

FUNDAMENTOS MATEMÁTICOS DE INVESTIGACIÓN OPERATIVA

Programación Lineal, Hipótesis y Series de Taylor

© Moisés Arreguín Sámano
Hernán Vinicio Villa Sánchez
Ángel Leyva Ovalle
Numa Inain Gaibor Velasco



FUNDAMENTOS MATEMÁTICOS DE INVESTIGACIÓN OPERATIVA

Programación Lineal, Hipótesis
y Series de Taylor

© Moisés Arreguín Sámano

Hernán Vinicio Villa Sánchez

Ángel Leyva Ovalle

Numa Inain Gaibor Velasco



© Autores

Moisés Arreguín Sámano, Docente de la Universidad Estatal de Bolívar, Bolívar, Ecuador.

Hernán Vinicio Villa Sánchez, Docente Investigador Universidad Técnica de Ambato, Facultad de Contabilidad y Auditoría; Ingeniero en Administración de Empresas; Magíster en Formulación, Evaluación y Gestión de Proyectos Sociales y Productivos; Ambato, Ecuador.

Ángel Leyva Ovalle, Docente Investigador Universidad Autónoma Chapingo (UACH-DiCiFo), México. Ingeniero forestal con orientación en economía y ordenación por la Universidad Autónoma Chapingo (UACH), con experiencia en la División de Ciencias Forestales (DiCiFo), específicamente en el Departamento de productos forestales.

Numa Inain Gaibor Velasco, Docente Investigador en la Escuela de Ingeniería en Riesgos de Desastres, Universidad Estatal de Bolívar (UEB), Ecuador. Ingeniero en Administración de Desastres y Gestión del Riesgo por la UEB (2019) y Magíster en Gestión de Riesgos con mención en Manejo de la Respuesta a Desastres por la Universidad Internacional SEK.



CASA EDITORA DEL POLO

Casa Editora del Polo - CASEDELPO CIA. LTDA.
Departamento de Edición

Editado y distribuido por:

Editorial: Casa Editora del Polo
Sello Editorial: 978-9942-816
Manta, Manabí, Ecuador. 2019
Teléfono: (05) 6051775 / 0991871420
Web: www.casadelpo.com
ISBN: 978-9942-684-24-0
DOI: <https://doi.org/10.23857/978-9942-684-24-0>

© Primera edición
© Marzo - 2025
Impreso en Ecuador

Revisión, Ortografía y Redacción:

Lic. Jessica M. Mero Vélez

Diseño de Portada:

Michael J. Suárez-Espinar

Diagramación:

Ing. Edwin A. Delgado-Veliz

Director Editorial:

Lic. Henry D. Suárez Vélez

Todos los libros publicados por la Casa Editora del Polo, son sometidos previamente a un proceso de evaluación realizado por árbitros calificados.

Este es un libro digital y físico, destinado únicamente al uso personal y colectivo en trabajos académicos de investigación, docencia y difusión del Conocimiento, donde se debe brindar crédito de manera adecuada a los autores.

© **Reservados todos los derechos.** Queda estrictamente prohibida, sin la autorización expresa de los autores, bajo las sanciones establecidas en las leyes, la reproducción parcial o total de este contenido, por cualquier medio o procedimiento, parcial o total de este contenido, por cualquier medio o procedimiento.

Comité Científico Académico

Dr. Lucio Noriero-Escalante
Universidad Autónoma de Chapingo, México

Dra. Yorkanda Masó-Dominico
Instituto Tecnológico de la Construcción, México

Dr. Juan Pedro Machado-Castillo
Universidad de Granma, Bayamo. M.N. Cuba

Dra. Fanny Miriam Sanabria-Boudri
Universidad Nacional Enrique Guzmán y Valle, Perú

Dra. Jennifer Quintero-Medina
Universidad Privada Dr. Rafael Bellosó Chacín, Venezuela

Dr. Félix Colina-Ysea
Universidad SISE. Lima, Perú

Dr. Reinaldo Velasco
Universidad Bolivariana de Venezuela, Venezuela

Dra. Lenys Piña-Ferrer
Universidad Rafael Bellosó Chacín, Maracaibo, Venezuela

Dr. José Javier Nuñez-Castillo
Universidad Cooperativa de Colombia, Santa Marta,
Colombia

Constancia de Arbitraje

La Casa Editora del Polo, hace constar que este libro proviene de una investigación realizada por los autores, siendo sometido a un arbitraje bajo el sistema de doble ciego (peer review), de contenido y forma por jurados especialistas. Además, se realizó una revisión del enfoque, paradigma y método investigativo; desde la matriz epistémica asumida por los autores, aplicándose las normas APA, Sexta Edición, proceso de anti plagio en línea Plagiarisma, garantizándose así la científicidad de la obra.

Comité Editorial

Abg. Néstor D. Suárez-Montes
Casa Editora del Polo (CASEDELPO)

Dra. Juana Cecilia-Ojeda
Universidad del Zulia, Maracaibo, Venezuela

Dra. Maritza Berenguer-Gouarnaluses
Universidad Santiago de Cuba, Santiago de Cuba, Cuba

Dr. Víctor Reinaldo Jama-Zambrano
Universidad Laica Eloy Alfaro de Manabí, Ext. Chone

Contenido

PROLOGO.....	9
INTRODUCCION.....	10
CAPITULO 1	
MODELADO CON PROGRAMACIÓN LINEAL (PL).....	13
1.1.Generalidades.....	14
1.2.Juegosdeempresas.....	15
1.3. Modelo Metaltec.....	16
1.3.1.Metaltec software.....	19
1.4.DesarrollodelmodelodeProgramaciónLineal.....	23
1.5. Diseño experimental.....	26
1.5.1. Análisis de varianza de una variable.....	26
1.5.2. Tipos de modelos.....	27
CAPITULO II	
HIPÓTESIS.....	40
2.1. Generalidades.....	41
2.2. Tipos de Hipótesis.....	44
2.2.1. De investigación o de trabajo.....	44
CAPITULO III	
SERIE DE TAYLOR.....	70
3.1. Generalidades.....	71
3.2.EjemplosdelaseriedeTaylor.....	71

3.2.1.Serie Geométrica.....	71
3.2.2.Serie Exponencial.....	72
3.2.3. Serie Seno.....	73
3.2.4. Representación matemática.....	73
3.4. Función Coseno.....	75
3.5. Vistas isométricas de un punto	86

CAPITULO IV

DERIVADAS PARCIALES DE ORDEN SUPERIOR	92
4.1. Definición.....	93
4.2. Representación	93
4.3.LaRegladelaCadena.....	94
4.4. Cálculo Integral por Partes con Regla de Cadena.....	97
4.4.1. Las trazas o cortes.....	97
4.4.2. Ecuación Paramétrica.....	98
4.5. Superficies Cuadráticas.....	101
BIBLIOGRAFÍA.....	121

La investigación operativa es un campo de la matemática aplicada que juega un papel crucial en la resolución de problemas complejos en la toma de decisiones. A lo largo de las últimas décadas, la necesidad de optimizar recursos, gestionar la incertidumbre y modelar fenómenos reales con una precisión rigurosa ha impulsado el desarrollo de herramientas matemáticas avanzadas. Este libro, *Fundamentos Matemáticos de Investigación Operativa: Programación Lineal, Hipótesis y Series de Taylor*, tiene como objetivo proporcionar una base sólida y comprensible sobre estos métodos, enfocados en su aplicación en problemas reales de investigación y análisis.

La obra está dirigida tanto a estudiantes que inician su camino en la investigación operativa como a profesionales que buscan ampliar su comprensión de las técnicas fundamentales que sustentan esta disciplina. A través de tres capítulos bien definidos, el lector explorará desde el modelado de problemas utilizando programación lineal hasta la construcción y aplicación de hipótesis estadísticas, y la formulación de la serie de Taylor como una herramienta poderosa en el análisis matemático.

A lo largo de este libro, se ha intentado que cada tema sea accesible y útil para aquellos que se adentran por primera vez en estos conceptos, pero también lo suficientemente detallado para satisfacer a quienes ya tienen un conocimiento previo. La matemática es una herramienta poderosa, y en este libro se pretende mostrar cómo puede ser utilizada para resolver problemas reales, optimizar procesos y hacer que las decisiones basadas en datos sean más efectivas.

En el mundo moderno, la toma de decisiones informadas es fundamental para el éxito de cualquier organización. Desde la distribución de recursos en una fábrica hasta la gestión de inventarios, pasando por la optimización de rutas logísticas y la planificación financiera, las técnicas matemáticas desempeñan un papel clave en la mejora de la eficiencia y la eficacia de los sistemas. La investigación operativa, como subcampo de las matemáticas aplicadas, se dedica a desarrollar modelos y métodos que permitan enfrentar estos problemas, utilizando herramientas cuantitativas y técnicas computacionales avanzadas.

El libro Fundamentos Matemáticos de Investigación Operativa: Programación Lineal, Hipótesis y Series de Taylor, tiene como objetivo proporcionar a los lectores una comprensión sólida de tres áreas fundamentales dentro de la investigación operativa y la estadística: la programación lineal, el análisis de hipótesis y las series de Taylor. A través de una combinación de teoría y ejemplos prácticos, buscamos ofrecer a los estudiantes y profesionales una introducción completa a estas herramientas, mostrándoles cómo aplicarlas en contextos reales.

La programación lineal (PL) es una de las técnicas más poderosas y ampliamente utilizadas en la investigación operativa. Este capítulo comienza con una introducción a los principios básicos de la PL, proporcionando una visión general de cómo se utiliza para modelar problemas de optimización donde los recursos son limitados. A lo largo del capítulo, se exploran conceptos clave como las restricciones, las funciones objetivo, y las soluciones óptimas. También se presenta el Modelo Metaltec, un caso de estudio práctico que ilustra cómo estas herramientas matemáticas pueden ser aplicadas a problemas reales, y se introduce el software Metaltec, que facilita la implementación de estos modelos. Además, se abordan técnicas de diseño experimental, permitiendo la realización de pruebas controladas para validar hipótesis y

mejorar la toma de decisiones en entornos dinámicos.

La formulación de hipótesis es esencial en cualquier proceso de investigación científica, ya que permite establecer conjeturas que luego serán probadas mediante experimentación y análisis estadístico. Este capítulo explora los diferentes tipos de hipótesis, diferenciando entre hipótesis de trabajo y estadísticas. Se analiza cómo las hipótesis sirven para orientar investigaciones y establecer bases para la comparación de tratamientos o variables, utilizando herramientas como el análisis de varianza (ANOVA). Los lectores aprenderán cómo desarrollar y probar hipótesis estadísticas, lo que es clave para el análisis de datos en diversos campos, desde la investigación clínica hasta los estudios de mercado.

La serie de Taylor es una técnica matemática que permite aproximar funciones complejas mediante polinomios, facilitando su análisis y computación. En este capítulo, se presenta la teoría de la serie de Taylor, mostrando cómo puede utilizarse para aproximar funciones geométricas, exponenciales y trigonométricas. La serie de Taylor es una herramienta esencial en el análisis numérico, la ingeniería y las ciencias aplicadas, ya que permite trabajar con funciones complicadas de una manera más manejable. A través de ejemplos y representaciones matemáticas, los lectores aprenderán cómo se construye y utiliza esta serie en diversas aplicaciones.

Este libro está diseñado para ser una herramienta educativa que sirva como base sólida en la comprensión de los fundamentos matemáticos que sustentan la investigación operativa moderna. A lo largo de los capítulos, el enfoque se mantiene práctico, con ejemplos aplicados que facilitan la comprensión de los temas tratados y muestran cómo las técnicas matemáticas pueden ser utilizadas en situaciones reales. Nuestro objetivo es que los lectores no solo adquieran

conocimiento teórico, sino también habilidades para aplicar estos métodos de manera efectiva en sus áreas de interés profesional o académico.



CAPITULO 1

MODELADO CON PROGRAMACIÓN LINEAL (PL)

1.1. Generalidades

Entre las herramientas y técnicas de enseñanza y aprendizaje en alumnos de pregrado y postgrado y la educación en ingeniería de la administración, los juegos de simulación de empresas están tomando un mayor protagonismo porque es una técnica que incluye no sólo el contenido temático de una disciplina, sino también contenido de motivación, de actitud y sociales que dan el animador una herramienta docente con mayores posibilidades de éxito porque los estudiantes o los jugadores participan en el juego de la iteración con sus colegas, la investigación de herramientas y de experimentar el proceso de toma de decisiones con menor riesgo, en realidad, pero ciertas características del mismo, que te hace aprender jugando.

Esta investigación continúa el trabajo iniciado en el proyecto Metaltec - Juego de empresas dedicadas a la cualificación de los directivos de las Micro y Pequeñas Industrias (Rigonazgo 2007) citado por (Adams, Benitez, Guidek, & Domínguez, 2016) donde se procedió a realizar una mejora o perfeccionamiento de las bases del juego, la incorporación de metodologías, la inclusión de temas en el modelo, la estandarización y la automatización de los procedimientos y mecanismos en el juego Metaltec con el objetivo de crear un módulo para el procesamiento de programación lineal (PL).

Los temas aplicados en esta investigación y en las grupos de jugadores sucesivas son de ingeniería economía, la gestión de la producción, la aplicación de la programación lineal, el análisis y proyección de la demanda, el uso de flujo de caja, costeo variable, las decisiones conjuntas y conceptos integrados que se prueban conjuntamente o por separado, aunque nos limitamos en este artículo a presentar los resultados específicos de un módulo de programación lineal que forma parte del modelo de juego de empresa Metaltec consistente en una metodología para enseñanza de programación lineal.

1.2. Juegos de empresas.

En la literatura la conceptualización de juegos de empresa tiene varios enfoques, entre ellas podemos decir que es una técnica de enseñanza y aprendizaje basados en la simulación, es decir, la utilización de modelos para el estudio de los problemas reales de naturaleza compleja (Ornellas et al 2008), una actividad de formación estructurada, con un objetivo de aprendizaje (Kirby, 1995), también se define como una actividad previamente planificada para afrontar los desafíos para los jugadores

(Gramigna 2000), un ejercicio de toma de decisiones en torno a un modelo de operación de negocios (Santos 2003), y también las técnicas de simulación que transporte a los participantes a las situaciones específicas de la empresa, la mejora de las habilidades técnicas, la comunicación y las relaciones personales (Adams, Benitez, Guidek, & Domínguez, 2016).

En el nivel más práctico u operativo, el concepto de juegos de negocios, está basado en modelos matemáticos desarrollados para simular ciertos ambientes de las variables clave de estos entornos (Kopittke 1992), "son abstracciones de desarrollos matemáticos simplificados relacionados con el mundo de los negocios" (Santos, 2003), y "se basa en un modelo matemático específico que abarca las características físicas, tecnológicas, financieras, económicas, de una empresa real en una forma simplificada" (Rabenschlag, 2005) citado por (Adams, Benitez, Guidek, & Domínguez, 2016).

Para Kirchhof (2006) un juego puede ser descrito como modelos matemáticos de simulación que ayuden en la formación de personas vinculadas a áreas que proporcionan pruebas de las estrategias empresariales, las decisiones y evaluación de las mismas.

Sobre la base de estas definiciones, juegos de negocios son

modelos matemáticos para simular una situación empresarial real expresado en forma simplificada, con el objetivo de proporcionar a los jugadores los conceptos técnicos y metodologías de aprendizaje a través de los mecanismos de aplicación en apoyo del partido (Adams, Benitez, Guidek, & Domínguez, 2016)..

En la literatura se observa que hay una importante producción intelectual vinculada a la empresa de desarrollo de juegos en Brasil, como en el mundo, hace tan poco de los juegos desarrollados en PPGEF / UFSM, el juego de empresa de basado en el método de la unidad la producción de Endeavor (Kirchhof 2006).

El juego JogABC, que ilustra a los jugadores como una empresa utiliza costeo ABC para mejorar sus procesos de producción (Rossato 2006), juego de simulación de inversiones en el mercado financiero (Portes, 2007), modelo matemático de Análisis de Inversiones de la Sociedad (De Paula Reis 2008) y también el modelo utilizado en este estudio Metaltec - Juego de empresa se centró en la calificación de los administradores de las industrias pequeñas y microempresas (Rigonazgo 2007) que proporciona un modelo de información de carácter general basado en el comportamiento real de la industria de aberturas y afines de la región de Santa Maria-RS (BR)se caracteriza por estar constituido por Pymes.

1.3. Modelo Metaltec

Metaltec es un juego de micro y pequeña empresa, desarrollado en el Programa de Posgrado en Ingeniería de Producción de la Universidad Federal de Santa María, basada en la red de herrerías PYMES Unimetal, situado en la Ciudad de Santa Maria, Rio Grande do Sul, Brasil. Metaltec simula una rama industrial de herrería compuesta por PYMES, la fabricación de productos se realiza bajo pedido.

Varias empresas operan en esta demanda del mercado

y compiten entre sí, a través de las decisiones adoptadas en cada corrida (Rigodanzo, 2007). La simulación tiene como objetivo reproducir las condiciones reales de la gestión de dichas empresas.

Las áreas de producción, las finanzas, recursos humanos y comercialización son considerados en el modelo, y todo jugador debe tomar decisiones racionales por lo tanto debe realizar un estudio detallado de los impactos de estas decisiones en los resultados empresariales que se lleva a cabo.

Existen 4 productos, de los cuales los jugadores no pueden diferenciarse por la calidad, los productos son: las ventanas, rejas, portones de contrapeso y puertas de seguridad. Los aspectos financieros contemplados en el juego para animar a los jugadores para optimizar los recursos para maximizar sus ingresos.

A medida que el juego se desarrolla la toma de decisiones está influenciada por dos factores principales (Adams, Benitez, Guidek, & Domínguez, 2016):

- Los resultados adoptados en períodos anteriores y que sean conocidos por todas las empresas;
- Las perspectivas de los resultados que sus propias decisiones o de otros jugadores pueden aparecer en el modelo de simulación en ejercicios futuros. En este segundo factor variable el riesgo está presente, y en la mayoría de los casos, los jugadores deben ejecutar o desarrollar alguna técnica o metodología apropiada a los riesgos a que están expuestos.

La dinámica del juego es sencilla, el animador (profesor) propone las condiciones del juego, estableciendo las reglas, fechas y componentes del juego, la figura siguiente presenta el flujo de trabajo del juego Metaltec.

El juego de empresas Metaltec, simula el ambiente de

una micro y pequeña empresa, donde los líderes (el equipo), dividen las tareas de gestión, cada miembro del equipo tiene un conjunto de tareas relacionadas y coordinar las decisiones con sus colegas.

Los equipos deben definir los objetivos para su empresa en una fase inicial del juego, planteando las estrategias de acción que deben adoptarse para alcanzar los objetivos propuestos. Ellos deberían evaluar su posición sobre los objetivos establecidos, debatir sobre el futuro de la compañía y su posicionamiento con relación a la competencia.

La toma de decisiones debe ser planificada, pues un mal diagnóstico de problema puede dar lugar a consecuencias inesperadas, incluso si está seguro de que la aplicación de soluciones aparentemente correcta.

Figura 1.1. Algoritmo Modelo Metaltec



Fuente: (Adams, Benitez, Guidek, & Domínguez, 2016)

El juego incluye una serie de decisiones (Figura 1.2) que abarcan la compra de materias primas (hierro, aluminio, cerraduras y accesorios), el mantenimiento y la inversión en activos fijos, la contratación y el despido de funcionarios, o la capacidad de producción de máquinas y RRHH y las inversiones.

El software entrega informes generales y particulares con información del mercado y las empresas, también revistas o periódicos con cambios en el mercado. Mientras que recibe en cada jugada las decisiones de cada equipo o empresa.

Figura 1.2. Juego Metaltec



Fuente: (Adams, Benitez, Guidek, & Domínguez, 2016)

1.3.1. Metaltec software

El juego Metaltec software, es un sistema de información que incorpora metodologías para la gestión del juego de empresas Metaltec con el objetivo fundamental de proporcionar al docente o coordinador, la planificación formal de la administración del juego, permitiendo la incorporación de instrumentos teóricos y prácticos de ingeniería de producción y, por último, la automatización de la dinámica del juego en los aspectos operacionales, tales como la recepción de las decisiones, informes y la transmisión de información y consulta con los jugadores, la siguiente figura muestra la interfase del juego de empresa.

Figura 1.3. Interfase del Juego Metaltec

The screenshot displays the 'Metaltec' game interface, which is a complex financial simulation tool. It features several data entry and display sections:

- Top Section:** Includes fields for 'ID de empresa', 'Periodo', and 'Finado'. There are also dropdown menus for 'Ejercicio' and 'Cuentas de balance'.
- Left Panel:** Contains a 'Cuentas de balance' section with a tree view showing assets, liabilities, and equity. Below this is a 'Cuentas de resultado' section with a similar tree view.
- Main Area:** Divided into several tables:
 - RECURSOS DE INVERSIÓN:** A table with columns for 'CANTIDAD', 'CANTIDAD', 'CANTIDAD', and 'CANTIDAD'.
 - FINANCIACIÓN:** A table with columns for 'CANTIDAD', 'CANTIDAD', 'CANTIDAD', and 'CANTIDAD'.
 - ESTABLECIMIENTO EN EL TIEMPO:** A table with columns for 'CANTIDAD', 'CANTIDAD', 'CANTIDAD', and 'CANTIDAD'.
 - RESULTADOS FINANCIEROS:** A table with columns for 'CANTIDAD', 'CANTIDAD', 'CANTIDAD', and 'CANTIDAD'.
 - PLANIFICACIÓN DE LA PRODUCCIÓN:** A table with columns for 'CANTIDAD', 'CANTIDAD', 'CANTIDAD', and 'CANTIDAD'.
- Bottom Section:** Includes a 'Modifica Cuentas' button and a 'Salir' button.

Fuente: (Adams, Benitez, Guidek, & Domínguez, 2016)

La metodología de este estudio es la descripción de tareas, técnicas, etapas y procedimientos realizados para elaborar, proponer, evaluar y validar el modelo juego de empresas Metaltec, por lo que este parte se presenta la caracterización de la investigación, los mecanismos y técnicas para crear el modelo, las fases de investigación y las aplicaciones posteriores del modelo.

La investigación se ordenó y desarrollo considerando las siguientes etapas (Adams, Benitez, Guidek, & Domínguez, 2016):

1. Preparación: fue la etapa inicial y en los que hemos aprendido el juego, fui primero en la disciplina, Temas Especiales en Ingeniería de Producción: los juegos de negocios PPGEF-UFSM, que se aplicó en el primer juego en la búsqueda de Rigonazgo, 2007. En este contexto, continúa con esta línea de investigación. Entre los problemas observados es la normalización y la mejora de las metodologías para un juego muy interesante como el Metaltec, a partir de la cual se define:

- Definición del Problema de Investigación
- Definición de los objetivos generales y específicos y los métodos para construir y poner en práctica el modelo.

- Definición de la hora de crear el juego y la disponibilidad de clases para su aplicación.

2. Definición: La definición de la investigación fue una revisión importante sobre todo en los juegos de las empresas existentes en Brasil y cómo se aplican.

- Revisión de la Literatura de Juegos de empresas, las metodologías utilizadas en la creación y normalización de los juegos.

- Revisión de la Literatura y las aplicaciones de software para juegos en la investigación académica.

3. Realización preliminar: fue la construcción del modelo, algunos conceptos y metodologías utilizadas.

- Definición del modelo propuesto, en particular la metodología aplicada en el modelo, el mecanismo de normalización y de su funcionamiento. -También se definen todas las cuestiones o conceptos en el modelo, tratando de incorporar o mejorar las herramientas simples y añadir más compleja.

- Comenzar la aplicación del modelo en diferentes clases, cada una de las características, especialmente en el aspecto del conocimiento, y cada uno de ellos con una solicitud por separado.

4. Producto final: el final de ejecución se analiza las solicitudes, cada una según la información disponible y la necesidad de cada uno de ellos.

- Sistema de Análisis Metaltec donde los problemas se consideraron en la demanda, el comportamiento de las variables, la observación de las necesidades y el comportamiento de los animales es también la opinión de los jugadores.

- Mejoras y propuestas de modelo: a partir de toda la

información disponible, y una vez que las aplicaciones se realizan algunas mejoras en el funcionamiento del modelo, tratando de incorporar una vez para cada aplicación.

Los procedimientos para la construcción del modelo son:

- Estudiar y analizar el modelo de herrería en el modelo de Microsoft Excel.
- Establecimiento del diagrama de flujo de cómo explicar la secuencia de instrucciones en algoritmos y programas.
- Construcción de tablas de la base de datos en Excel para analizar el comportamiento del sistema.
- Diseño del sistema visual. Establecimiento de principios básicos y generales, en lenguaje Visual Basic Project 6.0 y la base de datos y los formularios de solicitud están diseñados con Excel.
- Programación y definición de las variables.
- Definición de los procedimientos e instrucciones.
- Definiciones y redacción de manuales de usuario, los jugadores, las herramientas y el sistema de apoyo.
- La mejora continua y las pruebas de funcionamiento durante la ejecución del juego.

Debido a la complejidad del modelo y los elementos utilizados en esta investigación, modelado, metodologías y temáticas globales, las aplicaciones se realizan mediante el desarrollo de la construcción del modelo en el que se evalúan los distintos elementos y en parte a la discreción de los investigadores teniendo en cuenta la conveniencia y la disponibilidad de clases la solicitud.

La primera aplicación del modelo se evaluó el desempeño del modelo original que la metodología de estudio que representa el flujo de trabajo del juego, la decisión de examinar

las formas de entrada automática de datos y problemas de aplicación de la partida. Todo esto por la observación de los problemas y los elementos encontrados en el juego de ir y analizar el comportamiento de las variables.

El juego se llevó a cabo en la clase de posgraduados, segundo semestre de 2007 - PPGE - UFSM - la disciplina de los negocios de juegos con 6 equipos (uno es el regulador) y 14 jugadores (Adams, Benitez, Guidek, & Domínguez, 2016).

En la segunda aplicación del juego se centró en la incorporación de técnicas de investigación operativa específicamente para programación lineal, la construcción del sistema de apoyo a la decisión para determinar la demanda de herramientas para el análisis de tendendescuento, el flujo de caja, ingresos, inversiones.

El procedimiento es pedir a los jugadores para crear un modelo que puede ser implementado en el juego con el Metaltec y evaluar la posibilidad de incluir en el juego.

El juego se ha aplicado a una clase de la disciplina de la investigación operativa de la carrera de licenciatura en Administración de Empresas de la Universidad Nacional de Misiones - Argentina, el primer semestre de 2008 con 6 equipos (uno es el regulador) y 18 jugadores.

En esta segunda aplicación se procedió a evaluar a los alumnos en el aprendizaje de PL, presentando una clase previa de optimización mediante el análisis simbólico y la resolución de situaciones de varios casos de estudio (asignación, mezcla, transporte, fuerza de ventas) con el uso Solver de Excel. Con esta metodología se evaluó la aplicación de Solver en el modelo Metaltec.

1.4. Desarrollo del modelo de Programación Lineal:

En Investigación Operativa, los problemas de programación lineal (LP) son problemas de optimización en la función

objetivo y las restricciones son lineales.

La programación lineal es una técnica importante de la optimización y muchos de los problemas prácticos en investigación de operaciones se pueden expresar como problemas de programación lineal. Algunos casos especiales de programación lineal, tales como problemas de redes, transporte y asignación se aplican a problemas de producción lo que ha generado un gran desarrollo en algoritmos especializados para sus soluciones como también innumerables aplicaciones en la gestión de la producción.

La Programación Lineal es un procedimiento o algoritmo matemático mediante el cual se resuelve un problema indeterminado, formulado a través de ecuaciones lineales, optimizando la función objetivo, también lineal. Consiste en optimizar (minimizar o maximizar) una función lineal, denominada función objetivo, de tal forma que las variables de dicha función estén sujetas a una serie de restricciones que expresamos mediante un sistema de inecuaciones lineales (Adams, Benitez, Guidek, & Domínguez, 2016).

La opción Solver en Excel se puede utilizar para resolver la optimización lineal y no lineal, sujeta a restricciones lineales que contienen a las variables de decisión. Solver se puede utilizar para resolver problemas con un máximo de 200 variables de decisión, las limitaciones de 100 y 400 las limitaciones implícitas simple (límites inferior y superior y / o restricciones en las variables de decisión conjunto).

En cuanto al modelo general de programación lineal propuesto a los alumnos, se ha planteado un problema general donde pueden plantear tanto una maximización de ganancias como una minimización de costos ya que el modelo Metaltec permite calcular tanto los costos como las utilidades de la empresa y sujeta a las restricciones de recursos humanos (técnicos), maquinarias, materias primas disponibles, financiero, entre otros.

La siguiente figura muestra el modelo general propuesto, dejando la absoluta libertad a los mismos para plantear.

Figura 1.4. Modelo general propuesto

$$Z = \sum_{j=1}^m C_j \cdot X_j \quad \text{Función Objetivo}$$

Sujeto a las siguientes restricciones

$$\left\{ \sum_{j=1}^m a_{ij} \cdot X_j (\leq, \geq, =) b_i \right\}_{i=1}^n$$

$$X_j \geq 0, \text{ todo } j.$$

Fuente: (Adams, Benitez, Guidek, & Domínguez, 2016)

El modelo Metaltec de programación lineal encarado por los alumnos es un modelo con 4 variables “Xi” que permiten calcular la cantidad de cada producto a elaborar por la empresa y que para la mayoría de los grupos se representa con la figura siguiente, presentada más abajo (Adams, Benitez, Guidek, & Domínguez, 2016).

Figura 1.5. Modelo Metaltec de programación lineal

Función Objetivo	
	$\sum_{j=1}^4 C_j \cdot X_j$
Sujeto a las restricciones	
Cantidad max. Prod 1	$\sum_{j=1}^4 A1j \cdot X1j = b1 \leq B1$
Cantidad max. Prod 2	$\sum_{j=1}^4 A2j \cdot X2j = b2 \leq B2$
Cantidad max. Prod 2	$\sum_{j=1}^4 A3j \cdot X3j = b3 \leq B3$
Cantidad max. Prod 3	$\sum_{j=1}^4 A4j \cdot X4j = b4 \leq B4$
Cantidad max. Hierro	$\sum_{j=1}^4 A5j \cdot X5j = b5 \leq B5$
Cantidad max. Aluminio	$\sum_{j=1}^4 A6j \cdot X6j = b6 \leq B6$
Cantidad max. accesorios	$\sum_{j=1}^4 A7j \cdot X7j = b7 \leq B7$
CMO Horas	$\sum_{j=1}^4 A8j \cdot X8j = b8 \leq B8$
CMH Horas	$\sum_{j=1}^4 A9j \cdot X9j = b9 \leq B9$
No negatividad	$\sum_{j=1}^4 Xj \geq 0$

Fuente: (Adams, Benitez, Guidek, & Domínguez, 2016)

1.5. Diseño experimental

Según (Padrón, 1996), debido a la variabilidad de variables, para evitar que un tratamiento sea favorecido o puesto en desventaja en forma sistemática en sus repeticiones, Fisher ideó la técnica de aleatoriedad, cuya finalidad es dar una estimulación insesgada del error experimental. Es decir, la aleatoriedad tiende a destruir la correlación entre errores y hacer válidas las pruebas de significación. El ejemplo más palpable es la rifa de un objeto, pues si se colocan papeles o fichas numeradas en un recipiente y se supone que están completamente mezcladas, cualquier secuencia en que salgan se considera aleatoria. Cuando el investigador tiene pocos tratamientos recurre a esta técnica. No obstante, es recomendable usar una tabla de números aleatorios.

1.5.1. Análisis de varianza de una variable

El Análisis de Varianza, ANOVA, ADEVA, ANDEVA, ANVA, AVAR, ANOVA de Fisher, Análisis de Varianza de Fisher ó, según terminología inglesa, Analysis of Variance es una prueba utilizada para contrastar la hipótesis que varias medias son iguales debido al uso de la distribución F de Fisher como parte de su contraste.

Fueron ideadas por Sir Ronald Aylmer Fisher en 1925 y son una de las pruebas más importantes de la estadística actual. Es una extensión de la prueba T para dos muestras independientes. Por tanto, sirve para estudiar la incidencia de una variable nominal u ordinal (categórica) con más de las alternativas, en una variable de intervalo o de razón (cuantitativa).

Si las medias de las poblaciones son iguales, las medias de muestras de grupos serán parecidas, por lo que las únicas diferencias serán las atribuidas al azar y el valor del estimador F será 1. Si las medias poblacionales son diferentes, el valor será mayor que 1, tal que será más alto conforme mayor sea

la diferencia.

De acuerdo con (Vicéns, Herrarte, & Medina, 2005), el análisis de varianza puede contemplarse como un caso especial de modelización econométrica, donde el conjunto de variables explicativas o exógenas son ficticias y la variable dependiente o endógena es continua.

En estas situaciones, la estimación del modelo significa la realización de una ANOVA clásico, con amplia tradición y uso en estudios o diseños experimentales. Una ampliación a este planteamiento se presenta cuando se dispone de una variable de control que permita corregir el resultado del experimento mediante análisis de covarianza respecto a variable de estudio. Entonces, se calcula un análisis de covarianza o, también, llamada ANCOVA.

El ANOVA parte de algunos supuestos o hipótesis que han de cumplirse:

- La variable dependiente debe medirse al menos a nivel de intervalo.
- Independencia de las observaciones.
- La distribución de los residuales debe ser normal.
- Homocedasticidad (homogeneidad de varianzas).

1.5.2. Tipos de modelos

Los tipos de modelos que existen en ANOVA son:

• **Modelo I.** Efectos fijos, Sistemático o Paramétrico:- Se aplica a situaciones en las que el experimentador ha sometido al grupo o material analizado a varios factores, cada uno de los cuales le afecta sólo a la media, permaneciendo la “variable respuesta” con una distribución normal. Este modelo se supone cuando el investigador se interesa únicamente por los niveles del factor presentes en el experimento, por lo que cualquier variación observada en las puntuaciones se deberá al error

experimental. Puede tomar pocos o muchos valores o niveles, a cada uno de los cuales se asignan los grupos o muestras.

Si se toman K niveles del factor, a cada uno se asignan las muestras y las inferencias se refieren exclusivamente a los K niveles y no a otros que podrían haber sido incluidos, el ANOVA se llama de efectos fijos. El interés del diseño se centra en saber si esos niveles concretos difieren entre sí. Se aplican en investigaciones de Ciencias sociales.

• **Modelo II.** Efectos aleatorios o Componentes de Varianza:- Los modelos de efectos aleatorios se usan para describir situaciones en que ocurren diferencias incomparables en material o grupo experimental. El ejemplo más simple es el de estimar la media desconocida de una población compuesta de individuos diferentes y en el que esas diferencias se mezclan con los errores del instrumento de medición.

Este modelo se supone cuando el investigador está interesado en una población de niveles, teóricamente infinitos del factor de estudio, de los que únicamente una muestra al azar (t niveles) están presentes en el experimento. Cuando los niveles son muchos y se seleccionan al azar K niveles, pero las inferencias se desean hacer respecto al total de niveles.

La idea básica es que el investigador no tiene interés en niveles particulares del factor. Se aplican en diseños experimentales con muestras.

Modelo III. Efectos mixtos: Cuando se utilizan dos factores, cada uno con varios niveles, uno de efectos fijos y otro de efectos aleatorios. Se aplica en casos con uno o más factores de clases anteriores.

Las principales partes o fracciones de Análisis de Varianza de Fisher pueden variar, según el Diseño Experimental que se trabaje. Sin embargo, el ANOVA está particionada en los siguientes componentes explicativos:

Factor de Variación o Variation factor: indica el elemento o factor a estudiar dentro del ANOVA. Por ejemplo: Efecto total, Efecto tratamiento, Bloques (repeticiones), Factor A (Trat. 1rio), Parcela Grande, Sub Parcela, Error A, Factor B (Trat. 2rio), Error B, Factor C (Trat. 3rio), Error C, Interacciones doble AB, AC, BC, triple ABC, cuádruples ABCD, Prueba de Contrastes Ortogonales, Polinomios Ortogonales y/o Efecto Total.

Grados de libertad o Degrees of freedom: Son el número de Contrastes Ortogonales menos el número de restricciones impuestas, que pueden hacerse en un grupo de datos.

Los grados de libertad pueden descomponerse al igual que la suma de cuadrados en $GI\text{ total} = GI\text{ entre grupos o Tratamientos} + GI\text{ dentro de grupos o Error}$; aunque, sus divisiones pueden aumentar según el diseño experimental que se trabaje. Por ejemplo: Efecto total, Efecto tratamiento, Bloques (repeticiones), Factor A (Trat. 1rio), Parcela Grande, Sub Parcela, Error A, Factor B (Trat. 2rio), Error B, Factor C (Trat. 3rio), Error C, Interacciones doble AB, AC, BC, triple ABC, cuádruples ABCD, Prueba de Contrastes Ortogonales, Polinomios Ortogonales y/o Efecto Total. Simultáneamente, denota la significación entre tratamientos cuando se calcula o estima el cociente de dividir la suma de cuadrados entre número de grados de libertad.

Suma de cuadrados o Sum of square: indica la variación parcial o total de un conjunto de muestras de cada elemento o factor a estudiar dentro del ANOVA. Se puede descomponer en Suma de Cuadrados del Factor de estudio (SC), número real relacionado con varianza que mide la variación debida al "factor", "tratamiento" o tipo de situación estudiado, Suma de Cuadrados del Error de estudio (SC), número real relacionado con la varianza que mide la variación dentro de cada "factor", "tratamiento" o tipo de situación y/o Suma de Cuadrados del Total de estudio ($SC + SC_{\text{Error}}$) (Suma de Cuadrados

de Error α)), número real relacionado con la varianza que mide la variación total.

F Calculada o F value: El estadístico F de Snedecor está asociado a un nivel crítico de probabilidad de obtener valores, como el obtenido o mayores. Si éste es mayor al nivel de error asumido, habitualmente 95 % de confiabilidad estadística, error ó nivel de significancia (α) igual a 0.05, se interpreta “con base en resultados obtenidos de datos muestrales (poblacionales), no se acepta hipótesis nula, que indica igualdad de medias de tratamientos ($H_0: \mu_1 = \mu_2 = \mu_3 = \mu_4 = \dots = \mu_k$)” ($H_0: \mu_1 = \mu_2 = \mu_3 = \mu_4 = \dots = \mu_k$ = no se rechaza la alternativa, que indica diferencia de al menos una media de tratamiento respecto al resto ($H_a: \mu_i \neq \mu_j$ para algún $i \neq j$ o $H_a: \tau_i \neq 0$ para algún i)”

Con diferencia parcial respecto a (Gujarati & Porter, 2010) que afirman “con base en una prueba de significancia, como prueba t, se dice “aceptar” la hipótesis nula, todo lo que se afirma es que, con base en la evidencia dada por la muestra, no existe razón para rechazarla, no se sostiene que la hipótesis nula sea verdadera con absoluta certeza” y, de acuerdo con (Kmenta, 1971) y (González, 1985), “de la misma manera que un tribunal se pronuncia un veredicto de “no culpable” en vez de “inocente”, así la conclusión de una prueba estadística es “no rechazar” en vez de “aceptar”.

Esto se basa, en parte, en las siguientes razones:

(Wackerly, Mendenhall III, & Scheaffer, 2010) señalan “la probabilidad α recibe el nombre de nivel de significancia o, en forma más sencilla, nivel de prueba. Aun cuando se recomiendan con frecuencia pequeños valores de α , el valor real α para usar en un análisis es un tanto arbitrario.

Sin embargo, **p-value** es el nivel de significancia alcanzado, relacionado con una prueba y es un estadístico que representa el valor más pequeño de α el cual la información muestral indica que H_0 puede ser rechazada. En cierto sentido, p-value

permite al lector de la investigación evaluar la magnitud de la discrepancia entre datos observados e H_0 ".

(Navidi, 2006) afirma "si se rechaza H_0 se concluye que era falsa, en cambio si H_0 no se rechaza no se concluye que H_0 es verdadera, pues sólo se puede concluir que H_0 es factible, pero nunca se puede llegar a la conclusión H_0 es verdadera debido a que, por un lado, el estadístico de prueba es consistente con hipótesis H_a y está un poco en desacuerdo con H_0 tal que la única cuestión es si el nivel de desacuerdo medido con p-value es suficientemente grande para presentar H_0 como no factible y, por otro lado, una regla general conveniente, considerando p-value como menor que un umbral específico, indica rechazar o no aceptar H_0 cada vez que $p \leq 0.05$, interpretando que es "estadísticamente significativo"; aunque, este criterio no tiene ninguna base científica y, además, ésta es una mala práctica por varias razones:

1.No proporciona ninguna manera de decir si p-value era apenas menor que 0.05 o si era mucho menor,

2.Notificar que un resultado era estadísticamente significativo a un nivel de 5% implica que hay gran diferencia entre p-value justo debajo de 0.05 y uno justo arriba de 0.05, cuando efectivamente hay una diferencia pequeña,

3.Un trabajo así no permite al lector decidir por ellos mismos si p-value es suficientemente pequeño para rechazar o no aceptar H_0 . Por ejemplo: Si un lector cree que no debe rechazarse a menos que $p < 0.01$ entonces informar que solamente que $p < 0.05$ no permite al lector determinar si rechaza o no acepta o acepta o no rechaza H_0 ,

4.Valores pequeños de p-value señalan que H_0 es improbable que sea verdadera, tal que es tentador pensar que p-value representa la probabilidad que H_0 sea verdadera. No obstante, la verdad o falsedad de H_0 no se puede cambiar

mediante la repetición del experimento; por lo tanto, no es correcto hablar de “probabilidad” que H_0 sea verdadera,

5. La clase de probabilidad que analiza si no se rechaza o no se acepta H_0 , útil solamente cuando se aplica a resultados que pueden resultar en formas diferentes cuando se repiten experimentos debido a que para definir el p-value como probabilidad de observar un valor extremo de un estadístico como $\bar{X}$, pues su valor podría ser diferente si el experimento se repite, se llama probabilidad frecuentista (la frecuencia de un suceso en una muestra se define como el cociente entre número de veces que ha ocurrido el suceso en la muestra y el tamaño de la misma.

Empíricamente, se observa que al ir aumentando el tamaño de una muestra, la frecuencia de sucesos tiende a estabilizarse de un número fijo tal que se ha denominado como ley de estabilidad de frecuencias o ley única del azar y ese número ideal, límite que alcanzaría la frecuencia de un suceso si se obtuviera una muestra infinita del experimento es el primer concepto de probabilidad de un suceso tal que esta interpretación, frecuentista desarrollada a partir de conceptos de probabilidad y se centra en el cálculo de probabilidades y contrastes de hipótesis, sigue siendo la más ampliamente usada y ha inspirado el desarrollo formal de la teoría del cálculo de probabilidades tal que la probabilidad de un suceso va ser una medida del subconjunto del espacio muestral que corresponde a este suceso y las interpretaciones de esa medida constituirán distintos modelos probabilísticos de la realidad. Por último, (Triola, 2004) afirma “si bajo un supuesto dado, la probabilidad de un suceso observado particular es excepcionalmente pequeña, concluimos que el supuesto probablemente no sea correcto”),

6. La probabilidad subjetiva, que se basa en la creencias e ideas en que se realiza la evaluación de probabilidades y se define como en aquella que un evento asigna el individuo

basándose en la evidencia disponible con base en su experiencia; es decir, es una forma de cuantificar, por medio de factores de ponderación individuales, la probabilidad de que ocurra cierto evento cuando no es posible de cuantificar de otra manera más confiable, calcula la probabilidad que un enunciado, como H_0 , sea verdadero y es importante en la teoría de Estadística Bayesiana (es un subconjunto del campo de la estadística en la que la evidencia sobre el verdadero estado del mundo se expresa en términos de grados de creencia o, más específicamente, las probabilidades bayesianas).

Es sólo una de una serie de interpretaciones de la probabilidad y hay otras técnicas estadísticas que no se basan en “grados de creencia”. La inferencia bayesiana es un enfoque de la inferencia estadística, que es distinta de la inferencia frecuentista. Se basa específicamente en el uso de probabilidades bayesianas al resumir las pruebas.

La formulación de modelos estadísticos para su uso en la estadística bayesiana tiene la característica adicional, no está presente en otros tipos de técnicas estadísticas, que requiere la formulación de un conjunto de distribuciones previas para los parámetros desconocidos, sus consideraciones habituales en diseño de experimentos se extienden en el caso de diseño Bayesiano de experimentos para incluir la influencia de las creencias anteriores y requiere ciertas técnicas computacionales modernas para la inferencia bayesiana, como diversos tipos de técnicas Monte Carlo de cadenas de Markov, pues según (Wackerly, Mendenhall III, & Scheaffer, 2010), las pruebas de hipótesis se pueden abordar desde una perspectiva de Bayes con base en información previo a datos acerca de un parámetro para su distribución posterior.

En consecuencia, se está interesado en probar que el parámetro θ se encuentra en uno de dos conjuntos de valores Ω_0 y Ω_a . Cuando se pruebe $H_0: \theta \in \Omega_0$ contra $H_a: \theta \in \Omega_a$

, usando método de cálculo de probabilidades posteriores sería $P^*(\theta \in \Omega_0)$ y $P^*(\theta \in \Omega_a)$ aceptando la hipótesis con más alta probabilidad posterior; es decir, aceptar H_0 si $P^*(\theta \in \Omega_0) > P^*(\theta \in \Omega_a)$ o aceptar H_a si $P^*(\theta \in \Omega_a) > P^*(\theta \in \Omega_0)$ y 7) Si un resultado tiene un p-value pequeño se dice que es "estadísticamente significativo", representa "importante" y, en consecuencia, resulta tentador pensar que resultados estadísticamente significativos siempre son importantes.

No obstante, a veces este tipo de resultados no tienen importancia científica o práctica. Entonces, un resultado puede ser estadísticamente significativo sin ser lo suficientemente grande para tener importancia práctica, pues una diferencia es estadísticamente significativa cuando es grande comparada con su σ o s e, incluso, cuando σ o s es muy pequeña puede ser estadísticamente significativa.

Por lo tanto, p-value no mide la significancia práctica, sino el grado de confianza que se puede tener que el valor verdadero es muy diferente del valor especificado por H_0 tal que cuando p-value es pequeño se tiene la confianza que el valor verdadero es, en verdad, muy diferente, pero no implica que la diferencia sea lo bastante grande para que tenga importancia práctica".

-(Anderson, Sweeney, & Williams, 2006) sostienen "en una prueba de hipótesis sólo se puede tomar una de dos decisiones: aceptar o rechazar H_0 . En lugar de "aceptar" H_0 algunos investigadores prefieren enunciar la decisión como "No se rechaza H_0 ", "No se puede rechazar H_0 " o, también, "Los resultados muestrales no permiten rechazar H_0 ".

En pruebas de hipótesis p-value, proporciona la probabilidad de obtener, suponiendo que H_0 sea verdadera, un valor muestral estadístico de prueba tan extremo, por lo menos o más extremo, como el obtenido tal que p-value compara la probabilidad con el nivel de significancia".

Por lo tanto, si no se acepta H_0 , el siguiente paso es efectuar la prueba de significancia entre medias de tratamientos con el fin de conocer cuál o cuáles son estadísticamente mejores.

-(Triola, 2004) afirma “estadístico de prueba es un valor calculado a partir de datos muestrales, que se usa para tomar la decisión sobre rechazo de H_0 . Por lo tanto, sirve para determinar si existe evidencia significativa en contra de H_0 . ρ es la probabilidad de obtener un valor del estadístico de prueba que sea al menos tan extremo como el que representa a datos muestrales, suponiendo que H_0 es verdadera.

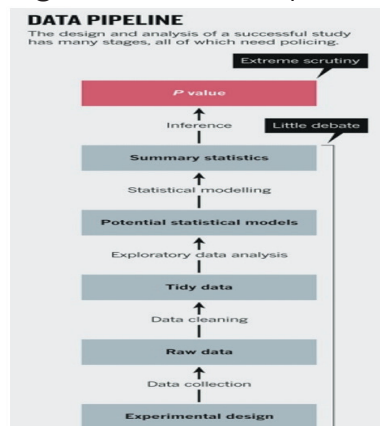
Algunos libros dicen “aceptar H_0 ” en vez de “no rechazar H_0 ”, pero no se está probando H_0 , sea que use el término “aceptar” o “no rechazar”, tal que únicamente se afirma que la evidencia muestral no es lo suficientemente fuerte como para justificar el rechazo de H_0 . El término “aceptar” es un poco confuso, pues parece indicar incorrectamente que H_0 ha sido aprobada tal que la frase “no rechazar” indica correctamente que la evidencia disponible no es lo suficientemente fuerte para justificar el rechazo de H_0 ”.

$\Pr(>F)$ o ρ -Value: Es la probabilidad de observar un valor muestral tan extremo o más extremo que el valor observado dado que H_0 es verdadera o es una manera de expresar la probabilidad H_0 no sea verdadera. Complementariamente, $\Pr(>F)$ es el nivel de probabilidad que H_a caiga en la zona de rechazo o, complementariamente, no hay ninguna diferencia entre los dos grupos o no hay correlación entre un par de características.

Según (Gujarati & Porter, 2010), ρ -Value, valor de probabilidad, nivel observado, exacto de significancia o probabilidad exacta de cometer error tipo I, estima la probabilidad real de obtener un valor del estadístico de prueba tan grande o mayor que el obtenido en un ejercicio o, en otras palabras, es el nivel de significancia más bajo al cual puede rechazarse una hipótesis nula, mientras que Leek

y Peng afirman que $\Pr(>F)$ o p -Value es la cima del iceberg:

Figura 1.6. Data Pipeline



Fuente: (Leek & Peng, 2015)

Sin embargo, afirman “no hay ninguna estadística más calumniada que el p -Value. Incluso, investigadores que supervisan el análisis de datos deberían exigir de sus instituciones completar la formación en comprensión de resultados y posibles problemas con un análisis e, incluso, la investigación estadística se centra en gran parte en Estadística matemática, exclusión de procesos involucrados en análisis de datos y comportamiento tal que para resolver este problema más profundo se estudiar cómo la gente realiza análisis de datos en el mundo real”.

De acuerdo con la publicación “Estadísticos emiten alerta sobre mal uso de $\Pr(>F)$ o p -Value” y (Wasserstein & Lazar, 2016) “el valor de P es una prueba común para juzgar la fuerza de la evidencia científica, contribuye al número de resultados de la investigación que no puede ser reproducido, la asociación estadística americana (ASA) advierte en un comunicado publicado hoy.

El grupo ha tomado el paso inusual de emitir principios para guiar el uso del p -Value, que dice no se puede determinar si una hipótesis es verdadera o si los resultados son importantes. En su comunicado, la ASA informa a investigadores evitar

conclusiones científicas o tomar decisiones de política basadas solo en valores p .

Los investigadores deben describir no sólo el análisis de datos que produjo resultados estadísticamente significativos, la sociedad dice, pero todas pruebas y decisiones en cálculos. De lo contrario, pueden parecer falsamente robustos resultados.

El más pequeño el valor de p -Value indica menos probabilidad que un conjunto observado de valores ocurra por casualidad, suponiendo que la hipótesis nula es verdadera. Un valor de $p \leq 0.05$ se toma, generalmente, para expresar que un resultado es estadísticamente significativo y garantiza la publicación o hay 95% de probabilidad que una hipótesis dada es correcta.

Pero no es necesariamente verdadero, según notas de instrucción de ASA. En cambio, significa si la hipótesis nula (H_0 según (Lind, Marchal, & Mason, 2006), H_0 se lee “H subíndice 0”, pues H significa hipótesis y el subíndice 0 señala “no hay diferencia”.

En cambio, H_1 o H_a , leído como “H subíndice 1 o a” señala que “al menos 1 es diferente respecto al resto de los tratamientos”. La demostración matemática que lo sustenta se encuentra en la siguiente nota al pie y da razón que exista H_0 e H_a en una prueba de dos colas.

Además, con base en una muestra de una variable aleatoria X de ley P_θ en decidir entre dos hipótesis: $H_0: \theta \in \theta_0$ (hipótesis nula o privilegiada) vs $H_1: \theta \in \theta_1$ (hipótesis alternativa o alterna) tal que $\theta_0 \cup \theta_1 = \theta$ y $\theta_0 \cap \theta_1 = \emptyset$. Si d_0 y d_1 representan, respectivamente, las decisiones de no rechazar H_0 o H_1 y $D = \{d_0, d_1\}$.

Entonces: sea $x: \Omega \rightarrow E \subseteq \mathbb{R}$ una variable aleatoria de ley P_θ . Se llama “Prueba de Hipótesis Pura” o test puro a toda aplicación: $\phi: E^n \rightarrow D$ tal que ϕ equivale a particionar E^n en dos conjuntos: $W = \phi^{-1}(d_0) = \{x \in E^n: \text{si se observa } x, \text{ no se acepta } H_0\}$

$Y W^c = \emptyset^{-1}(d_1) = \{x \in E^n: \text{si se observa } x, \text{ no se rechaza } H_0\}.$

El nombre de H_0 proviene que H_0 se asumirá como verdadera, salvo que datos muestrales indiquen su falsedad. Nula debe entenderse como neutra. H_0 nunca se considera probada o demostrada, salvo estudiando todos los datos de la población, puede diferir en un valor pequeño imperceptible para el muestreo, que puede ser imposible; aunque, puede no ser aceptada por los datos.

Con base en esto, no se debe afirmar “se acepta H_0 ” siendo lo correcto “no se rechaza H_0 ” y, por abuso de lenguaje, es común hallar la expresión “se acepta la H_0 ” en lugar de “no se rechaza H_0 ”.

Generalmente, H_0 se elige según con el principio de simplicidad científica, que establece que únicamente se abandona un modelo simple a favor de otro más complejo cuando la evidencia a favor de este último sea fuerte. Por el carácter dicotómico, si no se acepta H_0 , automáticamente no se rechaza H_1 .

Finalmente, se llama riesgo de primera especie a la probabilidad de rechazar H_0 cuando es verdadera: $\alpha_0(\emptyset) = P_0(W_{(\text{Región crítica de prueba})})$, mientras que el riesgo de segunda especie es la probabilidad de aceptar H_0 cuando es falsa: $\beta_\theta(\emptyset) = P_1(W_{(\text{Región de aceptación})}^c) = 1 - P_1(W)$.

Se llama “Potencia de una prueba” a la probabilidad de no aceptar H_0 cuando es falsa $-P_1(W) = 1 - \beta_\theta(\emptyset) -$, también, se denomina “Nivel de significación de una prueba” al valor $\alpha = \sup \alpha_\theta(\emptyset) \theta \in \theta_0$.

Cuando $\theta \in \mathbb{R}$ se estudian problemas del tipo:

$$P_0: \begin{cases} H_0: \theta = \theta_0, \\ H_1: \theta = \theta_1 \text{ con } \theta_0 \neq \theta_1, \end{cases}$$

$$P_1: \begin{cases} H_0: \theta \leq \theta_0, \\ H_1: \theta > \theta_0 \text{ (Prueba de cola derecha)}, \end{cases}$$

$$P_2: \begin{cases} H_0: \theta \geq \theta_0, \\ H_1: \theta < \theta_0 \text{ (Prueba de cola izquierda)}, \end{cases}$$

$$P_3: \begin{cases} H_0: \theta_0 \leq \theta \leq \theta_1, \\ H_1: \theta < \theta_0 \text{ (Prueba de cola izquierda)} \text{ o } \theta > \theta_1 \text{ (Prueba de cola derecha)} (\theta_0 < \theta_1) \end{cases}$$

$$P_4: \begin{cases} H_0: \theta < \theta_0 \text{ o } \theta > \theta_1, \\ H_1: \theta_0 \leq \theta \leq \theta_1, \end{cases} \text{ o } P_5: \begin{cases} H_0: \theta = \theta_0, \\ H_1: \theta \neq \theta_1. \end{cases}$$

El teorema Neyman-Pearson indica, en contexto de prueba propuesta: $\forall \alpha \in [0, 1]$ un test puro $\emptyset$, de nivel α , de potencia máxima, definido por la región crítica:

$W = \left\{ x \in E^n: \frac{L(\theta_0, x)}{L(\theta_1, x)} \leq k \right\} \rightarrow k$ se determina de la condición $\alpha = P_0(W)$
 L representa la función de verosimilitud y x una observación de la muestra. Otras pruebas importantes, de nivel α , son

A): $\begin{cases} H_0: \mu > \mu_0, \\ H_1: \mu \leq \mu_0 \end{cases}$ y B): $\begin{cases} H_0: \mu < \mu_0, \\ H_1: \mu \geq \mu_0 \end{cases}$ (Capa, 2015)) es verdadera y todos los otros supuestos son validados, tal que hay un 5% de probabilidad de obtener resultados al menos tan extrema como el observado. Un *p-Value* no puede indicar la importancia de un resultado. Vickers dice que “investigadores deben tratar la estadística como una ciencia y no una receta”.

Una vez que se han calculado grados de libertad, sumas de cuadrados, medias cuadráticas o cuadrados medios, F de Fisher o distribución F de Snedecor y Probabilidad que caiga en la región de rechazo ($\Pr(> F)$) Rho value o valor Ro (*p-Value*)



CAPITULO II

HIPÓTESIS

2.1. Generalidades

Son las guías para una investigación o estudio. Las hipótesis indican lo que tratamos de probar y se definen como explicaciones tentativas del fenómeno investigado. Las hipótesis son el centro, la médula o el eje del método deductivo cuantitativo y se puede tener una, dos o varias hipótesis. Se derivan de la teoría existente (Williams, 2003) citado por (Hernandez, Fernandez, & Baptista, 2006) y deben formularse a manera de proposiciones.

De hecho, son respuestas provisionales a las preguntas de investigación. Por ejemplo: si una pregunta de investigación es “¿le gustará a Paola?” y una hipótesis derivada sería “Le resulto atractivo a Paola”. Esta hipótesis es una explicación tentativa y está formulada como proposición. Después investigamos si se acepta o se rechaza la hipótesis, al cortejar a Paola y observar el resultado obtenido.

Es importante mencionar que no todas las investigaciones cuantitativas plantean hipótesis. El hecho de que formulemos o no hipótesis depende de un factor esencial: el alcance inicial del estudio.

Las hipótesis no necesariamente son verdaderas, pueden o no serlo, pueden o no comprobarse con datos y pueden ser más o menos generales o precisas e involucrar a dos o más variables, pero en cualquier caso son sólo proposiciones sujetas a comprobación empírica y a verificación en la realidad. Son explicaciones tentativas, no los hechos en sí.

Al formularlas, el investigador no está totalmente seguro de que vayan a comprobarse. Las investigaciones cuantitativas que formulan hipótesis son aquellas cuyo planteamiento define que su alcance será correlacional o explicativo, o las que tienen un alcance descriptivo, pero que intentan pronosticar una cifra o un hecho.

La formulación de hipótesis en estudios cuantitativos con

diferentes alcances es:

Alcance de estudio	Formulación de hipótesis
Exploratorio	No se formulan hipótesis.
Descriptivo	Sólo se formulan hipótesis cuando se pronostica un hecho o dato.
Correlacional	Se formulan hipótesis correlacionales.
Explicativo	Se formulan hipótesis causales.

Ejemplos:

1. “La proximidad geográfica entre los hogares de las parejas de novios está vinculada positivamente con el nivel de satisfacción que les proporciona su relación”.

2. “El índice de cáncer pulmonar es mayor entre los fumadores que entre los no fumadores”.

3. “A mayor variedad en el trabajo, habrá mayor motivación intrínseca hacia éste”.

Bajo el enfoque cuantitativo, es natural que las hipótesis surjan del planteamiento del problema que, se vuelve a evaluar y si es necesario se replantea después de revisar la literatura. Las hipótesis pueden surgir de un postulado de una teoría, del análisis de ésta, de generalizaciones empíricas pertinentes al problema de investigación, estudios revisados o antecedentes consultados o por analogía -relación de semejanza entre cosas distintas-, al aplicar cierta información a otros contextos, como la teoría del campo en psicología, que surgió de la teoría del comportamiento de los campos magnéticos.

Sin embargo, hipótesis útiles y fructíferas pueden originarse en planteamientos del problema cuidadosamente revisados, aunque el cuerpo teórico que las sustente no sea abundante. A veces la experiencia y la observación constante ofrecen

materia potencial para el establecimiento de hipótesis importantes y lo mismo se dice de la intuición.

No obstante, existe una relación muy estrecha entre el planteamiento del problema, la revisión de la literatura y las hipótesis. Mediante del proceso quizá se ocurra otras hipótesis que no estaban contempladas en el planteamiento original, producto de nuevas reflexiones, ideas o experiencias; discusiones con profesores, colegas o expertos en el área; incluso, “de analogías, al descubrir semejanzas entre la información referida a otros contextos y la poseída para el estudio” (Rojas & Rojas, 2000).

Lo que sí constituye una grave falla en la investigación es formular hipótesis sin haber revisado con cuidado la literatura, pues se cometen errores como sugerir hipótesis de algo bastante comprobado o algo que ha sido contundentemente rechazado. En conclusión, la calidad de las hipótesis está relacionada en forma positiva con el grado en que se haya revisado la literatura exhaustivamente.

Dentro del enfoque cuantitativo, para que una hipótesis sea digna de tomarse en cuenta, debe reunir algunos requisitos:

1. La hipótesis debe referirse a una situación “real”. Como argumenta (Rojas & Rojas, 2000), las hipótesis sólo pueden someterse a prueba en un universo y un contexto bien definidos. Por ejemplo, “a mayor satisfacción laboral mayor productividad”.

2. Las variables o términos de la hipótesis deben ser comprensibles, precisos y lo más concretos posible. Términos vagos o confusos no tienen cabida en una hipótesis. Así, globalización de la economía y sinergia organizacional son conceptos imprecisos y generales que deben sustituirse por otros más específicos y concretos.

3. La relación entre variables propuesta por una hipótesis debe ser clara y verosímil (lógica). Por ejemplo: “la disminución

del consumo del petróleo en Estados Unidos se relaciona con el grado de aprendizaje del álgebra por parte de niños que asisten a escuelas públicas en Ecuador”.

4. Los términos o variables de la hipótesis deben ser observables y medibles, así como la relación planteada entre ellos, o sea, tener referentes en la realidad. Por ejemplo: “los hombres más felices van al cielo” o “la libertad de espíritu está relacionada con la voluntad angelical”, implican conceptos o relaciones que no poseen referentes empíricos; por tanto, no son útiles como hipótesis para investigar científicamente ni se pueden someter a prueba en la realidad.

Las hipótesis deben estar relacionadas con técnicas disponibles para probarlas. Este requisito está estrechamente ligado con el anterior y se refiere a que al formular una hipótesis, tenemos que analizar si existen técnicas o herramientas de investigación para verificarla. Por ejemplo: Alguien quiere probar hipótesis referentes a la desviación presupuestal en el gasto gubernamental de un país latinoamericano o a la red de narcotraficantes en México, pero no dispone de formas eficaces para obtener sus datos.

2.2. Tipos de Hipótesis

2.2.1. De investigación o de trabajo

Son proposiciones tentativas sobre la(s) posibles relaciones entre dos o más variables y deben cumplir con los cinco requisitos mencionados. Se suelen simbolizar como H_1, H_2, H_3 , etcétera. Pueden ser:

Descriptivas de un valor o dato pronosticado. Se utilizan a veces en estudios descriptivos, para intentar predecir un dato o valor en una o más variables que se van a medir u observar. Pero cabe comentar que no en todas las investigaciones descriptivas se formulan hipótesis de esta clase o que sean afirmaciones más generales.

Ejemplo:

H_i: "El aumento del número de divorcios de parejas cuyas edades oscilan entre los 18 y 25 años, será de 20% el próximo año".

H_i: "La inflación del próximo semestre no será superior a 3%".

•**Correlacionales.** Especifican las relaciones entre dos o más variables y corresponden a los estudios correlacionales. No sólo pueden establecer que dos o más variables se encuentran vinculadas, sino también cómo están asociadas. Alcanzan el nivel predictivo y parcialmente explicativo.

En una hipótesis de correlación, el orden en que coloquemos las variables no es importante (ninguna variable antecede a la otra; no hay relación de causalidad).

Es lo mismo indicar "a mayor X, mayor Y"; que "a mayor Y, mayor X"; o "a mayor X, menor Y"; que "a menor Y, mayor X"). Es común que cuando en la investigación se pretende correlacionar diversas variables se tengan varias hipótesis y cada una de ellas relacione un par de variables. Estas hipótesis deben contextualizarse en su realidad (con qué parejas) y someterse a prueba empírica.

Ejemplos:

H_i : "A mayor exposición por parte de los adolescentes a videos musicales con alto contenido sexual, mayor manifestación de estrategias en las relaciones interpersonales para establecer contacto sexual".

H_i: "A mayor autoestima, habrá menor temor al éxito".

H_i: "Las telenovelas latinoamericanas muestran cada vez un mayor contenido sexual en sus escenas".

H_i: "Quienes logran más altas puntuaciones en el examen de estadística tienden a alcanzar las puntuaciones más

elevadas en el examen de economía” es igual a: “los que logran tener las puntuaciones más elevadas en el examen de economía son quienes tienden a obtener más altas puntuaciones en el examen de estadística”.

H₁: “A mayor atracción física, menor confianza”.

H₂: “A mayor atracción física, mayor proximidad física”.

H₃: “A mayor atracción física, mayor equidad”.

H₄: “A mayor confianza, mayor proximidad física”.

H₅: “A mayor confianza, mayor equidad”.

H₆: “A mayor proximidad, mayor equidad”.

• **Diferencia de grupos.** Estas hipótesis se formulan en investigaciones cuya finalidad es comparar grupos.

Por ejemplo:

H_i: “El efecto persuasivo para dejar de fumar no será igual en los adolescentes que vean la versión del comercial televisivo en colores, que el efecto en los adolescentes que vean la versión del comercial en blanco y negro”.

H_i: “Los adolescentes le atribuyen más importancia al atractivo físico en sus relaciones de pareja, que las adolescentes a las suyas”.

H_i: “El tiempo que tardan en desarrollar el SIDA las personas contagiadas por transfusión sanguínea, es menor que las que adquieren el VIH por transmisión sexual”.

H_i: “Las escenas de la telenovela La verdad de Paola presentarán un mayor contenido sexual que las de la telenovela Sentimientos de Christian, y éstas, a su vez, un mayor contenido sexual que las escenas de la telenovela Mi último amor: Mariana”

•**Causales.** Este tipo de hipótesis no solamente afirma la o las relaciones entre dos o más variables y la manera en que se manifiestan, sino que además propone un “sentido de entendimiento” de las relaciones.

Tal sentido puede ser más o menos completo, esto depende del número de variables que se incluyan, pero todas estas hipótesis establecen relaciones de causa-efecto.

Ejemplos:

H₁ : “La desintegración del matrimonio provoca baja autoestima en los hijos e hijas”.

H₁ : “Un clima organizacional negativo crea bajos niveles de innovación en los empleados”.

H₁ : “Un clima organizacional negativo crea bajos niveles de innovación en los empleados”.

H₁ : “La cohesión y la centralidad en un grupo sometido a una dinámica, así como el tipo de liderazgo que se ejerza dentro del grupo, determinan la eficacia de éste para alcanzar sus metas primarias”.

H₁ : “La variedad y la autonomía en el trabajo, así como la retroalimentación proveniente del desarrollo de éste, generan mayor motivación intrínseca y satisfacción laborales”.

H₁ : “La paga aumenta la motivación intrínseca de los trabajadores, cuando se administra de acuerdo con el desempeño”.

H₁ : “La paga incrementa la satisfacción laboral”.

H₂ : “La integración, la comunicación instrumental y la comunicación formal incrementan la satisfacción laboral”.

H₃ : “La centralización disminuye la satisfacción laboral”.

H₄ : “La satisfacción laboral influye en la reasignación de personal”.

H₅: “La oportunidad de capacitación mediatiza la vinculación entre la satisfacción laboral y la reasignación de personal”.

H₆: “La reasignación de personal afecta la integración, la efectividad organizacional, la formalización, la centralización y la innovación”.

• **Nulas.** Las hipótesis nulas son, en cierto modo, el reverso de las hipótesis de investigación. Son proposiciones que niegan o refutan la relación entre variables, sólo que sirven para refutar o negar lo que afirma la hipótesis de investigación.

Debido a que este tipo de hipótesis resulta la contrapartida de la hipótesis de investigación, hay prácticamente tantas clases de hipótesis nulas como de investigación. Las hipótesis nulas se simbolizan así: H_0 .

Ejemplos:

H_0 : “El aumento del número de divorcios de parejas cuyas edades oscilan entre los 18 y 25 años, no será de 20% el próximo año”.

H_0 : “No hay relación entre la autoestima y el temor al éxito”.

H_0 : “Las escenas de la telenovela La verdad de Paola no presentarán mayor contenido sexual que las de la telenovela Sentimientos de Christian, ni éstas tendrán mayor contenido sexual que las escenas de la telenovela Mi último amor: Mariana”. Esta hipótesis niega la diferencia entre grupos y también podría formularse así: “no existen diferencias en el contenido sexual entre las escenas de las telenovelas La verdad de Paola, Sentimientos de Christian y Mi último amor Mariana”. O bien, “el contenido sexual de las telenovelas La verdad de Paola, Sentimientos de Christian y Mi último amor Mariana es el mismo”.

H_0 : “La percepción de la similitud en religión, valores y

creencias no provoca mayor atracción" (hipótesis que niega la relación causal).

• **Alternativas.** Son posibilidades diferentes o "alternas" ante las hipótesis de investigación y nula: ofrecen otra descripción o explicación distinta de las que proporcionan estos tipos de hipótesis.

Las hipótesis alternativas se simbolizan como H_a y sólo pueden formularse cuando efectivamente hay otras posibilidades, además de las hipótesis de investigación y nula.

Ejemplos:

H_i : "El candidato A obtendrá en la elección para la presidencia del consejo escolar entre 50 y 60% de la votación total".

H_o : "El candidato A no obtendrá en la elección para la presidencia del consejo escolar entre 50 y 60% de la votación total".

H_a : "El candidato A obtendrá en la elección para la presidencia del consejo escolar más de 60% de la votación total".

H_a : "El candidato A obtendrá en la elección para la presidencia del consejo escolar menos del 50% de la votación total".

No existen reglas específicas para formular hipótesis, a menudo se sugiere la forma de hipótesis nula y alternativa; es decir, no hay reglas universales, ni siquiera consenso entre los investigadores.

En consecuencia, las expectativas teóricas o trabajo empírico previo o ambos pueden ser la base para formular hipótesis; aunque, sin importar la forma de postular hipótesis es extremadamente importante que el investigador las plantee antes de la investigación empírica, sino el (los) o la

investigadora (s) serán culpables de razonamiento circulares o profecías autocumplidas; es decir, si se formula la hipótesis después de examinar resultados empíricos puede presentarse la tentación de formular hipótesis de manera que justifique resultados.

Se puede leer en un artículo de alguna revista científica un reporte de investigación donde sólo se establezca la hipótesis de investigación; y, en otra, encontrar un artículo donde únicamente se plantea la hipótesis nula. Un artículo en una tercera revista, en el cual se puedan encontrar solamente las hipótesis de investigación y nula, pero no las alternativas.

En una cuarta publicación otro artículo que contenga la hipótesis de investigación y las alternativas. Y otro más donde aparezcan hipótesis de investigación, nulas y alternativas. Algunos investigadores sólo enuncian una hipótesis nula o de investigación presuponiendo que quien lea su reporte deducirá la hipótesis contraria. Sin embargo, la calidad de una investigación no necesariamente está relacionada con el número de hipótesis que contenga.

En este sentido, se debe tener el número de hipótesis necesarias para guiar el estudio, ni una más ni una menos. Cada investigación es diferente, pues algunas contienen gran variedad de hipótesis porque el problema de investigación es complejo (por ejemplo, pretenden relacionar 15 o más variables), mientras que otras contienen una o dos hipótesis.

Asimismo, en una investigación se pueden formular hipótesis descriptivas de un dato que se pronostica en una variable, hipótesis correlacionales, hipótesis de la diferencia de grupos, hipótesis causales e incluso hipótesis de cola derecha, cola izquierda o de ambas colas. Ejemplos:

Tabla 2.1. Tipos de Hipótesis

Preguntas de Investigación	Hipótesis
¿Cuál será a fin de año el nivel de desempleo en la Ciudad de Guaranda?	El nivel de desempleo en la Ciudad de Guaranda será de 7% a fin de año (H_1 : 7%).
¿Cuál es el nivel promedio de ingreso familiar mensual en la Ciudad de Guaranda?	El nivel promedio de ingreso familiar mensual en la Ciudad de Guaranda oscila entre 650 y 700 USD (H_1 : $649 < \bar{X} < 701$).
¿Existen diferencias entre los barrios de la Ciudad de Guaranda respecto a nivel de desempleo? (¿Hay barrios con mayores índices de desempleo?)	Existen diferencias respecto al nivel de desempleo entre barrios de la Ciudad de Guaranda (H_1 : Índice 1 $\neq$ Índice 2 $\neq$ Índice 3 $\neq$ Índice k).
¿Cuál es el nivel de escolaridad promedio de los jóvenes y las jóvenes que viven en la Ciudad de Guaranda? ¿Existen diferencias por género al respecto?	No se dispone de información, no se establecen hipótesis.
¿Está relacionado el desempleo con incrementos en la delincuencia de dicha ciudad?	A mayor desempleo, mayor delincuencia (H_1 : $r_{xy} \neq 0$).
¿Provoca el nivel de desempleo un rechazo contra la política fiscal gubernamental?	El desempleo provoca rechazo contra política fiscal gubernamental (H_1 : $X \rightarrow Y$).

Fuente: (Rojas & Rojas, 2000)

Estadísticas

Las hipótesis estadísticas son la hipótesis nula y la hipótesis alterna. En el campo de la utilización y aprovechamiento de la estadística, las decisiones se toman siempre sobre determinadas hipótesis.

La eficiencia de las campañas publicitarias o de los proceso de producción se fundan en criterios numéricos, y tales hipótesis se expresan en función de parámetros estadísticos. En el análisis de todo problema de investigación, la contrastación de una hipótesis dada se realiza aceptando o rechazando la hipótesis nula.

Funciones de las hipótesis

Las principales funciones de las hipótesis son:

En primer lugar, son las guías de una investigación en el

enfoque cuantitativo. Formularlas nos ayuda a saber lo que tratamos de buscar, de probar. Proporcionan orden y lógica al estudio (Selltiz et al., 1980) citado por (Hernandez, Fernandez, & Baptista, 2006).

En segundo lugar, tienen una función descriptiva y explicativa, según sea el caso. Cada vez que una hipótesis recibe evidencia empírica en su favor o en su contra, nos dice algo acerca del fenómeno con el que se asocia o hace referencia.

En tercer lugar, es probar teorías. Cuando varias hipótesis de una teoría reciben evidencia positiva, la teoría va haciéndose más robusta; cuanta más evidencia haya en favor de aquéllas, más evidencia habrá en favor de ésta.

En cuarto lugar, consiste en sugerir teorías. Diversas hipótesis no están asociadas con teoría alguna; pero llega a suceder que como resultado de la prueba de una hipótesis, se pueda construir una teoría o las bases para ésta.

Las hipótesis del proceso cuantitativo se someten a prueba o escrutinio empírico para determinar si son apoyadas o refutadas, de acuerdo con lo que el investigador observa. Ahora bien, en realidad no se puede probar que una hipótesis sea verdadera o falsa, sino argumentar que fue apoyada o no de acuerdo con ciertos datos obtenidos en una investigación particular.

Desde el punto de vista técnico, no se acepta una hipótesis a través de un estudio, sino que se aporta evidencia en su favor o en su contra. En el enfoque cuantitativo, las hipótesis se someten a prueba en la “realidad” cuando se aplica un diseño de investigación, se recolectan datos con uno o varios instrumentos de medición y se analizan e interpretan.

Se debe tener claro que no aceptar “rechazar” o no rechazar “aceptar” una hipótesis nula depende de α , nivel de significancia o probabilidad de cometer error tipo I,

probabilidad de rechazar la hipótesis cuando es verdadera. α se fija en niveles de 1, 5 o cuando mucho, 10%, pero según (Gujarati & Porter, 2010) no hay nada “sagrado” acerca de estos valores tal que cualquier otro valor sería por igual apropiado.

No obstante, comúnmente se fijan estos niveles de α para control de calidad, diseños experimentales o investigaciones sociales, respectivamente y, de igual manera, los econométricos tienen por costumbre fijar el valor α en estos niveles como máximo y escogen un estadístico de prueba que haga la probabilidad de cometer un error tipo II sea lo más pequeño posible.

Como uno menos la probabilidad de cometer error tipo II se conoce como potencia de prueba, este procedimiento equivale a maximizar la potencia de prueba.

Comparación de medias de tratamiento

Aunque en origen el Análisis de la Varianza (AOV, ANOVA, ADEVA, ANDEVA ó ANVA) fue introducido por Fisher para evaluar los efectos de los distintos niveles de un factor sobre una variable respuesta continua, desde un punto de vista puramente abstracto el ANOVA va a permitir generalizar el contraste de igualdad de medias de dos a k poblaciones (Arriaza, y otros, 2008).

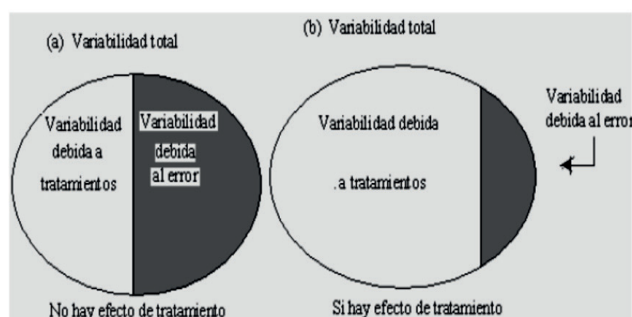
Es decir, al estudiar el comportamiento de los tratamientos de un factor, mediante un análisis de la varianza, el único objetivo es saber si globalmente dichos tratamientos difieren significativamente entre sí o conocer qué tratamientos concretos producen mayor efecto o cuáles son los tratamientos diferentes entre sí.

El nombre de análisis de varianza (ANOVA) viene del hecho de que se utilizan cocientes de varianzas para probar la hipótesis de igualdad de medias. La idea general de esta técnica es separar la variación total en dos partes: la

variabilidad debida a los tratamientos y la debida al error.

Cuando la primera predomina claramente sobre la segunda es cuando se concluye que los tratamientos tienen efecto; es decir, las medias son diferentes. Cuando los tratamientos contribuyen igual o menos que el error, se concluye que las medias son iguales (González, Métodos estadísticos y principios de diseño experimental, 2010):

Figura 2.1. Variación total de un DCA



Fuente: (González B. G., Métodos estadísticos y principios de diseño experimental, 2010)

No se propondrá pues ningún modelo teórico, sino que el objetivo se limitará a usar la técnica para contrastar la hipótesis

$$H_0: \mu_1 = \mu_2 = \mu_3 = \mu_4 = \dots = \mu_k = \mu(0) \text{ vs}$$

$$H_a: \mu_i \neq \mu_j \text{ para algún } i \neq j \text{ ó}$$

$$H_0: \tau_1 = \tau_2 = \tau_3 = \tau_4 = \dots \tau_k = 0 \text{ vs } H_a: \tau_i \neq 0 \text{ para algún } i.$$

Al igual que se ha hecho para una y dos poblaciones, se evaluarán las hipótesis previas relativas a la calidad de la muestra, a la estructura de probabilidad, normal o no, de la población y a si las distintas poblaciones tienen varianzas iguales o distintas, propiedad esta última conocida como

homocedasticidad (Arriaza, y otros, 2008).

Se explica algunas técnicas para analizar con mayor detalle los datos de un experimento, con posterioridad a realización del Análisis de Varianza. Si este análisis confirma que la prueba F del análisis de la varianza resultó al menos significativa al 5 % o "Fisher Protegido" es conveniente investigar qué medias son distintas, que indica que para calcular matrices de comparaciones de medias de tratamiento debe existir al menos una diferencia mínima significativa denotada por algunos softwares como R mediante *. Para ello, se usan diversas técnicas cuyo objeto es identificar qué tratamientos son estadísticamente diferentes y en cuánto oscila el valor de esas diferencias.

En algunos casos, el uso de estas técnicas está supeditado al resultado del análisis de varianza; en otros casos, las técnicas pueden emplearse directamente sin haber realizado previamente dicho análisis. Este conjunto de técnicas se engloba bajo la denominación de contrastes para comparaciones múltiples, pues su objetivo fundamental es comparar entre sí medias de tratamientos o grupos de ellas.

En primer lugar, se estudia un procedimiento intuitivo y cualitativo basado en la representación gráfica de datos del experimento. Después del método gráfico, se considera la técnica de comparación por parejas introducida por Fisher en 1935 llamada método de la diferencia mínima significativa o método LSD (Least Significant Difference), basada en la construcción de tests de hipótesis para la diferencia de cualquier par de medias.

Cuando el número posibles de comparaciones es elevado, la aplicación reiterada de este procedimiento, para un nivel de significación a dado, puede conducir a un número grande de rechazos de la hipótesis nula aunque no existan diferencias reales. El intento de disminuir el problema de falsos rechazos justifica la introducción de otros procedimientos

para comparaciones múltiples.

Entre estos métodos se estudia primero el basado en la desigualdad de Bonferroni. Enseguida, se aborda otra forma de solventar las deficiencias mencionadas anteriormente basada en el rango estudentizado, que da lugar al método de la diferencia significativa honesta propuesto por Tukey o método HSD (Honestly Significant Difference) y procedimientos de rangos múltiples de Newman-Keuls y Duncan.

Finalmente, se estudia el procedimiento general de Scheffé, basado en la construcción de intervalos de confianza simultáneos para todas las posibles diferencias de medias y que permite su extensión a comparaciones más generales denominadas contrastes.

Otro problema de índole diferente ligado a las comparaciones múltiples consiste en las comparaciones de distintos tratamientos con un control que suele emplearse en determinados campos de investigación; por ejemplo, en estudios epidemiológicos. Dicho problema se aborda mediante el contraste de Dunnett.

De acuerdo con (Padrón, 1996) y Monar (2017), Contrastes Ortogonales es la prueba de comparación de medias de tratamientos que registra un sólo grado de libertad debido a que representa un contraste con un grado de libertad y, a diferencia de pruebas con Factores Cualitativos, en análisis de comparaciones se usan totales de tratamientos en vez de medias, pues ahorra y evita errores por redondeo de cifras e incluye todos los datos en vez de un valor medio que puede ser sesgado en el cálculo.

Enseguida se presentan los Coeficientes y grados del polinomio para Comparaciones Ortogonales:

Tabla 2.2. Coeficientes y grados del polinomio

Número de niveles del factor	Grado del Polinomio	Totales de niveles del factor en estudio					
		t_1	t_2	t_3	t_4	t_5	$\sum C_i^2$
2	1	-1	+1				2
3	1	-1	0	+1			2
	2	+1	-2	+1		6	
4	1	-3	-1	+1	+3	20	
	2	+1	-1	-1	+1	4	
	3	-1	+3	-3	+1	20	
5	1	-2	-1	0	+1	+2	10
	2	+2	-1	-2	-1	+2	14
	3	-1	+2	0	-2	+1	10
	4	+1	-4	+6	-4	+1	10

Fuente: Modificado de (Padrón, 1996)

Es importante mencionar que en la Prueba de Contrastes Ortogonales los Coeficientes, grados del polinomio, símbolos, cantidades simbólicas de comparaciones positivas (+) y/o negativas (-) pueden cambiar sin influir en el resultado. Por ejemplo +3 -1 -1 -1 ó -3 +1 +1 +1, +6 -2 -2 -2 ó -6 +2 +2 +2, +12 -4 -4 -4 ó -12 +4 +4 +4, +18 -6 -6 -6 ó -18 +6 +6 +6 etcétera.

Con base en la experiencia del investigador, la Prueba de Contrastes Ortogonales será previamente conocida, planificada y orientada en conocer la diferencia de medias de tratamientos de interés específico.

Enseguida, se muestra un ejemplo de su estimación en Diseño Experimental de Parcela Sub-Sub Dividida o Doblemente Dividida con Bloques Completos al Azar (PSSD-DBCA) con dos niveles de factores A, B y tres de C:

Tabla 2.3. Análisis de Varianza

Análisis de Varianza (ANOVA-ANVA-ANDEVA) para Diseño de Parcelas Divididas con Tres Factores (($\cong 2^3$ o ABC)					
F. de V (Factor de Variación)	Gl (Grados de Libertad)	SC (suma de Cuadrados)	CM (Cuadrado Medio)	F (calculada)	F (calculada)
VF (Variation factor)	Df (Degrees of freedom)	Sum Sq (sum of square)	Mean Sq (Mean squares)	F value (Calculated)	
Bloques (repeticiones)	$r - 1$	$\sum_{j=1}^r \frac{Y_{.j}^2}{abc} - \frac{Y_{...}^2}{rabc}$	$\frac{SC_{\text{Bloques (repes)}}}{Gl_{\text{Bloques (repes)}}$	$\frac{CM_{\text{Bloques (repes)}}}{CM_{\text{Error A}}}$	
Factor A (Tratamiento primario)	$a - 1$	$\sum_{i=1}^a \frac{Y_{i..}^2}{rbc} - \frac{Y_{...}^2}{rabc}$	$\frac{SC_A}{Gl_A}$	$\frac{CM_A}{CM_{\text{Error A}}}$	
Error A	$(a - 1)(r - 1)$	$SC_{\text{Parcela Grande}} - SC_A - SC_{\text{Bloques (repes)}}$	$\frac{SC_{\text{Error A}}}{Gl_{\text{Error A}}}$		
Parcela Grande		$\sum_{i=1}^a \sum_{j=1}^r \frac{Y_{ij.}^2}{bc} - \frac{Y_{...}^2}{rabc}$			
Sub Parcela		$\sum_{i=1}^a \sum_{k=1}^b \frac{Y_{ijk}^2}{c} - \frac{Y_{...}^2}{rabc}$			
B (Tratamiento secundario)	$b - 1$	$\sum_{k=1}^b \frac{Y_{.k.}^2}{rac} - \frac{Y_{...}^2}{rabc}$	$\frac{SC_B}{Gl_B}$	$\frac{CM_B}{CM_{\text{Error B}}}$	

Tabla 2.3 Análisis de Varianza

Análisis de Varianza (ANOVA-ANVA-ANDEVA) para Diseño de Parcelas Divididas con Tres Factores ((≅ 2 ³ o ABC)					
F. de V _(Factor de Variación)	Gl _(Grados de Libertad)	SC _(Suma de Cuadrados)	CM _(Cuadrado Medio)	F _(calculada)	
VF _(Variation factor)	Df _(Degrees of freedom)	Sum Sq _(sum of square)	Mean Sq _(Mean squares)	F value _(Calculated)	
Interacción AB	(a - 1)(b - 1)	$\sum_{i=1}^a \sum_{k=1}^b \frac{Y_{ik}^2}{rc} - \frac{Y_{...}^2}{rabc} - SC_A - SC_B$	$\frac{SC_{AB}}{Gl_{AB}}$	$\frac{CM_{AB}}{CM_{Error B}}$	
Error B	a (b - 1)(r - 1)	$SC_{Sub Parcela} - SC_B - SC_{AB} - SC_{Parcela Grande}$	$\frac{SC_{Error B}}{Gl_{Error B}}$		
C (Tratamiento terciario)	c - 1	$\sum_{j=1}^c \frac{Y_{...j}^2}{rab} - \frac{Y_{...}^2}{rabc}$	$\frac{SC_C}{Gl_C}$	$\frac{CM_C}{CM_{Error C}}$	
PRUEBAS DE CONTRASTES ORTOGONALES					
$\overline{V_1}$ vs $\overline{V_2 V_3}$	1	$SC_{1, \overline{V_1} vs \overline{V_2 V_3}} = \frac{C_1^2 \text{ ó Factor } Q_1^2}{\sum_{j=1}^{n=4} r \cdot Q}$	$\frac{SC_{V_1 vs V_2 V_3}}{Gl_{V_1 vs V_2 V_3}}$	$\frac{CM_{V_1 vs V_2 V_3}}{CM_{Error C}}$	
$\overline{V_2}$ vs $\overline{V_3}$	1	$SC_{2, \overline{V_2} vs \overline{V_3}} = \frac{C_2^2 \text{ ó Factor } Q_2^2}{\sum_{j=2}^{n=4} r \cdot Q}$	$\frac{SC_{V_2 vs V_3}}{Gl_{V_2 vs V_3}}$	$\frac{CM_{V_2 vs V_3}}{CM_{Error C}}$	
POLINOMIOS ORTOGONALES					
Lineal	1	$SC_{Función Lineal} \cdot \frac{C^2 \text{ o Factor } Q^2}{r \cdot \sum 1^2}$	$\frac{SC_{Función Lineal}}{Gl_{Función Lineal}}$	$\frac{CM_{Función Lineal}}{CM_{Error C}}$	

Tabla 2.3 Análisis de Varianza

Análisis de Varianza (ANOVA-ANVA-ANDEVA) para Diseño de Parcelas Divididas con Tres Factores ($\cong 2^3$ o ABC)					
F. de V (Factor de Variación) 0	Gl (Grados de Libertad) 0	SC (Suma de Cuadrados) 0	CM (Cuadrado Medio) 0	F (calculada) 0	F (calculada) 0
VF (Variation factor)	Df (Degrees of freedom)	Sum Sq (Sum of square)	Mean Sq (Mean squares)	F value (calculated)	F value (calculated)
Cuadrático	1	SC Función Cuadrática: $\frac{Q^2 \text{ o Factor } Q^2}{r \cdot \Sigma I_i^2}$	$\frac{SC_{\text{Función Cuadrática}}}{Gl_{\text{Función Cuadrática}}}$	$\frac{CM_{\text{Función Cuadrática}}}{CM_{\text{Error C}}}$	$\frac{CM_{\text{Función Cuadrática}}}{CM_{\text{Error C}}}$
Interacción AC	$(a-1)(c-1)$	$\sum_{i=1}^a \sum_{k=1}^c \frac{Y_{ik}^2}{rb} - \frac{Y_{...}^2}{rabc} - SC_A - SC_C$	$\frac{SC_{AC}}{Gl_{AC}}$	$\frac{CM_{AC}}{CM_{\text{Error C}}}$	$\frac{CM_{AC}}{CM_{\text{Error C}}}$
Interacción BC	$(b-1)(c-1)$	$\sum_{k=1}^c \sum_{j=1}^b \frac{Y_{jk}^2}{ra} - \frac{Y_{...}^2}{rabc} - SC_B - SC_C$	$\frac{SC_{BC}}{Gl_{BC}}$	$\frac{CM_{BC}}{CM_{\text{Error C}}}$	$\frac{CM_{BC}}{CM_{\text{Error C}}}$
Interacción ABC	$(a-1)(b-1)(c-1)$	$\sum_{i=1}^a \sum_{j=1}^b \sum_{k=1}^c \frac{Y_{ijk}^2}{r} - \frac{Y_{...}^2}{rabc} - SC_A - SC_B - SC_C - SC_{AB} - SC_{AC} - SC_{BC}$	$\frac{SC_{ABC}}{Gl_{ABC}}$	$\frac{CM_{ABC}}{CM_{\text{Error C}}}$	$\frac{CM_{ABC}}{CM_{\text{Error C}}}$
Error C	$ab(c-1)(r-1)$	$SC_{\text{Total}} - SC_{\text{Bloques(repues)} - SC_A - SC_{\text{Error A}} - SC_B - SC_{\text{Error B}} - SC_C - SC_{\text{Error C}} - SC_{AC} - SC_{AB} - SC_{BC} - SC_{ABC}$	$\frac{SC_{\text{Error C}}}{Gl_{\text{Error C}}}$		

Tabla 2.3 Análisis de Varianza

Análisis de Varianza (ANOVA-ANVA-ANDEVA) para Diseño de Parcelas Divididas con Tres Factores ($\cong 2^3$ o ABC)				
F. de V _(Factor de Variación)	Gl _(Grados de Libertad)	SC _(Suma de Cuadrados)	CM _(Cuadrado Medio)	F _(calculada)
0	0	0	0	0
VF _(Variation factor)	Df _(Degrees of freedom)	Sum Sq _(Sum of square)	Mean Sq _(Mean squares)	F value _(Calculated)
Total	rabc - 1	$\sum_{i=1}^a \sum_{j=1}^r \sum_{k=1}^b \sum_{l=1}^c Y_{ijkl}^2 - \frac{Y_{....}^2}{rabc}$		

Tabla 2.4. Comparaciones de medias de tratamientos

Comparación	Variedades de Linaza (Linum usitatissimum)			Comparaciones					
	Var 1	Var 2	Var 3						
	$\sum \text{Var 1}$	$\sum \text{Var 2}$	$\sum \text{Var 3}$	C_i ó Factor Q	C_i^2 ó Factor Q_i^2	$\sum_{i=1}^{n=k} P_i$	$\sum_{i=1}^{n=k} P_i^2$	$\sum_{i=1}^{n=k} r$	$SC_i = \frac{C_i^2 \text{ ó Factor } Q_i^2}{\sum_{i=1}^{n=4} r * Q}$
$\overline{V_1}$ vs $\overline{V_2 V_3}$	-2	+1	+1	$C_1 \text{ o } Q_1$ $= (\sum \text{Var 1}$ $- \sum \text{Var 2})$ +	$C_1^2 \text{ ó } Q_1^2$ $= (\sum \text{Var 1}$ $- \sum \text{Var 2})^2$ +	-2 + 1 + 1 = 0	$(-2)^2$ + $(+1)^2$ + $(+1)^2$ = 6	$\sum_{i=1}^{n=4} P_i^2$ * rab	$C_1^2 \text{ ó Factor } Q_1^2$ $\sum_{i=1}^{n=4} r * Q$

Tabla 2.4 Comparaciones de medias de tratamientos

Comparación	Variedades de Linaza (Linum usitatissimum)			Comparaciones					
	Var 1	Var 2	Var 3						
	$\sum \text{Var 1}$	$\sum \text{Var 2}$	$\sum \text{Var 3}$	C_i ó Factor Q	C_i^2 ó Factor Q_i^2	$\sum_{i=1}^{n=k} P_i$ = 0	$\sum_{i=1}^{n=k} P_i^2$ = 0	$\sum_{i=1}^{n=k} r$ * Q	SC_i = $\frac{C_i^2 \text{ ó Factor } Q_i^2}{\sum_{i=1}^{n=4} r * Q}$
				$(\sum \text{Var 1}$ - $\sum \text{Var 2})$ ó $(-2$ + $\sum \text{Var 1})$ + $(+1$ + $\sum \text{Var 2})$ + $(+1 * \sum \text{Var 3})$	$(\sum \text{Var 1}$ - $\sum \text{Var 2})^2$ ó $(-2$ + $\sum \text{Var 1})^2$ + $(+1$ + $\sum \text{Var 2})^2$ + $(+1 * \sum \text{Var 3})^2$				

Tabla 2.4 Comparaciones de medias de tratamientos

Comparación	Variedades de Linaza (Linum usitatissimum)			Comparaciones					
	Var 1	Var 2	Var 3						
	$\sum \text{Var 1}$	$\sum \text{Var 2}$	$\sum \text{Var 3}$	C_i ó Factor Q	C_i^2 ó Factor Q_i^2	$\sum_{i=1}^{n-k} P_i$ $= 0$	$\sum_{i=1}^{n-k} P_i^2$ $= 0$	$\sum_{i=1}^{n-k} r$ $= + Q$	SC_i $= \frac{C_i^2 \text{ ó Factor } Q_i^2}{\sum_{i=1}^{n-k} r + Q}$
					$\left(+1 \right. \\ \left. + \sum \text{Var 3} \right)^2$				
$\overline{V_2}$ vs $\overline{V_3}$	0	-1	+1	C_2 ó Q_2 $= \left(\sum \text{Var 2} \right. \\ \left. - \sum \text{Var 3} \right)$ ó $\left(0 \right. \\ \left. + \sum \text{Var 1} \right)$	C_2^2 ó Q_2^2 $= \left(\sum \text{Var 2} \right. \\ \left. - \sum \text{Var 3} \right)^2$ ó $\left(0 \right. \\ \left. + \sum \text{Var 1} \right)^2$ +	$0 - 1$ $+ 1$ $= 0$	$(0)^2$ $+ (-1)^2$ $+ (+1)^2$ $= 2$	$\sum_{i=1}^{n=4} P_2^2$ $= \text{rab}$	$\frac{C_2^2 \text{ ó Factor } Q_2^2}{\sum_{i=2}^{n=4} r + Q}$

Tabla 2.4 Comparaciones de medias de tratamientos

Comparación	Variedades de Linaza (Linum usitatissimum)			Comparaciones					
	Var 1	Var 2	Var 3						
	$\sum \text{Var 1}$	$\sum \text{Var 2}$	$\sum \text{Var 3}$	C_i ó Factor Q	C_i^2 ó Factor Q_i^2	$\sum_{i=1}^{n=k} P_i$ $= 0$	$\sum_{i=1}^{n=k} P_i^2$ $= 0$	$\sum_{i=1}^{n=k} r$ $* Q$	SC_i $= \frac{C_i^2 \text{ ó Factor } Q_i^2}{\sum_{i=1}^{n=4} r * Q}$
				$\begin{pmatrix} -1 \\ * \sum \text{Var 2} \\ + \\ +1 \\ * \sum \text{Var 3} \end{pmatrix}$	$\begin{pmatrix} -1 \\ * \sum \text{Var 2} \\ + \\ +1 \\ * \sum \text{Var 3} \end{pmatrix}^2$				

Fuente: Elaboración propia con base en modificaciones de (González B. G., Métodos estadísticos y principios de diseño experimental, 2010) y (Hurtado & Gómez, 2010). Nota: r = Número de repeticiones, a = Niveles de Factor A y b = Niveles de Factor B.

De acuerdo con Monar (2017) y los métodos de estimación opcionales $C_i \text{ ó } Q_i$, con símbolo positivo (+) o $C_i * \sum \text{Var } i$, con símbolo negativo (-), del cálculo C_i ó Factor Q no es importante, pues la estimación de C_i^2 ó Factor C_i^2 transforma en valor positivo (+) cualquiera de ambas opciones.

Sin embargo, existen algunas Reglas de Contraste:

- Cualquier función de la forma $P_1T_1 + P_2T_2 + \dots + P_kT_k$ se llama contraste si $\sum_{i=1}^{n=k} P_i = 0$ es decir, si la suma de coeficientes es igual a 0. Por ejemplo: $-3 + 1 + 1 + 1 = 0$; $0 - 2 + 1 + 1 = 0$; $0 + \dots + 0 - 1 + 1 = 0$ etcétera. Entonces si $C_1 = T_1 - T_3$ ($+1 - 1 = 0$) ó $C_2 = T_1 + T_3 - 2T_2$ ($+1 + 1 - 2 = 0$) tal que C_1 y C_2 son contrastes.

- Si C es un posible contraste. El C_i Coeficiente es establecido por el investigador en función de comparaciones ortogonales que se haya establecido- son los mínimos números enteros que sumados algebraicamente dar $C_1 = (V_1 - V_2) + (V_1 - V_3) + (V_1 - V_4)$, $C_2 = (V_2 - V_3) + (V_2 - V_4)$ y $C_3 = (V_3 - V_4)$ ó, según Monar (2017) y Ojeda (2017),

- $Q_i = \sum T_i$ (Total/Tratamiento) * C_i (Coeficiente: $-3, +1$).

- Si C es un posible contraste, la cantidad $SC_i = \frac{C_i^2 \text{ ó Factor } Q_i^2}{\sum_{i=1}^{n=4} r * Q}$ es una componente de suma de cuadrados de tratamientos y representa un contraste con un grado de libertad.

- $\sum_{i=1}^{n=k} r * Q, r^*$ niveles de factor implica elementos faltantes o que no son estudiados -Factor A por ejemplo- y existen otros casos simples en que sólo se considera r. Por ejemplo: un DBCA con sólo 5 variedades de maíz con 4 repeticiones.

La comparación de medias de tratamientos se interpreta como cualquier ANOVA normal, tomando en cuenta sus seis criterios de lectura.

Es el experimento planeado para detectar cómo responde los niveles de un factor cuando se comparan con varios niveles de otro factor. Su metodología es un conjunto de técnicas matemático-estadísticas utilizadas para modelar y analizar problemas en que una variable de interés, dependiente o exógena es influenciada por otras, independientes o exógenas.

Su objetivo es optimizar la variable de interés que se logra al determinar las condiciones óptimas de operación del

sistema. Por ejemplo: aplicar varias dosis de un fungicida en diferentes niveles de humedad ambiental en un invernadero o usar niveles de nitrógeno en dosis de fósforo.

En estos experimentos el investigador busca el punto máximo de respuesta al factor estudiado, pues en general, el tipo de respuesta en experimentación no es Modelo Lineal a pesar que las respuestas de mayor interés son la lineal, cuadrática y/o cúbica, sino curvilínea empleando Modelos Cuadrático, Cúbico, Potencial, Exponencial, Logarítmico, Inverso, Sigmoidal, etcétera o usando transformación Box-Cox.

En esta situación es posible aplicar el análisis de regresión. Sin embargo, cuando el incremento o intervalo entre niveles de un factor son iguales se puede hacer análisis de datos en forma más sencilla mediante Coeficientes de Polinomios Ortogonales.

Si sólo la respuesta lineal es significativa se deduce que su respuesta a varios niveles de un factor es constante; es decir, de cada incremento del factor se obtiene un aumento adicional en otro factor. La respuesta puede ser positiva o negativa, según aumente o disminuya la respuesta del segundo factor en relación con el primero. Si la función cuadrática es significativa, la respuesta del segundo no es constante; es decir, los datos se ajustan a una parábola.

Respuesta es una cantidad medible, cuyo valor se afecta por el cambio de niveles de factores, entendidos como condiciones cuantitativas o cualitativas del proceso que influyen la variable de respuesta, dependiente o endógena.

Al mencionar que un valor de respuesta Y depende de niveles de variables independientes o exógenas ($X: X_1, X_2, X_3, X_4, \dots, X_k$) de K factores ($\varepsilon: \varepsilon_1, \varepsilon_2, \varepsilon_3, \varepsilon_4, \dots, \varepsilon_k$)... se entiende que existe una función matemática de $X_1, X_2, X_3, X_4, \dots, X_k$, cuyo valor para combinación de niveles de factores corresponde a

$$Y = f(X_1, X_2, X_3, X_4, \dots, X_k).$$

La función de respuesta se puede presentar con una ecuación polinomial. El éxito de una investigación de superficie de respuesta depende que se pueda ajustar a un polinomio de primer o segundo grado.

Por ejemplo: suponga que la función de respuesta para niveles de dos factores puede expresarse usando el polinomio de primer grado $Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2$ tal que $\beta_0, \beta_1, \beta_2$ son coeficientes de regresión a estimar, X_1 y X_2 representar niveles de ϵ_1 y ϵ_2 , respectivamente.

Suponga que se recolectan $N \geq 3$ valores respuesta (Y), con estimadores b_0, b_1 y b_2 , se obtienen b_0, b_1 y b_2 , respectivamente. Al sustituir coeficientes de regresión por estimadores se obtiene ecuación estimada $\hat{Y} = b_0 + b_1 X_1 + b_2 X_2$.

Sea una función $y = f(x)$ y $(R_f(\text{Rango de función y } \text{ó } R_f)) f(x)(D_f(\text{Dominio de función y } \text{ó } D_f))$ tal que el Dominio de una Función (D_f) es el conjunto de valores que toma X en la función ó Dominio de Función (variable explicativa, independiente o exógena), mientras que Rango de una Función (R_f) es el conjunto de valores que toma Y en F_x ó D_f (variable explicada, dependiente o endógena).

Ejemplos:

Tabla 2.5. Estimación de Dominio de Funciones Polinómicas

Cálculo o estimación de Dominio de Funciones Polinómicas			
Función	Ecuación matemática	Estimación de Dominio	Observaciones
Primer grado	$f(x) = 3x - 7$	$D_f: x \in \mathbb{R}$	En cualquier de estos casos el conjunto de Dominio sería el conjunto de números reales, representado por $\mathbb{R}$, pues la variable explicativa, independiente, exógena puede tomar cualquier valor del conjunto de números reales.
Cuadrática	$g(t) = 4t^2 - t + 1$	$D_g: t \in \mathbb{R}$	
Quinto grado	$r(u) = u^2 - 3u^5 + 10$	$D_r: u \in \mathbb{R}$	
Cálculo o estimación de Dominio de Funciones Racionales			
Racional (fracción o cocientes de expresiones)	$h(x) = \frac{x - 6}{x - 5}$	$D_h: x \in \mathbb{R} - \{5\}$ o $(-\infty, 5) \cup (5, +\infty)$	Es importante cuidar que su denominador no sea = 0; es decir, $x - 5 \neq 0$, pues sería una función indeterminada. Entonces, si $x - 5 \neq 0 \Rightarrow x \neq 5$. Por lo tanto, el dominio de la función h son x pertenecientes al conjunto de los reales ($\mathbb{R}$), exceptuando elemento 5 ($-\{5\}$).
	$l(m) = \frac{x - 1}{2x + 3}$	$D_l: m \in \mathbb{R} - \left\{-\frac{3}{2}\right\}$	Es importante cuidar que su denominador no sea = 0; es decir, $2x + 3 \neq 0$, pues sería una función indeterminada. Entonces, si $2x + 3 \neq 0 \Rightarrow x \neq -\frac{3}{2}$. Por lo tanto, el dominio

Tabla 2.5. Estimación de Dominio de Funciones Polinómicas

Cálculo o estimación de Dominio de Funciones Polinómicas			
Función	Ecuación matemática	Estimación de Dominio	Observaciones
		$\circ (-\infty, -\frac{3}{2}) \cup (-\frac{3}{2}, +\infty)$	de la función I son m pertenecientes al conjunto de los reales ($\mathbb{R}$), exceptuando elemento $-\frac{3}{2}$ ($-\{\frac{3}{2}\}$).
	$y = f(x) = \frac{x+3}{x-2}$	$D_y: x \in \mathbb{R} - \{2\}$ $\circ (-\infty, 2) \cup (2, +\infty)$	Es importante cuidar que su denominador no sea 0; es decir, $x - 2 \neq 0$, pues sería una función indeterminada. Entonces, si $x - 2 \neq 0 \Rightarrow x \neq 2$. Por lo tanto, el dominio de la función y son x pertenecientes al conjunto de los reales ($\mathbb{R}$), exceptuando elemento 2 ($-\{2\}$). Si a la función y se le dieran valores tanto por la izquierda ($\frac{1.99+3}{1.99-2} = -499$) como por la derecha de 2 ($\frac{2.01+3}{2.01-2} = 501$) implica que si x tiende a 2 por la izquierda ($x \rightarrow 2^-$) $\Rightarrow$ y tiende a menos infinito ($y \rightarrow -\infty$), pero si x tiende a 2 por la derecha ($x \rightarrow 2^+$) $\Rightarrow$ y tiende a más infinito ($y \rightarrow +\infty$), se presentaría una asíntota vertical

Tabla 2.5. Estimación de Dominio de Funciones Polinómicas

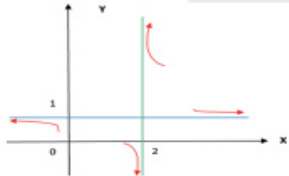
Cálculo o estimación de Dominio de Funciones Polinómicas			
Función	Ecuación matemática	Estimación de Dominio	Observaciones
			(recta a la que se aproxima, pero sin haber contacto) en 2, siendo su asíntota horizontal en azul ($y = \frac{1x+3}{1x-2} = \frac{1}{1} = 1$):
			
	$R(n) = \frac{n+1}{n^2-6n+8}$	$D_h: x \in \mathbb{R} - \{2; 4\}$ $\circ (-\infty, 2) \cup (2, 4) \cup (4, +\infty)$	Es importante cuidar que su denominador no sea 0; es decir, $n^2 - 6n + 8 \neq 0$, pues sería una función indeterminada. Entonces, si se factoriza su denominador sería $(n-4)(n-2) \neq 0 \Rightarrow n-4 \neq 0$ y $n-2 \neq 0 \therefore n \neq$

Tabla 2.5. Estimación de Dominio de Funciones Polinómicas

Cálculo o estimación de Dominio de Funciones Polinómicas			
Función	Ecuación matemática	Estimación de Dominio	Observaciones
			4 y $n \neq 2$. Por lo tanto, el dominio de la función R son n pertenecientes al conjunto de los reales ($\mathbb{R}$), exceptuando elementos 2 y 4 ($\{2; 4\}$).
	$Q(L) = \frac{L-2}{L^2+1}$	$D_Q: L \in \mathbb{R}$ $\circ (-\infty, +\infty)$	Es importante cuidar que su denominador no sea 0; es decir, $L^2 + 1 \neq 0$, pues sería una función indeterminada. Sin embargo, al ser cuadrado será una expresión positiva sin importan que valor tome L aunque sumando el valor 1 sería más positiva. Por lo tanto, su denominador jamás será = 0.
Cálculo o estimación de Dominio de Funciones con Raíces			
Radical par	$Z(v) = \sqrt[4]{v-9}$	$D_Z: v \in [9, \infty)$ $\circ (-\infty, +\infty)$	Si se tiene un radical -símbolo que indica que es una raíz cuadrada- par se debe garantizar que el radicando -número del que se obtiene la raíz cuadrada- no sea negativo. Entonces, $v-9 \geq 0 \Rightarrow v \geq 9$.

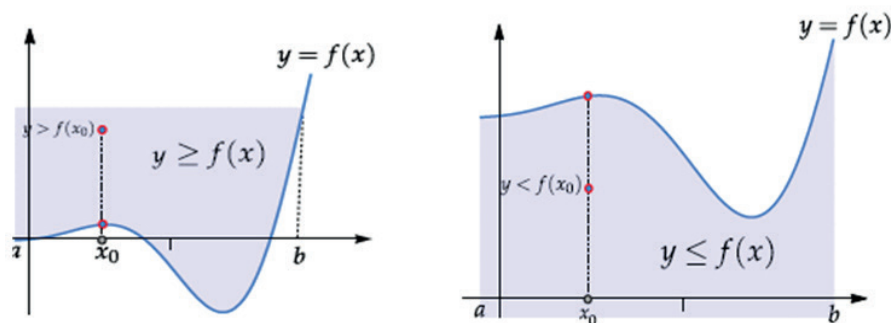
Fuente: Elaboración propia con base en modificaciones de

cursos de Diseño Experimental con Mat. Luís Ramón Guasgua
Amaguaña Ph. D.

Como en funciones de una variable, el dominio máximo de f es conjunto de puntos $(x,y) \in \mathbb{R}^2$ tal que $z=f(x,y)$ o $F(x,y,z)=0$ esté bien definida. Su representación gráfica en regiones definidas por desigualdades está dada por cualquier combinación de ecuaciones y desigualdades que pueden definir una región, de manera implícita.

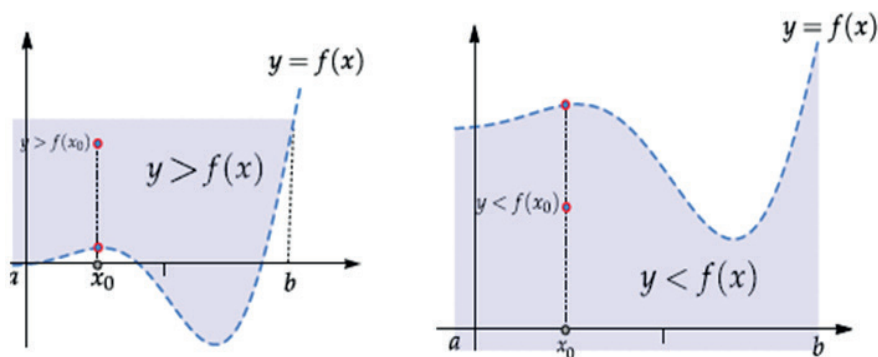
Los casos sencillos son:

• **Desigualdades tipo $y \geq f(x)$ o $y \leq f(x)$** en intervalo $[a,b]$ que consiste en representar pares (x,y) con $y=f(x)$ tal que los puntos (x,y) con $y \geq f(x)$ confirman una región por encima de la representación gráfica f , incluyéndola y, de manera análoga, puntos (x,y) con $y \leq f(x)$ conforman una región por debajo de la representación gráfica de f , incluyéndola.



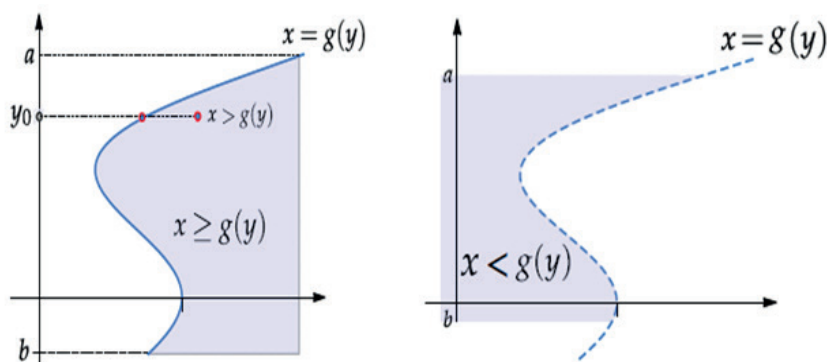
Fuente: (Mora, 2016)

• **Desigualdades tipo $y > f(x)$ o $y < f(x)$** en intervalo $[a,b]$ que consiste en representar pares (x,y) con $y=f(x)$ tal que los puntos (x,y) con $y > f(x)$ confirman una región por encima de la representación gráfica f , no incluyéndola y, de manera análoga, puntos (x,y) con $y < f(x)$ conforman una región por debajo de la representación gráfica de f , no incluyéndola.



Fuente: (Mora, 2016)

Desigualdades tipo $y \geq g(y)$ o $y \leq g(y)$ en intervalo $[a, b]$ que consiste en representar pares (x, y) con $x = g(y)$ tal que los puntos (x, y) con $y \geq g(x)$ confirman una región a la derecha de la representación gráfica de g , incluyéndola y, de manera análoga, los puntos (x, y) con $y \leq g(x)$ conforman la región a la izquierda de la representación gráfica de g , incluyéndola. Aunque, la región no incluye la curva cuando la desigualdad es estricta.

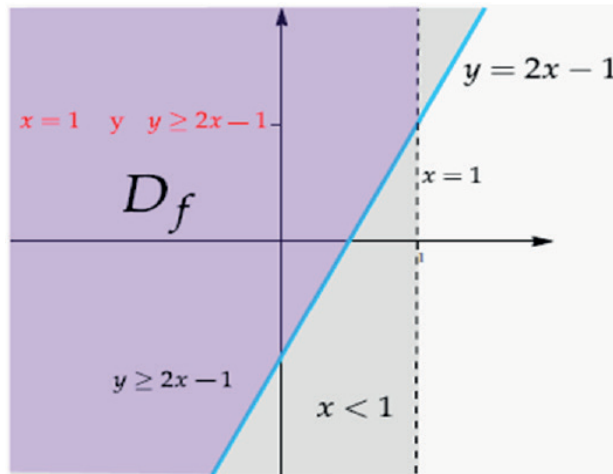


Fuente: (Mora, 2016)

Ejemplo a:

Determine y realice la representación gráfica del dominio de función $z = \sqrt{3y - 6x + 3} + \ln(1 - x) + 1$ entonces, Por lo tanto $\tan^3 y - 6x + 3 \geq 0 \Rightarrow y = \frac{6x-3}{3} = 2x - 1$ y $1 - x > 0$. Por lo tanto, $D_f: \{(x, y) \in \mathbb{R}^2 \text{ tal que } y \geq 2x - 1 \text{ y } x < 1\}$, el dominio de f es intersección de región $y \geq 2x - 1$ (región por encima de la recta $y = 2x - 1$, incluida) y región $x < 1$ (región a

izquierda de recta $x=1$, no incluida).



Fuente: (Mora, 2016)



CAPITULO III

SERIE DE TAYLOR

3.1. Generalidades

La Serie de Taylor, fue creada por el matemático británico Brook Taylor, Edmonton, Middlesex, Inglaterra, 18 de Agosto de 1685 - Somerset House, Londres, 29 de Diciembre de 1731.

En matemáticas, una serie de Taylor es una aproximación de funciones, entendida magnitud o cantidad es función de otra si el valor de la primera depende del valor de la segunda, como el área A de un círculo es función de su radio r tal que el valor del área es proporcional al cuadrado del radio ($A = \pi * r^2$) o duración T de un viaje en tren entre dos ciudades separadas por una distancia d de 150 km depende de la velocidad v a que se desplace el tren tal que la duración es inversamente proporcional a la velocidad, d/v) explicando que A es la primera magnitud (área o duración) se la denomina variable dependiente y la cantidad de la que depende (radio o velocidad) es la variable independiente, mediante una serie de potencias, entendida como una serie de forma $\sum_{n=0}^{\infty} a_n (x - c)^n = a_0 + a_1(x - c)^1 + a_2(x - c)^2 + a_3(x - c)^3 + a_4(x - c)^4 + \dots + a_n(x - c)^n$, alrededor de $c=x$ tal que C es centro y a_n sus coeficientes, entendidos como términos de una sucesión que usualmente corresponde con la serie de Taylor de alguna función conocida aunque si, en ocasiones, el centro $C=0$ la serie se denomina de MacLaurin, o suma de potencias enteras de polinomios como $(x - a)^n$.

3.2. Ejemplos de la serie de Taylor

Algunos ejemplos de la serie de Taylor son:

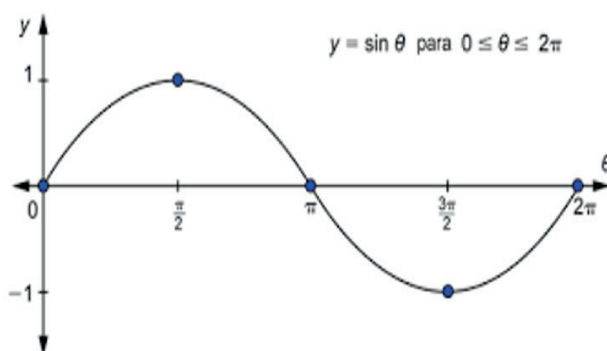
3.2.1. Serie Geométrica

Es serie en que la razón entre sus términos sucesivos permanece constante. Por ejemplo: $\frac{1}{2} + \frac{1}{4} + \frac{1}{8} + \frac{1}{16} + \dots + \frac{1}{2^n} + \dots$ se dice que es geométrica debido a que cada término sucesivo se obtiene de multiplicar el anterior valor por $\frac{1}{2}$ o cuando cada uno de los cuadrados púrpuras tiene $\frac{1}{4}$ del área del cuadrado anterior más grande ($\frac{1}{2} * \frac{1}{2} = \frac{1}{4}$, $\frac{1}{4} * \frac{1}{4} = \frac{1}{16}$) que la suma

72

3.2.3. Serie Seno

Tal que $\sin(x) = \sum_{n=0}^{\infty} \frac{(-1)^n x^{2n+1}}{(2n+1)!} = x - \frac{x^3}{3!} + \frac{x^5}{5!} - \frac{x^7}{7!} + \frac{x^9}{9!} + \dots + \frac{x^n}{n!}$



Fuente: (Zill & Wright, 2011)

3.2.4. Representación matemática

La representación matemática formal de serie de Taylor es:

$$f(x) = f(1) + f(2) + f(3) + f(4) + \dots + f(n)$$

$$f(x) \approx \frac{f(a)}{0!} + \frac{f'(a)}{1!} + \frac{f''(a)}{2!} + \frac{f'''(a)}{3!} + \frac{f^{IV}(a)}{4!} + \dots + \frac{f^n(a)}{n!}$$

$$f(x) \approx \frac{f(a)}{0!} + \frac{f'(a)}{1!}(x-a)' + \frac{f''(a)}{2!}(x-a)'' + \frac{f'''(a)}{3!}(x-a)''' + \frac{f^{IV}(a)}{4!}(x-a)'''' + \dots + \frac{f^n(a)}{n!}(x-a)^n$$

Donde, $n!$ es factorial de n y $f^n(a)$ denota n -ésima derivada de f para el valor a de variable respecto de la que se deriva.

Serie de MacLaurin ó Serie Taylor alrededor de 0. Fue creada por el matemático escocés Colin MacLaurin, Kilmodan, Febrero de 1698 - Edimburgo, 14 de Junio de 1746. En matemáticas, una serie de Taylor está centrada sobre el punto cero, $a=0$, es llamada Serie de McLaurin o, en otras palabras, es derivada de orden cero de f como la propia f y

tanto $(x - a)^0$ como $0!$ Son definidos como 1 ($0!=1$) aunque en caso que $a=0$ la serie es denominada McLaurin. Presenta las siguientes ventajas:

- La derivación e integración de una de estas series se puede realizar término a término, que resultan operaciones triviales.

- Se puede utilizar para calcular valores aproximados de funciones.

- Es posible calcular la optimidad de la aproximación.

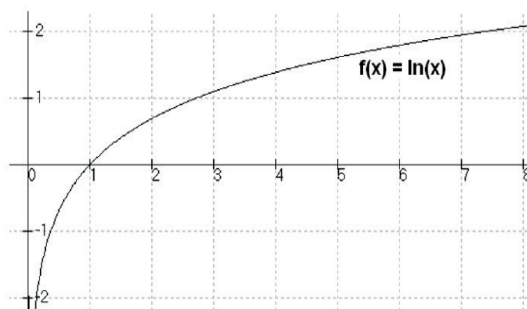
Algunos ejemplos de la serie de MacLaurin son:

Logaritmo natural

Donde $\text{Ln}(1 + x) = \sum_{n=1}^{\infty} \frac{(-1)^{n+1}}{n} x^n$. En matemáticas, se denomina logaritmo natural o informalmente logaritmo neperiano al logaritmo cuya base es el número e, un número irracional cuyo valor aproximado es 2,71828182845904523536 02874713527.

El logaritmo natural se suele denominar como $\ln(x)$ o a veces como $\log_e(x)$ e incluso en algunos contextos $\log(x)$, pues ese número cumple la propiedad que el logaritmo vale 1. El logaritmo natural de un número x es entonces el exponente a al que debe ser elevado el número e para obtener x.

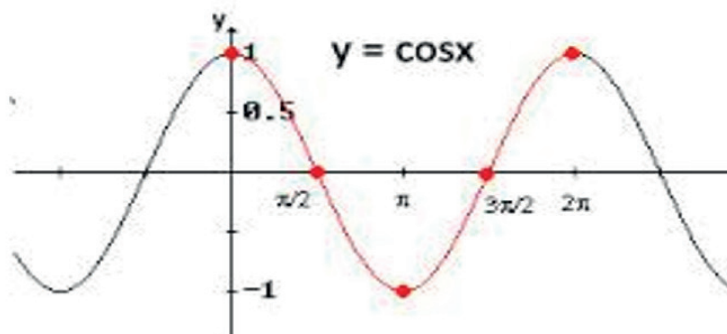
Por ejemplo, el logaritmo de 7,38905... e^2 , pues $e^2=7.38905$. El logaritmo de e es 1 debido a que $e^1=e$.



Fuente: (Zill & Wright, 2011)

3.4. Función Coseno

Donde $\cos(x) = \sum_{n=0}^{\infty} \frac{(-1)^n}{(2n)!} x^{2n}$



Fuente: (Zill & Wright, 2011)

La representación matemática formal de serie de MacLaurin es:

$$f(x) \approx \frac{f(0)}{0!} + \frac{f'(0)}{1!}(x)^1 + \frac{f''(0)}{2!}(x)^2 + \frac{f'''(0)}{3!}(x)^3 + \frac{f^{IV}(0)}{4!}(x)^4 + \dots$$

$$+ \frac{f^n(0)}{n!}(x)^n$$

Además de la obvia aplicación de utilizar funciones polinómicas en vez de funciones de mayor complejidad para analizar el comportamiento local de una función, las series de Taylor-MacLaurin tienen muchas otras aplicaciones.

Por ejemplo: análisis de límites, estudios paramétricos, estimación de números irracionales acotando su error, teorema de L'Hopital para la resolución de límites indeterminados, estudio de puntos estacionarios en funciones (máximos o mínimos relativos o puntos sillas de tendencia estrictamente creciente o decreciente), estimación de integrales, determinación de convergencia, suma de algunas series importantes, estudio de orden, parámetro principal de infinitésimos, estimación de puntos óptimos en Diseños Experimentales Diseño de Bloques Divididos Alpha Látices Simple, Bloques Divididos Alpha Látices Cuadrados usados en

mejoramiento genético, etcétera.

Ejemplo:

$$Y = f(x) = \cos x \text{ con } x = 0$$

$$f(x) = \cos x \quad f(0) = 1$$

$$f'(x) = -\sin x \quad f'(0) = 0$$

$$f''(x) = -\cos x \quad f''(0) = -1$$

$$f'''(x) = \sin x \quad f'''(0) = 0$$

$$f^{IV}(x) = \cos x \quad f^{IV}(0) = 1$$

$$f(x) \approx \frac{f(0)}{0!} + \frac{f'(0)}{1!}(x)^1 + \frac{f''(0)}{2!}(x)^2 + \frac{f'''(0)}{3!}(x)^3 + \frac{f^{IV}(0)}{4!}(x)^4 + \dots$$

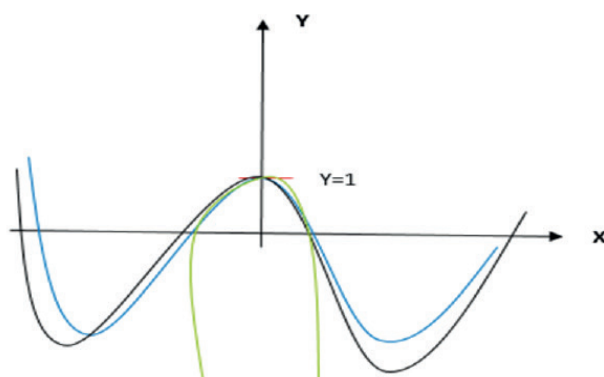
$$+ \frac{f^n(0)}{n!}(x)^n$$

$$f(x) \approx \frac{f(1)}{0!} + \frac{f'(0)}{1!}(x)^1 + \frac{f''(-1)}{2!}(x)^2 + \frac{f'''(0)}{3!}(x)^3 + \frac{f^{IV}(1)}{4!}(x)^4$$

$$+ \dots + \frac{f^n(0)}{n!}(x)^n$$

$$f(x) \approx 1 - \frac{(x)^2}{2!} + \frac{(x)^4}{4!} - \frac{(x)^6}{6!} + \frac{(x)^8}{8!} = \frac{(x)^n}{n!}$$

La representación gráfica de la convergencia de sus diferentes funciones sería:



Fuente: Elaboración propia usando Matlab con base en (Zill & Wright, 2011)

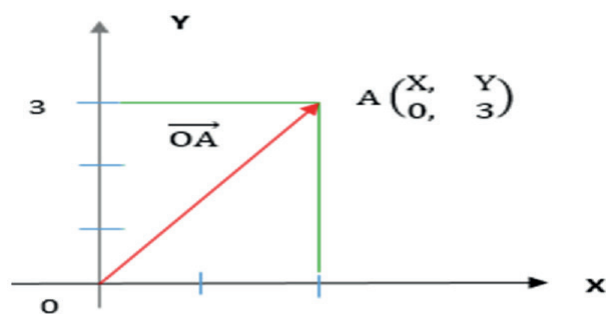
Una función escalar de dos variables $f: \mathbb{R}^2 \rightarrow \mathbb{R}$ con dominio $D \subseteq \mathbb{R}^2$ asignada a cada par $(x, y) \in D$, un único número real denotado con $f(x, y) = z$. El grafico de f es el conjunto $\{(x, y, z) : x, y, \in D \text{ y } z = f(x, y)\}$. El criterio (formula) que define a f puede ser explícito o implícito:

Función de dos variables	
Forma Explícita	Forma Implícita
$z = x^2 + y^2$	$F(x, y, z) = 0$
$f(x, y) = x^2 + y^2$	$\overbrace{F(x, y, z)}^{x^2 + y^2 + z^2 - 1 = 0}; z \geq 0$
Ejemplo: $f(1, 2) = 1^2 + 2^2 = 5$ y $f(0, 3) = 9$. Por lo tanto, los puntos $(1, 2, 5)$ y $(0, 3, 9)$ Si $x = 1$ y $y = 0$ entonces $z = 0$ y el punto $(1, 0, 0)$ está en la gráfica de la función. Si $x = 1$ y $y = 0$ entonces $1^2 + 1^2 + z^2 = 1 \Rightarrow z^2 = -1$; es decir, $(1, 1)$ no está en dominio de la función.	

Fuente: (Zill & Wright, 2011)

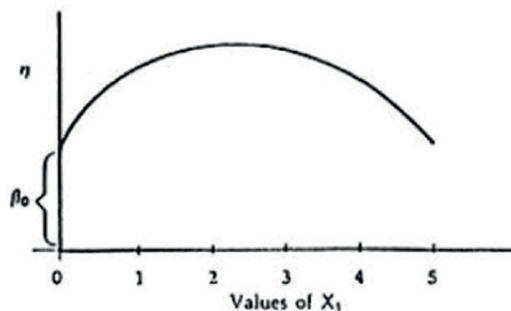
Paralelamente, existen magnitudes que quedan determinadas dando un solo número real, como longitud de regla, masa de un cuerpo, tiempo transcurrido entre dos sucesos, densidad, volumen, trabajo, potencia, aceleración, cantidad de movimiento, campo magnético, flujo de calor o materia, etcétera. Tales magnitudes se llaman escalares y pueden ser representadas sobre la recta real mediante un número que indica su medida. En otras magnitudes no es suficiente dar un número para determinarlas, sino que para conocer la magnitud o modulo –vector medida- en un punto ($\overrightarrow{OA}$: distancia entre puntos inicial O y final A o longitud del segmento), entendido como número que coincide con la “longitud” del vector en la representación gráfica o distancia euclídea, se requiere conocer dirección, medida del ángulo que hace la línea recta por la que está dada y sentido, entendido como uno de los sentidos posibles sobre recta que para por el módulo o vector –magnitudes vectoriales-, con

que el punto se mueve.



Fuente: (Mora, 2016)

Tal que, $\overrightarrow{OA} = (A_x * \vec{i}) + (A_y * \vec{j}) \Rightarrow \overrightarrow{OA} = (2\vec{i}) + (3\vec{j})$. Ahora bien, la relación $Y = f(X_1, X_2, X_3, X_4, \dots, X_k)$ entre Y y niveles de K factores ($\varepsilon: \varepsilon_1, \varepsilon_2, \varepsilon_3, \varepsilon_4, \dots, \varepsilon_k$) representa una superficie. Con K factores, la superficie está en $K+1$ dimensiones. Por ejemplo: si se tiene $f(X_1)$, la superficie está en dos dimensiones:



Fuente: (Zill & Wright, 2011)

Ejemplo α:

$$y = f(x) = \ln(x) \text{ con } x = 1$$

$$f(x) = \ln(x) \quad f(0) = 0$$

$$f'(x) = \frac{1}{x} = x^{-1} \quad f'(0) = 1$$

$$f''(x) = -1x^{-2} \quad f''(0) = -1$$

$$f'''(x) = 2x^{-3} \quad f'''(0) = 2$$

$$f^{IV}(x) = -6x^{-4} \quad f^{IV}(0) = -6$$

$$f(x) \approx \frac{f(a)}{0!} + \frac{f'(a)}{1!} + \frac{f''(a)}{2!} + \frac{f'''(a)}{3!} + \frac{f^{IV}(a)}{4!} + \dots + \frac{f^n(a)}{n!}$$

$$f(x) \approx \frac{f(a)}{0!} + \frac{f'(a)}{1!} (x-a)' + \frac{f''(a)}{2!} (x-a)'' + \frac{f'''(a)}{3!} (x-a)''' + \frac{f^{IV}(a)}{4!} (x-a)'''' + \dots + \frac{f^n(a)}{n!} (x-a)^n$$

$$f(x) \approx 0 + 1(x-1)^1 - \frac{1}{2!}(x-1)^2 + \frac{2}{3!}(x-1)^3 - \frac{6}{4!}(x-1)^4 + \dots + \frac{f^n(a)}{n!}(x-1)^n$$

$$f(x) \approx (x-1)^1 - \frac{1(x-1)^2}{2 * 1} + \frac{2(x-1)^3}{3 * 2 * 1} - \frac{6(x-1)^4}{4 * 3 * 2 * 1} + \dots + \frac{f^n(a)}{n!}(x-1)^n$$

$$f(x) \approx \frac{(x-1)^1}{1} - \frac{(x-1)^2}{2} + \frac{(x-1)^3}{3} - \frac{(x-1)^4}{4} + \dots + \frac{f^n(a)}{n!}(x-1)^n = \frac{(x-1)^n}{n!}$$

Las funciones tienen diferentes formas geométricas:

Tabla 3.1. Formas geométricas de diversas funciones

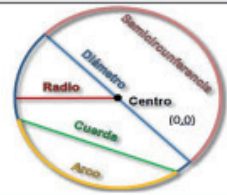
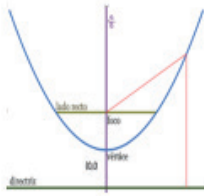
Nombre	Ecuación matemática	Equivalencia	Definición	Figura geométrica
Circunferencia	$(x-h)^2 + (y-k)^2 = r^2$	$x^2 + y^2 = r^2$ si x y y valen 0 el centro de circunferencia es un punto en el centro	Línea curva cerrada cuyos puntos equidistan de otro situado en el mismo plano que se llama centro	
Parábola	$(y-k)^2 = (x-h)^2$	$y = x^2$ y $x = y^2$	Sección cónica de excentricidad igual a 1, resultante de cortar un cono recto con un plano cuyo ángulo de inclinación respecto al eje de revolución del cono sea igual al presentado por su generatriz. Cuando $y = x^2$ las líneas van hacia arriba (+),	

Tabla 3.1. Formas geométricas de diversas funciones

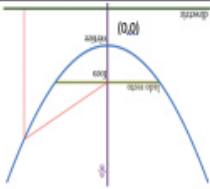
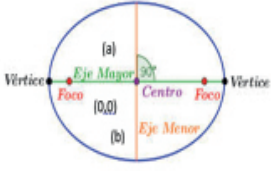
Nombre	Ecuación matemática	Equivalencia	Definición	Figura geométrica
			pero si $y = -x^2$ las líneas van hacia abajo (-).	
			Curva cerrada con dos ejes de simetría que resulta al cortar la superficie de un cono por un plano oblicuo al eje de simetría -con ángulo mayor que el de la generatriz respecto del eje de revolución. Si el denominador de x^2 es >	

Tabla 3.1. Formas geométricas de diversas funciones

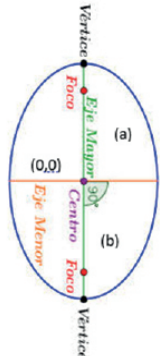
Nombre	Ecuación matemática	Equivalencia	Definición	Figura geométrica
Elipse	$\frac{(x-h)^2}{a^2} + \frac{(y-k)^2}{b^2} = 1$	$\frac{x^2}{a^2} + \frac{y^2}{b^2} = 1$	al denominador de y^2 , por ejemplo $\frac{x^2}{9} + \frac{y^2}{4} = \frac{x^2}{3^2} + \frac{y^2}{2^2} = 1$, se dice que el eje mayor es sobre eje X (elipse horizontal) y si el denominador de x^2 es < al denominador de y^2 , por ejemplo $\frac{x^2}{4} + \frac{y^2}{9} = \frac{x^2}{2^2} + \frac{y^2}{3^2} = 1$, se dice que el eje mayor es sobre eje Y (elipse vertical).	

Tabla 3.1. Formas geométricas de diversas funciones

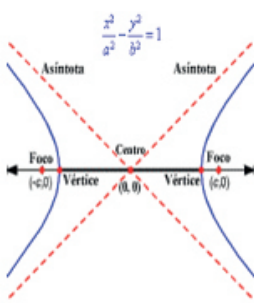
Nombre	Ecuación matemática	Equivalencia	Definición	Figura geométrica
Hipérbola	$\frac{(x-h)^2}{a^2} - \frac{(y-k)^2}{b^2} = 1$	$\frac{x^2}{a^2} - \frac{y^2}{b^2} = 1$	Sección cónica, una curva abierta de dos ramas obtenida cortando un cono recto por un plano oblicuo al eje de simetría, y con ángulo menor que el de la generatriz respecto del eje de revolución. También, se considera el lugar geométrico de los puntos de un plano tales que el valor absoluto de la diferencia de sus distancias a dos puntos fijos, llamados focos, es igual a la distancia entre los vértices, la cual es una	

Tabla 3.1. Formas geométricas de diversas funciones

Nombre	Ecuación matemática	Equivalencia	Definición	Figura geométrica
			constante positiva. Si el denominador de x^2 es $>$ al denominador de y^2, por ejemplo $\frac{x^2}{9} - \frac{y^2}{4} = \frac{x^2}{3^2} - \frac{y^2}{2^2} = 1$ se dice que las hipérbolas son sobre eje X (hipérbola horizontal) aunque si $x = y^2$ la hipérbola estará sobre lado positivo de eje X (+) en cambio si $x = -y^2$ la hipérbola estará sobre lado negativo de eje X (-), pero si el denominador de	

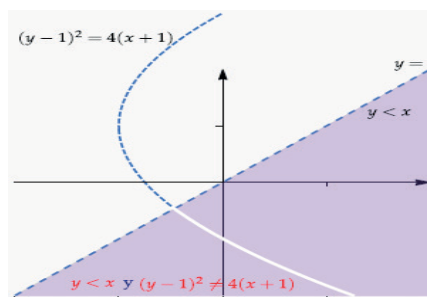
Tabla 3.1. Formas geométricas de diversas funciones

Nombre	Ecuación matemática	Equivalencia	Definición	Figura geométrica
			$\frac{x^2}{4} - \frac{y^2}{9} = \frac{x^2}{2^2} - \frac{y^2}{3^2} = 1$ se dice que las hipérbolas es sobre eje Y (hipérbola vertical).	

Fuente: Elaboración propia con base en curso de Cálculo Vectorial con Ing. Freddy Vinicio Lema Lema

Ejemplo b:

Determine y realice la representación gráfica del dominio de función $f(0,0) = \frac{1}{y^2-2y-4x-3} + \frac{1}{\sqrt{x-y}}$ tal que Dominio $\Leftrightarrow$ (si sólo si) $y^2 - 2y - 4x - 3 \neq 0 \cap x - y > 0 \Rightarrow y < x$. Entonces, $D_f = \{(x,y) \in \mathbb{R}^2: y^2 - 2y - 4x - 3 \neq 0 \Rightarrow (y-1)^2 \neq 4(x+1) \text{ o } (y-k)^2 \neq 4(x-h) \text{ (Parábola)} \cap x - y > 0 \Rightarrow x - y > 0, x > y \text{ tal que } y < x\}$.



Fuente: (Mora, 2016)

Ejemplo c:

Determine y realice la representación gráfica del dominio

de función $f(x, y) = \ln(x - y^2) + \sqrt{1 - y^2 - \frac{x^2}{2}}$ tal que Dominio

$\Leftrightarrow$ (si sólo si) $x - y^2 > 0 \Rightarrow x > y^2$ (Parábola) y $1 - y^2 - \frac{x^2}{2} \geq 0 \Rightarrow$

$$-y^2 - \frac{x^2}{2} \geq -1 \Rightarrow \frac{y^2}{1} + \frac{x^2}{2} \leq 1 \Rightarrow \frac{x^2}{2} + \frac{y^2}{1} \left(\frac{x^2}{\sqrt{2}(\sqrt{a})} + \frac{y^2}{\sqrt{1}(\sqrt{b})} \right) \leq$$

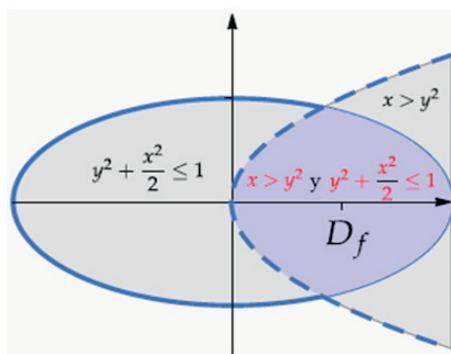
1(Elipse). Entonces, $D_f = \{(x, y) \in \mathbb{R}^2 \text{ tal que } x - y^2 > 0 \Rightarrow x >$

$y^2 \cap y^2 + \frac{x^2}{2} \leq 1\}$. Gráficamente el dominio de f es

intersección azul de región $x > y^2$ (región derecha de

parábola $x = y^2$, sin incluirla) y región $y^2 + \frac{x^2}{2} \leq 1$ (interior de

elipse $y^2 + \frac{x^2}{2} = 1$, incluyéndola).



Fuente: (Mora, 2016)

Ejemplo d:

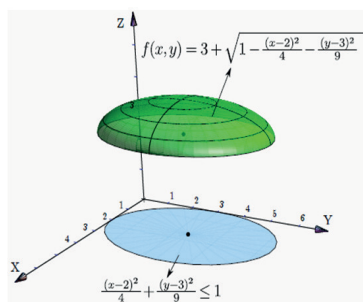
Considere la función $f(x, y) = 3 + \sqrt{1 - \frac{(x-2)^2}{4} - \frac{(y-3)^2}{9}}$; $x \geq 0$

para que la función esté definida. Entonces, el dominio

máximo es el conjunto $D_f = \{(x, y) \in \mathbb{R}^2 : \frac{(x-2)^2}{4} - \frac{(y-3)^2}{9} \leq 1\}$

(dominio de función en color celeste); es decir, D_f es región

encerrada por elipse $\frac{(x-2)^2}{4} - \frac{(y-3)^2}{9} = 1$ incluyendo el borde o límite.

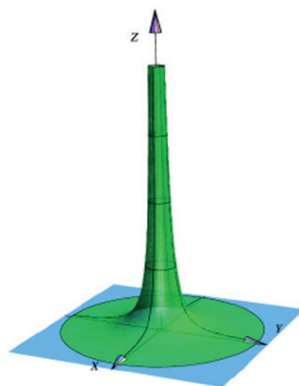


Fuente: (Mora, 2016)

Ejemplo e:

La función $z = \frac{1}{x^2+y^2}$ (dominio de función en color celeste)

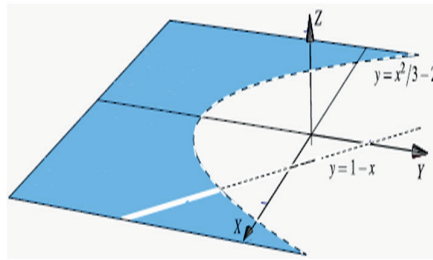
sólo se define en $(0,0)$ tal que el dominio máximo de esta función es el conjunto $D_f = \mathbb{R}^2 - \{(0,0)\}$



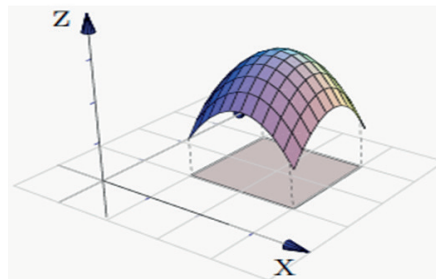
Fuente: (Mora, 2016)

Ejemplo f:

Considere la función $f(x, y) = \frac{\text{Log}(x^2-3(y+2))}{x+y-1}$ tal que $x + y - 1 \neq 0 \Rightarrow y \neq -x + 1$ (sus puntos no están en el dominio) y $x^2 - 3(y + 2) \Rightarrow y < \frac{x^2}{3} - 2$ (región por debajo de parábola $y = \frac{x^2}{3} - 2$, pero el dominio no la incluye, sino sólo el área por arriba de su línea "punteada"). Entonces, $D_f = \{(x, y) \in \mathbb{R}^2: y < \frac{x^2}{3} - 2 \cap y \neq -x + 1\}$.

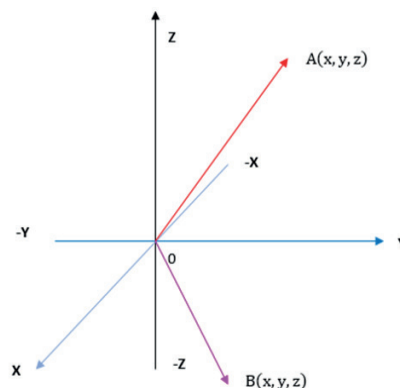
**Ejemplo g:**

Considere función $f(x,y) = 3 - (x-2)^2 - (y-2)^2 \in \mathbb{R}^2$. Su dominio máximo es $\mathbb{R}^2$ aunque se hace su representación gráfica de f sobre un dominio restringido. Por ejemplo: $D = [1, 3] * [1, 3]$.



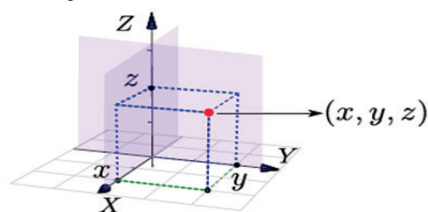
Fuente: (Mora, 2016)

En el caso de vectores $\mathbb{R}^3$ los vectores se representan gráficamente así:

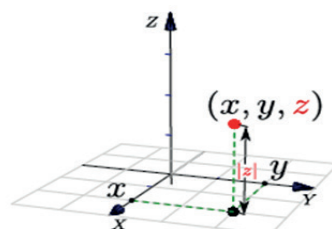


Fuente: (Mora, 2016)

Donde, $y = f(x) \Rightarrow z = f(x, y)$. De acuerdo con Mora (2016), otra forma de representar el Espacio Tridimensional con puntos $P(x, y, z) \Rightarrow x, y, z \in \mathbb{R}$ es:



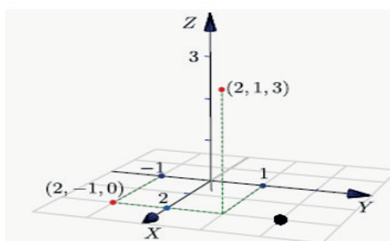
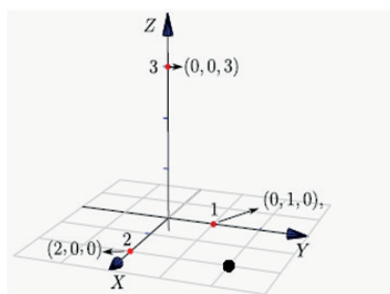
(a) Coordenadas cartesianas

(b) Punto $P = (x, y, z)$

Fuente: (Mora, 2016)

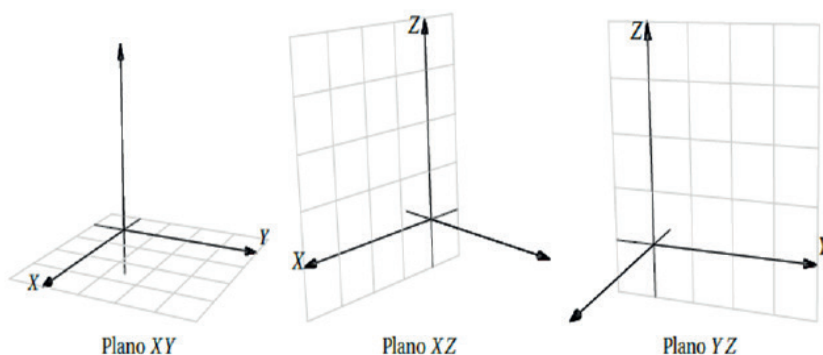
Ejemplo:

Los puntos en eje X tienen coordenadas $(x, 0, 0), x \in \mathbb{R}$, puntos en eje Y tienen coordenadas $(0, y, 0), y \in \mathbb{R}$ y puntos en eje Z tienen coordenadas $(0, 0, z), z \in \mathbb{R}$. Entonces:



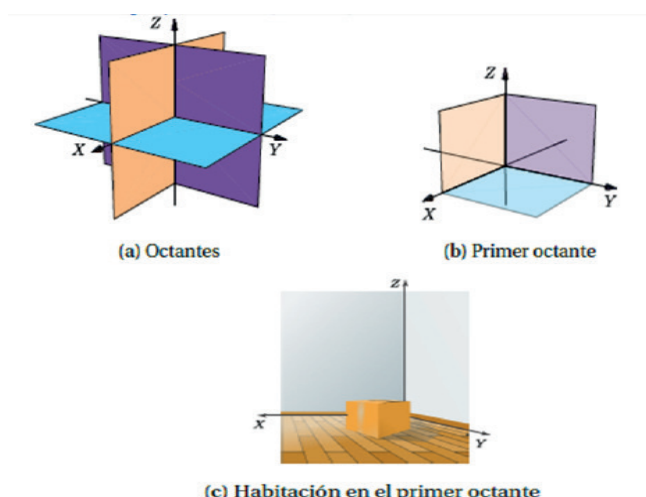
Fuente: (Mora, 2016)

Hay tres planos XY, XZ y YZ que contienen un par de coordenadas de ejes coordenados: XY contiene eje X y Y , XZ contiene X y Z y, finalmente, YZ contiene Y y Z .



Fuente: (Mora, 2016)

Hay tres planos XY, XZ y YZ dividen el espacio en ocho partes llamadas Octantes tal que el primer octante corresponde a la parte positiva de los ejes:

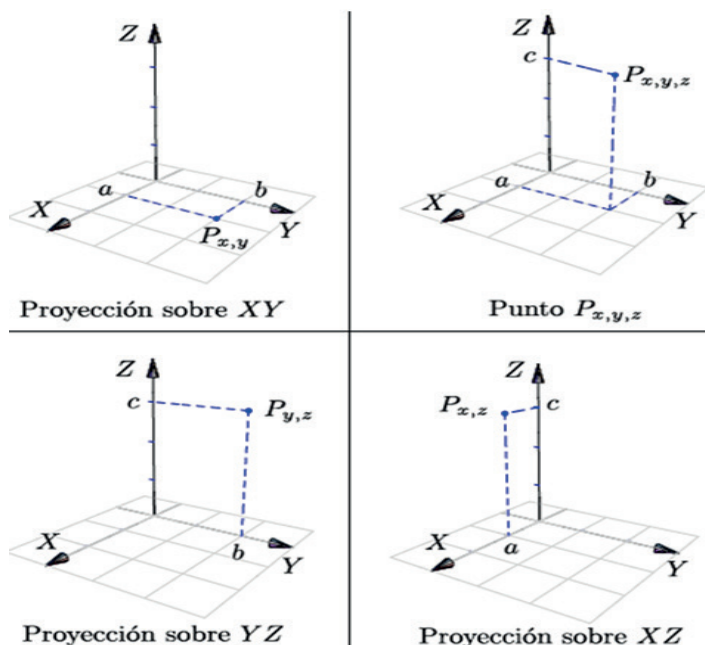


Fuente: (Mora, 2016)

3.5. Vistas isométricas de un punto

Para visualizar las vistas isométricas de un punto considere el punto $P_{(x,y,z)} = (a,b,c)$ en espacio tridimensional, se define su vista de este punto en plano XY como punto $P_{(x,y)} = (a,b,0)$. Entonces,

se define la vista en plano YZ como punto $P_{(y,z)} = (0, b, c)$ último, la vista del plano XZ como $P_{(x,z)} = (a, 0, c)$ Por lo tanto, estas vistas se denominan “Proyecciones Perpendiculares” del punto en plano respectivo:

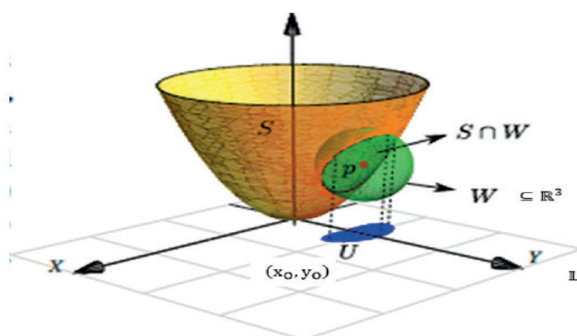


Fuente: (Mora, 2016)

Las superficies en $\mathbb{R}^3$ espacio son las superficies formadas en puntos (x,y,z) tal que su forma implícita es $z=f(x,y)$ ó, en forma explícita, $F(x,y,z)=0$. Sin embargo, no todas las superficies “suaves” en $\mathbb{R}^3$, con o sin desigualdades, se pueden describir con sólo la ecuación $F(x,y,z)=0$. En general, el estudio de superficies de manera general, se requiere estudiarlas “totalmente”. Así:

Una superficie $\mathbf{S}$ es un subconjunto de $\mathbb{R}^3$ que, en un entorno de cualquiera de sus puntos, luce como un “parche de $\mathbb{R}^2$ ”. Es decir, para cada $\mathbf{p} \in \mathbf{S} \ni$ un entorno abierto $\mathcal{U} \subseteq \mathbb{R}^2$ y un entorno $\mathcal{W} \subseteq \mathbb{R}^3$ que contiene a $\mathbf{p}$ tal que se puede establecer una biyección ($\Leftrightarrow$). continua (HOMEOMORFISMO). Tal que $r_p: \mathcal{U} \rightarrow \mathbf{S} \cap \mathcal{W}$ a cada homomorfismo r_p se llama “parche” o PARAMETRIZACIÓN del conjunto abierto $\mathbf{S} \cap \mathcal{W}$. Entonces, una

colección de “parches” que cubren S se llama “ATLAS DE S ”.



Fuente: (Mora, 2016)

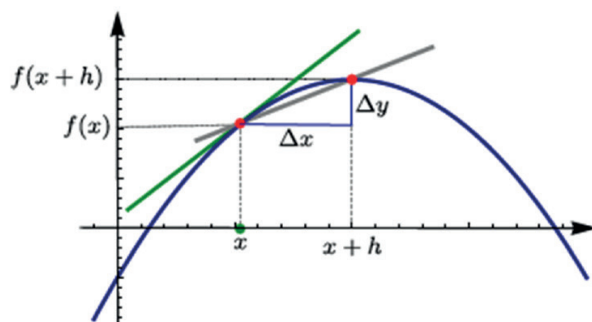
Si una superficie S tiene ecuación $z=f(x,y)$ con (x,y) ó (x,z) $(x,z) \in \mathbb{D} \subseteq \mathbb{R}^2$, entonces la superficie sería de un solo “parche” y una parametrización equivaldría a:

$$r(x,y) = x\vec{i} + y\vec{j} + f(x,y)\vec{k} \text{ ó, equivalente a } r(x,y) =$$

$$\begin{pmatrix} x \\ y \\ f(x,y) \end{pmatrix} (x,y) \in \mathbb{D}$$

Según Lema (2017) y Moral (2016), las Derivadas Parciales ($z=f(x,y)$) son conocidas en cálculo en una variable, expresadas en “Tasa de Cambio” como:

$$f'(x) = \lim_{\Delta x \rightarrow 0} \left(\frac{\Delta y}{\Delta x} \right) = \lim_{h \rightarrow 0} \left[\frac{f(x+h) - f(x)}{h} \right]$$

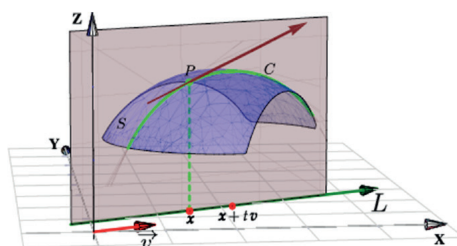


Fuente: (Mora, 2016)

Entonces, si el límite existe implica ($\Rightarrow$) que la derivada es la recta tangente a f en x . Por lo tanto, $f: \mathbb{R}^2 \rightarrow \mathbb{R}$, la derivada de f en $x = (x_0, y_0) \in \mathbb{R}^2$ en la dirección del vector unitario $v = (v_1, v_2) \in \mathbb{R}^2$ mide la tasa (instantánea) de cambio de f mediante recta $L_{(h)} = x + hv$ cuando $h=0$. Obteniendo la derivada en dirección de v con el límite. $\lim_{h \rightarrow 0} \left[\frac{f(x_0 + hv_1, y_0 + hv_2) - f(x_0, y_0)}{h} \right]$

Enseguida se tiene que el módulo de la recta como vector es $\|x - (x + hv)\| = \|x - x - hv\| = \|hv\| = h$. suponiendo que $\vec{v}$ es unitario (vector cuyo módulo es igual a 1 o vector normalizado).

Geométricamente, ésta derivada es la pendiente de la recta tangente a curva $C(h) = (x_0 + hv_1, y_0 + hv_2, f(x + hv))$ en $h=0$. Tal que, ésta curva es la intersección de superficie S de ecuación $z=f(x,y)$ con el plano generado por recta L :



Fuente: (Mora, 2016)

En cálculo diferencial, una Derivada Parcial de una función de diversas variables es la derivada respecto a cada una de esas variables manteniendo las otras como constantes ($A=f(x,y,z,...)$). Las derivadas parciales son útiles en cálculo vectorial, geometría diferencial, física matemática, etcétera. “ ∂ ” es letra “d” redondeada, conocida como “d de Jacobi”, con que se representa la “Derivada Parcial”.

El interés particular será respecto a las “Derivadas Parciales”:

$$\frac{\partial f}{\partial x} < \begin{array}{l} \text{Derivada en dirección de eje X} \\ \text{Derivada parcial respecto a X} \end{array}$$

$$\frac{\partial f}{\partial y} < \begin{array}{l} \text{Derivada en dirección de eje Y} \\ \text{Derivada parcial respecto a Y} \end{array}$$

Por definición, sea $u \subseteq \mathbb{R}^n$ un conjunto abierto y sea $f: u \rightarrow \mathbb{R}$.

Entonces, la derivada parcial $\frac{\partial f}{\partial x_i}$ de f respecto a variable X_i

en punto $x = (x_1, \dots, x_n)$ se define como:

$$\begin{aligned}\frac{\partial f}{\partial x_i} &= \lim_{h \rightarrow 0} \frac{f(x_1, x_2, \dots, x_i + h, \dots, x_n) - f(x_1, \dots, x_n)}{h} \\ &= \lim_{h \rightarrow 0} \left[\frac{f(x + he_i) - f(x)}{h} \right]\end{aligned}$$

Siempre y cuando el límite exista ($\exists$), donde $e_i = (0, \dots, 1, \dots, 0)$ con 1 en la i -ésima posición. El dominio de $\frac{\partial f}{\partial x_i}$ es el subconjunto de $\mathbb{R}^2$ en que este límite existe.

El Caso de Dos Variables implica que, siendo $z = f(x, y)$, puede denotarse de varias formas equivalentes: $\frac{\partial f}{\partial x}$, $\frac{\partial z}{\partial x}$, z_x o f_x . Entonces:

$$\begin{aligned}\frac{\partial f}{\partial x} &= \lim_{h \rightarrow 0} \left[\frac{f(x + h, y) - f(x, y)}{h} \right] \\ \frac{\partial f}{\partial y} &= \lim_{h \rightarrow 0} \left[\frac{f(x, y + h) - f(x, y)}{h} \right]\end{aligned}$$

El Cálculo Directo de las ∂ se puede comprender mediante los siguientes ejemplos:

$$1. f(x, y) = \sqrt[3]{X} * \sqrt[3]{Y} = X^{1/3} * Y^{1/3} \text{ tal que } \frac{\partial f}{\partial x} = \frac{1}{3X^{2/3}} * Y^{1/3} = \frac{\sqrt[3]{Y}}{3\sqrt[3]{X^2}}.$$

$$\text{Ahora bien, } \frac{\partial f}{\partial y} = \frac{1}{3Y^{2/3}} * X^{1/3} = \frac{\sqrt[3]{X}}{3\sqrt[3]{Y^2}}.$$

Para saber si la función es derivable en $(x, y) = (0, 0)$ se tiene que aplicar la definición:

$$\frac{\partial f}{\partial x}(0,0) = \lim_{h \rightarrow 0} \left[\frac{f(0+h,0) - f(0,0)}{h} \right] = \lim_{h \rightarrow 0} \left(\frac{0-0}{h} \right) = 0$$

$$\frac{\partial f}{\partial y}(0,0) = \lim_{h \rightarrow 0} \left[\frac{f(0,0+h) - f(0,0)}{h} \right] = \lim_{h \rightarrow 0} \left(\frac{0-0}{h} \right) = 0$$

2. $z = x^y \Rightarrow x > 0$ y, recordando regla de la cadena, es
 $(a^\mu)' = ? \Rightarrow a^\mu * \text{Ln}(a) * \mu'$

$$\frac{\partial z}{\partial x} = y * x^{y-1}$$

$$\frac{\partial z}{\partial y} = x^y * \text{Ln}(x) * y' = x^y * \text{Ln}(x)$$

Si $\mathbb{C}(r, \theta) = r^n * \text{Cos}(n\theta)$ $n \in \mathbb{N}$ en un Octante. Tal que:

$$\frac{\partial \mathbb{C}}{\partial r} = n * r^{n-1} * \text{Cos}(n\theta)$$

$$\frac{\partial \mathbb{C}}{\partial \theta} = r^n * (-n) * \text{Sen}(n\theta)$$

Si:

$$z = \frac{\text{Cos}(xy) + x\text{Sen}(2y)}{2} \Rightarrow \frac{\partial z}{\partial x} = \frac{-y \text{Sen}(xy) + \text{Sen}(2y)}{2}$$

$$\frac{\partial z}{\partial y} = \frac{-x \text{Sen}(xy) + 2x \text{Cos}(2y)}{2}$$



CAPITULO IV

DERIVADAS PARCIALES DE ORDEN SUPERIOR

4.1. Definición

Las Derivadas Parciales de Orden Superior es una derivada parcial de diversas variables, es su derivada respecto a una de esas variables manteniendo constantes las otras. La derivada de orden superior comprende las derivadas a partir de la segunda derivada o más y que se efectúa derivando tantas veces como se indique.

4.2. Representación

Las derivadas parciales son útiles en cálculo vectorial y geometría diferencial. Su representación es, si $f(x,y)$, $\frac{\partial z}{\partial x}$ y $\frac{\partial z}{\partial y}$, funciones de dos variables, se tienen $\frac{\partial^2 z}{\partial x^2}$ y $\frac{\partial^2 z}{\partial y^2}$ como segundas derivadas parciales.

También, se tienen las siguientes notaciones:

$$(f_x)_x = f_{xx} = f_{11} = \frac{\partial}{\partial x} \left(\frac{\partial f}{\partial x} \right) = \frac{\partial^2 f}{\partial x^2} = \frac{\partial^2 z}{\partial x^2}$$

$$(f_x)_y = f_{xy} = f_{12} = \frac{\partial}{\partial y} \left(\frac{\partial f}{\partial x} \right) = \frac{\partial^2 f}{\partial y \partial x} = \frac{\partial^2 z}{\partial y \partial x}$$

$$(f_y)_x = f_{yx} = f_{21} = \frac{\partial}{\partial x} \left(\frac{\partial f}{\partial y} \right) = \frac{\partial^2 f}{\partial x \partial y} = \frac{\partial^2 z}{\partial y \partial x}$$

$$(f_y)_y = f_{yy} = f_{22} = \frac{\partial}{\partial y} \left(\frac{\partial f}{\partial y} \right) = \frac{\partial^2 f}{\partial y^2} = \frac{\partial^2 z}{\partial y^2}$$

Si se deriva respecto a x , la variable y se tomará como una constante y no como una variable, así como viceversa.

Ejemplo a: Sea $f(x,y) = x^3 + x^2y^2 + y^3$ ó $z = x^3 + x^2y^2 + y^3$. Por lo tanto:

$$\frac{\partial z}{\partial x} = 3x^2 + 2xy^2 + 0 \quad \left| \quad \frac{\partial z}{\partial y} = 0 + 2x^2y + 3y^2 \right.$$

$$\frac{\partial^2 z}{\partial x^2} = 6x + 2y^2 \quad \left| \quad \frac{\partial^2 z}{\partial y^2} = 2x^2 + 6y \right.$$

Ejemplo b: Se tiene, sea $f: \mathbb{R} \rightarrow \mathbb{R}$ una función dos veces derivable y $z = f(\mu)$ con $\mu = x^3y^4$. Entonces:

$$\frac{\partial z}{\partial x} = f'(\mu) * 3x^2y^4$$

$$\frac{\partial^2 z}{\partial x^2} = f''(\mu) * 3x^2y^4 * 3x^2y^4 + 6xy^4 * f'(\mu)$$

$$\frac{\partial^2 z}{\partial y \partial x} = f''(\mu) * 4x^3y^3 * 3x^2y^4 + 4 * 3x^2y^3 * f'(\mu)$$

(multiplican coeficientes de $\partial y \partial x$ con exponentes menores de ambas variables)

$$\frac{\partial z}{\partial x} = f'(\mu) * 4x^3y^3$$

$$\frac{\partial^2 z}{\partial y^2} = f''(\mu) * 4x^3y^3 * 4x^3y^2 + 12x^3y^2 * f'(\mu)$$

$$\frac{\partial^2 z}{\partial x \partial y} = f''(\mu) * 3x^2y^4 * 4x^3y^3 + 3 * 4x^2y^3 * f'(\mu)$$

(multiplican coeficientes de $\partial x \partial y$ con exponentes menores de ambas variables)

4.3. La Regla de la Cadena

ggvariable se deriva, según ORDEN pedido: $\partial y \partial x$ o $\partial y \partial x$.

La Regla de la Cadena es una fórmula para la derivada de la composición de dos funciones. Tiene aplicaciones en el cálculo algebraico de derivadas cuando existe composición de funciones.

En términos intuitivos, si una variable y depende de una segunda variable u , que a la vez depende de una tercera variable x ; entonces, la razón de cambio de y con respecto a x puede ser calculada con el producto de la razón de cambio de y con respecto a u multiplicado por la razón de cambio de u con respecto a x .

Al recordar el cálculo en una variable, Método de Integración por partes por ejemplo, se tienen $f(u)$ y $u(x)$ derivables, tal que si $\mu = x^2$, $y = \text{Ln}(x^2) \Rightarrow y = \text{Ln}(\mu) \therefore y' = \frac{1}{x^2} * 2x$ (derivada interna de μ). Por lo tanto:

$$\frac{\delta f}{\delta x} = \frac{\delta f}{\delta \mu} - \frac{\delta \mu}{\delta x}$$

Ejemplo a: $\int x\sqrt{x+3} \delta x$ tal que $\mu = x$, variable más fácil de derivar, $\rightarrow \mu' = \delta x$ mientras que, la variable más fácil de integral, $v' = \sqrt{x+3}$ es $\rightarrow$ su integral es $\frac{(x+3)^{\frac{1}{2}+1}}{\frac{1}{2}+1} = \frac{2}{3} * (x+3)^{\frac{3}{2}}$.

Por lo tanto, $\int \mu \delta v = \mu v - \int v \delta \mu = \frac{2}{3} * (x+3)^{\frac{3}{2}} * x - \frac{2}{3} \int (x+3)^{\frac{3}{2}} \delta x = \frac{2}{3} * (x+3)^{\frac{3}{2}} * x - \frac{2}{3} * \frac{2}{5} x^{\frac{5}{2}} + c \therefore \int x\sqrt{x+3} \delta x = \frac{2}{3} \sqrt{(x+3)^3} * x - \frac{4}{15} * \sqrt{x^5} + c$

Ejemplo b: $\int \frac{x}{\sqrt{2x-5}} \delta x$ al que $\mu = x$, variable más fácil de derivar, $\rightarrow \mu' = \delta x$ mientras que, la variable más fácil de integral, $v' = (2x-5)^{-\frac{1}{2}}$ es $\rightarrow$ su integral es $\frac{(2x-5)^{-\frac{1}{2}+1}}{-\frac{1}{2}+1} = 2\sqrt{2x-5}$. Si: $\mu = 2x-5 \Rightarrow \delta \mu = 2 \delta x \therefore \delta x = \frac{\delta \mu}{2}$.

Entonces, $\int \mu \delta v = \mu v - \int v \delta \mu = x * (2x-5)^{\frac{1}{2}} - 2 \int \frac{\mu^{\frac{1}{2}}}{2} \delta \mu = x * (2x-5)^{\frac{1}{2}} - \frac{2}{3} * \left(\mu^{\frac{3}{2}}\right) + c$, pues $\int \sqrt{2x-5} \delta x = \frac{(2x-5)^{\frac{1}{2}+1}}{\frac{1}{2}+1} = \frac{(2x-5)^{\frac{3}{2}}}{\frac{3}{2}} = \frac{2}{3} * \left(\mu^{\frac{3}{2}}\right)$. Por lo tanto, $\int \frac{x}{\sqrt{2x-5}} \delta x = x(2x-5)^{\frac{1}{2}} - \frac{2}{3} \left[(2x-5)^{3/2}\right] + c$

Sin embargo, la regla de la cadena en varias variables presenta los siguientes casos particulares:

CASO I: Sean $x = x(t)$ y $y = y(t)$ funciones derivables, $z = f(x, y)$ diferenciable en $(x, y) = [x(t), y(t)]$. Entonces, $z = f(x(t), y(t))$ es derivable en y . Por lo tanto:

$$\frac{\delta z}{\delta t} = \frac{\partial f}{\partial x} x'(t) + \frac{\partial f}{\partial y} y'(t) = \frac{\partial f}{\partial x} * \frac{\delta x}{\delta t} + \frac{\partial f}{\partial y} * \frac{\delta y}{\delta t}$$

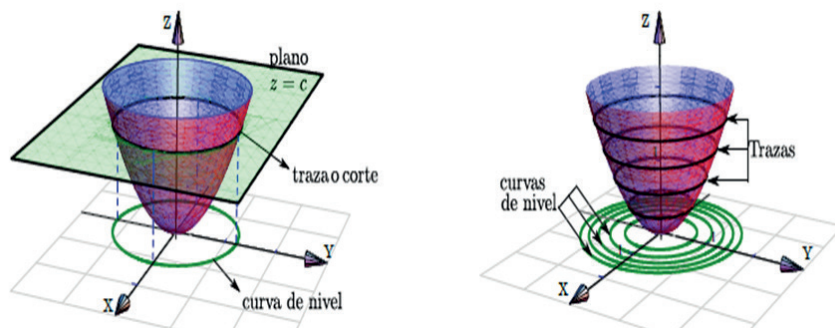
CASO II: Sean funciones $\mu = \mu(x, y)$ y $\nu = \nu(x, y)$ con derivadas parciales en (x, y) . Si $z = f(\mu, \nu)$ diferenciable en $(\mu, \nu) = \left(\begin{smallmatrix} \mu(x, y) \\ \nu(x, y) \end{smallmatrix} \right)$ tal que $z = f(\mu, \nu)$ tiene derivadas parciales de primer orden en (x, y) y:

$$\frac{\partial z}{\partial x} = \frac{\partial f}{\partial \mu} * \frac{\partial \mu}{\partial x} + \frac{\partial f}{\partial \nu} * \frac{\partial \nu}{\partial x}$$

$$\frac{\partial z}{\partial y} = \frac{\partial f}{\partial \mu} * \frac{\partial \mu}{\partial y} + \frac{\partial f}{\partial \nu} * \frac{\partial \nu}{\partial y}$$

4.4. Cálculo Integral por Partes con Regla de Cadena

El bosquejo gráfico consiste en un conjunto de curvas, llamadas trazas o cortes verticales y horizontales:



Fuente: (Mora, 2016)

En las curvas de nivel y trazas o cortes se tiene los siguientes elementos:

- Sea S una superficie en el espacio: $F(x, y, z) = 0$.
- Sea $(x, y) \in \mathbb{R}$ que satisfacen: $F(x, y, c) = 0$.
- La misma que define una curva en plano XY conocida

como Curva de Nivel de la superficie S.

-Geométricamente es el corte entre proyección XY y plano $z=c$.

-A las curvas $F(x,y,c)=0$; $z=c$ (si existen) se les llama trazas o cortes de superficie S.

4.4.1. Las trazas o cortes

Las trazas o cortes tienen como fin realizar el dibujo de una superficie S de ecuación implícita $F(x,y,z)=0$ o explícita $z=f(x,y)$. Los cortes a esta superficie S mediante planos paralelos (Π) a planos coordenados, que tienen apariencia de “alambres” se llaman trazas o cortes.

Ejemplo a: Superficie (S), plano (Π), traza $z=f(c,y)$; $x=c$ o $F(c,y,c)=0$; $x=c$. Se llama superficie al conjunto de puntos del espacio euclidiano tridimensional cuyas coordenadas satisfacen la ecuación $F(x,y,z)=0$; aunque, esta ecuación expresa una relación entre tres variables, pero no siempre es así.

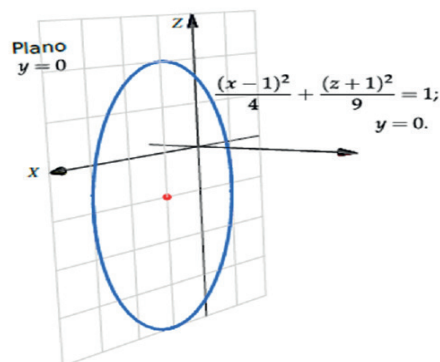
Su objetivo principal es estudiar superficies simples, como planos, superficies cilíndricas y superficies cuadráticas. Por lo tanto, las curvas en el espacio se describen por medio de una ecuación cartesiana $F(x,y)=c$. Por ejemplo, circunferencia de radio a, $x^2 + y^2 = a^2$. Por lo tanto, la curva $\mathcal{C}$ como un conjunto de puntos se representa así:

$$\mathcal{C} = \{(x,y) \in \mathbb{R}^2 / F(x,y) = c\}$$

Así, las curvas $\mathbb{R}^3$ podrían ser definidas por un par de ecuaciones, como intersección de dos superficies:

$$F_1(x,y,z) = c_1; F_2(x,y,z) = c_2$$

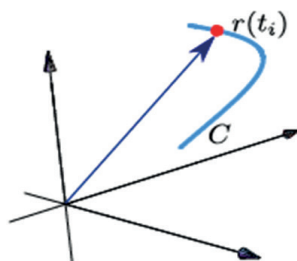
Ejemplo b: Elipse $\cdot \frac{(x-1)^2}{4} + \frac{(z+1)^2}{9} = 1$ se gráfica en plano (Π) XZ con $Y=0$.



Fuente: (Mora, 2016)

4.4.2. Ecuación Paramétrica

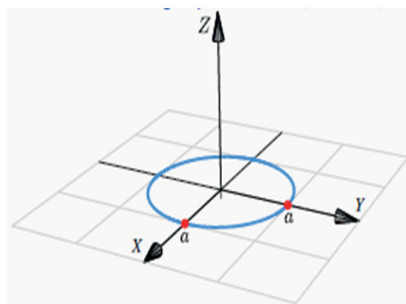
La ecuación paramétrica es otra manera de definir una curva, es como el lugar geométrico de un punto en movimiento. Donde $r(t)$ es la posición del punto en el instante t . Siendo las ecuaciones en Curvas Planas, $r: \mathbb{R} \rightarrow \mathbb{R}^2$ tal que $r(t) = x(t)\vec{i} + y(t)\vec{j}$ y Curvas en el Espacio, $r: \mathbb{R} \rightarrow \mathbb{R}^3$ tal que $r(t) = x(t)\vec{i} + y(t)\vec{j} + z(t)\vec{k}$, respetivamente.



Fuente: (Mora, 2016)

Ejemplo c: Una circunferencia en plano (Π) (x,y), de radio a con centro en el origen. Su gráfica cartesiana sería $x^2 + y^2 = a^2; z = 0$ En cambio, una ecuación paramétrica sería

$$r(t) = a \cos(t)\vec{i} + a \sin(t)\vec{j} + 0\vec{k} \text{ ó } r(t) = \begin{pmatrix} a \cos(t) \\ a \sin(t) \\ 0 \end{pmatrix}$$



Fuente: (Mora, 2016)

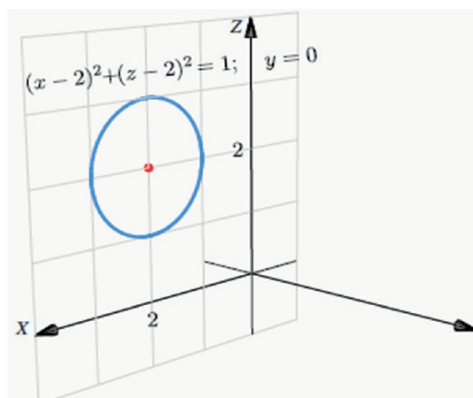
Donde:

-Curva en plano XY $\rightarrow F(x,y)=0 \quad z=0$.

-Curva en plano XZ $\rightarrow F(x,z)=0 \quad y=0$.

-Curva en plano YZ $\rightarrow F(y,z)=0 \quad x=0$

Ejemplo d: Se grafica $\mathcal{C} = (x - 2)^2 + (z - 2)^2 = 1; y = 0$. Tal que, si $\mathcal{C} = (x - 2)^2 + (z - 2)^2 = 1; y = 0$ implica plano (Π) x,y con radio = 1 y centro $(2,0,2)$. En cambio, si $\mathcal{C}: r(t) = \begin{pmatrix} 2 + \cos(t) \\ 0 \\ 2 + \sin(t) \end{pmatrix}$ tal que $t \in [0, 2\pi]$ ó $r(t) = (2 + \cos(t))\vec{i} + 0\vec{j} + (2 + \sin(t))\vec{k}$

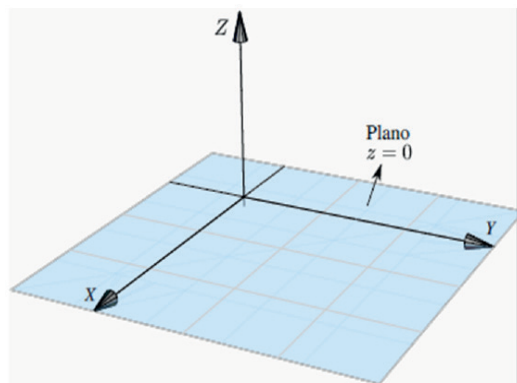


Fuente: (Mora, 2016)

Ejemplo e: El $\Pi: z = 0 \rightarrow (x, y, z)/y \in \mathbb{R} \Rightarrow z = 0$ es el plano XY.

Paramétricamente esto es: $\Pi: r(t, s) = \begin{pmatrix} t \\ s \\ 0 \end{pmatrix} = \begin{pmatrix} \vec{i} \\ \vec{j} \\ \vec{k} \end{pmatrix} = t \vec{i} + s$

$\vec{j} + 0 \vec{k}$ tal que $(t, s) \in \mathbb{R} * \mathbb{R}$:

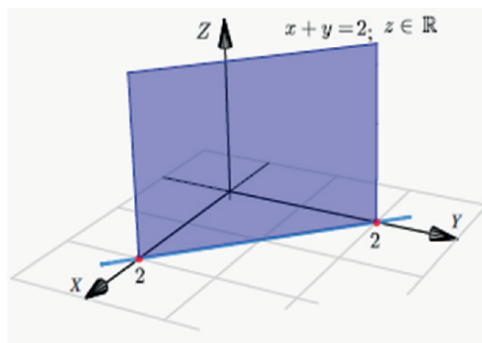


Fuente: (Mora, 2016)

Los planos de ecuación cartesiana con una variable ausente significan que una variable tiene coeficiente nulo. Ejemplo f: $\Pi: 0x+by+cz=d$, Ejemplo g: $\Pi: x+y=2 \rightarrow \{(x, y, z): x+y=2; z \in \mathbb{R}\}$ tal que $x+y=2$ es recta con $z=0$ y el plano se levanta sobre la recta.

Parametrizada esta ecuación sería $x+y=2 \rightarrow t+y=2$. Donde

$\Pi: r(t, s) = \begin{pmatrix} t \\ z-t \\ s \end{pmatrix} = \begin{pmatrix} \vec{i} \\ \vec{j} \\ \vec{k} \end{pmatrix}$ tal que $t, s \in \mathbb{R} * \mathbb{R}$:



Fuente: (Mora, 2016)

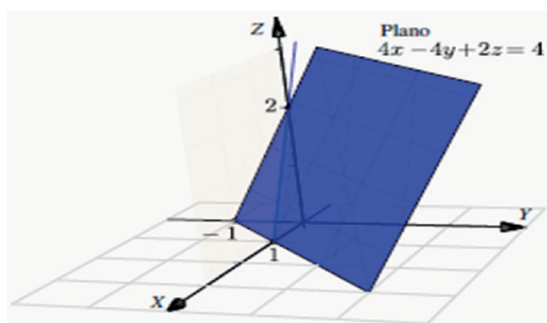
Los planos de ecuación cartesiana sin variables ausentes significan que $ax+by+cz=d$ tal que $a,b,c \neq 0$. Entonces, se determina intersecciones con cada eje coordenado, se traza segmentos y se extiende con paralelas. Ejemplo h: $\Pi: 4x-4y+2z=4$ implica

x	y	z
0	0	2
1	0	0
0	-1	0

Siendo los puntos de intersección con ejes de coordenadas $x=1, y=-1$ y $z=2$ tal que $A(1,0,0), B(0,-1,0)$ y $C(0,0,2)$, parametrizado

$$\Pi: r(t, s) = A + (B - A)t + (C - A)s = \begin{pmatrix} t \\ t + s - 1 \\ zs \end{pmatrix} =$$

$$\begin{pmatrix} \vec{i} \\ \vec{j} \\ \vec{k} \end{pmatrix} \text{ tal que } t, s \in \mathbb{R}:$$



Fuente: (Mora, 2016)

4.5. Superficies Cuadráticas

Las superficies cuadráticas se producen al rotar una cónica alrededor de su eje focal. Estas superficies satisfacen una ecuación de segundo grado en x, y, z llamadas superficies cuadráticas o cuádricas. Se considera cuadrática en posición estándar (sin rotación) si $AX^2 + BY^2 + CZ^2 + Dx + Ey + Fz + G = 0$. Existen

17 tipos estándar cuadráticas:

Tabla 3.2. Tipos estándar cuadráticas o superficies canónicas centradas en el origen. (Cont.).

Tipo	Ecuación	Observaciones	Figura geométrica
Paraboloide	$\left(\frac{X}{a}\right)^2 \pm \left(\frac{Y}{b}\right)^2 = \frac{Z}{c}$	Es una cuádrica, un tipo de superficie tridimensional que se describe mediante ecuaciones. Pueden ser elípticos $\left(\left(\frac{X}{a}\right)^2 + \left(\frac{Y}{b}\right)^2 = \frac{Z}{c}\right)$ o hiperbólicos $\left(\left(\frac{X}{a}\right)^2 - \left(\frac{Y}{b}\right)^2 = \frac{Z}{c}\right)$, según sea que sus términos cuadráticos (los que contienen variables elevadas al cuadrado, aquí indicadas	

Tabla 3.2. Tipos estándar cuadráticas o superficies canónicas centradas en el origen. (Cont.).

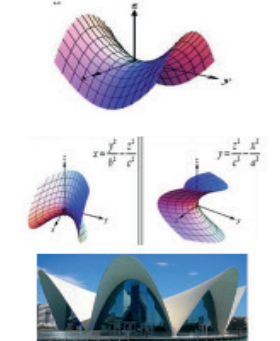
Tipo	Ecuación	Observaciones	Figura geométrica
		como x e y) tengan igual o distinto signo, respectivamente.	

Tabla 3.2. Tipos estándar cuadráticas o superficies canónicas centradas en el origen. (Cont.).

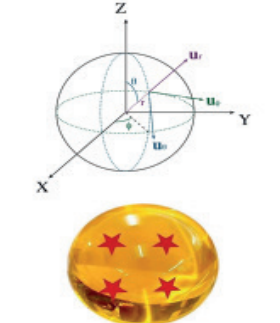
Tipo	Ecuación	Observaciones	Figura geométrica
Esfera	$\frac{X^2}{a^2} + \frac{Y^2}{b^2} + \frac{Z^2}{c^2} - 1 = 0$	Es una cuádrica, caso particular de esferoide y un tipo de superficie tridimensional que se describe mediante una ecuación.	

Tabla 3.2. Tipos estándar cuadráticas o superficies canónicas centradas en el origen. (Cont.).

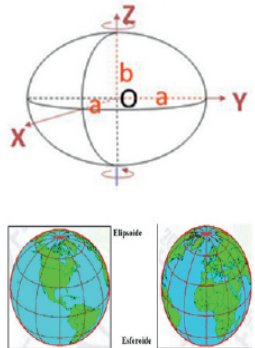
Tipo	Ecuación	Observaciones	Figura geométrica
Esferoide	$\frac{X^2}{a^2} + \frac{Y^2}{b^2} + \frac{Z^2}{c^2} - 1 = 0$	Es un caso particular de elipsoide. Es un elipsoide de revolución; es decir, la superficie que se obtiene al girar una elipse alrededor de uno de sus ejes principales. Por convenio, el eje de simetría se denomina c y se sitúa en el eje de coordenadas cartesianas z ; el eje perpendicular al de simetría se denomina a .	

Tabla 3.2. Tipos estándar cuadráticas o superficies canónicas centradas en el origen. (Cont.).

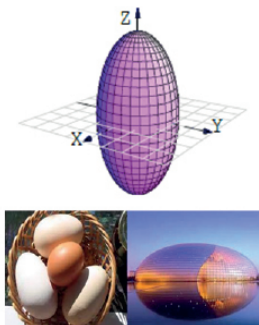
Tipo	Ecuación	Observaciones	Figura geométrica
Elipsoide	$\frac{X^2}{a^2} + \frac{Y^2}{b^2} + \frac{Z^2}{c^2} = 1$	Es simétrico respecto a cada uno de los tres planos coordenados y tiene intersección con ejes coordenados en $(\pm a, 0, 0)$, $(0, \pm b, 0)$ y $(0, 0, \pm c)$. La traza del elipsoide sobre cada uno de los planos coordenados es un único punto o una elipse.	

Tabla 3.2. Tipos estándar cuadráticas o superficies canónicas centradas en el origen. (Cont.).

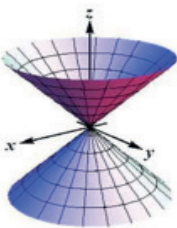
Tipo	Ecuación	Observaciones	Figura geométrica
Cono	$\frac{Z^2}{c^2} = \frac{X^2}{a^2} + \frac{Y^2}{b^2}$	Cuádricas formadas por desplazamiento de una recta (generatriz) que pasa siempre por un punto fijo llamado vértice a lo largo de una curva plana (cerrada o abierta) que se halla en un plano diferente al del vértice, denominada directriz del cono. Es decir, es un sólido de revolución generado por el giro de un triángulo rectángulo alrededor de uno de sus catetos. Al círculo conformado por el	

Tabla 3.2. Tipos estándar cuadráticas o superficies canónicas centradas en el origen. (Cont.).

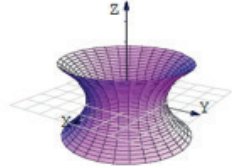
Tipo	Ecuación	Observaciones	Figura geométrica
		otro cateto se denomina base y al punto donde confluyen las generatrices se llama vértice o cúspide. Análogamente se tendrán conos: $\frac{y^2}{b^2} = \frac{x^2}{a^2} + \frac{z^2}{c^2}$ y $\frac{x^2}{a^2} = \frac{y^2}{b^2} + \frac{z^2}{c^2}$.	
Hiperboloide de una hoja	$\frac{X^2}{a^2} + \frac{Y^2}{b^2} - \frac{Z^2}{c^2} = 1$	Hiperboloide de una hoja tiene trazas sobre planos horizontales $z = k$, que son elipses $\frac{x^2}{a^2} + \frac{y^2}{b^2} = 1 + \frac{k^2}{c^2}$. Sus trazas sobre planos verticales son hipérbolas o un	

Tabla 3.2. Tipos estándar cuadráticas o superficies canónicas centradas en el origen. (Cont.).

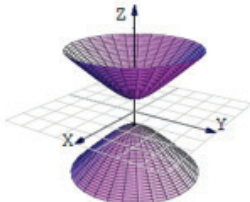
Tipo	Ecuación	Observaciones	Figura geométrica
		par de rectas que se intersecan, mientras que un hiperboloide de dos hojas es una superficie con dos <i>hojas</i> o mantos separados. Sus trazas sobre planos horizontales $z = k$ son elipses y sobre planos verticales son hipérbolas.	
Hiperboloide de dos hoja	$\frac{Z^2}{c^2} - \frac{Y^2}{b^2} - \frac{X^2}{a^2} = 1$		

Tabla 3.2. Tipos estándar cuadráticas o superficies canónicas centradas en el origen. (Cont.).

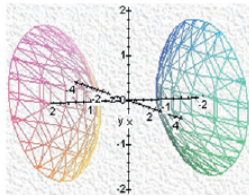
Tipo	Ecuación	Observaciones	Figura geométrica
Hiperboloide elíptico	$\frac{X^2}{a^2} - \frac{Y^2}{b^2} - \frac{Z^2}{c^2} = 1$	Es una superficie cuádrica, como elipsoide, hiperboloide elíptico de dos hojas, cono elíptico, hiperboloide elíptico de una hoja, paraboloide elíptico y el paraboloide hiperbólico. En otras palabras, es el cuerpo formado por la superficie que	

Tabla 3.2. Tipos estándar cuadráticas o superficies canónicas centradas en el origen. (Cont.).

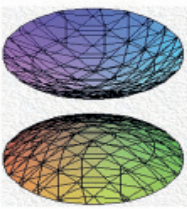
Tipo	Ecuación	Observaciones	Figura geométrica
	$-\frac{X^2}{a^2} - \frac{Y^2}{b^2} + \frac{Z^2}{c^2} = 1$	recubre dos hipérbolas ortogonales entre sí y que comparten un mismo eje. Dependiendo de qué eje comparten se obtienen el hiperboloide elíptico de dos hojas (si el eje común es el mayor o real) y el hiperboloide elíptico de una hoja (cuando el eje común es el menor o imaginario).	 

Tabla 3.2. Tipos estándar cuadráticas o superficies canónicas centradas en el origen. (Cont.).

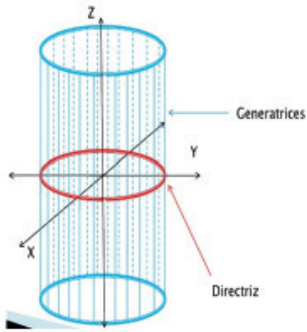
Tipo	Ecuación	Observaciones	Figura geométrica
Cilindro	$X^2 + Y^2 = r^2$	Se origina del griego "kylindros" que significa "rodillo", es un cuerpo sólido geométrico que posee dos extremos planos y redondos u ovalados muy similares y con un lado curvo, cuyo desarrollo es un rectángulo. Es una superficie de las denominadas cuádricas formada por el desplazamiento paralelo de una recta llamada generatriz a lo largo de una curva plana, denominada directriz del	

Tabla 3.2. Tipos estándar cuadráticas o superficies canónicas centradas en el origen. (Cont.).

Tipo	Ecuación	Observaciones	Figura geométrica
		cilindro. Si la directriz es un círculo y la generatriz es perpendicular a él, entonces la superficie obtenida, llamada cilindro circular recto, será de revolución y tendrá por lo tanto todos sus puntos situados a una distancia fija de una línea recta, el eje del cilindro. El sólido encerrado por esta superficie y por dos planos perpendiculares al eje también es llamado cilindro. Este sólido es utilizado como una superficie <u>Gausiana</u> .	

Tabla 3.2. Tipos estándar cuadráticas o superficies canónicas centradas en el origen. (Cont.).

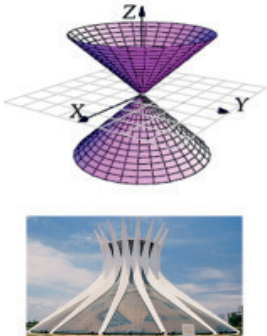
Tipo	Ecuación	Observaciones	Figura geométrica
Cono elíptico	$\frac{x^2}{a^2} + \frac{y^2}{b^2} = \frac{z^2}{c^2}$	Sus trazas sobre planos horizontales $z = k$ son elipses. Sus trazas sobre planos verticales corresponden a hipérbolas o un par de rectas	

Tabla 3.2. Tipos estándar cuadráticas o superficies canónicas centradas en el origen. (Cont.).

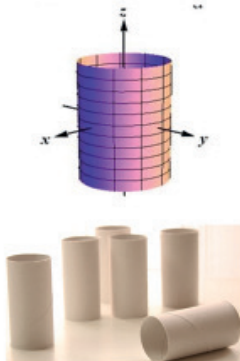
Tipo	Ecuación	Observaciones	Figura geométrica
Cilindro elíptico	$\frac{x^2}{a^2} + \frac{y^2}{b^2} = 1$	Es una cuádrica formada por desplazamiento paralelo de una recta llamada generatriz a lo largo de una curva plana, que puede ser cerrada o abierta, denominada directriz del cilindro. Se construye desplazando una elipse a través de un eje que pasa por su centro y que es perpendicular al plano de la elipse.	

Tabla 3.2. Tipos estándar cuadráticas o superficies canónicas centradas en el origen. (Cont.).

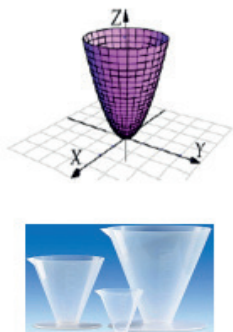
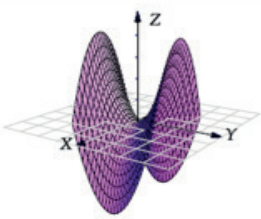
Tipo	Ecuación	Observaciones	Figura geométrica
Paraboloide elíptico	$\frac{x^2}{a^2} + \frac{y^2}{b^2} = \frac{z}{c}$	Sus trazas sobre planos horizontales $z = k$ son elipses $\frac{x^2}{a^2} + \frac{y^2}{b^2} = \frac{k}{c}$. Sus trazas sobre planos verticales, sean $x = k$ o $y = k$ son parábolas.	

Tabla 3.2. Tipos estándar cuadráticas o superficies canónicas centradas en el origen. (Cont.).

Tipo	Ecuación	Observaciones	Figura geométrica
Paraboloide hiperbólico	$\frac{Y^2}{b^2} - \frac{X^2}{a^2} = \frac{z}{c}$	Sus trazas sobre planos horizontales $z = k$ son hipérbolas o dos rectas ($z = 0$). Sus trazas sobre planos verticales paralelos al plano x son parábolas que abren hacia abajo, mientras que las trazas sobre planos verticales paralelos al plano YZ son parábolas que abren hacia arriba. Su grafica tiene la forma de "una silla de montar".	 

Fuente: Elaboración propia con base (Zill & Wright, 2011), cursos de Cálculo Vectorial con Ing. Freddy Vinicio Lema Lema (2015), (Mora, 2016)

Generalmente, la relación entre respuesta y variables independientes se desconoce. En consecuencia, se requiere un modelo que aproxime esta relación funcional entre Y con variables independientes o exógenas $\dots\dots(X_1, X_2, X_3, X_4, \dots, X_k)$; modelo provee bases para un nuevo experimento que lleva a un nuevo modelo y el ciclo se repite. Si la respuesta se describe adecuadamente por una función lineal de variables exógenas usa el modelo de primer orden $Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_k X_k + \varepsilon$.

Los parámetros del modelo se estiman mediante el método de Mínimos Cuadrados Ordinarios (MCO). Una vez obtenidos los estimadores se sustituyen en ecuación y se obtiene el modelo ajustado $\hat{Y} = b_0 + b_1 X_1 + b_2 X_2 + \dots + b_k X_k + \varepsilon$. Este modelo se usa cuando se estudia el comportamiento de variable de respuesta, dependiente o endógena únicamente en la región y cuando se conoce la forma de la superficie.

Para estimar los coeficientes se requiere $N \geq k+1$ valores respuesta (Y). El análisis de datos de corridas se presenta en una tabla de Análisis de Varianza con las siguientes fuentes de variación que contribuyen a su variación total:

Tabla 3.3. Análisis de varianza o ANOVA

F. de V _(Factor de Variación) o VF (Variation factor)	Gl _(Grados de Libertad) o Df (Degrees of freedom)	SC _(Suma de Cuadrados) o Sum Sq (Sum of squares)	CM _(Cuadrado Medio) o Mean Sq (Mean square)	F _(calculada) o F value (Calculated)	F _(Tablas)	R ²
Total	N - 1					
Regresión	p - 1 (Numerador)	SSR $= \sum_{u=1}^N (\hat{Y}_u - \bar{Y})^2$	SCR/p - 1	$\frac{C. M. Regresión}{C. M. Residuo}$	$\frac{Gl_{(Numerador)}}{Gl_{(Denominador)}}$	$\frac{SSR}{SST}$
Residuo	N - p (Denominador)	SSE $= \sum_{u=1}^N (Y_u - \bar{Y}_u)^2$	SCR/N - p			

Fuente: Elaboración propia con base en modificaciones de (González B. G., Métodos estadísticos y principios de diseño experimental, 2010) y (Hurtado & Gómez, 2010). Y_u : es el valor observado en la u-ésima corrida y p representa el número de términos del modelo ajustado.

La prueba de significancia de esta ecuación de regresión ajustada tiene por H_0 : β_s , menos $\beta_0=0$ vs H_a : β_s , menos $\beta_0 \neq 0$. La prueba supone que el error se comporta normalmente y se usa el estadístico de F de Fisher Snedecor, tal que si el valor **F_{Calculada}** es mayor que **F_{tablas}** (α ; p-1 o numerador y N-p o denominador) indica que la variación explicada por el modelo es significativamente mayor que la variación no explicada o inexplicable.

Se presenta por la no planaridad o curvatura de superficie de respuesta, pues no se detecta debido a la exclusión de términos cuadráticos o cúbicos, como $b_{ii}x_i^2$, $b_{ii}x_i^3$ o de términos de producto cruzado $b_{ii}x_i x_j x_k$ que referentes al efecto interacción entre factores.

Esta prueba requiere que el diseño del experimento satisfaga los siguientes elementos:

-Número de distintos puntos de diseño n debe exceder el número de términos en modelo ajustado ($n > k+1$).

Al menos 2 réplicas deben colectarse en uno o más puntos del diseño para estimar la varianza del error.

Valores del error aleatorio (ε_{lu}) deben asumir una distribución normal e independiente con una varianza común σ^2 (homocedasticidad).

Si se cumplen condiciones 1 y 2, la suma de cuadrados residual se compone de dos fuentes de variación. La primera es la falta de ajuste del modelo ajustado debido a exclusión de términos de mayor orden y la segunda es la variación del error puro. Su cálculo requiere la suma de cuadrados calculada de réplicas que recibe el nombre “error puro de suma de cuadrados” y sustraer la suma de cuadrados residual, para obtener la suma de cuadrados de la falta de ajuste

$$SS \text{ Error Puro} = \sum_{l=1}^n \sum_{u=1}^n (Y_{lu} - \bar{Y}_l)^2$$

Donde, Y_{lu} es u -ésima observación del l -ésimo punto del diseño, $u=1,2,3,4,\dots,r_l$, $l=1,2,3,4,\dots,n$ y $\bar{Y}_l$ es promedio de r_l observaciones del l -ésimo punto del diseño.

$$SS \text{ Falta de Ajuste} = SSE - SS \text{ Error Puro} = \sum_{l=1}^n r_l (\hat{Y}_l - \bar{Y}_l)^2$$

Donde, $\hat{Y}_l$ es valor predicho de respuesta en l -ésimo punto del diseño. La prueba de adecuación del modelo ajustado es:

$$F_{\text{Calculada}} = \frac{SS \text{ Falta de Ajuste}/n - p}{SS \text{ Error Puro}/N - n} = \frac{\sum_{l=1}^n r_l (\hat{Y}_l - \bar{Y}_l)^2 / n - p}{\sum_{l=1}^n \sum_{u=1}^n (Y_{lu} - \bar{Y}_l)^2 / N - n}$$

La prueba usa el estadístico de F de Fisher Snedecor, tal que si el valor $F_{\text{Calculada}}$ es mayor que $F_{\text{tablas}}(\alpha; n-p \text{ o numerador y } N-n \text{ o denominador})$ Cuando $F_{\text{Calculada}} \leq CME_{\text{Residual}}$ es usado para estimar σ^2 y, también, se usa para probar la significancia del modelo ajustado tal que cuando la hipótesis de suficiencia de ajuste no se acepta

se debe elevar el grado del modelo aumentando términos de producto cruzado y/o términos de mayor grado en $X_1, X_2, X_3, X_4, \dots, X_k$. Si se requieren puntos adicionales para estimar todos los coeficientes, se añaden. Se colectan datos y se vuelve a realizar el análisis. Sin embargo, si no se rechaza la hipótesis se puede inferir que la superficie es plana. Una vez que se tiene la ecuación y se ha probado el ajuste se buscan niveles que mejoren los valores de respuesta.

Frecuentemente, la estimación inicial de condiciones de operación óptimas está alejada del óptimo real, pues en este caso se desea mover rápidamente a la vecindad del óptimo. El método de máxima pendiente en ascenso es un procedimiento para recorrer secuencialmente la trayectoria de máxima pendiente, que conduce al máximo aumento de respuesta. Cuando desea minimización se trata de mínima pendiente en descenso. De acuerdo con (Montgomery, 2012), la dirección de ascenso máximo es donde $\hat{Y}$ aumenta más rápido y es paralela a la normal de la superficie de respuesta ajustada. Los incrementos a lo largo de su trayectoria son proporcionales a coeficientes de regresión $\beta_0, \beta_1, \beta_2, \dots, \beta_k$. Los experimentos se hacen hasta que deje de observarse un incremento en la respuesta. Entonces, se ajusta un nuevo modelo de primer orden con que se determina una nueva trayectoria y se sigue con el procedimiento. Por último, se consigue llegar a la cercanía del óptimo, que ocurre cuando existe falta de ajuste del modelo de primer orden.

Un algoritmo propuesto por Montgomery (2012) es, suponiendo que punto $X_1 = X_2 = X_3 = X_4 = \dots = X_k = 0$:

1. Se elige un tamaño de incremento o “escalón” en una de las variables del proceso, como ΔX_j , se elige usualmente la variable que más se sabe o la que tiene mayor coeficiente de regresión absoluto $|\beta_j|$.

2. El tamaño de incremento en otras variables es:

$$\Delta X_j = \frac{\hat{\beta}_i}{\hat{\beta}_j / \Delta X_j}$$

Donde, $i=1,2,3,4,\dots,k$ e $i \neq j$

Se convierte ΔX_i de variables codificadas a variables naturales.

Ejemplo (Montgomery, 2012): Un Ingeniero Agroindustrial desea determinar condiciones de operación que maximicen rendimiento de una acción. Tiempo y temperatura son dos variables controlables influyen. Actualmente, el proceso opera con un tiempo de reacción de 35 minutos y temperatura de 155°F , produciendo un rendimiento de 40 %. El ingeniero decide que la región de exploración sea (30,40) minutos de reacción y $(150,160)^{\circ} \text{F}$. Para simplificar los cálculos, las variables independientes o exógenas se codifican en intervalo $(-1,1)$. Las variables codificadas son:

$$X_1 = \frac{(\epsilon_1 - 35)}{5}$$

$$X_2 = \frac{(\epsilon_2 - 155)}{5}$$

Donde, ϵ_1 es variable natural tiempo y ϵ_2 es variable natural temperatura. En la siguiente tabla se muestran los datos, se usa un diseño factorial 2^k aumentando en cinco puntos centrales. Las observaciones centrales sirven para estimar el error experimental y permiten probar la adecuación del modelo de primer orden:

Tabla 3.4. Datos del proceso para ajustar a un modelo de primer orden

Variables Naturales		Variables Codificadas		Respuesta
ϵ_1	ϵ_2	X_1	X_2	Y
30	150	-1	-1	39.3
30	160	-1	1	40.0
40	150	1	-1	40.9
40	160	1	1	41.5
35	155	0	0	40.3
35	155	0	0	40.5
35	155	0	0	40.7
35	155	0	0	40.2
35	155	0	0	40.6

Fuente: Elaboración propia con base en (Hurtado & Gómez, 2010)

Con el método de Mínimos Cuadrados Ordinarios se obtiene:

$$\hat{Y} = 40.44 + 0.775 X_1 + 0.325 X_2$$

Enseguida se muestra el Análisis de Varianza, siendo F de regresión global altamente significativa:

Tabla 3.5. Análisis de Varianza o ANOVA

F. de V (Factor de Variancia)	Gl (Grados de Libertad)	SC (Suma de Cuadrados)	CM (Cuadrado Medio)	F (calculada)	R ²
VF (Variation factor)	Df (Degrees of freedom)	Sum Sq (Sum of square)	Mean Sq (Mean squares)	F value (Calculated F)	
Total	$N - 1 = 8$	$SST = \sum_{u=1}^N (Y_u - \bar{Y})^2 = 3.0022$			
Regresión (β_1, β_2)	$p - 1 = 2$	$SSR = \sum_{u=1}^N (\hat{Y}_u - \bar{Y})^2 = 2.8250$	$SCR/p - 1 = 1.4125$	$\frac{C. M. Regresión}{C. M. Residuo} = 47.83$	$\frac{SSR}{SST} = 94.09$
Residuo (Falta de ajuste)	$N - p = 6$	$SSE = \sum_{u=1}^N (Y_u - \hat{Y}_u)^2 = 0.1772$	$SCR/N - p = 0.0026$	$= 0.06$	
(Error puro)	4	$= 0.0052$ $= 0.1720$	$= 0.0430$		

Fuente: Elaboración propia con base en modificaciones de (González B. G., Métodos estadísticos y principios de diseño experimental, 2010) y (Hurtado & Gómez, 2010)

El modelo indica desplazamiento de 0.775 unidades en dirección X_1 por cada 0.325 unidades en dirección X_2 . Se sabe que la trayectoria pasa por el punto $X_1=0, X_2=0$ y tiene pendiente $0.325/0.775=41.94\%$. En este ejemplo, se decide usar 5 minutos como incremento en tiempo de reacción, equivalente a variable codificada $\Delta X_1 = 1$. Los incrementos a lo largo de la trayectoria son $\Delta X_1 = 1$. $\Delta X_2 = 0.325/0.775 = 0.42$.

El ingeniero calcula punto en su trayectoria y observa el rendimiento en cada uno hasta notar un decremento en su respuesta. Los resultados aparecen a continuación. Los incrementos se muestran tanto para variables codificadas como naturales. Esto se debe a que las codificadas son más fáciles de manejar matemáticamente y las naturales son las usadas en llevar a cabo el proceso:

Tabla 3.6. Experimento de máximo ascenso

Incrementos	Variables Codificadas		Variables Naturales		Respuesta
	X_1	X_2	ϵ_1	ϵ_2	Y
Origen	0	0	35	155	
Δ	1.00	0.42	5	2	
Origen + Δ	1.00	0.42	40	157	41.00
Origen + 2Δ	2.00	0.84	45	159	42.90
Origen + 3Δ	3.00	1.26	50	161	47.10
Origen + 4Δ	4.00	1.68	55	163	49.70
Origen + 5Δ	5.0	2.10	60	165	53.80
Origen + 6Δ	6.0	2.52	65	167	59.90
Origen + 7Δ	7.0	2.94	70	169	65.00

Tabla 3.6. Experimento de máximo ascenso

Incrementos	Variables Codificadas		Variables Naturales		Respuesta
	X_1	X_2	ϵ_1	ϵ_2	Y
Origen + 8Δ	8.0	3.36	75	171	70.40
Origen + 9Δ	9.00	3.78	80	173	77.60
Origen + 10Δ	10.00	4.23	85	175	80.30
Origen + 11Δ	11.00	4.62	90	177	76.20
Origen + 12Δ	12.00	5.04	95	179	75.10

Fuente: Elaboración propia con base en (Hurtado & Gómez, 2010)

Se observa un aumento en respuesta hasta el décimo incremento. A partir del undécimo se produce un decremento en rendimiento. Por lo tanto, se debe ajustar otro modelo de primer orden en su cercanía del punto $(\epsilon_1 = 85, \epsilon_2 = 175)$. La región explorada para ϵ_1 es (80,90) y para ϵ_2 es (170,180). Las variables codificadas son:

$$X_1 = \frac{(\epsilon_1 - 85)}{5}$$

$$X_2 = \frac{(\epsilon_2 - 175)}{5}$$

Nuevamente, se usa un diseño factorial 2^k . Los datos son:

Tabla 3.7. Datos del proceso para segundo modelo de primer orden

Variables Naturales		Variables Codificadas		Respuesta
ϵ_1	ϵ_2	X_1	X_2	Y
80	170	-1	-1	76.5
80	180	-1	1	77.0
90	170	1	-1	78.0
90	180	1	1	79.5
85	175	0	0	79.9
85	175	0	0	80.3
85	175	0	0	80.0
85	175	0	0	79.7
85	175	0	0	79.8

Fuente: Elaboración propia con base en (Hurtado & Gómez, 2010)

Con el método de Mínimos Cuadrados Ordinarios, el modelo de primer orden ajustado es:

$$\hat{Y} = 78.97 + 1.00 X_1 + 0.500 X_2$$

Enseguida se presenta su ANOVA:

Tabla 3.8. Análisis de Varianza o ANOVA

F. de V. (Factor de Variación)	GL (Grados de Libertad)	SC (Suma de Cuadrados)	CM (Cuadrado Medio)	F (calculada)	R ²
0	0	0	0	0	
VF (Variation factor)	Df (Degrees of freedom)	Sum Sq (Sum of square)	Mean Sq (Mean square)	F value (Calculated F)	
Total	$N - 1 = 8$	$SST = \sum_{u=1}^N (Y_u - \bar{Y})^2$ $= 16.1200$			
Regresión (β_1, β_2)	$p - 1 = 2$	$SSR = \sum_{u=1}^N (\hat{Y}_u - \bar{Y})^2$ $= 5.0000$	$SCR/p - 1$ $= 2.5000$	$\frac{C.M. Regresión}{C.M. Residuo}$ $= 1.35$	$\frac{SSR}{SST}$ $= 31.02\%$
Residuo (Falta de ajuste)	$N - p = 6$	$SSE = \sum_{u=1}^N (Y_u - \bar{Y}_u)^2$ $= 11.1200$	$SCR/N - p$ $= 5.507$	$= 102.91$	
(Error puro)	2 4	$= 10.9080$ $= 0.2120$	$= 5.4540$ $= 0.0530$		

Fuente: Elaboración propia con base en modificaciones de (González B. G., Métodos estadísticos y principios de diseño experimental, 2010) y (Hurtado & Gómez, 2010)

El resultado de prueba de falta de ajuste implica que el modelo de primer orden no es una aproximación adecuada, por lo que se trata de una superficie con curvatura y se logra llegar a la cercanía del óptimo.

El modelo de segundo orden es:

$$Y = \beta_0 + \sum_{i=1}^k \beta_i X_i + \sum_{i=1}^k \beta_{ii} X_i^2 + \sum_{i < j}^k \sum_{j=1}^k \beta_{ij} X_i X_j + \varepsilon$$

En este modelo, β_i son coeficientes de regresión para términos de primer orden, β_{ii} son coeficientes para términos cuadráticos puros, β_{ij} son coeficientes para términos de producto cruzado o cruz y ε es término error aleatorio. Los términos cuadráticos puros y cruzados son de segundo orden. El número de términos en ecuación está dado por:

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

Los parámetros del modelo se estiman mediante Método de Mínimos Cuadrados. Una vez que se tienen estimadores se sustituyen en ecuación para obtener el modelo ajustado en el vecindario del valor óptimo de respuesta:

$$\hat{Y} = b_0 + \sum_{i=1}^k b_i X_i + \sum_{i=1}^k b_{ii} X_i^2 + \sum_{i < j}^k \sum_{j=1}^k b_{ij} X_i X_j$$

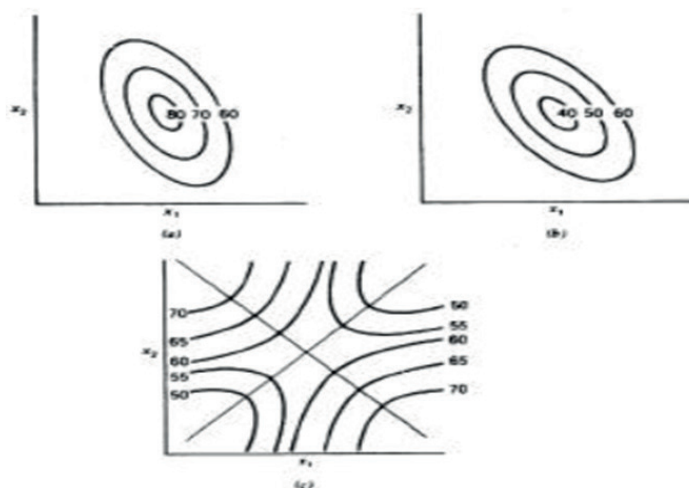
La significancia de coeficientes estimados y su ajuste del modelo se prueban con estadístico F de Fisher Snedecor. Una vez verificado si el modelo tiene suficiencia de ajuste y los coeficientes son significativos, se procede a localizar las coordenadas del “punto estacionario” y se analiza detalladamente su sistema de respuesta.

De acuerdo con (Montgomery, 2012) y en caso que exista valor máximo, suponiendo que se desea maximizar la respuesta, será el conjunto $X_1 = X_2 = X_3 = X_4 = \dots = X_k$. Tal que sus derivadas parciales son:

$$\frac{\partial \hat{y}}{\partial X_1} = \frac{\partial \hat{y}}{\partial X_2} = \frac{\partial \hat{y}}{\partial X_3} = \frac{\partial \hat{y}}{\partial X_4} = \dots = \frac{\partial \hat{y}}{\partial X_k} = 0$$

Donde, el punto $(X_{1,0}, X_{2,0}, X_{3,0}, X_{4,0}, \dots, X_{k,0})$ se denomina “punto estacionario” y puede ser:

- Un punto de respuesta máxima.
- Un punto de respuesta mínima.
- Un punto silla (se ubica punto máximo y mínimo simultáneamente)



Fuente: (Zill & Wright, 2011)

Se puede obtener este punto usando notación matricial para modelo de segundo orden:

$$\hat{Y} = \hat{\beta}_0 + X'b + X'BX$$

Donde:

$$x = \begin{bmatrix} X_1 \\ X_2 \\ M \\ X_k \end{bmatrix} \quad b = \begin{bmatrix} \hat{\beta}_1 \\ \hat{\beta}_2 \\ M \\ \hat{\beta}_k \end{bmatrix} \quad y \quad B = \begin{bmatrix} \hat{\beta}_{11}, & \frac{\hat{\beta}_{12}}{2}, & \square & \frac{\hat{\beta}_{1k}}{2} \\ \frac{\hat{\beta}_{21}}{2}, & \hat{\beta}_{22}, & K, & \frac{\hat{\beta}_{2k}}{2} \\ M & \square & \square & \square \\ \frac{\hat{\beta}_{k1}}{2}, & \frac{\hat{\beta}_{k2}}{2}, & K & \hat{\beta}_{kk} \end{bmatrix}$$

Donde, b es vector ($k \times 1$) de coeficientes de regresión de primer orden, B es una matriz simétrica ($k \times k$), cuya diagonal principal está formada por coeficientes de términos cuadráticos puros ($\hat{\beta}_{ii}$) y elementos fuera de esta corresponden a un medio del valor de coeficientes cuadráticos mixtos ($\hat{\beta}_{ij}, i \neq j$). La derivada de $\hat{Y}$ respecto a vector X igualada a cero es:

$$\frac{\delta \hat{Y}}{\delta X} = b + 2BX = 0$$

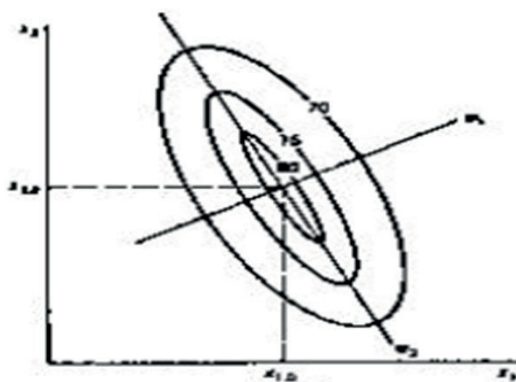
Tal que, el punto estacionario es la solución de esta ecuación:

$$x_0 = -\frac{1}{2}B^{-1}b$$

Se sustituye esta en ecuación matricial para modelo de segundo orden:

Hallado el punto estacionario, se caracteriza la superficie de respuesta. Es decir, se determina si se trata de un punto de respuesta máximo, mínimo o silla (se ubica punto máximo y mínimo simultáneamente). La forma directa de hacerlo es mediante la gráfica de contornos del modelo ajustado; sin embargo, un análisis más formal es útil.

Como alternativa se puede expresar la forma de superficie de respuesta usando un nuevo conjunto de variables $W_1, W_2, W_3, W_4, \dots, W_k$, cuyos ejes representan ejes principales de superficie de respuesta, interpretados en el punto estacionario como:



Fuente: (Zill & Wright, 2011)

Esta da como resultado el modelo ajustado llamado “**forma canónica**”:

$$\hat{Y} = \hat{Y}_0 + \lambda_1 W_1^2 + \lambda_2 W_2^2 + \lambda_3 W_3^2 + \lambda_4 W_4^2 + \dots + \lambda_k W_k^2$$

Donde, w_i son variables independientes transformadas y λ_k son constantes o valores propios también conocidos como raíces características, auto valores o eigenvalores tomados de matriz B.

La naturaleza de superficie de respuesta puede determinarse a partir del punto estacionario, el signo y magnitud de λ_i . Si todas las λ_i son positivas se puede tomar como un punto de respuesta mínima, si todas las λ_i son negativas es un punto de respuesta máxima y si las λ_i tienen signos distintos es un punto de respuesta silla (se ubica punto máximo y mínimo simultáneamente).

BIBLIOGRAFÍA

Adams, J. A., Benitez, M. D., Guidek, R. C., & Domínguez, G. A. (2016). Enseñanza de Programación Lineal y Juegos de Empresa. Asunción, Paraguay: Universidad Autónoma de Paraguay.

Adhikari, R., & Dandekar, R. (2012). Optimization models for control of parasitic diseases. *Mathematical and Computer Modelling*, 56(1-2), 233-250.

Ahuja, R. K., Magnanti, T. L., & Orlin, J. B. (1993). *Network Flows. Theory, Algorithms and Applications*. Upper Saddle River. New Jersey, USA.: Prentice-Hall, Inc.

Alarcón, S. (1998). La programación estocástica discreta como instrumento de planificación de plantaciones de olivos en Castilla-La Mancha. *Economía Agraria* N0. 183, 173-200.

Alidaee, A., & Aneja, Y. P. (2018). Non-predetermined goal programming. In *Operations research: A model-based approach*. New Jersey, U.S.A.: Springer, Cham.

Ancco, P. R., Chalco, C. D., Clavijo, T. P., Delgado, H. P., Huarac, S. Y., & Ugarte, F. D. (2020). Métodos para el procesamiento de imágenes digitales. Universidad Nacional de San Antonio Abad del Cusco, Cusco Perú, 1-8.

Anderson, D. R., Sweeney, D. J., & Williams, T. A. (2006). *Estadística para administración y economía*. Col. Cruz Manca, Santa Fe, México, D. F.: Cengage Learning Editores, S. A. de C. V.

Arévalo, J. J., Castro, A., & Villa, É. (2002). Un análisis del ciclo económico en competencia imperfecta. *Revista de Economía Institucional*, vol. 4, núm. 7, segundo semestre, 11-39 Pp.

Arriaza, G. A., Fernández, P. F., López, S. M., Muñoz, M. M., Pérez, P. S., & Sánchez, N. A. (2008). Estadística Básica con R y R-Commander. Cádiz, España: Servicio de Publicaciones de la Universidad de Cádiz.

Atucha, A. J., & Gualdoni, P. (2018). El funcionamiento de los mercados. Mar de Plata, Argentina: Facultad de Ciencias Económicas y Sociales, Universidad Nacional de Mar del Plata.

Avdoshin, S. M., & Pesotskaya, E. Y. (2011). Software risk management. National Research University Higher School of Economics , 1-6 Pp.

Baquela, E. G., & Redchuk, A. (2013). Optimización Matemática con R. Volumen I: Introducción al modelado y resolución de problemas. Madrid, España: Bubok Publishing S. L.

Bazaraa, M. S., Sherali, H. D., & Shetty, C. M. (2009). Nonlinear programming: Theory and applications. Hoboken, New Jersey, U.S.A.: John Wiley & Sons, Inc, Publication.

Bertsimas, D., & Tsitsiklis, J. N. (1997). Introduction to Linear Optimization. Belmont, Mass. U.S.A.: Athena Scientific, Belmont, Massachusetts.

Bertsimas, D., Teo, C., & Vohra, R. (1999). On dependent randomized rounding algorithms. ELSEVIER, 105 –114.

Bhatt, A., Das, V., & Sundaram, V. (2011). Economic analysis of malaria control programs: A review of the literature. Malaria Journal, 10(1), 202.

Birge, J. R., & Louveaux, F. (2011). Introduction to Stochastic Programming. New York, USA: Springer Science+Business Media, .

Brambila, P. J. (2011). Bioeconomía: Instrumentos para su análisis económico. Texcoco, estado de México. México: Secretaría de Agricultura, Ganadería, Desarrollo Rural, Pesca

y Alimentación (SAGARPA) y Colegio de Postgraduados (COLPOS).

Broz, D., Mac Donagh, P., Arce, J., & Yapura, P. (2017). La Investigación Operativa, la Ingeniería Forestal y los Problemas Sectoriales: Ante la Necesidad de un Cambio de Paradigma. *Yvyrareta*, 64-72.

Burbano, R., Espinosa, A. M., & Horna, L. (2018). Economía para matemáticos. Quito, Ecuador: Escuela Politécnica Nacional (EPN).

Burkard, R., Dell'Amico, M., & Martello, S. (2012). Operations research: Theory and practice. New Jersey, U.S.A.: Springer Science & Business Media.

Bustos, F. E. (2023). La importancia del teorema de suficiencia de Kuhn-Tucker en la tarea de decisiones organizacionales. Caracas, Venezuela: Escuela Superior de Cómputo de Venezuela.

Cabrera, G. E., & Díaz, G. E. (2021). Manual de uso de Jupyter Notebook para aplicaciones docentes. Madrid, España: Universidad Complutense de Madrid.

Camps, P. R., Casillas, S. L., Costal, C. D., Gibert, G. M., Martín, E. C., & Pérez, M. O. (2005). Software libre. Bases de datos. Av. Tibidado, 39-43, Barcelona España: Fundació per a la Universitat Oberta de Catalunya. Eureka Media, SL.

Capa, S. H. (2015). Probabilidades y estadística para una gestión científica de la información. Quito, Ecuador: Escuela Politécnica Nacional .

Carballo, R. E., & Peñate, P. E. (2014). Benchmarking de Herramientas Forenses Móviles. Managua: COMPDES, UNAN, Managua.

Carro, P. R. (2014). Investigación de operaciones en administración. 1ra Edición. Mar de Plata, Provincia de Buenos

Aires, Argentina: Facultad de Ciencias Económicas y Sociales.
Universidad Nacional de Mar de Plata.

Cestero, E. V., & Caballero, A. M. (2017). Data Science y Redes Complejas. Métodos y aplicaciones. Tomás Bretón, 21-28045, Madrid, España: Editorial Centro de Estudios Ramón Areces, S. A.

Challenger, P. I., Díaz, R. Y., & Becerra, G. R. (Abril-Junio, 2014). El lenguaje de programación Python. Ciencias Holguín, vol. XX, núm. 2., 1-13.

Dantzig, G. B. (1963). Linear Programming and Extensions. Chichester, West Sussex, U.S.A.: Princeton University Press.

Dantzig, G. B., & Wolfe, P. (1954). The decomposition principle for linear programming. Operations Research, 2(1), 76-81.

Daza, P. J. (2006). Estadística Aplicada con Excel. Lima, Perú: Grupo Editorial Megabyte s.a.c.

De los Reyes García, M. M., & Romero, C. J. (2004). Investigación de operaciones I. UAM Unidad Azcapotzalco. Av. San Pablo 180. Col. Reynosa Tamaulipas. Delegación Azcapotzalco. México, D. F.: División de Ciencias Básicas e Ingeniería. Departamento de Sistemas. Universidad Autónoma Metropolitana-Azcapotzalco.

De Mendiburu, F. (2007). Análisis estadístico con R. Av. Túpac Amaru 210, Rímac 15333, Lima, Perú: Sección de extensión universitaria y proyección social. Centro de estudiantes de facultad de ingeniería económica y ciencias sociales. Universidad Nacional de Ingeniería.

De Mendiburu, F. (2017). Tutorial de agricolae (Versión 1.2-8). La Molina-Perú: Departamento Académico de Estadística e Informática de la Facultad de Economía y Planificación. Universidad Nacional Agraria La Molina-Perú.

Delgado, Q. S. (2023). Aprende Python. Ciudad de México, México: Mundi Prensa.

Deng, Z. X., & Wang, X. (2022). A survey on non-predetermined goal programming. *European Journal of Operational Research*, 297(2), 1115-1133.

Devia, J., Azacon, J., López, N., Villarroel, Z., Carvajal, A., Rodríguez, B., . . . Suárez, F. (2012). Software TORA. Maturín, Venezuela: Núcleo de Monagas. Universidad de Oriente de Venezuela.

Di Perro, M. (2017). What is the Blocchaing? SChool of Computing at DePaul University, Chicago, IL, USA, 92-95.

Díaz, C. M., & Gómez, M. P. (2011). Comportamiento de volatilidad del tipo de México 1970 a 2010. Michoacán, México: Universidad Michoacana de San Nicolás de Hidalgo.

DiLorenzo, T. J. (1992). El Mito del Monopolio Natural. *The Review of Austrian Economics* Vol. 9, No. 2 , 1-15 Pp.

Eiselt, H. A., & Sandblom, C. L. (2009). Integer programming: Theory and practice. New Jersey, U.S.A.: Springer.

Engler, P. A. (2009). Estrategías del manejo del riesgo. Quilamapu, Chillán, Chile: TRAMA Impresores S. A.

Galindo, D. T. (2006). Estadística. Métodos y Aplicaciones. Quito, Ecuador: Prociencia Editores.

Garcés, R. A. (2020). Optimización convexa. Aplicaciones en operación y dinámica de sistemas de potencia. Pereira, Colombia: Universidad Tecnológica de Pereira, Pereira, Colombia.

García, M. A., & García, S. (2015). Investigación de operaciones en el control de enfermedades parasitarias. *Revista de Investigación Académica*, 26, 1-12.

García, M. M., & Romero, C. J. (2004). Investigación de

operaciones 1. 1ra parte. Delegación Azcapotzalco, D. F.: Universidad Autónoma Metropolitana (UAM-Azcapotzalco).

Gen, M., & Cheng, R. (2000). Genetic algorithms and engineering optimization. New York, U.S.A.: Wiley-Interscience Publication.

Gómez, G. M., Danglot, B. C., & Vega, F. L. (2003). Sinopsis de pruebas estadísticas no paramétricas. Cuándo usarlas. Revista Mexicana de Pediatría, 10.

Gómez, P. M. (2006). Introducción a la microeconomía. Barcelona, España: Universidad de Barcelona.

González, B. G. (1985). Métodos Estadísticos y Principios de Diseño Experimental. Quito, Ecuador: Facultad de Ciencias Agrícolas. Universidad Central del Ecuador.

González, B. G. (2010). Métodos estadísticos y principios de diseño experimental. Quito, Ecuador: Universidad Central del Ecuador.

González, J. A., García, S. J., Chalita, T. L., Matus, G. J., Cruz, G. B., Sangerman, J. D., . . . Fortis, H. M. (2012). Modelo de equilibrio espacial para determinar costos de transporte en la distribución de durazno en México. Revista Mexicana de Ciencias Agrícolas Vol.3 Núm.4 , 701-712 Pp.

Grossman, S. S., & Flores, G. J. (2012). Álgebra Lineal. México D. F.: Colonia Santa Fe, Delegación Álvaro Obregón, México D. F.

Guerrero, S. H. (2009). Programación Lineal Aplicada. Bogotá, Colombia: Ecoe Ediciones.

Gujarati, D. N., & Porter, D. C. (2010). Econometría. Delegación Álvaro Obregón, México, D. F.: McGraw-Hill/ Interamericana Editorres, S.A. de C.V.

Hernández, N., Soto, F., & Caballero, A. (2009). Modelos de simulación de cultivos, características y usos. Cultivos

Tropicales, vol. 30, núm. 1, 73-82 Pp.

Hernandez, S. R., Fernandez, C. C., & Baptista, L. P. (2006). Metodología de la investigación. Colonia Desarrollo Santa Fe, Delegación Álvaro Obregón, México, D. F. : McGraw Hill/ Interamericana Editores, S.A. de C. V.

Herrera, A. L. (2009). Perspecctivas para la carna de bovino en México aplicando un sistema de trazabilidad. Campus Montecillo, Texcoco, estado de México: Colegio de Postgraduados. Tesis de maestría.

Hillier, F. S., & Lieberman, G. J. (2010). Introducción a la investigación de operaciones. Delegación Álvaro Obregón, México D. F.: McGraw-Hill Companies, Inc.

Hillier, F. S., & Lieberman, G. J. (2015). Operations Reseach. New York, U.S.A.: McGraw-Hill Education.

Hung, M. T., & Hung, C. C. (2019). A review of non-predetermined goal programming: Methods, applications, and challenges. European Journal of Operational Research, 274(3), 909-924.

Hurtado, M. J., & Gómez, F. R. (2010). Diseño Experimental. Valencia, España: Universidad de Valencia.

Iusem, A. N. (2003). On the convergence properties of the projectedgradient method for convex optimization. Computational and Applied Mathematics, 37–52.

Kaiser, H. M., & Messer, K. D. (2011). Mathematical Programming for Agricultural, Enviromental, and Resource Economics. Rosewood Drive, Danvers, United States of America: John Wiley & Sons, Inc.

Kmenta, J. (1971). Elements of econometrics. Collier Macmillan Canada, Inc.: Macmillan Publishing Ccompany.

Larrañaga, L. J., Zulueta, G. E., Elizagarate, U. F., & Alzola, B. J. (2020). Algoritmos Meméticos en problemas de

Investigación Operativa. Escuela Universitaria de Ingeniería de Vitoria-Gasteiz. Universidad del País Vasco Euskal Herriko Unibertsitatea (UPV-EHU)., 1-19.

Leek, J. T., & Peng, R. D. (2015). Statistics: P values are just the tip of the iceberg. Maryland, USA.: School of Public Health in Baltimore.

Li, H., Zhang, J., & Zhang, X. (2019). Randomized rounding for general integer programs. *Operations Research Letters*, 47(2), 164-169.

Li, W., & Yang, L. (2022). A novel non-predetermined goal programming model for renewable energy planning. *Energy*, 220, 122772.

Lind, D. A., Marchal, W. G., & Mason, R. D. (2004). Estadística para Administración y Economía. Col. del Valle, México D. F.: Alfaomega Grupo Editor S. A. de C. V.

Lind, D. A., Marchal, W. G., & Mason, R. D. (2006). Estadística para Administración y Economía. Col. del Valle, México D. F.: Alfa Omega Grupo Editor S. A. de C. V.

López, B. E., & González, R. B. (2014). Diseño y Análisis de Experimentos. Fundamentos y Aplicaciones en Agronomía. Ciudad Universitaria zona 12. Ciudad de Guatemala: Universidad de San Carlos de Guatemala, Facultad de Agronomía.

Loubet, B. (2020). Excel: Herramienta Solver. Centro Universitario, Ciudad de Mendoza. Provincia de Mendoza, Argentina.: Facultad de ciencias económicas. Universidad Nacional de Cuyo.

Ipízar, M. J. (12 de Octubre de 2023). Universidad Latina. Obtenido de La organización industrial: <https://cisprocr.com/cispro/system/files/Clase%209%20Organizaci%C3%B3n%20Industrial.pdf>

Luenberger, D. G. (1984). *Linear and Nonlinear Programming*. New Jersey, U.S.A.: Springer.

Machado, D. E., & Coto, F. H. (2017). Sistema de adquisición de datos con Python y Arduino. *Revista, Ciencia, Ingeniería y Desarrollo Tec Lerdo*, 1-6.

Madrid, C. C. (2020). Filosofía de la ciencia del cambio climático: Modelos, problemas e incertidumbres. *Revista Colombiana de Filosofía de la Ciencia*, vol. 20, núm. 41, 201-234 Pp.

Martínez, Á. (26/10/2023 de Octubre 26 de 2023). Optimización global. Obtenido de Optimización global: http://www.dma.uvigo.es/~aurea/Optimizaci%C3%B3n_global.

Martínez, A. L. (2013). *Modelo de Programación Cuadrática y Ratios Financieros para minimizar el riesgo de las inversiones en la Bolsa de Valores de Lima*. Lima, Perú: Facultad de Ciencias Matemáticas, E.A.P. de Investigación Operativa. Universidad Nacional Mayor de San Marcos.

Mehrotra, S. (2022). Non-predetermined goal programming: A review of recent advances. *Journal of the Operational Research Society*, 73(5), 1018-1032.

Mendiburu, F. (2007). *Análisis Estadístico con "R"*. Lima, Perú: Sección de Extensión Universitaria y Proyección Social. Centro de Estudios de la Facultad de Ingeniería Económica y Ciencias Sociales. Universidad Nacional de Ingeniería.

Mendiburu, F. (2015). *Agricolae tutorial (Version 1.2-3)*. La Molina, Perú: Departamento Académico de Estadística e Informática de la Facultad de Economía y Planificación. Universidad Agraria La Molina.

Mijangos, L. J. (2019). *Investigación de Operaciones II*. Tuxtla Guitiérrez, Chiapas, México: Instituto Tecnológico de Tuxtla Guitiérrez.

Moghaddam, S. A., Arabshahi, A., Yazdani, D., & Dehshibi, M. M. (2012). A novel Hybrid Algorithm for Optimization in Multimodal Dynamic Enviroments. ResearchGate, 143-148.

Montgomery, D. C. (2004). Diseño y análisis de experimentos. México, D. F.: Limusa S. A. de C. V.

Montgomery, D. C. (2012). Design and analysis of experiments. Arizona, USA: John Wiley & Sons, Inc.

Muñoz, C. R., Ochoa, H. M., & Morales, G. M. (2011). Investigación de operaciones. Delegación Álvaro Obregón, México D. F.: McGraw-Hill/Interamericana Editores, S.A. DE C.V.

Murillo, F. E. (2016). Riesgo agropecuario. Revista de la Carrera de Ingeniería Agronómica – UMSA, 103-127 Pp.

Navidi, W. (2006). Estadística para ingenieros y científicos. Colonia Desarrollo Santa Fe, Delegación Álvaro Obregón, México,D. F.: McGrawHill McGraw-Hill/Interamericana Editores, S. A. de C. V.

Nocedal, J., & Wright, S. J. (2006). Numerical optimization. New Jersey, U.S.A.: Springer Science & Business Media.

Nolasco, V. J. (2018). Python. Apliacaciones prácticas. Calle Jarama, 3A, Poligono Industrial Igarsa. Paracuellos de Jarama, Madrid, España: Ra-Ma Editorial.

Oliveira, W., Rocha, R. S., & Katz, N. (2010). Esquistossomose mansônica: situação atual e perspectivas de controle. Revista da Sociedade Brasileira de Medicina Tropical, 43(2), 123-130.

O'Neill, P., & Thomas, J. (2010). Mathematical modelling of infectious diseases. Nature Reviews Microbiology, 8(11), 803-814.

Padrón, C. E. (1996). Diseños Experimentales: con aplicaciones a la agricultura y la ganadería. Ciudad de México, México: Trillas. Mex.

Parra, R. F. (2016). Curso de Estadística con R. . Santander, Cantabria: Instituto Cántabro de Estadística.

Pérez, A. E. (2017). Relajación Lagrangiana, métodos heurísticos y metaheurísticos en algunos modelos de Localización e Inventarios. Valladolid, España: Máster en investigación en Matemáticas, Universidad de Valladolid.

Prawda, W. J. (2004). Métodos y modelos de investigación de operaciones Vol. I: Modelos determinísticos. Ciudad de México, México: Limusa. Noriega Editores.

Quintana, S. F. (octubre 2016/marzo 2017). Dinámica, escalas y dimensiones del cambio climático. Tla-Melaua, revista de Ciencias Sociales. Facultad de Derecho y Ciencias Sociales, año 10, núm. 41, 180-200 Pp.

Ravindran, A., & Garg, D. P. (2018). Meta-heuristics for optimization. New Jersey, U.S.A.: Springer.

Rodríguez, C. E. (11 de Octubre de 2023). La competencia imperfecta. Obtenido de Biblioteca digital de la Universidad Católica Argentina: <http://bibliotecadigital.uca.edu.ar/repositorio/contribuciones/competencia-imperfecta-carlos-rodriguez.pdf>

Rodríguez, V. A. (2007). Cambio climático, agua y agricultura. Desarrollo Rural Sostenible, N° 1, II Etapa, 13-23 Pp.

Rojas, C. L., & Rojas, C. L. (2000). Exploración al diseño experimental. Granada, España: Ciencia e Ingeniería Neogranadina, 9, 51–59. Universidad Militar Nueva Granada.

Romero, C., & Ventura, S. (2013). Metaheuristics: Applications to optimization and machine learning. New Jersey, U.S.A.: Springer.

Rosales, A. J. (2001). Análisis y diseño en el espacio de estado utilizando Matlab. quito, Ecuador: Escuela de Ingeniería, Escuela Politécnica Nacional .

Rosell, E. K. (2022). Optimización con programación dinámica. Barcelons, España: Facultad de Matemáticas e Informática, Universidad de Barcelona.

Rossum, V. G. (2009). Tutorial Python versión 2.5.2. Ciudad de México, México: Mundi Prensa.

Salas, F. V. (2009). Modelos de negocio y nueva economía industrial. *Universia Business Review* | Tercer Trimestre, 1-24 Pp.

Sampayo, M. S. (2019). Apuntes de economía del medio ambiente. Iztacala, D. F., México: Universidad Nacional Autónoma de México (UNAM).

Sarasa, C. A. (2017). Gestión de la información web usando Python. *Rambla del Poblenou*, 156. Barcelona, España.: Editorial UOC (Oberta UOC Publishing, SL).

SEMARNAT, & CONAFOR. (2019). Estrategia Nacional de Sanidad Forestal 2019-2024. Ciudad de México, México.: Secretaría de Medio Ambiente y Recursos Naturales; Comisión Nacional Forestal.

Szigeti, F., Cardillo, J., Henet, J. C., & Calvet, J. L. (2014). El Método de Relajación Aplicado a Optimización de Sistemas Discretos. *ResearchGate*, 1-9.

Taha, H. A. (2004). Investigación de operaciones 7ma Edición. Atlacomulco No. 500-5º piso. Col. Industrial Atoto. Naucalpan de Juárez, Edo. de México: Pearson Educación de México, S.A. de C.V.

Taha, H. A. (2012). Investigación de Operaciones. Fayetteville, Arkansas, USA: PEARSON EDUCACIÓN DE MÉXICO, S.A. de C.V.

Taha, H. A. (2012). Investigación de Operaciones. Fayetteville, Arkansas, USA: PEARSON EDUCACIÓN DE MÉXICO, S.A. de C.V.

Toledo, T. R. (2009). El riesgo en la agricultura. Quilamapu,

Chillán, Chile: TRAMA Impresores S. A.

Triola, M. F. (2004). Estadística. Col. Industrial Atoto, Naucalpan de Juárez, Edo. de México: Pearson Educación de México, S. A. de C. V.

Trívez, B. F. (2004). Economía Espacial: una disciplina en auge. Estudios de Economía Aplicada, Vol. 22-3, 409-429 Pp.

Urquía, M. A., & Martín, V. C. (2016). Métodos de simulación y modelado. Madrid, España: Universidad Nacional de Educación a Distancia.

Valencia, N. E. (2018). Investigación Operativa. Programación lineal, problemas resueltos con soluciones detalladas. San Juan Bautista de Ambato, Provincia Tungurahua, Ecuador: Universidad Técnica de Ambato (UTA).

Vanegas, C. J. (2018). Metodología para la planeación de requerimientos de materiales-MRP. Bogotá, Colombia: Facultad de Ingeniería. Especialización en Gerencia Integral de Proyectos. Universidad Militar Nueva Granada.

Vásquez, A. V. (2014). Diseños Experimentales con SAS. Cajamarca, Perú: CONCYTEC FONDECYT.

Vega, G. C. (05 de Mayo de 2021). IESIP. Obtenido de iesip.edu.ve: <https://virtual.iesip.net/mod/page/view.php?id=5933>

Velásquez, S., & Velásquez, R. (2012). Modelado con variables aleatorias en simulink utilizando simulación Montecarlo. Universidad, Ciencia y Tecnología, Vol. 16, No. 64, 203-211 Pp.

Vercher, G. E. (2015). Soft Computing approaches to portfolio selection. Boletín de Estadística e Investigación Operativa. Vol. 31. No. 1, 23-46.

Vicéns, O. J., Herrarte, S. A., & Medina, M. E. (2005). Análisis de la varianza (ANOVA). Distrito Federal, México: Trillas, S. A.

Villota, C. E. (2010). Control moderno y óptimo (MT 227C). Lima, Perú: Departamento Académico de Ingeniería Aplicada, Facultad de Ingeniería Mecánica, Universidad Nacional de Ingeniería.

Adams, J. A., Benítez, M. D., Guidek, R. C., & Domínguez, G. A. (2016). Enseñanza de Programación Lineal y Juegos de Empresa. Asunción, Paraguay: Universidad Autónoma de Paraguay.

Adhikari, R., & Dandekar, R. (2012). Optimization models for control of parasitic diseases. *Mathematical and Computer Modelling*, 56(1-2), 233-250.

Ahuja, R. K., Magnanti, T. L., & Orlin, J. B. (1993). *Network Flows. Theory, Algorithms and Applications*. Upper Saddle River. New Jersey, USA.: Prentice-Hall, Inc.

Alarcón, S. (1998). La programación estocástica discreta como instrumento de planificación de plantaciones de olivos en Castilla-La Mancha. *Economía Agraria* N0. 183, 173-200.

Alidaee, A., & Aneja, Y. P. (2018). Non-predetermined goal programming. In *Operations research: A model-based approach*. New Jersey, U.S.A.: Springer, Cham.

Ancco, P. R., Chalco, C. D., Clavijo, T. P., Delgado, H. P., Huarac, S. Y., & Ugarte, F. D. (2020). Métodos para el procesamiento de imágenes digitales. Universidad Nacional de San Antonio Abad del Cusco, Cusco Perú, 1-8.

Anderson, D. R., Sweeney, D. J., & Williams, T. A. (2006). *Estadística para administración y economía*. Col. Cruz Manca, Santa Fe, México, D. F.: Cengage Learning Editores, S. A. de C. V.

Arévalo, J. J., Castro, A., & Villa, É. (2002). Un análisis del ciclo económico en competencia imperfecta. *Revista de Economía Institucional*, vol. 4, núm. 7, segundo semestre, 11-39 Pp.

Arriaza, G. A., Fernández, P. F., López, S. M., Muñoz, M. M., Pérez, P. S., & Sánchez, N. A. (2008). Estadística Básica con R y R-Commander. Cádiz, España: Servicio de Publicaciones de la Universidad de Cádiz.

Atucha, A. J., & Gualdoni, P. (2018). El funcionamiento de los mercados. Mar de Plata, Argentina: Facultad de Ciencias Económicas y Sociales, Universidad Nacional de Mar del Plata.

Avdoshin, S. M., & Pesotskaya, E. Y. (2011). Software risk management. National Research University Higher School of Economics , 1-6 Pp.

Baquela, E. G., & Redchuk, A. (2013). Optimización Matemática con R. Volumen I: Introducción al modelado y resolución de problemas. Madrid, España: Bubok Publishing S. L.

Bazaraa, M. S., Sherali, H. D., & Shetty, C. M. (2009). Nonlinear programming: Theory and applications. Hoboken, New Jersey, U.S.A.: John Wiley & Sons, Inc, Publication.

Bertsimas, D., & Tsitsiklis, J. N. (1997). Introduction to Linear Optimization. Belmont, Mass. U.S.A.: Athena Scientific, Belmont, Massachusetts.

Bertsimas, D., Teo, C., & Vohra, R. (1999). On dependent randomized rounding algorithms. ELSEVIER, 105 –114.

Bhatt, A., Das, V., & Sundaram, V. (2011). Economic analysis of malaria control programs: A review of the literature. Malaria Journal, 10(1), 202.

Birge, J. R., & Louveaux, F. (2011). Introduction to Stochastic Programming. New York, USA: Springer Science+Business Media, .

Brambila, P. J. (2011). Bioeconomía: Instrumentos para su análisis económico. Texcoco, estado de México. México: Secretaría de Agricultura, Ganadería, Desarrollo Rural, Pesca

y Alimentación (SAGARPA) y Colegio de Postgraduados (COLPOS).

Broz, D., Mac Donagh, P., Arce, J., & Yapura, P. (2017). La Investigación Operativa, la Ingeniería Forestal y los Problemas Sectoriales: Ante la Necesidad de un Cambio de Paradigma. *Yvyrareta*, 64-72.

Burbano, R., Espinosa, A. M., & Horna, L. (2018). Economía para matemáticos. Quito, Ecuador: Escuela Politécnica Nacional (EPN).

Burkard, R., Dell'Amico, M., & Martello, S. (2012). Operations research: Theory and practice. New Jersey, U.S.A.: Springer Science & Business Media.

Bustos, F. E. (2023). La importancia del teorema de suficiencia de Kuhn-Tucker en la tarea de decisiones organizacionales. Carácas, Venezuela: Escuela Superior de Cómputo de Venezuela.

Cabrera, G. E., & Díaz, G. E. (2021). Manual de uso de Jupyter Notebook para aplicaciones docentes. Madrid, España: Universidad Complutense de Madrid.

Camps, P. R., Casillas, S. L., Costal, C. D., Gibert, G. M., Martín, E. C., & Pérez, M. O. (2005). Software libre. Bases de datos. Av. Tibidado, 39-43, Barcelona España: Fundació per a la Universitat Oberta de Catalunya. Eureka Media, SL.

Capa, S. H. (2015). Probabilidades y estadística para una gestión científica de la información. Quito, Ecuador: Escuela Politécnica Nacional .

Carballo, R. E., & Peñate, P. E. (2014). Benchmarking de Herramientas Forenses Móviles. Managua: COMPDES, UNAN, Managua.

Carro, P. R. (2014). Investigación de operaciones en administración. 1ra Edición. Mar de Plata, Provincia de Buenos

Aires, Argentina: Facultad de Ciencias Económicas y Sociales. Universidad Nacional de Mar de Plata.

Cestero, E. V., & Caballero, A. M. (2017). Data Science y Redes Complejas. Métodos y aplicaciones. Tomás Bretón, 21-28045, Madrid, España: Editorial Centro de Estudios Ramón Areces, S. A.

Challenger, P. I., Díaz, R. Y., & Becerra, G. R. (Abril-Junio, 2014). El lenguaje de programación Python. Ciencias Holguín, vol. XX, núm. 2., 1-13.

Dantzig, G. B. (1963). Linear Programming and Extensions. Chichester, West Sussex, U.S.A.: Princeton University Press.

Dantzig, G. B., & Wolfe, P. (1954). The decomposition principle for linear programming. Operations Research, 2(1), 76-81.

Daza, P. J. (2006). Estadística Aplicada con Excel. Lima, Perú: Grupo Editorial Megabyte s.a.c.

De los Reyes García, M. M., & Romero, C. J. (2004). Investigación de operaciones I. UAM Unidad Azcapotzalco. Av. San Pablo 180. Col. Reynosa Tamaulipas. Delegación Azcapotzalco. México, D. F.: División de Ciencias Básicas e Ingeniería. Departamento de Sistemas. Universidad Autónoma Metropolitana-Azcapotzalco.

De Mendiburu, F. (2007). Análisis estadístico con R. Av. Túpac Amaru 210, Rímac 15333, Lima, Perú: Sección de extensión universitaria y proyección social. Centro de estudiantes de facultad de ingeniería económica y ciencias sociales. Universidad Nacional de Ingeniería.

De Mendiburu, F. (2017). Tutorial de agricolae (Versión 1.2-8). La Molina-Perú: Departamento Académico de Estadística e Informática de la Facultad de Economía y Planificación. Universidad Nacional Agraria La Molina-Perú.

Delgado, Q. S. (2023). Aprende Python. Ciudad de México, México: Mundi Prensa.

Deng, Z. X., & Wang, X. (2022). A survey on non-predetermined goal programming. *European Journal of Operational Research*, 297(2), 1115-1133.

Devia, J., Azacon, J., López, N., Villarroel, Z., Carvajal, A., Rodríguez, B., . . . Suárez, F. (2012). Software TORA. Maturín, Venezuela: Núcleo de Monagas. Universidad de Oriente de Venezuela.

Di Perro, M. (2017). What is the Blocchaing? SChool of Computing at DePaul University, Chicago, IL, USA, 92-95.

Díaz, C. M., & Gómez, M. P. (2011). Comportamiento de volatilidad del tipo de México 1970 a 2010. Michoacán, México: Universidad Michoacana de San Nicolás de Hidalgo.

DiLorenzo, T. J. (1992). El Mito del Monopolio Natural. *The Review of Austrian Economics* Vol. 9, No. 2 , 1-15 Pp.

Eiselt, H. A., & Sandblom, C. L. (2009). Integer programming: Theory and practice. New Jersey, U.S.A.: Springer.

Engler, P. A. (2009). Estrategías del manejo del riesgo. Quilamapu, Chillán, Chile: TRAMA Impresores S. A.

Galindo, D. T. (2006). Estadística. Métodos y Aplicaciones. Quito, Ecuador: Prociencia Editores.

Garcés, R. A. (2020). Optimización convexa. Aplicaciones en operación y dinámica de sistemas de potencia. Pereira, Colombia: Universidad Tecnológica de Pereira, Pereira, Colombia.

García, M. A., & García, S. (2015). Investigación de operaciones en el control de enfermedades parasitarias. *Revista de Investigación Académica*, 26, 1-12.

García, M. M., & Romero, C. J. (2004). Investigación de

operaciones 1. 1ra parte. Delegación Azcapotzalco, D. F.: Universidad Autónoma Metropolitana (UAM-Azcapotzalco).

Gen, M., & Cheng, R. (2000). Genetic algorithms and engineering optimization. New York, U.S.A.: Wiley-Interscience Publication.

Gómez, G. M., Danglot, B. C., & Vega, F. L. (2003). Sinopsis de pruebas estadísticas no paramétricas. Cuándo usarlas. Revista Mexicana de Pediatría, 10.

Gómez, P. M. (2006). Introducción a la microeconomía. Barcelona, España: Universidad de Barcelona.

González, B. G. (1985). Métodos Estadísticos y Principios de Diseño Experimental. Quito, Ecuador: Facultad de Ciencias Agrícolas. Universidad Central del Ecuador.

González, B. G. (2010). Métodos estadísticos y principios de diseño experimental. Quito, Ecuador: Universidad Central del Ecuador.

González, J. A., García, S. J., Chalita, T. L., Matus, G. J., Cruz, G. B., Sangerman, J. D., . . . Fortis, H. M. (2012). Modelo de equilibrio espacial para determinar costos de transporte en la distribución de durazno en México. Revista Mexicana de Ciencias Agrícolas Vol.3 Núm.4 , 701-712 Pp.

Grossman, S. S., & Flores, G. J. (2012). Álgebra Lineal. México D. F.: Colonia Santa Fe, Delegación Álvaro Obregón, México D. F.

Guerrero, S. H. (2009). Programación Lineal Aplicada. Bogotá, Colombia: Ecoe Ediciones.

Gujarati, D. N., & Porter, D. C. (2010). Econometría. Delegación Álvaro Obregón, México, D. F.: McGraw-Hill/ Interamericana Editorres, S.A. de C.V.

Hernández, N., Soto, F., & Caballero, A. (2009). Modelos de simulación de cultivos, características y usos. Cultivos

Tropicales, vol. 30, núm. 1, 73-82 Pp.

Hernandez, S. R., Fernandez, C. C., & Baptista, L. P. (2006). Metodología de la investigación. Colonia Desarrollo Santa Fe, Delegación Álvaro Obregón, México, D. F. : McGraw Hill/ Interamericana Editores, S.A. de C. V.

Herrera, A. L. (2009). Perspecctivas para la carna de bovino en México aplicando un sistema de trazabilidad. Campus Montecillo, Texcoco, estado de México: Colegio de Postgraduados. Tesis de maestría.

Hillier, F. S., & Lieberman, G. J. (2010). Introducción a la investigación de operaciones. Delegación Álvaro Obregón, México D. F.: McGraw-Hill Companies, Inc.

Hillier, F. S., & Lieberman, G. J. (2015). Operations Reseach. New York, U.S.A.: McGraw-Hill Education.

Hung, M. T., & Hung, C. C. (2019). A review of non-predetermined goal programming: Methods, applications, and challenges. European Journal of Operational Research, 274(3), 909-924.

Hurtado, M. J., & Gómez, F. R. (2010). Diseño Experimental. Valencia, España: Universidad de Valencia.

Iusem, A. N. (2003). On the convergence properties of the projectedgradient method for convex optimization. Computational and Applied Mathematics, 37–52.

Kaiser, H. M., & Messer, K. D. (2011). Mathematical Programming for Agricultural, Enviromental, and Resource Economics. Rosewood Drive, Danvers, United States of America: John Wiley & Sons, Inc.

Kmenta, J. (1971). Elements of econometrics. Collier Macmillan Canada, Inc.: Macmillan Publishing Ccompany.

Larrañaga, L. J., Zulueta, G. E., Elizagarate, U. F., & Alzola, B. J. (2020). Algoritmos Meméticos en problemas de

Investigación Operativa. Escuela Universitaria de Ingeniería de Vitoria-Gasteiz. Universidad del País Vasco Euskal Herriko Unibertsitatea (UPV-EHU)., 1-19.

Leek, J. T., & Peng, R. D. (2015). Statistics: P values are just the tip of the iceberg. Maryland, USA.: School of Public Health in Baltimore.

Li, H., Zhang, J., & Zhang, X. (2019). Randomized rounding for general integer programs. *Operations Research Letters*, 47(2), 164-169.

Li, W., & Yang, L. (2022). A novel non-predetermined goal programming model for renewable energy planning. *Energy*, 220, 122772.

Lind, D. A., Marchal, W. G., & Mason, R. D. (2004). Estadística para Administración y Economía. Col. del Valle, México D. F.: Alfaomega Grupo Editor S. A. de C. V.

Lind, D. A., Marchal, W. G., & Mason, R. D. (2006). Estadística para Administración y Economía. Col. del Valle, México D. F.: Alfa Omega Grupo Editor S. A. de C. V.

López, B. E., & González, R. B. (2014). Diseño y Análisis de Experimentos. Fundamentos y Aplicaciones en Agronomía. Ciudad Universitaria zona 12. Ciudad de Guatemala: Universidad de San Carlos de Guatemala, Facultad de Agronomía.

Loubet, B. (2020). Excel: Herramienta Solver. Centro Universitario, Ciudad de Mendoza. Provincia de Mendoza, Argentina.: Facultad de ciencias económicas. Universidad Nacional de Cuyo.

Ipízar, M. J. (12 de Octubre de 2023). Universidad Latina. Obtenido de La organización industrial: <https://cisprocr.com/cispro/system/files/Clase%209%20Organizaci%C3%B3n%20Industrial.pdf>

Luenberger, D. G. (1984). Linear and Nonlinear Programming . New JErsey, U.S.A.: Springer.

Machado, D. E., & Coto, F. H. (2017). Sistema de adquisición de datos con Pyhton y Arduino. Revista, Ciencia, Ingeniería y Desarrollo Tec Lerdo, 1-6.

Madrid, C. C. (2020). Filosofía de la ciencia del cambio climático: Modelos, problemas e incertidumbres. Revista Colombiana de Filosofía de la Ciencia, vol. 20, núm. 41, 201-234 Pp.

Martínez, Á. (26/10/2023 de Octubre26 de 2023). Optimización global. Obtenido de Optimización global: http://www.dma.uvigo.es/~aurea/Optimizaci%C3%9Bn_global.

El libro “Fundamentos Matemáticos de Investigación Operativa” explora tres áreas clave: programación lineal, formulación de hipótesis y series de Taylor. A través de teoría y ejemplos prácticos, enseña cómo aplicar estas herramientas matemáticas en problemas reales, mejorando la eficiencia y toma de decisiones. La programación lineal optimiza recursos limitados, la formulación de hipótesis guía la investigación y el análisis de datos, y las series de Taylor aproximan funciones complejas, fundamentales en el análisis numérico y las ciencias aplicadas.



Moisés Arreguín Sámano

Profesor e investigador de la Universidad Estatal de Bolívar (UEB), Ecuador.



Hernán Vinicio Villa Sánchez

Docente Investigador Universidad Técnica de Ambato, Facultad de Contabilidad y Auditoría; Ingeniero en Administración de Empresas; Magíster en Formulación, Evaluación y Gestión de Proyectos Sociales y Productivos; Ambato, Ecuador.



Ángel Leyva- Ovalle

Docente Investigador Universidad Autónoma Chapingo (UACH-DiCiFo), México. Ingeniero forestal con orientación en economía y ordenación por la Universidad Autónoma Chapingo (UACH), con experiencia en la División de Ciencias Forestales (DiCiFo), específicamente en el Departamento de Productos Forestales.



Numa Inain Gaibor Velasco

Graduado con un título de Ingeniero en Administración para Desastres y Gestión del Riesgo en la Universidad Estatal de Bolívar, Ecuador, con un grado de Maestría en Gestión del Riesgo con mención en Preparación para la Respuesta en la Universidad Internacional SEK, Quito, Ecuador. Su experiencia como profesor-investigador en la Universidad Estatal de Bolívar, imparte cursos de pregrado relacionados con la gestión del riesgo de desastres.

ISBN: 978-9942-684-24-0



9 789942 684240