flujo

Tubo Dall

El *tubo Dall* (figura 5.13a) es otra versión del tubo Venturi, con una longitud de casi dos diámetros del tubo donde se monta. Una forma más corta es el *orificio Dall* cuya longitud es de 0.3 veces el diámetro del tubo donde se monta (figura 5.13b).

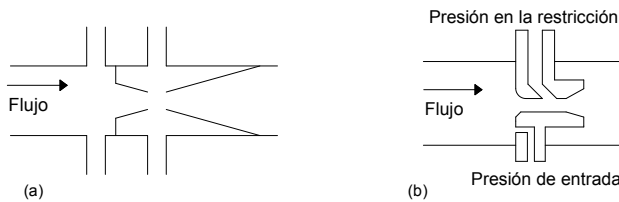


Figura 5.13 Medidor de flujo Dall. (a) Tubo de Dall; (b) Orificio Dall

Placa de orificio

La placa de orificio es un disco con un orificio (figura 5.14). El efecto de introducirla es restringir el flujo a la apertura del orificio y el canal de flujo a una región aún mas angosta corriente abajo (después) del orificio. La diferencia de presión se mide entre un punto de diámetro igual al tubo corriente arriba (antes) del orificio y un punto de diámetro igual a la mitad corriente abajo. La placa de orificio es sencilla y confiable, más económica que el Venturi, pero menos exacta, alrededor de $\pm 1.5\%$.

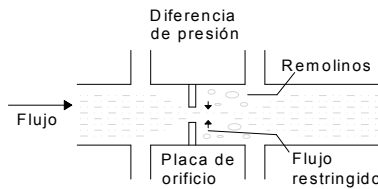


Figura 5.14 Medidor de flujo basado en una placa con orificio

Instrumentos de corrección

Motores de CD

El **motor de corriente directa**, consta de bobinas de alambre montadas en ranuras sobre un cilindro de material ferromagnético. Éste se conoce como *devanado de armadura* y está montado a su vez sobre rodamientos de modo que gira con una libertad en un campo magnético que se produce mediante una corriente que pasa a través de las bobinas de alambre (**figura 5.15**) conocidas como *bobinas de campo* o *devanados de campo*.

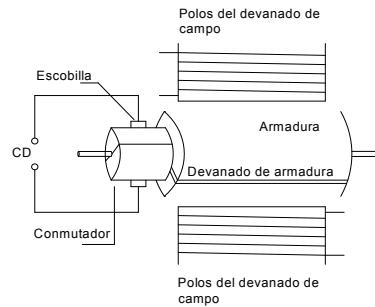


Figura 5.15 Elementos básicos de un motor de CD

Cuando una corriente pasa a través del devanado de la armadura, sobre el devanado actúan fuerzas que hacen que este gire. Para invertir la corriente, cada medio giro se usan escobillas y el conmutador a fin de mantenerlo girando. Con el motor controlado por la armadura, la velocidad se modifica al cambiar la magnitud de la corriente del devanado de armadura. Con frecuencia esto se hace usando fuentes de alimentación de voltaje fijas al devanado y se varía el tiempo para el cual la fuente está encendida. De esta manera se controla el valor promedio del voltaje y, por lo tanto, el de la corriente. Esto se conoce como modulación por ancho de pulso, PWM (*Pulse Width Modulation*).

Otra forma de motor de CD es el **motor de CD sin escobillas**. Este usa un imanes permanente para proveer el campo magnético, y gira dentro de un devanado estacionario. La **figura 5.16** describe el principio básico donde sólo se muestra una

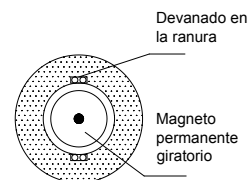


Figura 5.16 Motor de CD sin escobillas

bobina del devanado. El **motor de CD sin escobillas**, es en esencia un motor de CD convencional "*invertido*": el campo magnético es generado por imanes permanentes en el rotor, mientras que los devanados de la armadura están en el estator, eliminando la necesidad de escobillas y conmutador. La velocidad de rotación se puede controlar usando modulación por ancho de pulso. Se usan circuitos electrónicos para invertir la corriente y proveer la conmutación.

Motor de CA

Los **motores de corriente alterna** constan de dos partes básicas: un cilindro giratorio llamado *rotor* y una parte estacionaria denominada *estator*, el cual circunda al rotor y tienen los devanados que producen un campo magnético giratorio en el espacio que ocupa el rotor. Este campo magnético giratorio hace que gire el rotor.

La **figura 5.17** ilustra un motor de inducción monofásico de jaula de ardilla. El rotor consta de barras de cobre o aluminio que se fijan en las ranuras de dos anillos de los extremos para formar un conjunto de conductores paralelos conectados, la llamada jaula de ardilla.

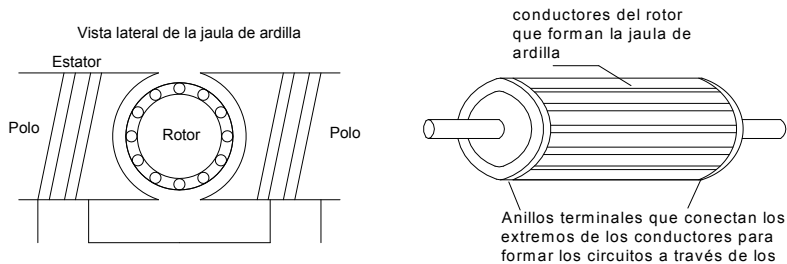


Figura 5.17 Motor de inducción jaula de ardilla

Cuando una corriente alterna pasa a través del devanado de estator, éste produce un campo magnético alterno; se induce una fuerza electromotriz

(fem) en los conductores del rotor y fluye corriente a través de ellos. En estos conductores que transportan corriente actúan fuerzas.

El rotor gira a una velocidad que determina la frecuencia de la corriente alterna aplicada al estator. Una forma de variar la velocidad de rotación es usar un circuito electrónico para controlar la frecuencia de la corriente alimentada al estator.

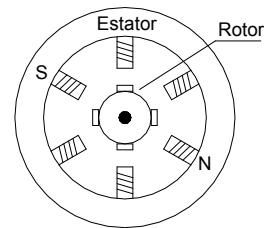
Motor de pasos

El **motor paso a paso** produce una rotación angular igual al ángulo de paso de cada pulso digital aplicado a su entrada que puede usar para posicionamiento angular exacto. Por ejemplo. Un motor de pasos puede dar una rotación de 1.8° por cada pulso de entrada, por lo tanto una entrada de 200 pulsos daría una rotación de 360° .

La **figura 5.18** muestra le principio básico del **motor paso a paso de reluctancia variable**. Su rotor está hecho de acero suave y tiene varios dientes; el número de dientes es menor al número de polos sobre el estator. El estator tiene pares de polos y cada par se energiza mediante la corriente que pasa a través de los devanados que los envuelven.

Cuando un par de polos se activa, se produce un campo magnético que atrae el par de dientes más cercano del rotor, de manera que los polos y dientes se alinean. Así mediante la conmutación de la corriente al siguiente par de polos, se puede lograr que gire el rotor.

Figura 5.18 Motor de pasos de reluctancia variable de 4 dientes y 6 polos



De esta manera el rotor puede girar en pasos mediante la conmutación secuencial de la corriente de un par de polos al siguiente.

La entrada de una secuencia de pulsos debe proporcionar las salidas a cada par del devanado del estator en la secuencia correcta. El sistema de manejo se puede obtener como un circuito integrado.

La **figura 5.19** ilustra el **circuito integrado SAA1027** y sus conexiones que se usa con motores de pasos que tienen cuatro pares de polos en el estator. La entrada para accionar un solo paso de giro del motor de pasos es una transición de voltaje de bajo a alto para la entrada en la terminal 15.

Las salidas del circuito integrado son corrientes en secuencia a lo largo de las conexiones café, negra, verde y amarilla a los devanados del estator. El motor se moverá en el sentido de las manecillas del reloj cuando en la terminal 3 se tiene un voltaje menor a 4.5V, y en el sentido contrario de las manecillas del reloj, cuando es mayor a 7.5V. Cuando la terminal 2 está en bajo, la salida se restaura a su posición normal.

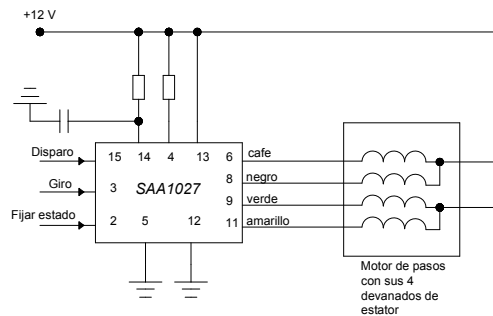


Figura 5.19 Circuito integrado SAA1027 para el motor de pasos

Relevador

En algunos sistemas de control se necesita que la salida del controlador conmute de encendido a apagado una corriente mucho más grande que la que requiere el elemento de corrección final, por ejemplo la corriente que requiere un calefactor eléctrico en un sistema de control de temperatura.

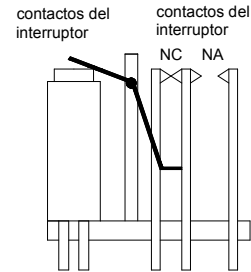


Figura 5.20 Relevador

Esto se puede lograr mediante un relevador. Los relevadores son interruptores que operan en forma eléctrica y en los que al cambiar una corriente en un circuito eléctrico conmutan una corriente en otro circuito.

En el relevador mostrado en la **figura 5.20**, cuando hay una corriente a través del solenoide, se produce un campo magnético que atrae a la armadura de hierro, mueve el rodillo de empuje y, de esta manera, cierra los contactos del interruptor normalmente abierto, NO (*Normally Open*) y abre los contactos del interruptor normalmente cerrado, NC (*Normally Closed*).

Válvula solenoide simple

Un **solenoid** es un núcleo de hierro suave, parcialmente en una bobina. Cuando la corriente pasa a través de la bobina se establece un campo magnético que atrae al núcleo de hierro un poco más hacia el interior de la bobina.

La **figura 5.21** muestra los elementos básicos de una **válvula solenoide simple neumática**. Cuando se aplica corriente al solenoide, la lanzadera se desplaza, redirigiendo el flujo de aire a través de los puertos de la válvula. En reposo (sin corriente), el aire de suministro fluye hacia el puerto de salida. Al activarse el solenoide, el suministro se interrumpe y el puerto de salida se conecta al escape, invirtiendo el estado de la salida neumática.

Cuando ya no se aplica corriente al solenoide, el resorte regresa la válvula a su posición original. De esta manera, en la válvula que se muestra, la corriente al solenoide conmuta la presión de aire al puerto de salida de encendido/apagado.

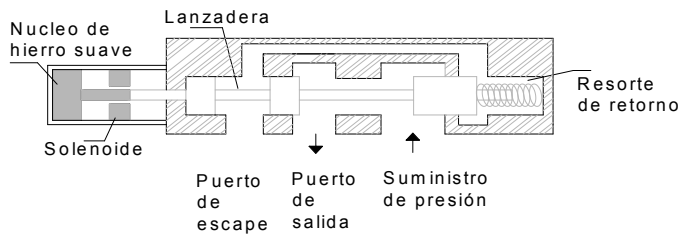


Figura 5.21 Válvula solenoide simple hidráulica

Instrumentos de medición utilizados en procesos industriales

Estos instrumentos permiten **supervisar**, **registrar** y **controlar** condiciones de operación en equipos, máquinas y sistemas industriales. La descripción de cada instrumento considera su función principal, aplicaciones comunes y consideraciones de ingeniería relacionadas con su selección e integración en sistemas de medición y control.

Los instrumentos considerados se agrupan con base en la variable física o eléctrica medida:

(1) Instrumentos que miden cantidades mecánicas o termodinámicas.

Incluyen dispositivos empleados para medir variables como temperatura, presión, flujo, nivel, velocidad y composición de fluidos o gases.

(2) Instrumentos que miden cantidades eléctricas. Incluyen dispositivos empleados para medir variables eléctricas como corriente, tensión, potencia y condiciones de operación de sistemas eléctricos.

Instrumentos que miden cantidades mecánicas o termodinámicas

Medición de temperaturas

1. Termómetros de mercurio con tubo de vidrio. Miden temperatura mediante la expansión térmica del mercurio, con indicación local. Se utilizan en tuberías, depósitos y equipos industriales para verificar condiciones térmicas de operación. **Consideración:** Son adecuados para mediciones locales donde no se requiere transmisión remota de la señal.

2. Termómetros de bulbo lleno de gas. Miden temperatura mediante la expansión de un gas encerrado en un bulbo, transmitiendo mecánicamente la señal al indicador. Se utilizan en procesos industriales donde se requiere

lectura remota de temperatura en gases o líquidos. **Consideración:** Su selección depende del rango de temperatura y la distancia entre el punto de medición y el indicador.

3. Termómetros de vapor. Miden temperatura mediante la expansión de un fluido contenido en un bulbo conectado a un tubo capilar. Se utilizan en aplicaciones industriales de temperatura moderada. **Consideración:** Deben seleccionarse considerando el rango térmico y las condiciones del proceso.

4. Termómetro de resistencia eléctrica. Miden temperatura mediante la variación de la resistencia eléctrica de un sensor (RTD) al cambiar la temperatura. Se utilizan en sistemas industriales donde se requiere precisión y comunicación con sistemas de control. **Consideración:** Son recomendables cuando se requiere integración con PLC, registradores o sistemas SCADA.

5. Termopares o pirómetros. Miden temperaturas elevadas mediante el efecto termoeléctrico (voltaje generado por diferencia de temperatura entre dos metales). Se utilizan en hornos, calderas, turbinas y sistemas de combustión. **Consideración:** Su selección depende del rango de temperatura, materiales del proceso y condiciones ambientales.

Medición de presión

6. Manómetro de tubo Bourdon. Miden presión mediante la deformación de un tubo curvado (tubo Bourdon) que se endereza al aplicar presión. Se utilizan en sistemas hidráulicos, neumáticos, líneas de vapor y recipientes presurizados, como cilindros de oxígeno, acetileno, argón, nitrógeno y otros gases industriales. **Consideración:** Deben seleccionarse considerando el rango de presión, precisión requerida y compatibilidad con el fluido medido.

7. Manómetros de tubo helicoidal o diafragma. Miden presiones bajas o moderadas mediante elementos sensibles como tubos helicoidales o

diafragmas. **Consideración:** Se utilizan donde se requiere mayor sensibilidad que la ofrecida por el tubo Bourdon.

8. Vacuómetros y manómetros combinados. Miden presiones por debajo y por encima de la atmosférica (vacío y presión positiva). Se utilizan en condensadores, bombas y sistemas de vacío. **Consideración:** Su selección depende del rango de presión negativa o positiva requerido.

9. Manómetros de tiro. Miden pequeñas diferencias de presión (tiro) en sistemas de combustión y ventilación. Se utilizan para supervisar hornos, calderas y extractores. **Consideración:** Permiten verificar condiciones adecuadas de operación y eficiencia del sistema.

Medición de flujo o gasto

10. Medidores de vapor. Miden el flujo másico o volumétrico de vapor en procesos industriales. Se utilizan en sistemas de generación, distribución de vapor y equipos térmicos. **Consideración:** Deben seleccionarse considerando presión, temperatura y rango de flujo esperado.

11. Medidores de agua. Miden el caudal volumétrico de agua en tuberías y sistemas hidráulicos. Se utilizan en circuitos de alimentación, enfriamiento y procesos industriales. **Consideración:** Su selección depende del tipo de fluido, precisión requerida y condiciones de instalación.

12. Medidores de aire. Miden el flujo volumétrico de aire en sistemas de ventilación, combustión y procesos neumáticos. **Consideración:** En muchas aplicaciones utilizan medición diferencial de presión para determinar el flujo.

Registradores de nivel

13. Indicadores y registradores de nivel. Miden o registran el nivel de líquidos (agua, aceite, químicos) o sólidos (polvo de carbón, graneles) en

tanques, calderas y silos. **Consideración:** Su selección depende del tipo de material, rango de medición y condiciones del recipiente.

Medición de Velocidad

14. Tacómetros.

Miden la velocidad angular (de rotación) instantánea de máquinas y equipos mecánicos, generalmente expresada en revoluciones por minuto (rpm). Se utilizan en motores, turbinas, generadores, bombas y otros sistemas rotativos industriales para supervisar condiciones de operación.

Existen diferentes tipos de tacómetros: *tacómetro de lengüeta vibratoria*, *tacómetro eléctrico*, *tacómetro de tipo reloj*, *tacómetro centrífugo* y *estroboscopio*.

Consideración: Su selección depende del rango de velocidad, precisión requerida, condiciones de operación y método de transmisión de la señal hacia sistemas de supervisión o control

15. Contadores de revoluciones. Registran el número total de vueltas realizadas por un eje o mecanismo durante un periodo determinado. Se utilizan en equipos rotativos para seguimiento de ciclos de operación, mantenimiento y evaluación de desgaste. **Consideración:** Son útiles cuando se requiere conocer la cantidad acumulada de ciclos o revoluciones, independientemente de la velocidad instantánea del equipo.

Medición de combustible

16. Básculas y medidores de carbón. Miden la cantidad de carbón utilizada en procesos industriales mediante sistemas de pesaje o medición de volumen. Se utilizan en bandas transportadoras, alimentadores de carbón, tolvas, silos y sistemas de alimentación de calderas u hornos industriales.

Consideración: Deben seleccionarse según el método de alimentación, capacidad requerida y precisión de medición.

17. Contadores de gas. Miden el consumo o flujo volumétrico de gases en procesos industriales. Pueden ser de desplazamiento positivo o de carga diferencial. Se aplican en sistemas de combustión, distribución y generación de energía. **Consideración:** Deben seleccionarse considerando presión, composición del gas y rango de operación.

18. Medidores de aceite. Miden el consumo de combustibles líquidos (aceite) mediante desplazamiento positivo. Se utilizan en motores, quemadores y procesos térmicos. **Consideración:** Su selección depende de viscosidad del fluido y condiciones de operación.

Análisis de gases

19. Analizadores de Gases. Miden la composición de gases mediante diferentes métodos de análisis. Se utilizan en procesos de combustión, control ambiental y supervisión de eficiencia industrial. **Consideración:** Deben seleccionarse según el gas a medir, precisión requerida y condiciones del proceso.

20. Analizadores de humo y gases de combustión. Miden la concentración de partículas y componentes presentes en los gases de escape, como dióxido de carbono (CO_2) y oxígeno (O_2). Se utilizan para evaluar la eficiencia de combustión, controlar emisiones y supervisar condiciones de operación en equipos térmicos. **Consideración:** Los instrumentos modernos emplean sensores electrónicos, como celdas fotoeléctricas, y deben seleccionarse considerando el rango de medición y las características del gas analizado.

21. Analizadores de CO_2 o O_2 . Determinan la concentración de oxígeno y dióxido de carbono en gases de proceso. Se utilizan principalmente en

sistemas de combustión, calderas, control de emisiones y supervisión de eficiencia energética. **Consideración:** Pueden utilizar diferentes métodos de medición, como principios químicos, eléctricos o mecánicos, y deben seleccionarse según el gas analizado, precisión requerida y condiciones del proceso.

Calorímetros

22. Calorímetros para vapor y combustible. Miden propiedades relacionadas con la calidad energética del vapor o combustible. Se utilizan principalmente en pruebas de rendimiento y evaluación de eficiencia. **Consideración:** Generalmente se emplean en análisis específicos y no como instrumentos de supervisión continua.

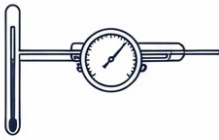
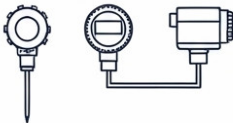
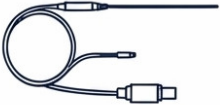
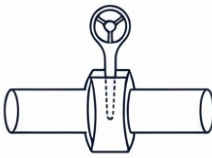
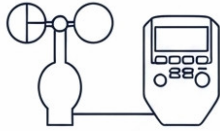
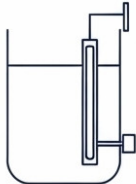
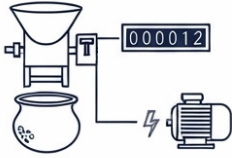
Medición atmosférica

23. Barómetros, higrómetros y termómetros ambientales. Miden condiciones ambientales: presión atmosférica (barómetros), humedad (higrómetros) y temperatura ambiente (termómetros). Se utilizan para monitoreo ambiental y corrección de condiciones de proceso. **Consideración:** Su aplicación depende de la precisión requerida y condiciones del entorno.

Alarmas industriales

24. Alarmas sonoras y anunciadores, Generan señales de advertencia (auditivas o visuales) al detectar condiciones anormales: temperaturas elevadas, niveles críticos, fallas eléctricas, etc. **Consideración:** Deben integrarse con sistemas de protección y supervisión del proceso.

INSTRUMENTOS DE MEDICIÓN MECÁNICA Y TERMODINÁMICA

1. TERMÓMETRO DE MERCURIO  Indicación Local	2. TERMÓMETRO DE BULBO LLENO  Transmisión Mecánica	3. TERMÓMETRO RTD (PT100)  Señal Eléctrica
4. TERMOPAR  Efecto Termoeléctrico	5. MANÓMETRO DE TUBO BOURDON ACETILENO  MANÓMETRO	6. MANÓMETRO DE DIAFRAGMA  Mayor Sensibilidad (Presión Baja)
7. VACUÓMETRO  Presión Negativa	8. MEDIDOR DE VAPOR / AGUA  Flujo Volumétrico	9. MEDIDOR DE AIRE (ANEMÓMETRO DE COPA)  Velocidad del Aire
10. INDICADOR DE NIVEL  Nivel de Líquidos/Sólidos	11. TACÓMETRO (ESTROBOSCOPIO)  Velocidad Angular (RPM)	12. CONTADOR DE REVOLUCIONES  RPM

FSD-FMTC

Instrumentos que miden cantidades eléctricas

Medición de tensión

1. **Voltímetros.** Miden la diferencia de potencial eléctrico en circuitos y equipos eléctricos. Se utilizan para verificar niveles de tensión, regulación de voltaje y condiciones de operación. **Consideración:** Su rango de medición y aislamiento deben ser adecuados para la tensión del sistema.

Medición de corriente

2. **Amperímetros.** Miden la corriente eléctrica en circuitos de alimentación, generadores, motores y sistemas auxiliares. Se utilizan para supervisar carga eléctrica y detectar condiciones anormales de operación. **Consideración:** Deben seleccionarse de acuerdo con el rango de corriente y el sistema de medición utilizado.

Medición de potencia

3. **Vatímetros.** Miden la potencia activa eléctrica consumida o generada por un sistema. Se utilizan en circuitos de alimentación, generadores, motores y equipos industriales para supervisar el consumo energético. **Consideración:** Deben seleccionarse considerando el nivel de potencia, tipo de alimentación y compatibilidad con sistemas de monitoreo.

Medición del factor de potencia

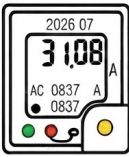
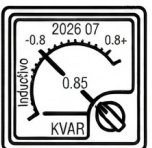
4. **Indicadores de factor de potencia.** Miden el factor de potencia de sistemas eléctricos. Se utilizan en generadores, tableros y procesos industriales para evaluar la eficiencia del consumo eléctrico. **Consideración:** Permiten identificar la necesidad de compensación de potencia reactiva.

Sincronización de sistemas eléctricos

5. Sincronoscopios. Indican las condiciones de sincronización (*frecuencia, ángulo de fase y voltaje*) entre fuentes eléctricas antes de conectarlas en paralelo. Se utilizan principalmente en generadores y sistemas de distribución eléctrica. **Consideración:** Son necesarios para garantizar una conexión segura entre fuentes de energía.

Protección y diagnóstico eléctrico

6. Detectores de tierra. Identifican fallas de aislamiento o conexiones no deseadas a tierra en sistemas eléctricos. Se utilizan en instalaciones industriales para protección y diagnóstico. **Consideración:** Su aplicación depende del esquema de puesta a tierra y protección eléctrica utilizado.

INSTRUMENTOS DE MEDICIÓN ELÉCTRICA		
1. VOLTÍMETRO AC	2. AMPERÍMETRO AC	3. VATÍMETRO
		
4. INDICADOR DE FACTOR DE POTENCIA	5. SINCRONOSCOPIO	6. DETECTOR DE FALLA A TIERRA
		
FSD-FMTC		

Anexos

Ecuaciones y Leyes Matemáticas y Físicas

FSD-FMTC

*Este anexo fue desarrollado en $\text{\LaTeX}$ e incorporado en la
Versión 2.0, Release 2 (v2.0-r2c1), como parte de la
actualización y complementación de esta obra (2026).*

Propósito

Estos anexos están diseñados para que todo lector, desde el nivel preuniversitario hasta el nivel avanzado, pueda consultar rápidamente la notación, definiciones, fórmulas y leyes utilizadas a lo largo de la obra. No sustituyen el desarrollo conceptual de los capítulos, pero permiten seguir el hilo del libro sin recurrir a materiales externos.

Índice de anexos

Anexo	Contenido
A	Notación matemática, lógica, conjuntos, funciones, tiempo
B	Axiomas y principios matemáticos (números reales)
C	Álgebra: operaciones elementales, productos notables, factorización, exponentes, radicales, ecuaciones, logaritmos, límites, álgebra booleana
D	Geometría plana, del espacio y analítica: fórmulas de perímetros, áreas, volúmenes, teoremas y geometría analítica
E	Trigonometría: razones, identidades, leyes de senos y cosenos
F	Vectores: operaciones, propiedades y aplicaciones
G	Números complejos: formas, operaciones, Euler, De Moivre
H	Matrices y sistemas de ecuaciones lineales: operaciones, determinantes, inversa, rango, métodos de solución (Gauss, Cramer, inversa), teorema de Rouché-Frobenius
I	Leyes físicas fundamentales: mecánica, oscilaciones, electricidad, fluidos, termodinámica
J	Fundamentos de cálculo: derivada, integral, teorema fundamental, ecuaciones diferenciales, transformada de Laplace
K	Conceptos esenciales de control: función de transferencia, bloques, PID, estabilidad, evolución temporal de sistemas
L	Ecuaciones de Maxwell: las cuatro ecuaciones, forma vectorial (diferencial), ecuación de onda

Anexo A

NOTACIÓN MATEMÁTICA, LÓGICA Y RELACIONAL

A.1. Alfabeto griego

Símbolo	Nombre	Símbolo	Nombre
1. $A \alpha$	alfa	14. $N \nu$	nu
2. $B \beta$	beta	15. $\Xi \xi$	xi
3. $\Gamma \gamma$	gamma	16. $O o$	ómicron
4. $\Delta \delta$	delta	17. $\Pi \pi$	pi
5. $E \epsilon$	épsilon	18. $P \rho$	rho
6. $Z \zeta$	zeta	19. $\Sigma \sigma$	sigma
7. $H \eta$	eta	20. $T \tau$	tau
8. $\Theta \theta$	theta	21. $\Upsilon \upsilon$	ípsilon
9. $I \iota$	iota	22. $\Phi \phi$	fi
10. $K \kappa$	kappa	23. $X \chi$	ji
11. $\Lambda \lambda$	lambda	24. $\Psi \psi$	psi
12. $M \mu$	mu	25. $\Omega \omega$	omega
13. $\Xi \xi$	xi		

A.2. Símbolos lógicos fundamentales

Símbolo	Nombre	Significado	Ejemplo
$\wedge$	Conjunción	“y” $\rightarrow$ verdadero si ambas proposiciones lo son	$p \wedge q$
$\vee$	Disyunción	“o” $\rightarrow$ verdadero si al menos una lo es	$p \vee q$
$\neg$	Negación	“no” $\rightarrow$ invierte el valor de verdad	$\neg p$
$\implies$	Implicación	“si ... entonces ...”	$p \implies q$
$\iff$	Doble implicación	“si y solo si” $\rightarrow$ equivalencia lógica	$p \iff q$
$\forall$	Cuantificador universal	“para todo”	$\forall x \in \mathbb{R}$
$\exists$	Cuantificador existencial	“existe al menos uno”	$\exists x \in \mathbb{N}$
$\exists!$	Existencia única	“existe exactamente uno”	$\exists! x$
$\therefore$	Por lo tanto	conclusión	...
$\because$	Porque	justificación	...

A.3. Conjuntos y operaciones entre conjuntos

Símbolo	Nombre	Significado
$\{\}$	Conjunto	Colección de elementos
$\in$	Pertenencia	“pertenecer a”
$\notin$	No pertenencia	“no pertenece a”
$\subset$	Subconjunto propio	$A \subset B$: todos los elementos de A están en B , pero $A \neq B$
$\subseteq$	Subconjunto	$A \subseteq B$: A está contenido en B o es igual
$\supset$	Superconjunto propio	$A \supset B$: A contiene a B pero no son iguales
$\supseteq$	Superconjunto	$A \supseteq B$: A contiene a B o es igual
$\cup$	Unión	$A \cup B$: elementos que están en A , en B , o en ambos
$\cap$	Intersección	$A \cap B$: elementos que están simultáneamente en A y en B
$\setminus$	Diferencia	$A \setminus B$: elementos que están en A pero no en B
$\emptyset$	Conjunto vacío	Conjunto sin elementos
$\mathbb{N}$	Números naturales	$\{1, 2, 3, \dots\}$ o $\{0, 1, 2, \dots\}$
$\mathbb{Z}$	Números enteros	$\{\dots, -2, -1, 0, 1, 2, \dots\}$
$\mathbb{Q}$	Números racionales	Cocientes de enteros: p/q con $q \neq 0$
$\mathbb{R}$	Números reales	Todos los números en la recta numérica
$\mathbb{C}$	Números complejos	$a + jb$ con $a, b \in \mathbb{R}$
$ A $	Cardinalidad	Número de elementos del conjunto A
$\times$	Producto cartesiano	$A \times B = \{(a, b) \mid a \in A, b \in B\}$

A.4. Relaciones de orden y comparación

Símbolo	Significado
$=$	Igual a
$\neq$	Diferente de
$<$	Menor que
$\leq$	Menor o igual que
$>$	Mayor que
$\geq$	Mayor o igual que
$\approx$	Aproximadamente igual
$\propto$	Directamente proporcional a
$\equiv$	Idéntico o equivalente por definición
$\cong$	Congruente (en geometría)
$\sim$	Semejante (en geometría)
$\perp$	Perpendicular
$\parallel$	Paralelo

A.5. Operadores matemáticos comunes

Símbolo	Nombre	Significado
$+, -, \times, \div$	Operaciones aritméticas	Suma, resta, multiplicación, división
$\sqrt{x}$	Raíz cuadrada	$\sqrt{x}$ es el número no negativo que al cuadrado da x
$\sqrt[n]{x}$	Raíz n-ésima	$\sqrt[n]{x}$ es el número que elevado a n da x
$ x $	Valor absoluto	Distancia de x al origen: $ x = \begin{cases} x & \text{si } x \geq 0 \\ -x & \text{si } x < 0 \end{cases}$
$\sum$	Sumatoria	$\sum_{i=1}^n a_i = a_1 + a_2 + \cdots + a_n$
$\prod$	Productoria	$\prod_{i=1}^n a_i = a_1 \cdot a_2 \cdot \cdots \cdot a_n$
$!$	Factorial	$n! = n \cdot (n-1) \cdot \cdots \cdot 2 \cdot 1$
Δ	Incremento	Variación finita: $\Delta x = x_2 - x_1$
∂	Derivada parcial	Derivada con respecto a una variable manteniendo fijas las demás
d	Diferencial	Cantidad infinitesimal
$\int$	Integral	Operación inversa a la derivada; área bajo la curva
$\oint$	Integral cerrada	Integral sobre un camino cerrado
∞	Infinito	Cantidad sin límite finito

A.6. Notación de funciones y mapeos

Símbolo	Significado
$f : A \rightarrow B$	f es una función que va del conjunto A al conjunto B
$x \mapsto f(x)$	El elemento x se mapea al valor $f(x)$
f^{-1}	Función inversa de f
$f \circ g$	Composición: $(f \circ g)(x) = f(g(x))$
$f(x)$	Valor de la función f evaluada en x
$\text{dom}(f)$	Dominio de f : valores donde está definida
$\text{ran}(f)$	Rango o imagen de f : valores que toma

A.7. Notación para variables dependientes del tiempo y señales

Símbolo	Significado
$x(t)$	Variable que depende del tiempo t
$\dot{x}$	Primera derivada respecto al tiempo (notación de Newton)
$\ddot{x}$	Segunda derivada respecto al tiempo
$x^{(n)}$	Derivada de orden n respecto al tiempo
x_0	Condición inicial (valor en $t = 0$)
x_e	Estado de equilibrio (valor constante en reposo)
$\mathcal{L}\{f(t)\}$	Transformada de Laplace de $f(t)$
$\mathcal{L}^{-1}\{F(s)\}$	Transformada inversa de Laplace
s	Variable compleja de Laplace: $s = \sigma + j\omega$
j	Unidad imaginaria ($j^2 = -1$)
ω	Frecuencia angular (rad/s)
f	Frecuencia (Hz), $f = \omega/(2\pi)$
$\mathcal{F}\{f(t)\}$	Transformada de Fourier de $f(t)$
$u(t)$	Señal de entrada o escalón unitario de Heaviside
$y(t)$	Señal de salida
$r(t)$	Señal de referencia (consigna)
$e(t)$	Señal de error: $e(t) = r(t) - y(t)$
$U(s)$	Transformada de Laplace de la entrada $u(t)$
$Y(s)$	Transformada de Laplace de la salida $y(t)$
$G(s)$	Función de transferencia: $G(s) = \frac{Y(s)}{U(s)}$
$H(s)$	Función de transferencia del sensor o de retroalimentación

Anexo B

PRINCIPIOS Y AXIOMAS MATEMÁTICOS FUNDAMENTALES

B.1. Propiedades de la igualdad

- **Reflexiva:** $a = a$ (todo número es igual a sí mismo).
- **Simétrica:** si $a = b$, entonces $b = a$.
- **Transitiva:** si $a = b$ y $b = c$, entonces $a = c$.
- **Sustitución:** si $a = b$, entonces cualquier expresión que contenga a puede reemplazar a a por b y viceversa.

B.2. Propiedades de los números reales

Propiedad	Expresión
Conmutatividad de la suma	$a + b = b + a$
Conmutatividad del producto	$a \cdot b = b \cdot a$
Asociatividad de la suma	$(a + b) + c = a + (b + c)$
Asociatividad del producto	$(a \cdot b) \cdot c = a \cdot (b \cdot c)$
Distributividad	$a \cdot (b + c) = a \cdot b + a \cdot c$
Elemento neutro de la suma	$a + 0 = a$
Elemento neutro del producto	$a \cdot 1 = a$
Inverso aditivo	$a + (-a) = 0$
Inverso multiplicativo	$a \cdot a^{-1} = 1$ (si $a \neq 0$)

B.3. Propiedades de orden

- **Ley de tricotomía:** para dos números reales a y b , exactamente una de las siguientes es verdadera: $a < b$, $a = b$, $a > b$.
- **Transitividad:** si $a < b$ y $b < c$, entonces $a < c$.
- **Monotonía de la suma:** si $a < b$, entonces $a + c < b + c$ para cualquier c .
- **Monotonía del producto:** si $a < b$ y $c > 0$, entonces $a \cdot c < b \cdot c$. Si $c < 0$, la desigualdad se invierte: $a \cdot c > b \cdot c$.

Anexo C

ÁLGEBRA

C.I. Operaciones elementales

Reglas de los signos

- **Suma y resta:** signos iguales se suman, signos diferentes se restan y se coloca el signo del mayor.
- **Multiplicación:** $(+) \cdot (+) = +$, $(+) \cdot (-) = -$, $(-) \cdot (-) = +$.
- **División:** $\frac{(+)}{(+)} = +$, $\frac{(+)}{(-)} = -$, $\frac{(-)}{(-)} = +$.
- **Potencia par:** $(-a)^n = a^n$ si n es par.
- **Potencia impar:** $(-a)^n = -a^n$ si n es impar.

Jerarquía de las operaciones (orden de ejecución)

1. **Paréntesis:** se resuelven primero los paréntesis, corchetes y llaves (de adentro hacia afuera).
2. **Potencias y raíces:** se resuelven antes que multiplicaciones y divisiones.
3. **Multiplicaciones y divisiones:** se resuelven de izquierda a derecha.
4. **Sumas y restas:** se resuelven de izquierda a derecha.

Recordatorio: “Paréntesis, Potencias, Multiplicación/División, Suma/Resta” (PPMD/SR).

Operaciones con fracciones

- **Suma y resta** (mismo denominador): $\frac{a}{c} \pm \frac{b}{c} = \frac{a \pm b}{c}$.
- **Suma y resta** (distinto denominador): $\frac{a}{b} \pm \frac{c}{d} = \frac{ad \pm bc}{bd}$.
- **Multiplicación:** $\frac{a}{b} \cdot \frac{c}{d} = \frac{a \cdot c}{b \cdot d}$.
- **División:** $\frac{a}{b} \div \frac{c}{d} = \frac{a}{b} \cdot \frac{d}{c} = \frac{a \cdot d}{b \cdot c}$.
- **Simplificación:** $\frac{a \cdot c}{b \cdot c} = \frac{a}{b}$ (si $c \neq 0$).

C.2. Productos notables

Expresión	Resultado
$(a + b)^2$	$a^2 + 2ab + b^2$
$(a - b)^2$	$a^2 - 2ab + b^2$
$(a + b)^3$	$a^3 + 3a^2b + 3ab^2 + b^3$
$(a - b)^3$	$a^3 - 3a^2b + 3ab^2 - b^3$
$(a + b + c)^2$	$a^2 + b^2 + c^2 + 2ab + 2ac + 2bc$
$(a + b)(a - b)$	$a^2 - b^2$
$(a + b)(a^2 - ab + b^2)$	$a^3 + b^3$
$(a - b)(a^2 + ab + b^2)$	$a^3 - b^3$

C.3. Factorización

Expresión	Factorización
$ax + ay$	$a(x + y)$
$x^2 - y^2$	$(x - y)(x + y)$
$x^3 - y^3$	$(x - y)(x^2 + xy + y^2)$
$x^3 + y^3$	$(x + y)(x^2 - xy + y^2)$
$x^2 + (a + b)x + ab$	$(x + a)(x + b)$
$ax^2 + bx + c$	$(px + q)(rx + s)$, con $pr = a$, $ps + qr = b$, $qs = c$

C.4. Leyes de exponentes

Ley	Expresión
Producto de potencias	$a^m \cdot a^n = a^{m+n}$
Cociente de potencias	$\frac{a^m}{a^n} = a^{m-n}$ (si $a \neq 0$)
Potencia de una potencia	$(a^m)^n = a^{m \cdot n}$
Potencia de un producto	$(ab)^n = a^n \cdot b^n$
Potencia de un cociente	$\left(\frac{a}{b}\right)^n = \frac{a^n}{b^n}$ (si $b \neq 0$)
Exponente cero	$a^0 = 1$ (si $a \neq 0$)
Exponente negativo	$a^{-n} = \frac{1}{a^n}$ (si $a \neq 0$)
Exponente fraccionario	$a^{m/n} = \sqrt[n]{a^m}$

C.5. Leyes de radicales

Ley	Expresión
Raíz de un producto	$\sqrt[n]{ab} = \sqrt[n]{a} \cdot \sqrt[n]{b}$
Raíz de un cociente	$\sqrt[n]{\frac{a}{b}} = \frac{\sqrt[n]{a}}{\sqrt[n]{b}}$ (si $b \neq 0$)
Raíz de una raíz	$\sqrt[m]{\sqrt[n]{a}} = \sqrt[mn]{a}$
Potencia de una raíz	$(\sqrt[n]{a})^m = \sqrt[n]{a^m}$
Simplificación	$\sqrt[n]{a^n} = a$ si $a \geq 0$

C.6. Ecuación lineal

$$ax + b = 0 \quad (a \neq 0) \implies x = -\frac{b}{a}$$

C.7. Ecuación cuadrática

$$ax^2 + bx + c = 0 \quad (a \neq 0)$$

$$x = \frac{-b \pm \sqrt{b^2 - 4ac}}{2a}$$

Discriminante: $\Delta = b^2 - 4ac$

- Si $\Delta > 0$: dos raíces reales distintas.
- Si $\Delta = 0$: una raíz real doble.
- Si $\Delta < 0$: dos raíces complejas conjugadas.

Relaciones entre raíces:

$$x_1 + x_2 = -\frac{b}{a}, \quad x_1 \cdot x_2 = \frac{c}{a}$$

C.8. Sistema de dos ecuaciones lineales

$$\begin{cases} a_1x + b_1y = c_1 \\ a_2x + b_2y = c_2 \end{cases}$$

Regla de Cramer (si $\Delta = a_1b_2 - a_2b_1 \neq 0$):

$$x = \frac{c_1b_2 - c_2b_1}{\Delta}, \quad y = \frac{a_1c_2 - a_2c_1}{\Delta}$$

C.9. Logaritmos

Propiedad	Expresión
Definición	$\log_b(a) = c \iff b^c = a$
Producto	$\log_b(xy) = \log_b x + \log_b y$
Cociente	$\log_b\left(\frac{x}{y}\right) = \log_b x - \log_b y$
Potencia	$\log_b(x^n) = n \log_b x$
Cambio de base	$\log_b a = \frac{\log_c a}{\log_c b}$
Logaritmo natural	$\ln x = \log_e x (e \approx 2,71828)$
Logaritmo decimal	$\log x = \log_{10} x$
Identidad fundamental	$b^{\log_b a} = a$
Logaritmo de 1	$\log_b 1 = 0$
Logaritmo de la base	$\log_b b = 1$

C.10. Introducción a límites

Concepto de límite

El límite de una función $f(x)$ cuando x tiende a un valor a es el valor al que se aproxima $f(x)$ a medida que x se acerca a a .

$$\lim_{x \rightarrow a} f(x) = L$$

Propiedades básicas de los límites

- **Límite de una constante:** $\lim_{x \rightarrow a} c = c$.
- **Límite de x :** $\lim_{x \rightarrow a} x = a$.
- **Suma:** $\lim_{x \rightarrow a} [f(x) + g(x)] = \lim_{x \rightarrow a} f(x) + \lim_{x \rightarrow a} g(x)$.
- **Resta:** $\lim_{x \rightarrow a} [f(x) - g(x)] = \lim_{x \rightarrow a} f(x) - \lim_{x \rightarrow a} g(x)$.
- **Producto:** $\lim_{x \rightarrow a} [f(x) \cdot g(x)] = \lim_{x \rightarrow a} f(x) \cdot \lim_{x \rightarrow a} g(x)$.
- **Cociente:** $\lim_{x \rightarrow a} \frac{f(x)}{g(x)} = \frac{\lim_{x \rightarrow a} f(x)}{\lim_{x \rightarrow a} g(x)}$ (si $\lim_{x \rightarrow a} g(x) \neq 0$).
- **Potencia:** $\lim_{x \rightarrow a} [f(x)]^n = \left[\lim_{x \rightarrow a} f(x) \right]^n$.

- **Raíz:** $\lim_{x \rightarrow a} \sqrt[n]{f(x)} = \sqrt[n]{\lim_{x \rightarrow a} f(x)}$ (si $f(x) \geq 0$ para n par).

Límites notables (fundamentales)

- $\lim_{x \rightarrow 0} \frac{\operatorname{sen} x}{x} = 1.$
- $\lim_{x \rightarrow 0} \frac{1 - \cos x}{x} = 0.$
- $\lim_{x \rightarrow \infty} \left(1 + \frac{1}{x}\right)^x = e.$

Límites al infinito

- $\lim_{x \rightarrow \infty} \frac{1}{x} = 0.$
- $\lim_{x \rightarrow \infty} \frac{1}{x^n} = 0$ (para $n > 0$).
- $\lim_{x \rightarrow \infty} \frac{a_n x^n + \dots}{b_m x^m + \dots} = \begin{cases} \infty & \text{si } n > m \\ \frac{a_n}{b_m} & \text{si } n = m \\ 0 & \text{si } n < m \end{cases}$

C.II. Álgebra booleana y tablas de verdad

Postulados del álgebra booleana

- Los valores posibles son 0 y 1, es decir falso (F) y verdadero (V).
- Operadores básicos: AND ($\cdot$), OR ($+$), NOT ($\bar{A}$ o $\neg A$).
- Elemento neutro para AND: $A \cdot 1 = A$.
- Elemento neutro para OR: $A + 0 = A$.
- Elemento complemento: $A \cdot \bar{A} = 0$, $A + \bar{A} = 1$.

Leyes fundamentales del álgebra booleana

- **Idempotencia:** $A + A = A$, $A \cdot A = A$.
- **Conmutatividad:** $A + B = B + A$, $A \cdot B = B \cdot A$.

- **Asociatividad:** $(A + B) + C = A + (B + C)$, $(A \cdot B) \cdot C = A \cdot (B \cdot C)$.
- **Distributividad:** $A \cdot (B + C) = A \cdot B + A \cdot C$, $A + B \cdot C = (A + B) \cdot (A + C)$.
- **Absorción:** $A + A \cdot B = A$, $A \cdot (A + B) = A$.
- **De Morgan:** $\overline{A \cdot B} = \overline{A} + \overline{B}$, $\overline{A + B} = \overline{A} \cdot \overline{B}$.
- **Doble negación:** $\overline{\overline{A}} = A$.

Compuertas lógicas y sus tablas de verdad

Compuerta	A	B	S	A*	B*	S*
AND ($A \cdot B$)	0	0	0	F	F	F
AND ($A \cdot B$)	0	1	0	F	V	F
AND ($A \cdot B$)	1	0	0	V	F	F
AND ($A \cdot B$)	1	1	1	V	V	V
OR ($A + B$)	0	0	0	F	F	F
OR ($A + B$)	0	1	1	F	V	V
OR ($A + B$)	1	0	1	V	F	V
OR ($A + B$)	1	1	1	V	V	V
NOT ($\overline{A}$)	0	--	1	F	--	V
NOT ($\overline{A}$)	1	--	0	V	--	F
XOR ($A \oplus B$)	0	0	0	F	F	F
XOR ($A \oplus B$)	0	1	1	F	V	V
XOR ($A \oplus B$)	1	0	1	V	F	V
XOR ($A \oplus B$)	1	1	0	V	V	F
NAND ($\overline{A \cdot B}$)	0	0	1	F	F	V
NAND ($\overline{A \cdot B}$)	0	1	1	F	V	V
NAND ($\overline{A \cdot B}$)	1	0	1	V	F	V
NAND ($\overline{A \cdot B}$)	1	1	0	V	V	F
NOR ($\overline{A + B}$)	0	0	1	F	F	V
NOR ($\overline{A + B}$)	0	1	0	F	V	F
NOR ($\overline{A + B}$)	1	0	0	V	F	F
NOR ($\overline{A + B}$)	1	1	0	V	V	F

Álgebra booleana en la práctica

- Los circuitos digitales se construyen combinando compuertas lógicas.
- Las tablas de verdad permiten verificar el comportamiento de cualquier circuito.
- El álgebra booleana es la base de la electrónica digital y de la lógica de programación.

GEOMETRÍA PLANA Y DEL ESPACIO

Principios fundamentales de la geometría

Principios generales del razonamiento matemático, formulados por Euclides:

Nota: En la formulación de estas “Nociones”, Euclides empleaba el término “magnitud” en un sentido abstracto, para referirse a cualquier entidad (números, entidades geométricas, figuras, etc.) capaz de ser comparada, igualada u operada dentro del razonamiento matemático.

Noción Común 1 (propiedad transitiva de la igualdad): magnitudes que son iguales a una misma magnitud, son iguales entre sí.

Noción Común 2 (uniformidad de la suma): si a magnitudes iguales se añaden magnitudes iguales, los totales son iguales.

Noción Común 3 (uniformidad de la resta): si de magnitudes iguales se restan magnitudes iguales, los restos son iguales.

Axiomas elementales adicionales:

1. El espacio tiene infinitos puntos, rectas y planos.
2. El plano tiene infinitos puntos y rectas.
3. La recta tiene infinitos puntos.
4. Por un punto pasan infinitas rectas.
5. Por una recta pasan infinitos planos.
6. Por tres puntos no alineados pasa un único plano.
7. Si dos puntos pertenecen a un plano, la recta que pasa por esos dos puntos también se encuentra en el mismo plano.

Axioma de existencia de la recta: por dos puntos distintos pasa una única recta.

Axioma de existencia de la circunferencia: dado un punto y una distancia, existe una circunferencia con ese centro y radio.

Axioma de medición de segmentos: todo segmento puede ser medido y su longitud es un número real positivo.

Axioma de la desigualdad triangular: en todo triángulo, la suma de las longitudes de dos lados cualesquiera es mayor que la longitud del lado restante.

Postulado de las paralelas (Euclides): por un punto exterior a una recta pasa exactamente una recta paralela a ella.

Los postulados de Euclides son los principios fundamentales a partir de los cuales se construye la forma sistemática de razonamiento que caracteriza al pensamiento matemático como ciencia abstracta.

D.I. Figuras planas: perímetros, áreas y elementos

Figura	Perímetro	Área	Elementos / Observaciones
Cuadrado (lado a)	$4a$	a^2	Diagonal: $d = a\sqrt{2}$
Rectángulo (base b , altura h)	$2(b + h)$	$b \cdot h$	Diagonal: $d = \sqrt{b^2 + h^2}$
Triángulo (base b , altura h)	$a + b + c$	$\frac{b \cdot h}{2}$	Semiperímetro: $s = \frac{a + b + c}{2}$
Triángulo equilátero (lado a)	$3a$	$\frac{\sqrt{3}}{4}a^2$	Altura: $h = \frac{\sqrt{3}}{2}a$
Triángulo rectángulo	$a + b + c$	$\frac{a \cdot b}{2}$	Pitágoras: $c^2 = a^2 + b^2$
Paralelogramo (base b , altura h)	$2(a + b)$	$b \cdot h$	Lado oblicuo: a
Rombo (lado a , diagonales d_1, d_2)	$4a$	$\frac{d_1 \cdot d_2}{2}$	$a = \frac{1}{2}\sqrt{d_1^2 + d_2^2}$
Trapezio (bases B, b , altura h)	$B + b + l_1 + l_2$	$\frac{(B + b) \cdot h}{2}$	l_1, l_2 : lados oblicuos
Trapezio isósceles (bases B, b , altura h)	$B + b + 2l$	$\frac{(B + b) \cdot h}{2}$	$l = \sqrt{h^2 + \left(\frac{B-b}{2}\right)^2}$
Polígono regular (n lados, lado a , apotema a_p)	$n \cdot a$	$\frac{n \cdot a \cdot a_p}{2}$	Ángulo central: $\theta = \frac{360^\circ}{n}$
Círculo (radio r)	$2\pi r$	πr^2	Diámetro: $D = 2r$
Sector circular (radio r , ángulo θ en rad)	$2r + r\theta$	$\frac{\theta}{2}r^2$	Longitud de arco: $L = r\theta$

D.2. Figuras del espacio: áreas, volúmenes y elementos

Figura	Área total	Volumen	Elementos / Observaciones
Cubo (arista a)	$6a^2$	a^3	Diagonal: $d = a\sqrt{3}$
Prisma recto (A_b : área base, h : altura)	$2A_b + A_{\text{lateral}}$	$A_b \cdot h$	Volumen = área base $\times$ altura
Prisma regular (n lados, lado a , altura h)	$2A_b + n \cdot a \cdot h$	$A_b \cdot h$	$A_b = \frac{n \cdot a \cdot a_p}{2}$
Cilindro (radio r , altura h)	$2\pi r^2 + 2\pi r h$	$\pi r^2 h$	Área lateral: $A_L = 2\pi r h$
Cilindro hueco (radios R, r , altura h)	$2\pi(R^2 - r^2) + 2\pi(R + r)h$	$\pi(R^2 - r^2)h$	R : exterior, r : interior
Cono (radio r , altura h , generatriz g)	$\pi r^2 + \pi r g$	$\frac{1}{3}\pi r^2 h$	$g = \sqrt{r^2 + h^2}$
Cono truncado (radios R, r , altura h , generatriz g)	$\pi(R^2 + r^2 + (R + r)g)$	$\frac{1}{3}\pi h(R^2 + Rr + r^2)$	$g = \sqrt{h^2 + (R - r)^2}$
Esfera (radio r)	$4\pi r^2$	$\frac{4}{3}\pi r^3$	Área del círculo máximo: $A = \pi r^2$
Pirámide (A_b : área base, h : altura)	$A_b + A_{\text{lateral}}$	$\frac{1}{3}A_b \cdot h$	Volumen = $\frac{1}{3}$ área base $\times$ altura
Pirámide regular (n lados, lado a , altura h)	$A_b + \frac{n \cdot a \cdot a_p}{2}$	$\frac{1}{3}A_b \cdot h$	a_p : apotema de la pirámide
Tetraedro regular (arista a)	$\sqrt{3}a^2$	$\frac{\sqrt{2}}{12}a^3$	4 caras triangulares equiláteras
Octaedro regular (arista a)	$2\sqrt{3}a^2$	$\frac{\sqrt{2}}{3}a^3$	8 caras triangulares equiláteras
Icosaedro regular (arista a)	$5\sqrt{3}a^2$	$\frac{5(3 + \sqrt{5})}{12}a^3$	20 caras triangulares equiláteras

D.3. Teoremas geométricos fundamentales

Teorema de Pitágoras: en un triángulo rectángulo, el cuadrado de la hipotenusa es igual a la suma de los cuadrados de los catetos: $c^2 = a^2 + b^2$.

Teorema de Tales: si varias rectas paralelas cortan a dos rectas secantes, los segmentos correspondientes son proporcionales.

Teorema de la bisectriz: la bisectriz de un ángulo divide al lado opuesto en segmentos proporcionales a los lados adyacentes.

Teorema de la altura: en un triángulo rectángulo, la altura sobre la hipotenusa cumple: $h^2 = m \cdot n$, donde m y n son las proyecciones de los catetos sobre la hipotenusa.

Teorema del cateto: en un triángulo rectángulo, cada cateto es media proporcional entre la hipotenusa y su proyección sobre ella: $a^2 = c \cdot m$, $b^2 = c \cdot n$.

Teorema de Menelao: en un triángulo ABC , una recta transversal corta a los lados AB , BC y CA en los puntos D , E y F respectivamente, entonces: $\frac{AD}{DB} \cdot \frac{BE}{EC} \cdot \frac{CF}{FA} = 1$.

Teorema de Ceva: en un triángulo ABC , las cevianas AD , BE y CF son concurrentes si y solo si: $\frac{BD}{DC} \cdot \frac{CE}{EA} \cdot \frac{AF}{FB} = 1$.

Ley de cosenos (Pitágoras generalizado): en cualquier triángulo, $c^2 = a^2 + b^2 - 2ab \cos C$.

Suma de ángulos internos de un triángulo: $\alpha + \beta + \gamma = 180^\circ = \pi$ radianes.

Suma de ángulos internos de un polígono de n lados: $S = (n - 2) \cdot 180^\circ$.

Suma de ángulos externos de un polígono convexo: $S_e = 360^\circ$, constante para todo polígono convexo.

Número de diagonales de un polígono de n lados: $D = \frac{n(n-3)}{2}$.

Ángulo central de un polígono regular de n lados: $\theta = \frac{360^\circ}{n}$.

Ángulo interior de un polígono regular de n lados: $\alpha = \frac{(n-2) \cdot 180^\circ}{n}$.

Teorema de los ángulos inscritos: un ángulo inscrito en una circunferencia es igual a la mitad del arco que subtiende.

Teorema de las cuerdas: si dos cuerdas AB y CD se cortan en un punto P interior a la circunferencia, entonces: $PA \cdot PB = PC \cdot PD$.

Teorema de la tangente y la secante: desde un punto exterior P , si se trazan una tangente y una secante: $PT^2 = PA \cdot PB$.

Teorema de las secantes: desde un punto exterior P , si se trazan dos secantes que cortan a la circunferencia en A , B y C , D , entonces: $PA \cdot PB = PC \cdot PD$.

Fórmula de Euler para poliedros convexos: $C + V = A + 2$, donde C es el número de caras, V el de vértices y A el de aristas.

D.4. Geometría analítica

Distancia entre dos puntos:

$$d = \sqrt{(x_2 - x_1)^2 + (y_2 - y_1)^2}$$

Punto medio de un segmento:

$$M = \left(\frac{x_1 + x_2}{2}, \frac{y_1 + y_2}{2} \right)$$

Pendiente de una recta:

$$m = \frac{y_2 - y_1}{x_2 - x_1}$$

Ecuación de la recta (pendiente-intercepto):

$$y = mx + b$$

Ecuación de la recta (punto-pendiente):

$$y - y_1 = m(x - x_1)$$

Ecuación general de la recta:

$$Ax + By + C = 0$$

Condición de paralelismo: $m_1 = m_2$.

Condición de perpendicularidad: $m_1 \cdot m_2 = -1$.

Circunferencia con centro (h, k) y radio r :

$$(x - h)^2 + (y - k)^2 = r^2$$

Elipse con centro (h, k) :

$$\frac{(x - h)^2}{a^2} + \frac{(y - k)^2}{b^2} = 1$$

Parábola con vértice (h, k) :

$$(y - k)^2 = 4p(x - h) \quad (\text{horizontal})$$

$$(x - h)^2 = 4p(y - k) \quad (\text{vertical})$$

Hipérbola con centro (h, k) :

$$\frac{(x - h)^2}{a^2} - \frac{(y - k)^2}{b^2} = 1$$

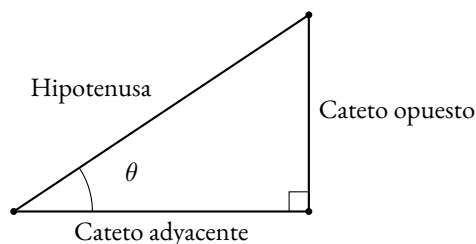
Anexo E

TRIGONOMETRÍA

E.1. Razones trigonométricas fundamentales

En el triángulo rectángulo que se muestra a la derecha, los lados que forman el ángulo recto se denominan *catetos*. Para el ángulo θ , el lado opuesto es el *cateto opuesto* y el lado contiguo es el *cateto adyacente*. Con esta nomenclatura, definimos:

Función	Definición (razón de lados)
$\text{sen } \theta$	$\frac{\text{cateto opuesto}}{\text{hipotenusa}}$
$\text{cos } \theta$	$\frac{\text{cateto adyacente}}{\text{hipotenusa}}$
$\text{tan } \theta$	$\frac{\text{cateto opuesto}}{\text{cateto adyacente}}$
$\text{cot } \theta$	$\frac{\text{cateto adyacente}}{\text{cateto opuesto}}$
$\text{sec } \theta$	$\frac{\text{hipotenusa}}{\text{cateto adyacente}}$
$\text{csc } \theta$	$\frac{\text{hipotenusa}}{\text{cateto opuesto}}$



E.2. Identidades pitagóricas

$$\text{sen}^2 \theta + \text{cos}^2 \theta = 1 \quad (\text{E.1})$$

$$1 + \text{tan}^2 \theta = \text{sec}^2 \theta \quad (\text{E.2})$$

$$1 + \text{cot}^2 \theta = \text{csc}^2 \theta \quad (\text{E.3})$$

E.3. Paridad (simetría)

$$\text{sen}(-\theta) = -\text{sen } \theta \quad (\text{E.4})$$

$$\text{cos}(-\theta) = \text{cos } \theta \quad (\text{E.5})$$

$$\text{tan}(-\theta) = -\text{tan } \theta \quad (\text{E.6})$$

E.4. Fórmulas de suma y diferencia de ángulos

$$\operatorname{sen}(a + b) = \operatorname{sen} a \cos b + \cos a \operatorname{sen} b \quad (\text{E.7})$$

$$\operatorname{sen}(a - b) = \operatorname{sen} a \cos b - \cos a \operatorname{sen} b \quad (\text{E.8})$$

$$\cos(a + b) = \cos a \cos b - \operatorname{sen} a \operatorname{sen} b \quad (\text{E.9})$$

$$\cos(a - b) = \cos a \cos b + \operatorname{sen} a \operatorname{sen} b \quad (\text{E.10})$$

$$\tan(a + b) = \frac{\tan a + \tan b}{1 - \tan a \tan b} \quad (\text{E.11})$$

$$\tan(a - b) = \frac{\tan a - \tan b}{1 + \tan a \tan b} \quad (\text{E.12})$$

E.5. Fórmulas del ángulo doble

$$\operatorname{sen}(2\theta) = 2 \operatorname{sen} \theta \cos \theta \quad (\text{E.13})$$

$$\cos(2\theta) = \cos^2 \theta - \operatorname{sen}^2 \theta = 2 \cos^2 \theta - 1 = 1 - 2 \operatorname{sen}^2 \theta \quad (\text{E.14})$$

$$\tan(2\theta) = \frac{2 \tan \theta}{1 - \tan^2 \theta} \quad (\text{E.15})$$

E.6. Fórmulas del ángulo medio

$$\operatorname{sen}^2 \left(\frac{\theta}{2} \right) = \frac{1 - \cos \theta}{2} \quad (\text{E.16})$$

$$\cos^2 \left(\frac{\theta}{2} \right) = \frac{1 + \cos \theta}{2} \quad (\text{E.17})$$

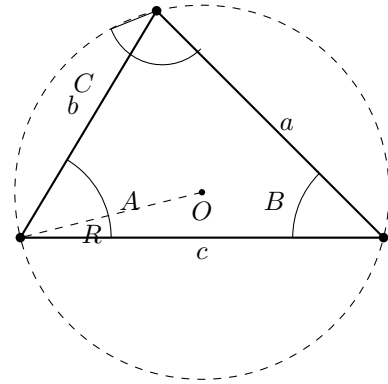
$$\tan \left(\frac{\theta}{2} \right) = \frac{1 - \cos \theta}{\operatorname{sen} \theta} = \frac{\operatorname{sen} \theta}{1 + \cos \theta} \quad (\text{E.18})$$

E.7. Ley de senos

Para cualquier triángulo con lados a , b , c y ángulos opuestos A , B , C , se cumple que las razones entre cada lado y el seno de su ángulo opuesto son iguales al diámetro de la circunferencia circunscrita:

$$\frac{a}{\operatorname{sen} A} = \frac{b}{\operatorname{sen} B} = \frac{c}{\operatorname{sen} C} = 2R$$

donde R es el radio de la circunferencia circunscrita al triángulo.



E.8. Ley de cosenos

Considerando el mismo triángulo ABC de la figura de la ley de senos, se cumplen las siguientes relaciones:

$$c^2 = a^2 + b^2 - 2ab \cos C$$

$$a^2 = b^2 + c^2 - 2bc \cos A$$

$$b^2 = a^2 + c^2 - 2ac \cos B$$

La ley de cosenos generaliza el teorema de Pitágoras: si el ángulo $C = 90^\circ$, entonces $\cos C = 0$ y la primera fórmula se reduce a $c^2 = a^2 + b^2$.

E.9. Identidades trigonométricas adicionales

Suma a producto (factorización)

$$\operatorname{sen} A + \operatorname{sen} B = 2 \operatorname{sen} \left(\frac{A+B}{2} \right) \cos \left(\frac{A-B}{2} \right) \quad (\text{E.19})$$

$$\operatorname{sen} A - \operatorname{sen} B = 2 \cos \left(\frac{A+B}{2} \right) \operatorname{sen} \left(\frac{A-B}{2} \right) \quad (\text{E.20})$$

$$\cos A + \cos B = 2 \cos \left(\frac{A+B}{2} \right) \cos \left(\frac{A-B}{2} \right) \quad (\text{E.21})$$

$$\cos A - \cos B = -2 \operatorname{sen} \left(\frac{A+B}{2} \right) \operatorname{sen} \left(\frac{A-B}{2} \right) \quad (\text{E.22})$$

Producto a suma

$$\operatorname{sen} A \cos B = \frac{1}{2}[\operatorname{sen}(A + B) + \operatorname{sen}(A - B)] \quad (\text{E.23})$$

$$\cos A \operatorname{sen} B = \frac{1}{2}[\operatorname{sen}(A + B) - \operatorname{sen}(A - B)] \quad (\text{E.24})$$

$$\cos A \cos B = \frac{1}{2}[\cos(A + B) + \cos(A - B)] \quad (\text{E.25})$$

$$\operatorname{sen} A \operatorname{sen} B = \frac{1}{2}[\cos(A - B) - \cos(A + B)] \quad (\text{E.26})$$

Ángulo triple

$$\operatorname{sen} 3\theta = 3 \operatorname{sen} \theta - 4 \operatorname{sen}^3 \theta \quad (\text{E.27})$$

$$\cos 3\theta = 4 \cos^3 \theta - 3 \cos \theta \quad (\text{E.28})$$

$$\tan 3\theta = \frac{3 \tan \theta - \tan^3 \theta}{1 - 3 \tan^2 \theta} \quad (\text{E.29})$$

Reducción de potencia

$$\operatorname{sen}^2 \theta = \frac{1 - \cos 2\theta}{2} \quad (\text{E.30})$$

$$\cos^2 \theta = \frac{1 + \cos 2\theta}{2} \quad (\text{E.31})$$

$$\tan^2 \theta = \frac{1 - \cos 2\theta}{1 + \cos 2\theta} \quad (\text{E.32})$$

Ángulos complementarios y suplementarios

$$\operatorname{sen} \left(\frac{\pi}{2} - \theta \right) = \cos \theta \quad (\text{E.33})$$

$$\cos \left(\frac{\pi}{2} - \theta \right) = \operatorname{sen} \theta \quad (\text{E.34})$$

$$\operatorname{sen}(\pi - \theta) = \operatorname{sen} \theta \quad (\text{E.35})$$

$$\cos(\pi - \theta) = -\cos \theta \quad (\text{E.36})$$

$$\operatorname{sen}(\pi + \theta) = -\operatorname{sen} \theta \quad (\text{E.37})$$

$$\cos(\pi + \theta) = -\cos \theta \quad (\text{E.38})$$

Anexo F

VECTORES

F.1. Notación

- Un vector se denota como $\mathbf{v}$ o $\vec{v}$.
- En el plano: $\mathbf{v} = (v_x, v_y)$.
- En el espacio: $\mathbf{v} = (v_x, v_y, v_z)$.
- Vectores unitarios: $\mathbf{i} = (1, 0, 0)$, $\mathbf{j} = (0, 1, 0)$, $\mathbf{k} = (0, 0, 1)$.
- Vector unitario: $\hat{\mathbf{v}} = \frac{\mathbf{v}}{\|\mathbf{v}\|}$.

F.2. Operaciones fundamentales

Operación	Expresión
Suma	$\mathbf{u} + \mathbf{v} = (u_x + v_x, u_y + v_y, u_z + v_z)$
Resta	$\mathbf{u} - \mathbf{v} = (u_x - v_x, u_y - v_y, u_z - v_z)$
Multipliación por escalar	$c\mathbf{v} = (cv_x, cv_y, cv_z)$
Magnitud	$\ \mathbf{v}\ = \sqrt{v_x^2 + v_y^2 + v_z^2}$
Producto escalar	$\mathbf{u} \cdot \mathbf{v} = u_x v_x + u_y v_y + u_z v_z = \ \mathbf{u}\ \ \mathbf{v}\ \cos \theta$
Producto vectorial	$\mathbf{u} \times \mathbf{v} = (u_y v_z - u_z v_y, u_z v_x - u_x v_z, u_x v_y - u_y v_x)$
Módulo del producto cruz	$\ \mathbf{u} \times \mathbf{v}\ = \ \mathbf{u}\ \ \mathbf{v}\ \sin \theta$

F.3. Propiedades importantes

- **Producto escalar:** conmutativo, distributivo y se anula si los vectores son perpendiculares ($\mathbf{u} \cdot \mathbf{v} = 0$).
- **Producto vectorial:** anticonmutativo ($\mathbf{u} \times \mathbf{v} = -\mathbf{v} \times \mathbf{u}$) y perpendicular a ambos.
- Dos vectores son paralelos si $\mathbf{u} \times \mathbf{v} = \mathbf{0}$.

F.4. Aplicaciones físicas de vectores

Cinemática

$$\text{Posición: } \mathbf{r}(t) = x(t)\mathbf{i} + y(t)\mathbf{j} + z(t)\mathbf{k} \quad (\text{F.1})$$

$$\text{Velocidad: } \mathbf{v}(t) = \frac{d\mathbf{r}}{dt} = \dot{x}\mathbf{i} + \dot{y}\mathbf{j} + \dot{z}\mathbf{k} \quad (\text{F.2})$$

$$\text{Aceleración: } \mathbf{a}(t) = \frac{d\mathbf{v}}{dt} = \frac{d^2\mathbf{r}}{dt^2} = \ddot{x}\mathbf{i} + \ddot{y}\mathbf{j} + \ddot{z}\mathbf{k} \quad (\text{F.3})$$

Dinámica

$$\text{Segunda ley de Newton: } \mathbf{F} = m\mathbf{a} \quad (\text{F.4})$$

$$\text{Fuerza gravitacional: } \mathbf{F}_g = -\frac{Gm_1m_2}{r^3}\mathbf{r} \quad (\text{F.5})$$

$$\text{Fuerza eléctrica: } \mathbf{F}_e = k_e \frac{q_1q_2}{r^3}\mathbf{r} \quad (\text{F.6})$$

$$\text{Fuerza de Lorentz: } \mathbf{F} = q(\mathbf{E} + \mathbf{v} \times \mathbf{B}) \quad (\text{F.7})$$

$$\text{Fuerza de resorte: } \mathbf{F}_s = -k\mathbf{x} \quad (\text{F.8})$$

Cantidades de movimiento

$$\text{Momento lineal: } \mathbf{p} = m\mathbf{v} \quad (\text{F.9})$$

$$\text{Impulso: } \mathbf{J} = \int \mathbf{F} dt = \Delta\mathbf{p} \quad (\text{F.10})$$

$$\text{Momento angular: } \mathbf{L} = \mathbf{r} \times \mathbf{p} \quad (\text{F.11})$$

$$\text{Torque: } \boldsymbol{\tau} = \frac{d\mathbf{L}}{dt} = \mathbf{r} \times \mathbf{F} \quad (\text{F.12})$$

Energía y trabajo

$$\text{Trabajo: } W = \int \mathbf{F} \cdot d\mathbf{r} \quad (\text{F.13})$$

$$\text{Potencia: } P = \frac{dW}{dt} = \mathbf{F} \cdot \mathbf{v} \quad (\text{F.14})$$

$$\text{Energía cinética: } E_c = \frac{1}{2}mv^2 \quad (\text{F.15})$$

Tabla resumen de aplicaciones vectoriales

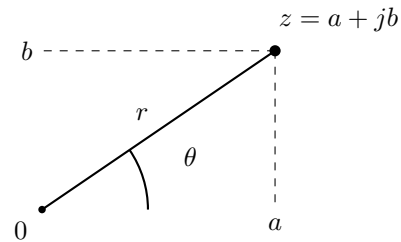
Cantidad física	Expresión vectorial	Unidades (SI)
Posición	$\mathbf{r} = x\mathbf{i} + y\mathbf{j} + z\mathbf{k}$	m
Velocidad	$\mathbf{v} = \frac{d\mathbf{r}}{dt}$	m/s
Aceleración	$\mathbf{a} = \frac{d\mathbf{v}}{dt}$	m/s ²
Cantidad de movimiento	$\mathbf{p} = m\mathbf{v}$	kg·m/s
Fuerza	$\mathbf{F} = m\mathbf{a}$	N
Momento angular	$\mathbf{L} = \mathbf{r} \times \mathbf{p}$	kg·m ² /s
Torque	$\boldsymbol{\tau} = \mathbf{r} \times \mathbf{F}$	N·m
Trabajo	$W = \mathbf{F} \cdot \Delta\mathbf{r}$	J
Potencia	$P = \mathbf{F} \cdot \mathbf{v}$	W
Campo eléctrico	$\mathbf{E} = \frac{\mathbf{F}}{q}$	N/C

Anexo G

NÚMEROS COMPLEJOS

G.1. Representaciones

Un número complejo $z = a + jb$ puede representarse como un punto en el plano complejo, donde el eje horizontal es la parte real y el vertical la parte imaginaria. El módulo r es la distancia al origen y el argumento θ es el ángulo con el eje real positivo.



Forma	Expresión
Rectangular	$z = a + jb$, con $j = \sqrt{-1}$
Polar	$z = r(\cos \theta + j \sin \theta)$
Exponencial	$z = r e^{j\theta}$
Módulo	$r = \ z\ = \sqrt{a^2 + b^2}$
Argumento	$\theta = \arg(z) = \text{atan2}(b, a)$

G.2. Operaciones básicas

Operación	Expresión
Suma	$(a + jb) + (c + jd) = (a + c) + j(b + d)$
Resta	$(a + jb) - (c + jd) = (a - c) + j(b - d)$
Multiplicación	$(a + jb)(c + jd) = (ac - bd) + j(ad + bc)$
División	$\frac{a + jb}{c + jd} = \frac{(a + jb)(c - jd)}{c^2 + d^2}$
Conjugado	$z^* = a - jb$
Producto por conjugado	$z \cdot z^* = a^2 + b^2 = \ z\ ^2$

G.3. Fórmulas fundamentales

A continuación se presentan las fórmulas más importantes que relacionan las distintas representaciones de los números complejos y sus operaciones.

Concepto / Fórmula	Expresión matemática
Fórmula de Euler	$e^{j\theta} = \cos \theta + j \sin \theta$
Identidades de Euler	$\cos \theta = \frac{e^{j\theta} + e^{-j\theta}}{2}, \quad \sin \theta = \frac{e^{j\theta} - e^{-j\theta}}{2j}$
Teorema de De Moivre	$(re^{j\theta})^n = r^n e^{jn\theta} = r^n (\cos n\theta + j \sin n\theta)$
Potencia de un complejo	$z^n = r^n e^{jn\theta} = r^n (\cos(n\theta) + j \sin(n\theta))$
Raíces n -ésimas	$z_k = r^{1/n} e^{j(\theta+2\pi k)/n}, \quad k = 0, 1, \dots, n-1$
Raíces de la unidad	$\omega_k = e^{j2\pi k/n}, \quad k = 0, 1, \dots, n-1$
Logaritmo complejo	$\ln z = \ln r + j(\theta + 2\pi k), \quad k \in \mathbb{Z}$
Conjugado en forma polar	$\bar{z} = r e^{-j\theta}$

Nota: En el logaritmo complejo, el valor principal se obtiene con $k = 0$ y suele denotarse $\text{Log } z = \ln r + j\theta$, donde θ es el argumento principal ($-\pi < \theta \leq \pi$).

G.4. Ejemplos de operaciones con números complejos

Los siguientes ejemplos muestran paso a paso las operaciones más comunes y la conversión entre formas.

Ejemplo 1: Paso de rectangular a polar

Dado $z = 3 + 4j$:

$$r = \|z\| = \sqrt{3^2 + 4^2} = \sqrt{25} = 5$$

$$\theta = \arg(z) = \text{atan2}(4, 3) \approx 0,9273 \text{ rad} \approx 53,13^\circ$$

Por tanto:

$$z = 5(\cos 0,9273 + j \sin 0,9273) = 5 e^{j 0,9273}$$

Ejemplo 2: Paso de polar a rectangular

Dado $z = 2 e^{j\pi/3}$:

$$z = 2 \left(\cos \frac{\pi}{3} + j \sin \frac{\pi}{3} \right) = 2 \left(\frac{1}{2} + j \frac{\sqrt{3}}{2} \right) = 1 + j\sqrt{3}$$

Su conjugado es $\bar{z} = 1 - j\sqrt{3}$.

Ejemplo 3: Multiplicación en forma polar y rectangular

Calcular $(3 + 4j)(1 + 2j)$:

■ **Forma rectangular:**

$$(3+4j)(1+2j) = 3 \cdot 1 + 3 \cdot 2j + 4j \cdot 1 + 4j \cdot 2j = 3 + 6j + 4j + 8j^2 = 3 + 10j - 8 = -5 + 10j$$

- **Forma polar:** Primero convertimos cada factor: $3 + 4j = 5e^{j0,9273}$, $1 + 2j = \sqrt{5} e^{j1,1071}$. El producto es:

$$5 \cdot \sqrt{5} e^{j(0,9273+1,1071)} = 5\sqrt{5} e^{j2,0344}$$

Comprobamos que esta forma coincide con $-5 + 10j$.

Ejemplo 4: Raíces complejas

Calcular las dos raíces cuadradas de $z = -1$ (es decir, $\sqrt{-1}$).

Escribimos $-1 = 1 \cdot e^{j\pi}$. Aplicando la fórmula de raíces:

$$z_k = 1^{1/2} e^{j(\pi+2\pi k)/2}, \quad k = 0, 1$$

Es decir:

$$z_0 = e^{j\pi/2} = j, \quad z_1 = e^{j(\pi/2+\pi)} = e^{j3\pi/2} = -j$$

Por supuesto, las soluciones son j y $-j$.

Ejemplo 5: Raíz cúbica de la unidad

Las tres raíces cúbicas de 1 son:

$$\omega_k = e^{j2\pi k/3}, \quad k = 0, 1, 2$$

que explícitamente son:

$$1, \quad -\frac{1}{2} + j\frac{\sqrt{3}}{2}, \quad -\frac{1}{2} - j\frac{\sqrt{3}}{2}.$$

Nota sobre el argumento

En la práctica, para calcular el argumento de un número complejo se usa la función $\text{atan2}(b, a)$, que devuelve el ángulo en el cuadrante correcto, a diferencia de $\arctan(b/a)$ que solo da el cuadrante I o IV.

Anexo H

MATRICES Y SISTEMAS DE ECUACIONES LINEALES

Este anexo presenta los conceptos fundamentales del álgebra matricial y su aplicación directa en la resolución de sistemas de ecuaciones lineales. Las matrices son una herramienta indispensable en el análisis de sistemas dinámicos, control, circuitos eléctricos, mecánica estructural y prácticamente cualquier área que requiera manejar múltiples variables interrelacionadas.

H.1. Definición y tipos de matrices

Definición

Una **matriz** es un arreglo rectangular de números dispuestos en filas y columnas. Se denota generalmente con una letra mayúscula en negrita:

$$\mathbf{A} = \begin{pmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{m1} & a_{m2} & \cdots & a_{mn} \end{pmatrix}$$

donde a_{ij} es el elemento que ocupa la fila i y la columna j . Se dice que $\mathbf{A}$ es de tamaño $m \times n$ (m filas, n columnas).

Tipos particulares

Tipo	Característica
Matriz fila	$1 \times n$ (una sola fila).
Matriz columna	$m \times 1$ (una sola columna).
Matriz cuadrada	$m = n$ (mismo número de filas y columnas).
Matriz diagonal	Cuadrada con $a_{ij} = 0$ para $i \neq j$.
Matriz identidad	Matriz diagonal con todos los elementos de la diagonal iguales a 1. Se denota $\mathbf{I}$.
Matriz nula	Todos sus elementos son cero. Se denota $\mathbf{0}$.
Matriz triangular superior	$a_{ij} = 0$ para $i > j$.
Matriz triangular inferior	$a_{ij} = 0$ para $i < j$.
Matriz simétrica	Cuadrada y $a_{ij} = a_{ji}$ para todo i, j .
Matriz antisimétrica	Cuadrada y $a_{ij} = -a_{ji}$ (implica $a_{ii} = 0$).
Matriz ortogonal	Cuadrada y $\mathbf{A}^T \mathbf{A} = \mathbf{A} \mathbf{A}^T = \mathbf{I}$.

H.2. Operaciones con matrices

Suma y resta

Dos matrices **A** y **B** del mismo tamaño se suman (o restan) elemento a elemento:

$$(\mathbf{A} \pm \mathbf{B})_{ij} = a_{ij} \pm b_{ij}$$

Multipliación por un escalar

$$(k\mathbf{A})_{ij} = k \cdot a_{ij}$$

Producto de matrices

El producto **AB** está definido si el número de columnas de **A** es igual al número de filas de **B**. Si **A** es $m \times p$ y **B** es $p \times n$, entonces **C** = **AB** es $m \times n$ y su elemento c_{ij} es:

$$c_{ij} = \sum_{k=1}^p a_{ik} b_{kj}$$

Propiedades (cuando las operaciones están definidas):

- Asociativa: $(\mathbf{AB})\mathbf{C} = \mathbf{A}(\mathbf{BC})$.
- Distributiva: $\mathbf{A}(\mathbf{B} + \mathbf{C}) = \mathbf{AB} + \mathbf{AC}$.
- **No conmutativa**: en general $\mathbf{AB} \neq \mathbf{BA}$.

Transpuesta de una matriz

La transpuesta de **A** (denotada $\mathbf{A}^T$) se obtiene intercambiando filas por columnas:

$$(\mathbf{A}^T)_{ij} = a_{ji}$$

Propiedades:

$$(\mathbf{A}^T)^T = \mathbf{A}, \quad (\mathbf{A} + \mathbf{B})^T = \mathbf{A}^T + \mathbf{B}^T, \quad (k\mathbf{A})^T = k\mathbf{A}^T, \quad (\mathbf{AB})^T = \mathbf{B}^T \mathbf{A}^T$$

H.3. Determinante de una matriz cuadrada

El determinante es un número asociado a toda matriz cuadrada que proporciona información sobre su invertibilidad y sobre la independencia lineal de sus filas/columnas.

Determinante de orden 2

$$\det \begin{pmatrix} a & b \\ c & d \end{pmatrix} = ad - bc$$

Determinante de orden 3 (regla de Sarrus)

$$\det \begin{pmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{pmatrix} = a_{11}a_{22}a_{33} + a_{12}a_{23}a_{31} + a_{13}a_{21}a_{32} - a_{13}a_{22}a_{31} - a_{12}a_{21}a_{33} - a_{11}a_{23}a_{32}$$

Propiedades fundamentales del determinante

- Si una fila (o columna) es combinación lineal de las demás, el determinante es cero.
- Si se intercambian dos filas (o columnas), el determinante cambia de signo.
- Si se multiplica una fila (o columna) por un escalar k , el determinante queda multiplicado por k .
- $\det(\mathbf{AB}) = \det(\mathbf{A}) \cdot \det(\mathbf{B})$.
- $\det(\mathbf{A}^T) = \det(\mathbf{A})$.
- $\det(\mathbf{I}) = 1$.

H.4. Matriz inversa

Una matriz cuadrada $\mathbf{A}$ es **invertible** (o no singular) si existe una matriz $\mathbf{A}^{-1}$ tal que:

$$\mathbf{A}\mathbf{A}^{-1} = \mathbf{A}^{-1}\mathbf{A} = \mathbf{I}$$

La inversa existe si y solo si $\det(\mathbf{A}) \neq 0$.

Cálculo de la inversa (método de la adjunta)

$$\mathbf{A}^{-1} = \frac{1}{\det(\mathbf{A})} \cdot \text{adj}(\mathbf{A})$$

donde $\text{adj}(\mathbf{A})$ es la matriz adjunta, formada por los cofactores transpuestos. Para una matriz 2×2 :

$$\mathbf{A} = \begin{pmatrix} a & b \\ c & d \end{pmatrix} \Rightarrow \mathbf{A}^{-1} = \frac{1}{ad - bc} \begin{pmatrix} d & -b \\ -c & a \end{pmatrix}$$

Propiedades de la inversa

- $(\mathbf{A}^{-1})^{-1} = \mathbf{A}$.
- $(\mathbf{AB})^{-1} = \mathbf{B}^{-1}\mathbf{A}^{-1}$ (si ambos son invertibles).
- $(\mathbf{A}^T)^{-1} = (\mathbf{A}^{-1})^T$.
- $\det(\mathbf{A}^{-1}) = \frac{1}{\det(\mathbf{A})}$.

H.5. Rango de una matriz

El **rango** de una matriz $\mathbf{A}$ es el número máximo de filas (o columnas) linealmente independientes. Coincide con el orden del mayor menor no nulo que se puede extraer de la matriz. Se denota $\text{rg}(\mathbf{A})$.

Propiedades

- $\text{rg}(\mathbf{A}) \leq \min(m, n)$.
- El rango no cambia al realizar operaciones elementales de fila (intercambiar filas, multiplicar una fila por un escalar no nulo, sumar a una fila un múltiplo de otra).
- Una matriz cuadrada de tamaño n es invertible si y solo si $\text{rg}(\mathbf{A}) = n$.

H.6. Sistemas de ecuaciones lineales

Un sistema de m ecuaciones lineales con n incógnitas $x_1, x_2, \dots, x_n$ se escribe como:

$$\begin{cases} a_{11}x_1 + a_{12}x_2 + \dots + a_{1n}x_n = b_1 \\ a_{21}x_1 + a_{22}x_2 + \dots + a_{2n}x_n = b_2 \\ \vdots \\ a_{m1}x_1 + a_{m2}x_2 + \dots + a_{mn}x_n = b_m \end{cases}$$

En forma matricial:

$$\mathbf{AX} = \mathbf{B}$$

donde

$$\mathbf{A} = \begin{pmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{m1} & a_{m2} & \dots & a_{mn} \end{pmatrix}, \quad \mathbf{X} = \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix}, \quad \mathbf{B} = \begin{pmatrix} b_1 \\ b_2 \\ \vdots \\ b_m \end{pmatrix}$$

H.7. Métodos de solución

Eliminación de Gauss

Consiste en transformar la matriz ampliada $[\mathbf{A} \mid \mathbf{B}]$ en una matriz triangular superior (o escalonada) mediante operaciones elementales de fila. Luego se resuelve el sistema por sustitución regresiva.

Eliminación de Gauss-Jordan

Continúa el proceso hasta obtener la forma escalonada reducida (matriz identidad en la parte de $\mathbf{A}$), de modo que la solución se lee directamente en la última columna.

Regla de Cramer (para sistemas con $\mathbf{A}$ cuadrada e invertible)

Cada incógnita x_i se calcula como:

$$x_i = \frac{\det(\mathbf{A}_i)}{\det(\mathbf{A})}$$

donde $\mathbf{A}_i$ es la matriz que se obtiene al reemplazar la columna i de $\mathbf{A}$ por el vector $\mathbf{B}$.

Uso de la matriz inversa

Si $\mathbf{A}$ es cuadrada e invertible, la solución única es:

$$\mathbf{X} = \mathbf{A}^{-1}\mathbf{B}$$

H.8. Clasificación de sistemas (Teorema de Rouché-Frobenius)

Sea $\mathbf{A}$ la matriz de coeficientes (tamaño $m \times n$) y $[\mathbf{A} \mid \mathbf{B}]$ la matriz ampliada. Entonces:

Condición	Tipo de sistema
$\text{rg}(\mathbf{A}) = \text{rg}([\mathbf{A} \mid \mathbf{B}]) = n$	Sistema compatible determinado (solución única).
$\text{rg}(\mathbf{A}) = \text{rg}([\mathbf{A} \mid \mathbf{B}]) < n$	Sistema compatible indeterminado (infinitas soluciones, con $n - \text{rg}$ grados de libertad).
$\text{rg}(\mathbf{A}) < \text{rg}([\mathbf{A} \mid \mathbf{B}])$	Sistema incompatible (no tiene solución).

H.9. Ejemplos ilustrativos

Ejemplo 1: Sistema 2×2

$$\begin{cases} 2x + 3y = 8 \\ x - y = -1 \end{cases} \Rightarrow \mathbf{A} = \begin{pmatrix} 2 & 3 \\ 1 & -1 \end{pmatrix}, \quad \mathbf{B} = \begin{pmatrix} 8 \\ -1 \end{pmatrix}$$

$$\det(\mathbf{A}) = 2(-1) - 3(1) = -5 \neq 0$$

$$\mathbf{A}^{-1} = \frac{1}{-5} \begin{pmatrix} -1 & -3 \\ -1 & 2 \end{pmatrix} = \begin{pmatrix} 1/5 & 3/5 \\ 1/5 & -2/5 \end{pmatrix}$$

$$\mathbf{X} = \mathbf{A}^{-1}\mathbf{B} = \begin{pmatrix} 1/5 & 3/5 \\ 1/5 & -2/5 \end{pmatrix} \begin{pmatrix} 8 \\ -1 \end{pmatrix} = \begin{pmatrix} (8/5) - (3/5) \\ (8/5) + (2/5) \end{pmatrix} = \begin{pmatrix} 1 \\ 2 \end{pmatrix}$$

Solución: $x = 1, y = 2$.

Ejemplo 2: Sistema 3×3 (regla de Cramer)

$$\begin{cases} x + y + z = 6 \\ 2x - y + z = 3 \\ x + 2y - z = 3 \end{cases}$$

$$\mathbf{A} = \begin{pmatrix} 1 & 1 & 1 \\ 2 & -1 & 1 \\ 1 & 2 & -1 \end{pmatrix}, \quad \mathbf{B} = \begin{pmatrix} 6 \\ 3 \\ 3 \end{pmatrix}$$

$$\begin{aligned} \det(\mathbf{A}) &= 1((-1)(-1) - 1 \cdot 2) - 1(2(-1) - 1 \cdot 1) + 1(2 \cdot 2 - (-1) \cdot 1) \\ &= 1(1 - 2) - 1(-2 - 1) + 1(4 + 1) \\ &= -1 + 3 + 5 = 7 \end{aligned}$$

$$\mathbf{A}_1 = \begin{pmatrix} 6 & 1 & 1 \\ 3 & -1 & 1 \\ 3 & 2 & -1 \end{pmatrix}, \quad \mathbf{A}_2 = \begin{pmatrix} 1 & 6 & 1 \\ 2 & 3 & 1 \\ 1 & 3 & -1 \end{pmatrix}, \quad \mathbf{A}_3 = \begin{pmatrix} 1 & 1 & 6 \\ 2 & -1 & 3 \\ 1 & 2 & 3 \end{pmatrix}$$

Calculando los determinantes:

$$\det(\mathbf{A}_1) = 7, \quad \det(\mathbf{A}_2) = 14, \quad \det(\mathbf{A}_3) = 21$$

Por lo tanto:

$$x = \frac{7}{7} = 1, \quad y = \frac{14}{7} = 2, \quad z = \frac{21}{7} = 3$$

Ejemplo 3: Cálculo de determinantes

Determinante de orden 2:

$$\det \begin{pmatrix} 2 & -1 \\ 3 & 4 \end{pmatrix} = 2 \cdot 4 - (-1) \cdot 3 = 8 + 3 = 11$$

Determinante de orden 3 (regla de Sarrus):

$$\begin{aligned} \det \begin{pmatrix} 1 & 2 & 3 \\ -1 & 0 & 4 \\ 2 & 1 & -2 \end{pmatrix} &= 1 \cdot 0 \cdot (-2) + 2 \cdot 4 \cdot 2 + 3 \cdot (-1) \cdot 1 - 3 \cdot 0 \cdot 2 - 2 \cdot (-1) \cdot (-2) - 1 \cdot 4 \cdot 1 \\ &= 0 + 16 - 3 - 0 - 4 - 4 = 5 \end{aligned}$$

Propiedad importante: Si una fila (o columna) es combinación lineal de las demás, el determinante es cero. Por ejemplo, en la matriz

$$\begin{pmatrix} 1 & 2 & 3 \\ 2 & 4 & 6 \\ -1 & 0 & 1 \end{pmatrix},$$

la segunda fila es el doble de la primera, por lo que su determinante es 0.

H.10. Aplicaciones en ingeniería y sistemas dinámicos

- **Análisis de circuitos eléctricos:** las leyes de Kirchhoff dan lugar a sistemas lineales que se resuelven por métodos matriciales.
- **Resolución de ecuaciones diferenciales lineales:** los sistemas de EDO lineales se representan en forma matricial $\dot{\mathbf{x}} = \mathbf{Ax} + \mathbf{Bu}$, y la solución involucra exponenciales de matrices.
- **Control moderno:** las ecuaciones de estado y la controlabilidad/observabilidad se estudian mediante el álgebra de matrices.
- **Mecánica estructural:** el método de rigideces utiliza matrices para calcular desplazamientos y fuerzas en estructuras.
- **Procesamiento de señales e imágenes:** las transformadas lineales (Fourier, wavelet, etc.) se apoyan en operaciones matriciales.

Anexo I

LEYES Y PRINCIPIOS FÍSICOS FUNDAMENTALES

I.1. Mecánica clásica

Ley / Principio	Fórmula	Descripción
Primera ley de Newton	$\sum \mathbf{F} = 0$	Un cuerpo en reposo o MRU permanece así si la fuerza neta es cero.
Segunda ley de Newton	$\sum \mathbf{F} = m\mathbf{a}$	La fuerza neta es igual a masa por aceleración.
Tercera ley de Newton	$\mathbf{F}_{12} = -\mathbf{F}_{21}$	Acción y reacción: fuerzas iguales y opuestas.
Ley de gravitación universal	$F = G \frac{m_1 m_2}{r^2}$	Fuerza de atracción entre dos masas.
Cantidad de movimiento	$\mathbf{p} = m\mathbf{v}$	Producto de masa por velocidad.
Conservación del momento	$\sum \mathbf{p}_{\text{inicial}} = \sum \mathbf{p}_{\text{final}}$	En sistema aislado, el momento total se conserva.
Trabajo mecánico	$W = \int \mathbf{F} \cdot d\mathbf{r}$	Energía transferida por una fuerza.
Energía cinética	$E_c = \frac{1}{2}mv^2$	Energía asociada al movimiento.
Energía potencial gravitacional	$E_p = mgh$	Energía por altura en campo gravitacional.
Energía potencial elástica	$E_e = \frac{1}{2}kx^2$	Energía almacenada en un resorte.
Potencia	$P = \frac{dW}{dt} = \mathbf{F} \cdot \mathbf{v}$	Rapidez de transferencia de energía.
Fricción (cinética)	$F_k = \mu_k N$	Fuerza de rozamiento entre superficies en movimiento relativo.
Fricción (estática)	$F_s \leq \mu_s N$	Fuerza máxima de rozamiento antes de que el movimiento comience.
Momento angular	$\mathbf{L} = \mathbf{r} \times \mathbf{p}$	Medida de la inercia rotacional de un cuerpo.
Teorema del trabajo-energía	$W_{\text{neto}} = \Delta E_c$	El trabajo total realizado sobre un cuerpo es igual al cambio en su energía cinética.
Fuerza centrípeta	$F_c = \frac{mv^2}{r}$	Fuerza dirigida hacia el centro de una trayectoria circular.

I.2. Oscilaciones y dinámica

Ley / Principio	Fórmula	Descripción
Ley de Hooke	$F = -kx$	Fuerza proporcional a la deformación.
Amortiguamiento viscoso	$F_d = -c\dot{x}$	Fuerza resistente proporcional a la velocidad.
Masa-resorte-amortiguador	$m\ddot{x} + c\dot{x} + kx = F(t)$	Modelo de segundo orden.
Dinámica rotacional	$\sum \tau = I\alpha$	Análogo rotacional de la segunda ley de Newton.
Torque	$\boldsymbol{\tau} = \mathbf{r} \times \mathbf{F}$	Capacidad de una fuerza para producir rotación.
Energía cinética rotacional	$E_r = \frac{1}{2}I\omega^2$	Energía asociada a la rotación.

I.3. Electricidad y circuitos

Ley / Principio	Fórmula	Descripción
Ley de Ohm	$V = IR$	Tensión proporcional a corriente y resistencia.
Potencia eléctrica	$P = VI = \frac{I^2 R}{V^2} = \frac{R}{V^2}$	Tasa de disipación de energía.
Ley de Kirchhoff (corrientes)	$\sum I_{\text{entra}} = \sum I_{\text{sale}}$	Conservación de la carga en un nodo.
Ley de Kirchhoff (voltajes)	$\sum V = 0$	Conservación de la energía en una malla.
Capacitor	$Q = CV$	Carga proporcional a la tensión.
Corriente en capacitor	$I = C \frac{dV}{dt}$	Corriente proporcional a la variación de tensión.
Inductor	$V = L \frac{dI}{dt}$	Tensión proporcional a la variación de corriente.
Ley de Coulomb	$F = k_e \frac{ q_1 q_2 }{r^2}$	Fuerza entre cargas puntuales.
Ley de Faraday	$\mathcal{E} = -\frac{d\Phi_B}{dt}$	Fem inducida por variación de flujo magnético.
Ley de Lenz	- -	La corriente inducida se opone al cambio de flujo.

I.4. Fluidos y termodinámica

Ley / Principio	Fórmula	Descripción
Presión	$P = \frac{F}{A}$	Fuerza por unidad de área.
Principio de Pascal	$\frac{F_1}{A_1} = \frac{F_2}{A_2}$	La presión se transmite en todas direcciones.
Principio de Arquímedes	$F_B = \rho_f g V_d$	Empuje igual al peso del fluido desalojado.
Ecuación de continuidad	$A_1 v_1 = A_2 v_2$	Conservación del caudal.
Ecuación de Bernoulli	$P + \frac{1}{2} \rho v^2 + \rho g h = \text{constante}$	Conservación de la energía en fluidos.
Primera ley de termodinámica	$\Delta U = Q - W$	Conservación de la energía en sistemas térmicos.
Segunda ley de termodinámica	$\Delta S \geq 0$	La entropía de un sistema aislado nunca disminuye.

I.5. Elasticidad y materiales

Relación	Fórmula	Descripción
Ley de Hooke (elasticidad)	$\sigma = E \varepsilon$	Esfuerzo proporcional a la deformación unitaria.
Modelo lineal de medición	$y = Kx + b$	Relación lineal entre entrada y salida.
Modelo calibrado ideal	$y = Kx$	Modelo sin desplazamiento.

Anexo J

FUNDAMENTOS DE CÁLCULO

Este anexo presenta los conceptos fundamentales del cálculo diferencial e integral, así como sus conexiones con las ecuaciones diferenciales y la transformada de Laplace. Su propósito es ofrecer una visión general que oriente al lector sobre las herramientas matemáticas que subyacen al análisis de sistemas dinámicos.

J.1. ¿Qué es el cálculo?

El cálculo es la rama de las matemáticas que estudia el cambio y la acumulación. Se divide en dos áreas principales:

- **Cálculo diferencial:** estudia la tasa de cambio instantánea (derivadas).
- **Cálculo integral:** estudia la acumulación de cantidades (integrales).

Ambas están unificadas por el Teorema Fundamental del Cálculo, que establece que la derivación y la integración son operaciones inversas.

El cálculo es la base de la física, la ingeniería, la economía y prácticamente todas las ciencias que modelan fenómenos continuos.

J.2. La derivada: cambio instantáneo

Definición intuitiva

La derivada de una función $f(x)$ en un punto x mide cómo cambia f cuando x varía ligeramente. Es la pendiente de la recta tangente a la curva $y = f(x)$ en ese punto.

Definición formal

$$f'(x) = \frac{df}{dx} = \lim_{\Delta x \rightarrow 0} \frac{f(x + \Delta x) - f(x)}{\Delta x}$$

Interpretación física

- La derivada de la posición respecto al tiempo es la **velocidad**: $v = \frac{dx}{dt}$.

- La derivada de la velocidad respecto al tiempo es la **aceleración**: $a = \frac{dv}{dt}$.
- En general, la derivada describe cualquier tasa de cambio: crecimiento poblacional, flujo de calor, corriente eléctrica, etc.

Derivadas de funciones elementales (referencia)

- $\frac{d}{dx}(c) = 0$ (derivada de una constante)
- $\frac{d}{dx}(x^n) = nx^{n-1}$ (regla de la potencia)
- $\frac{d}{dx}(\sin x) = \cos x$
- $\frac{d}{dx}(\cos x) = -\sin x$
- $\frac{d}{dx}(e^x) = e^x$
- $\frac{d}{dx}(\ln x) = \frac{1}{x}$

Reglas de derivación

- **Regla de la suma:** $\frac{d}{dx}[f(x) + g(x)] = f'(x) + g'(x)$
- **Regla del producto:** $\frac{d}{dx}[f(x) \cdot g(x)] = f'(x)g(x) + f(x)g'(x)$
- **Regla del cociente:** $\frac{d}{dx} \left[\frac{f(x)}{g(x)} \right] = \frac{f'(x)g(x) - f(x)g'(x)}{[g(x)]^2}$
- **Regla de la cadena:** $\frac{d}{dx}[f(g(x))] = f'(g(x)) \cdot g'(x)$

J.3. La integral: acumulación

Definición intuitiva

La integral de una función $f(x)$ en un intervalo $[a, b]$ representa el área bajo la curva $y = f(x)$ desde $x = a$ hasta $x = b$.

Definición formal (integral de Riemann)

$$\int_a^b f(x) dx = \lim_{n \rightarrow \infty} \sum_{i=1}^n f(x_i^*) \Delta x$$

donde la suma se toma sobre particiones del intervalo $[a, b]$.

Interpretación física

- La integral de la velocidad respecto al tiempo es el **desplazamiento**: $\Delta x = \int v(t) dt$.
- La integral de la aceleración es el **cambio de velocidad**: $\Delta v = \int a(t) dt$.
- La integral de una tasa de flujo es la **cantidad acumulada**.

Integral indefinida (antiderivada)

La integral indefinida es la operación inversa de la derivada:

$$\int f(x) dx = F(x) + C \quad \text{donde} \quad F'(x) = f(x)$$

La constante C aparece porque la derivada de cualquier constante es cero.

J.4. El Teorema Fundamental del Cálculo

Este teorema es el puente entre el cálculo diferencial y el integral. Consta de dos partes:

Primera parte (existencia de antiderivadas)

Si f es continua en $[a, b]$, entonces la función

$$F(x) = \int_a^x f(t) dt$$

es derivable en (a, b) y $F'(x) = f(x)$. Es decir, **toda función continua tiene una antiderivada**.

Segunda parte (evaluación de integrales)

Si F es cualquier antiderivada de f , entonces

$$\int_a^b f(x) dx = F(b) - F(a)$$

Importancia del teorema

- Permite calcular integrales definidas encontrando antiderivadas.
- Establece que la derivación y la integración son operaciones inversas.
- Es la base de la mayoría de las aplicaciones del cálculo en ciencias e ingeniería.

J.5. ¿Qué es una ecuación diferencial?

Definición general

Una ecuación diferencial es una ecuación que relaciona una función desconocida con una o más de sus derivadas.

Clasificación

- **Ecuación diferencial ordinaria (EDO):** involucra derivadas respecto a una sola variable independiente.
- **Ecuación diferencial parcial (EDP):** involucra derivadas respecto a dos o más variables independientes.

Orden de una ecuación diferencial

El orden es el de la derivada más alta que aparece en la ecuación.

Ejemplos clásicos

- **Crecimiento exponencial:** $\frac{dy}{dt} = ky$ (solución: $y(t) = y_0 e^{kt}$)
- **Oscilador armónico:** $\frac{d^2x}{dt^2} + \omega^2 x = 0$ (solución: $x(t) = A \cos(\omega t) + B \sin(\omega t)$)
- **Ley de enfriamiento de Newton:** $\frac{dT}{dt} = -k(T - T_{\text{amb}})$
- **Ecuación logística:** $\frac{dP}{dt} = rP \left(1 - \frac{P}{K}\right)$
- **Ecuaciones de Maxwell:** gobiernan el electromagnetismo (EDP).
- **Ecuación de Schrödinger:** gobierna la mecánica cuántica (EDP).

Importancia de las ecuaciones diferenciales

- Modelan prácticamente todos los fenómenos dinámicos en la naturaleza.
- Son la base de la física, la ingeniería, la biología y la economía.
- Resolver una ecuación diferencial es encontrar la función que la satisface, lo que generalmente implica integrar.

J.6. La relación entre derivada, integral y ecuación diferencial

Concepto	Relación
Derivada	Describe el cambio instantáneo.
Integral	Describe la acumulación.
Teorema Fundamental	Conecta derivada e integral como operaciones inversas.
Ecuación diferencial	Iguala una función con sus derivadas.
Resolver una EDO	Implica integrar (encontrar la función original).

En este sentido, la derivada y la integral son operaciones inversas, y las ecuaciones diferenciales combinan ambas para modelar sistemas complejos.

J.7. ¿Qué es la transformada de Laplace?

Definición

La transformada de Laplace es una herramienta matemática que convierte funciones del dominio del tiempo (t) al dominio de la frecuencia compleja (s):

$$\mathcal{L}\{f(t)\} = F(s) = \int_0^{\infty} f(t)e^{-st} dt$$

Propiedades fundamentales

- **Linealidad:** $\mathcal{L}\{af(t) + bg(t)\} = aF(s) + bG(s)$
- **Derivada:** $\mathcal{L}\{f'(t)\} = sF(s) - f(0)$
- **Derivada segunda:** $\mathcal{L}\{f''(t)\} = s^2F(s) - sf(0) - f'(0)$
- **Integral:** $\mathcal{L}\left\{\int_0^t f(\tau) d\tau\right\} = \frac{F(s)}{s}$

Transformadas de funciones elementales (referencia)

- $\mathcal{L}\{1\} = \frac{1}{s}$
- $\mathcal{L}\{t^n\} = \frac{n!}{s^{n+1}}$
- $\mathcal{L}\{e^{at}\} = \frac{1}{s-a}$
- $\mathcal{L}\{\sin(at)\} = \frac{a}{s^2 + a^2}$
- $\mathcal{L}\{\cos(at)\} = \frac{s}{s^2 + a^2}$
- $\mathcal{L}\{\delta(t)\} = 1$ (impulso unitario o delta de Dirac)
- $\mathcal{L}\{u(t)\} = \frac{1}{s}$ (escalón unitario de Heaviside)

¿Para qué sirve?

- Convierte ecuaciones diferenciales en ecuaciones algebraicas.
- Facilita la resolución de sistemas lineales con condiciones iniciales.
- Es la base del análisis de sistemas de control y circuitos eléctricos.
- Permite trabajar con funciones discontinuas (usando la función escalón).

Proceso general de resolución

1. Aplicar la transformada de Laplace a la ecuación diferencial.
2. Resolver la ecuación algebraica resultante para $Y(s)$.
3. Aplicar la transformada inversa $\mathcal{L}^{-1}\{Y(s)\}$ para obtener $y(t)$.

J.8. Teoremas importantes del cálculo

Teorema del valor medio para derivadas

Si f es continua en $[a, b]$ y derivable en (a, b) , entonces existe $c \in (a, b)$ tal que

$$f'(c) = \frac{f(b) - f(a)}{b - a}$$

Teorema del valor medio para integrales

Si f es continua en $[a, b]$, entonces existe $c \in [a, b]$ tal que

$$\int_a^b f(x) dx = f(c)(b - a)$$

Teorema de Rolle

Si f es continua en $[a, b]$, derivable en (a, b) y $f(a) = f(b)$, entonces existe $c \in (a, b)$ tal que $f'(c) = 0$.

Regla de L'Hôpital

Si $\lim_{x \rightarrow a} \frac{f(x)}{g(x)}$ es una forma indeterminada $\frac{0}{0}$ o $\frac{\infty}{\infty}$, entonces

$$\lim_{x \rightarrow a} \frac{f(x)}{g(x)} = \lim_{x \rightarrow a} \frac{f'(x)}{g'(x)}$$

siempre que el límite del cociente de las derivadas exista.

Teorema de Taylor

Una función suficientemente diferenciable puede aproximarse localmente por un polinomio:

$$f(x) = f(a) + f'(a)(x - a) + \frac{f''(a)}{2!}(x - a)^2 + \frac{f'''(a)}{3!}(x - a)^3 + \dots$$

Teorema de la función inversa

Si $f'(a) \neq 0$, entonces f tiene una función inversa local alrededor de a , y

$$(f^{-1})'(f(a)) = \frac{1}{f'(a)}$$

Teorema de la función implícita

Si $F(x, y) = 0$ define implícitamente a y como función de x , y $\partial F / \partial y \neq 0$, entonces

$$\frac{dy}{dx} = -\frac{\partial F / \partial x}{\partial F / \partial y}$$

J.9. Resumen conceptual integrado

Concepto	Idea fundamental
Derivada	Mide la rapidez de cambio instantáneo (pendiente de la tangente).
Integral	Mide la acumulación de una cantidad (área bajo la curva).
Teorema Fundamental	Conecta derivada e integral como operaciones inversas.
Ecuación diferencial	Relaciona una función con sus derivadas; modela sistemas dinámicos.
Transformada de Laplace	Convierte ecuaciones diferenciales en algebraicas.
Teorema del valor medio	Garantiza la existencia de un punto con pendiente promedio.
Teorema de Taylor	Permite aproximar funciones mediante polinomios.

J.10. Conexión con el resto de la obra

Los conceptos presentados en este anexo son la base matemática de gran parte del desarrollo del libro. La derivada, la integral y las ecuaciones diferenciales aparecen en:

- **Física:** leyes de Newton, circuitos eléctricos, termodinámica.
- **Control:** funciones de transferencia, estabilidad, respuesta temporal.
- **Sistemas dinámicos:** modelado y análisis.
- **Instrumentación:** respuesta de sensores, acondicionamiento de señales.

La transformada de Laplace es especialmente relevante en el análisis de sistemas lineales invariantes en el tiempo (SLIT), como se estudia en los capítulos de control.

Este anexo sirve como guía conceptual; las fórmulas detalladas y aplicaciones específicas se desarrollan en los capítulos correspondientes.

Anexo K

CONCEPTOS ESENCIALES DE CONTROL

K.1. Función de transferencia

$$G(s) = \frac{Y(s)}{U(s)}$$

K.2. Sistemas elementales

Sistema	Función de transferencia
Ganancia constante	K
Integrador puro	$\frac{K}{s}$
Derivador ideal	Ks
Primer orden	$\frac{K}{\tau s + 1}$
Segundo orden	$\frac{K\omega_n^2}{s^2 + 2\zeta\omega_n s + \omega_n^2}$

K.3. Conexiones de bloques

Configuración	Función de transferencia equivalente
Serie	$G_{eq} = G_1 G_2$
Paralelo (suma)	$G_{eq} = G_1 + G_2$
Retroalimentación negativa	$G_{eq} = \frac{G}{1 + GH}$
Retroalimentación positiva	$G_{eq} = \frac{G}{1 - GH}$

K.4. Lazo cerrado y error

$$T(s) = \frac{Y(s)}{R(s)} = \frac{G(s)}{1 + G(s)}, \quad \frac{E(s)}{R(s)} = \frac{1}{1 + G(s)}$$

K.5. Controladores PID

Tipo	Función de transferencia
P	$C(s) = K_p$
I	$C(s) = \frac{K_i}{s}$
D	$C(s) = K_d s$
PI	$C(s) = K_p + \frac{K_i}{s}$
PD	$C(s) = K_p + K_d s$
PID	$C(s) = K_p + \frac{K_i}{s} + K_d s$

$$u(t) = K_p \left[e(t) + \frac{1}{T_i} \int e(t) dt + T_d \frac{de(t)}{dt} \right]$$

K.6. Ecuación característica y estabilidad

$$1 + G(s)H(s) = 0$$

$$\operatorname{Re}(s_i) < 0$$

K.7. Comportamiento dinámico en el tiempo

Concepto	Relación
Escalón unitario	$R(s) = \frac{1}{s}$
Impulso unitario	$\mathcal{L}\{\delta(t)\} = 1$
Valor final	$\lim_{t \rightarrow \infty} f(t) = \lim_{s \rightarrow 0} sF(s)$
Valor inicial	$\lim_{t \rightarrow 0^+} f(t) = \lim_{s \rightarrow \infty} sF(s)$

Anexo L

ECUACIONES DE MAXWELL

Las ecuaciones de Maxwell son cuatro ecuaciones fundamentales que unifican la electricidad y el magnetismo en una sola teoría: el electromagnetismo. Constituyen la base de la teoría electromagnética clásica y describen el comportamiento de los campos eléctricos y magnéticos en la materia y en el vacío.

L.1. Las cuatro ecuaciones de Maxwell

Nombre	Ecuación	Significado físico
Ley de Gauss (eléctrica)	$\oint \mathbf{E} \cdot d\mathbf{A} = \frac{Q_{\text{enc}}}{\epsilon_0}$	Las cargas eléctricas son la fuente del campo eléctrico.
Ley de Gauss (magnética)	$\oint \mathbf{B} \cdot d\mathbf{A} = 0$	No existen monopolos magnéticos; el flujo magnético es siempre cero.
Ley de Faraday	$\oint \mathbf{E} \cdot d\mathbf{l} = -\frac{d\Phi_B}{dt}$	Un campo magnético variable crea un campo eléctrico.
Ley de Ampère-Maxwell	$\oint \mathbf{B} \cdot d\mathbf{l} = \mu_0 I + \mu_0 \epsilon_0 \frac{d\Phi_E}{dt}$	Las corrientes y los campos eléctricos variables crean campos magnéticos.

L.2. Forma vectorial (diferencial) de las ecuaciones

Ley	Forma vectorial (diferencial)
Gauss (eléctrica)	$\nabla \cdot \mathbf{E} = \frac{\rho}{\epsilon_0}$
Gauss (magnética)	$\nabla \cdot \mathbf{B} = 0$
Faraday	$\nabla \times \mathbf{E} = -\frac{\partial \mathbf{B}}{\partial t}$
Ampère-Maxwell	$\nabla \times \mathbf{B} = \mu_0 \mathbf{J} + \mu_0 \epsilon_0 \frac{\partial \mathbf{E}}{\partial t}$

L.3. Ecuación de onda electromagnética

De las ecuaciones de Maxwell se deduce que los campos eléctrico y magnético se propagan como ondas:

$$\nabla^2 \mathbf{E} = \mu_0 \varepsilon_0 \frac{\partial^2 \mathbf{E}}{\partial t^2}, \quad \nabla^2 \mathbf{B} = \mu_0 \varepsilon_0 \frac{\partial^2 \mathbf{B}}{\partial t^2}$$

donde $c = \frac{1}{\sqrt{\mu_0 \varepsilon_0}}$ es la velocidad de la luz en el vacío.

L.4. Explicación de las ecuaciones de Maxwell

- Unifican la electricidad y el magnetismo en una sola teoría.
- Son compatibles con la teoría de la relatividad.
- Establecen que la velocidad de la luz es constante en el vacío.
- Predicen la existencia de ondas electromagnéticas (luz, radio, microondas, etc).
- Son la base de toda la tecnología moderna: comunicaciones, electrónica, óptica, etc.

L.5. Conexión con el resto de la obra

Las ecuaciones de Maxwell:

- Se expresan en términos de campos vectoriales (Anexo F).
- Involucran derivadas parciales y ecuaciones diferenciales (Anexo J, antiguo I).
- Son la base de la teoría electromagnética (Anexo I, antiguo H).
- Utilizan constantes fundamentales como ε_0 , μ_0 y c .

Las ecuaciones de Maxwell representan una de las culminaciones de la física clásica, el punto en el que las matemáticas y la naturaleza convergen en un mismo lenguaje. Con ellas cerramos este recorrido, desde los fundamentos matemáticos hasta las leyes que describen los fenómenos físicos que rigen la naturaleza. Un recorrido en el que cada ecuación ha sido un puente entre la abstracción y la realidad.



Epílogo 2026

Has concluido este libro.

Al llegar a estas páginas, el autor espera haber cumplido el propósito que dio origen a la obra: presentar los fundamentos de los sistemas dinámicos mediante una secuencia lógica y coherente, demostrando que las matemáticas, las distintas ramas de la física, la instrumentación y la teoría de control no son saberes aislados, sino partes de una misma realidad.

Con frecuencia, estas disciplinas se estudian por separado, como si pertenecieran a mundos independientes. Como consecuencia de esa fragmentación, suelen percibirse como conocimientos excesivamente abstractos o completamente inconexos. Sin embargo, la práctica de la ingeniería demuestra lo contrario: **los problemas no conocen barreras teóricas**. Comprender un sistema exige integrar conceptos, hilar ideas y reconocer que la variedad de métodos, técnicas y procedimientos convergen en un propósito común: **entender para actuar**.

Si eres estudiante, tal vez este libro te haya permitido advertir, antes de terminar tu formación, que las asignaturas no son compartimentos estancos, sino etapas de un mismo proceso intelectual. Si no eres estudiante, pero buscas entender la física, las matemáticas y sus aplicaciones sin incurrir en tecnicismos innecesarios, y al concluir este libro has alcanzado tus objetivos de comprensión, asociación y reflexión, entonces en ambos casos este libro habrá cumplido su propósito.

Aprender matemáticas y física no consiste en memorizar fórmulas, sino en entender sus interrelaciones y principios. Una ecuación aislada parece un mero artificio algebraico; una ecuación comprendida e integrada dentro de un marco conceptual se convierte en un instrumento para interpretar la realidad.

La ingeniería es, ante todo, una disciplina de aplicación del conocimiento orientada a la resolución de problemas. Toda idea que pretende ser una solución se genera a partir de la facultad de observar, razonar y comprender. Visto así, el conocimiento no debe ser solamente una acumulación de datos, sino una estructura sistemática que resulte en un pensamiento de mayor claridad y discernimiento, a fin de entender un mundo donde el orden y el caos coexisten; pero cuyos sistemas pueden ser intencionalmente transformados.

Cada conocimiento adquirido altera la manera en que observamos el entorno y, a partir de ese momento, comprender deja de constituir un ejercicio meramente contemplativo para convertirse en una herramienta de responsabilidad, pues toda herramienta tecnológica lleva consigo una dimensión ética que ninguna ciencia puede ignorar.

A lo largo de este libro se ha presentado el concepto de **transformación como un principio fundamental presente en la física, las matemáticas y la teoría de control**. Una transformación, en su sentido más general, es el proceso mediante el cual un estado o sistema pasa de una representación a otra (por ejemplo, del dominio del tiempo al dominio de la frecuencia) para revelar propiedades que permanecían ocultas en su formulación original.

En este contexto, las representaciones matemáticas del sistema, como la función de transferencia obtenida mediante la transformada de Laplace, permiten describir el comportamiento de un proceso mediante una formulación formal que conserva las dependencias causales relevantes para su análisis, al tiempo que simplifica el tratamiento matemático. Esta simplificación no es arbitraria; responde a la necesidad de hacer manejable la complejidad inherente a los sistemas físicos, los procesos industriales, los sistemas biológicos y los modelos económicos y financieros.

Cuando el sistema ha sido definido, modelado y comprendido, puede ser monitoreado, controlado, mejorado, modificado o evolucionado. Las perturbaciones, el ruido de medición y las incertidumbres paramétricas pueden ser identificados, mitigados o tolerados mediante estrategias de control adaptativas y resilientes. El comportamiento del sistema queda entonces controlado cuando estas variables son anticipadas y ajustadas adecuadamente.

Al concluir este libro, es posible asentar una enseñanza de la ingeniería: **toda representación del sistema no es el fin, sino el medio para intervenir en la realidad con rigor y transformar el sistema con criterio.**

Finalmente, el autor te anima a seguir aprendiendo, pues ningún conocimiento posee carácter definitivo y ninguna explicación es absoluta. Cada nuevo aprendizaje genera nuevas interrogantes y cada solución encuentra siempre nuevos desafíos por resolver.

Quizá el futuro de la vida en la Tierra no dependa únicamente de la capacidad para crear tecnologías incrementalmente más sofisticadas, sino de la aplicación de estas tecnologías bajo preceptos tales como la responsabilidad, la sensatez y el sentido común, de modo que contribuyan a la construcción de un mundo donde el progreso se mida también por su capacidad para preservar la vida, la libertad, la dignidad y la razón.

Para futuras revisiones de este trabajo, así como para consultar otras publicaciones sobre tecnología, ingeniería y filosofía, puede visitar mi blog:

<https://codingforwar.home.blog>

[Página intencionalmente en blanco]

Bibliografía

Las referencias bibliográficas corresponden a las fuentes consultadas durante la elaboración original de este trabajo, entre 2006 y 2007. Se ha incorporado el manual de referencia de LaTeX2e, utilizado en la actualización de esta versión (2026).

Matemáticas

- (1) Dennis G. Zill. *Ecuaciones diferenciales con aplicaciones de modelado*. 6a edición. Thomson. México, 2003.
- (2) Dennis G. Zill. *Cálculo con geometría analítica*. Grupo Editorial Iberoamérica. México, 1987.
- (3) Edwards C. H, Penney E. D. *Ecuaciones diferenciales elementales*. 3a edición. Prentice-Hall. México, 1994.
- (4) Granville. *Cálculo diferencial e integral*. Limusa. México, 2001.
- (5) Grimaldi P. Ralph. *Discrete and combinatorial mathematics*. 3a edición. Addison-Wesley. USA, 1994.
- (6) Marsden E. J, Tromba J. A. *Cálculo vectorial*. 4a edición. Prentice-Hall. México, 1998.
- (7) Sean Mauch. *Advanced mathematical methods for Scientists and Engineers*. 2004.

Física

- (8) Halliday D, Resnick R, Krane S. *Física, volumen 1 y 2*. 4a edición. CECSA. México 1998.
- (9) Young D. H, Zemansky, Sears. *University Physics*. 8a edición. Addison-Wesley. USA, 1992.

Ingeniería eléctrica y de control

- (10) Bolton W. *Ingeniería de control*. 2a edición. Alfaomega. México, 2001.
- (11) Enriquez H. G. *Transformadores y motores de inducción*. 4a edición. Limusa. México, 2001.
- (12) Glover. J. D, Sarma S. M. *Sistemas de potencia análisis y diseño*. 3a edición. Thomson. México 2003.
- (13) Grainger J. J, Stevenson D. W. *Análisis de sistemas de potencia*. McGraw-Hill. México, 1996.
- (14) Guru B. S.; Hizirolu H. R. *Máquinas eléctricas y transformadores*. 3a edición. Oxford. México, 2003.

- (15) Gwyther H. F. G. *Potencia electrónica y electrónica de potencia*. Alfaomega. México 1997.
- (16) Hayt H. W, Kemmerly E. J. *Análisis de circuitos en ingeniería*. 5a edición. McGraw-Hill. México, 2001.
- (17) Hayt W. *Teoría electromagnética*. 5a edición. McGraw-Hill. México, 2003.
- (18) Matthew N. O. *Elementos de electromagnetismo*. 3a edición. Oxford. México, 2003.
- (19) Ogata Katsuhiko. *Ingeniería de control moderna*. 3a edición. Prentice-Hall. México, 1998.
- (20) Smith A. C, Corripio B. A. *Control automático de procesos*. Limusa. México, 1991.

Programación y computación

- (21) Jose A, Javier L. *Autocad V.14. Manual de actualización*. McGraw-Hill. España, 1998.
- (22) Morris M. M. *Arquitectura de computadoras*. 3a edición. Pearson. México 1993.
- (23) Scott Meyers. *Effective C++*. Addison-Wesley. USA, 1996.
- (24) The Programming Research Group. *High-integrity C++ coding standard manual - version 2.2*. USA, 2004.
- (25) Oetiker T, Serwin M, Partl H, Hyna I, Schlegl E. *The Not So Short Introduction to LaTeX2e or LaTeX in 280 minutes*. Nightly version 7.0. May 12, 2025.

```

#include <string>
#include <iostream>

#define TX1 "N~#vr{r!-p#vqn |{-qr-zv9-uvpvr |{-}!!voyr-zv-s|
znpv|{-'nyr{\\"n |{-zv-v{\\"r r!-}| -p|z} r{qr -ry-z#{q|G-"
#define TX2 "NO_3TUZ"
#define TX3 "QO_9ROT9ROU9YOU9ZNOU9R_O9Y_O9T_O9ZON3NON"
#define TX4 "Nqrzn!-n~#vr{r!-} |q#pr{9-} r!r ${9-}\" n{!zv\"r{-
'-\" n{!s| zn{-ry-p|{|pvzv{\\"|9-'-p#n!-nppv|{r!-} |z#r${r{-ry-}
|t r!|- r!}|{|noyr9-yn-$vnovyvqng-qry-s#\\"# |-'-yn-|o!r ${n{pvn-
qr-yr'r!-s#{qnzr{\\"nqn!-r{-} v{pv}v|!-qr-w#!\\"vpvn;"

char fx(char x) { return ((x - 32 - 13 + 95) % 95) + 32; }

std::string T(const char* S) {
    std::string R = S;
    for (char& x : R) if (x >= 32 && x <= 126) x = fx(x);
    return R;
}

int main() {
    std::cout << "--AGRADECIMIENTOS--" << std::endl;
    std::cout << T(TX1) << T(TX2) << std::endl;
    std::cout << "--DEDICATORIA--" << std::endl;
    std::cout << "A:" << T(TX3) << std::endl;
    std::cout << T(TX4) << std::endl;
    std::cout << "--ATTE--" << std::endl;
    std::cout << "Arcenio Brito H." << std::endl;
    return 0;
} //-- C++

/****
Algoritmo: Desplazamiento modular de 40 unidades sobre el
conjunto de caracteres ASCII imprimibles (95 símbolos)

f(x) = ((x - 40) mod 95) + 32
T(S) = f(x); x ∈ [32, 126]

donde:
    S = (c1, c2, ..., cn) es la cadena de texto de entrada
    [32, 126] -> códigos numéricos ASCII imprimibles
    ASCII imprimibles -> (' '=32,'A'=65,...,'a'=97,... '~'=126)
****/

```

Simbolismo de la portada

Q

uetzalcoatl emerge del entramado matemático no como teotl, sino como huella de una verdad antigua: toda civilización, toda casa y todo linaje se alzan sobre la raíz del conocimiento, como el árbol sobre la tierra que lo sustenta.

Toda ciudad de obsidiana y pluma oscila entre el orden que la erige y el caos que la erosiona. Cuando las sombras eclipsan la luz, cuando el canto enmudece y la lanza se vuelve contra la piedra que la sostiene, el equilibrio se quiebra y la tierra tiembla bajo nuestros pies. Solo el conocimiento, flecha que no se desvía; solo la razón, escudo que no se quiebra; solo la disciplina, jade que no se astilla, pueden devolver la piedra a su sitio y restaurar la armonía.

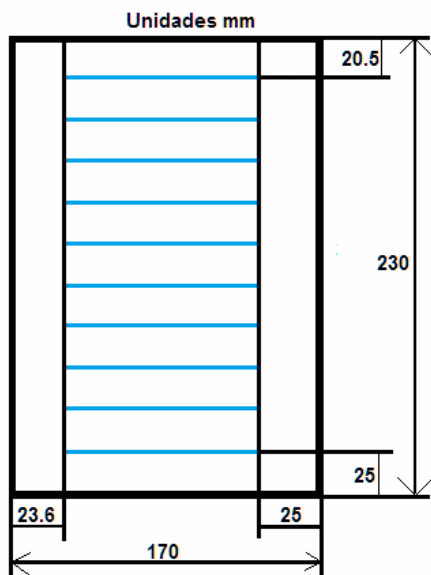
Este libro nace de esa cosmovisión, de esa palabra que los abuelos confiaron al fuego: el orden no se hereda, se cincela sobre la piedra del tiempo. No perdura por inercia, sino por el temple de la voluntad de aquellos que elevan la lanza, levantan el escudo y sostienen la casa del mundo mientras la noche reclama sus muros, para que el amanecer vuelva a encontrarla en pie, firme como el maíz que resiste la sequía y reverdece con la lluvia.

Apéndice E — Configuración de Impresión

Se recomienda imprimir el contenido en papel tamaño carta (216 × 279 mm).

Una vez impreso, las hojas deberán recortarse a un tamaño final de **230 × 170 mm** (alto × ancho), como se muestra en la **Figura E1** y en las capturas de pantalla de configuración.

Figura E1. Configuración de la hoja de impresión

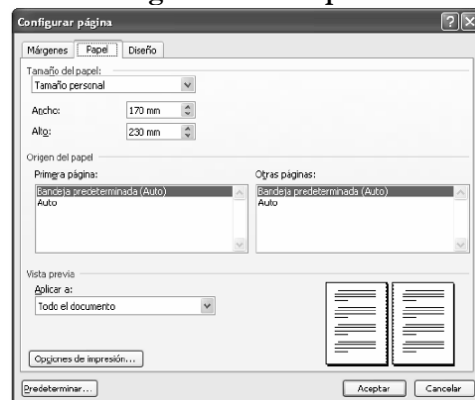


FSD-FMTC • v2.0 • 2026

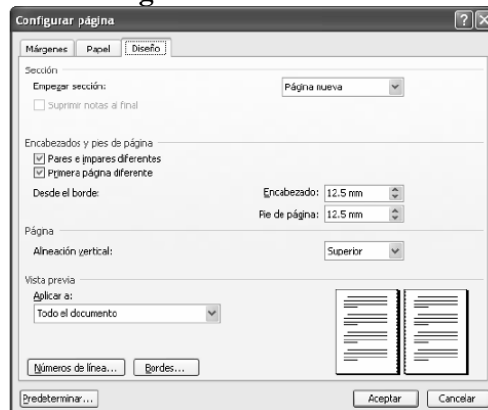
Configuración — Márgenes



Configuración — Papel



Configuración — Diseño



[Página intencionalmente en blanco]

FSD-FMTC • v2.0 • 2026

2026	
2025	
2024	
2023	
2022	
2021	
2020	
2019	
2018	
2017	
2016	
2015	
2014	
2013	
2012	
2011	
2010	
2009	
2008	
2007	
2006	

[Página intencionalmente en blanco]

Ensamblado

Fecha	Versión	Release	Filename
2026-08-04	2.0	r2	FSD-FMTC-v2.0-r2c1.pdf
2026-07-31	1.1	r1	FSD-FMTC-v1.1-r1c1.pdf
2007-05	1.0	r0	PFMTC-v1.0.pdf

FSD-FMTC
Componentes del ensamblado
Actual: v2.0 • r2 • 2026

Documento	A	B	Total	Versión	Settings
A1 - Portada.doc	1	2	2	1.1	S, LC
A2 - Presentacion.doc	1	2	2	1.1	N
A3 - Prologo.doc	1	4	4	1.1	N
A4 - Changelog.doc	1	2	2	1.1	N
A5 - Prefacio.doc	1	2	2	1.0	S, LB
A6 - Portada-2007.doc	1	2	2	1.0	S, LB
A7 - Notacion.doc	1	2	2	1.0++	N
B1 - Tabla de Contenido	1	4	4	1.0++	S, LB
C1 - Metodos Matematicos de Transformacion.doc	1	110	110	1.0	S, NL
C2 - Bloques Funcionales de Sistemas Fisicos.doc	111	146	36	1.0	N
C3 - Modelado de Sistemas Dinamicos.doc	147	172	26	1.0	S, NL
C4 - Sistemas de Contro.doc	173	196	24	1.0	N
C5 - Instrumentos de Control.doc	197	224	28	1.0	N
C6 -Ecuaciones y Leyes Matematicas y Fisicas.tex	1	50	50	2.0	N
D1 - Epilogo.doc	1	4	4	1.1	S, LB
D2 - Bibliografia.doc	1	2	2	1.0++	N
D3 - Dedicatoria.doc	1	2	2	1.0++	--
D4 - PageSetup.doc	1	2	2	1.0	S, LB
D5 - Ensamblado.doc	1	4	4	1.1++	N, RINF
D6 - Contraportada.doc	1	2	2	1.1	S, LC

Todo cuanto existe se manifiesta en el cambio; y en el cambio se revela el principio que permanece. Comprender ese principio es el primer paso para conocer la naturaleza.

$$\begin{pmatrix} 1 & 2 & 3 \\ 0 & 5 & 6 \\ 0 & 3 & 1 \end{pmatrix} = \sum_{n=0}^{\infty} A$$
$$\begin{pmatrix} 1 \\ \hbar \\ \vdots \\ n \end{pmatrix} = \begin{pmatrix} 1 \\ \hbar \\ \vdots \\ 1 \end{pmatrix} \sum_{i=0}^{\infty} \sum_{fnxv} \int \int \frac{dx}{dt}$$



Revisión
v2.0 · 2026
Original
v1.0 · 2007

FSD-FMTC

ClusterBR