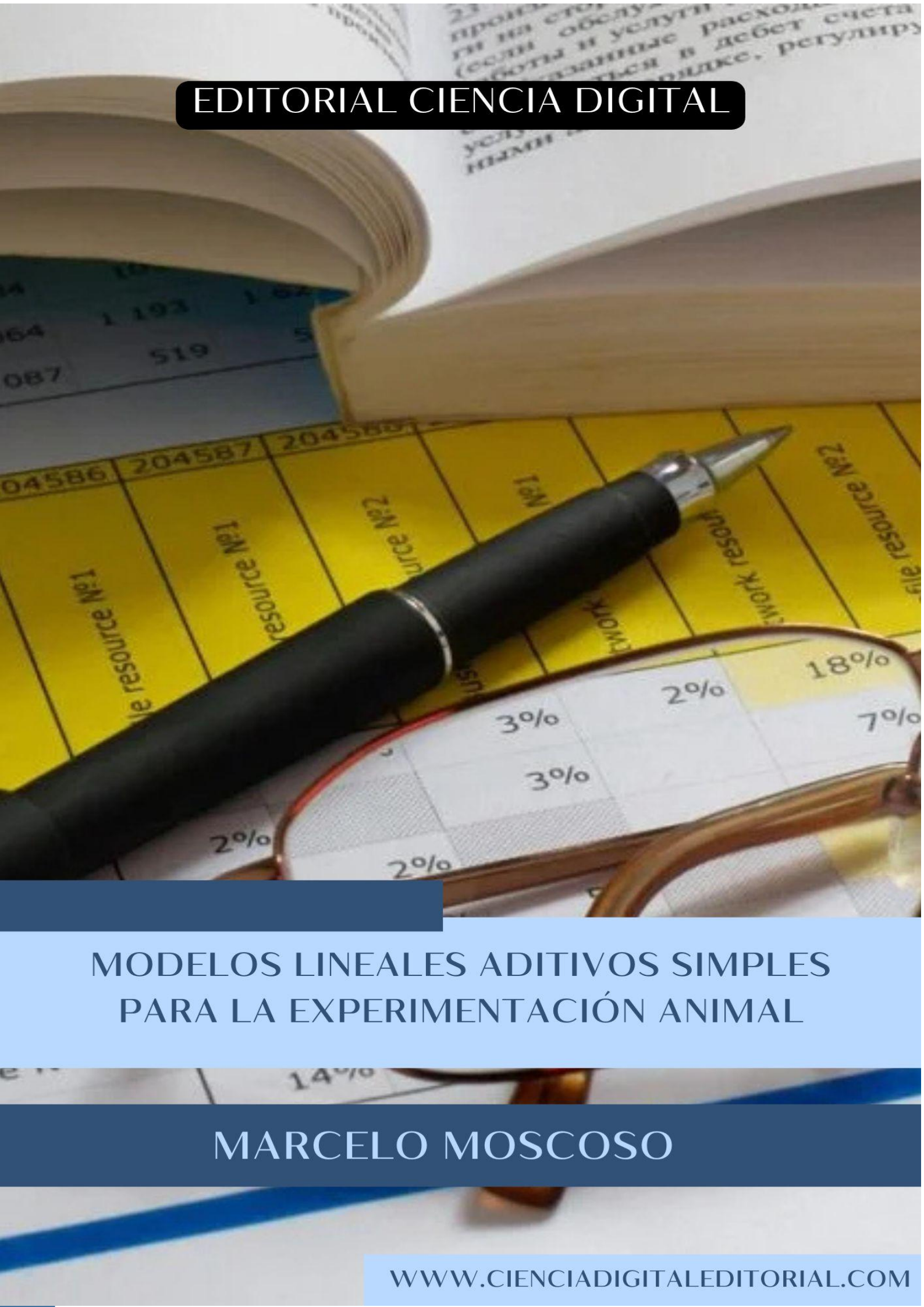




EDITORIAL CIENCIA DIGITAL

MODELOS LINEALES ADITIVOS SIMPLES
PARA LA EXPERIMENTACIÓN ANIMAL

MARCELO MOSCOSO

WWW.CIENCIADIGITALEEDITORIAL.COM

El libro **Modelos lineales aditivos simples para la experimentación animal** está avalado por un sistema de evaluación por pares doble ciego, también conocido en inglés como sistemas “*double-blind paper review*” registrados en la base de datos de la **EDITORIAL CIENCIA DIGITAL** con registro en la Cámara Ecuatoriana del Libro No.663 para la revisión de libros, capítulos de libros o compilación.

ISBN_978-9942-7437-2-5

Primera edición, septiembre 2025

Edición con fines didácticos

Coeditado e impreso en Ambato - Ecuador

El libro que se publica es de exclusiva responsabilidad de los autores y no necesariamente reflejan el pensamiento de la **Editorial Ciencia Digital**.

El libro queda en propiedad de la editorial y por tanto su publicación parcial y/o total en otro medio tiene que ser autorizado por el director de la **Editorial Ciencia Digital**.



Jardín Ambateño, Ambato, Ecuador

Teléfono: 0998235485 – 032-511262

Publicación:

w: www.cienciadigitaleditorial.com

w: <http://libros.cienciadigital.org/index.php/CienciaDigitalEditorial>

e: luisefrainvelastegui@cienciadigital.org

AUTORES

AUTORES



Marcelo Moscoso Gómez

Escuela Superior Politécnica de Chimborazo



La **Editorial Ciencia Digital**, creada por Dr.C. Efraín Velasteguí López PhD. en 2017, está inscrita en la Cámara Ecuatoriana del Libro con registro editorial No. 663.

El **objetivo** fundamental de la **Editorial Ciencia Digital** es un observatorio y lugar de intercambio de referencia en relación con la investigación, la didáctica y la práctica artística de la escritura. Reivindica a un tiempo los espacios tradicionales para el texto y la experimentación con los nuevos lenguajes, haciendo de puente entre las distintas sensibilidades y concepciones de la literatura.

El acceso libre y universal a la cultura es un valor que promueve Editorial Ciencia Digital a las nuevas tecnologías esta difusión tiene un alcance global. Muchas de nuestras actividades están enfocadas en este sentido, como la biblioteca digital, las publicaciones digitales, a la investigación y el desarrollo.

Desde su creación, la Editorial Ciencia Digital ha venido desarrollando una intensa actividad abarcando las siguientes áreas:

- Edición de libros y capítulos de libros
- Memoria de congresos científicos
- Red de Investigación

Editorial de las revistas indexadas en Latindex 2.0 y en diferentes bases de datos y repositorios: **Ciencia Digital** (ISSN 2602-8085), **Visionario Digital** (ISSN 2602-8506), **Explorador Digital** (ISSN 2661-6831), **Conciencia Digital** (ISSN 2600-5859), **Anatomía Digital** (ISSN 2697-3391), **Alfa Publicaciones** (ISSN 2773-7330) & **Ciencia & Turismo** (ISSN 3028-8665).



ISBN: 978-9942-7437-2-5



ISBN: 978-9942-7437-2-5 Versión Electrónica

-  Los aportes para la publicación de esta obra, está constituido por la experiencia de los investigadores

EDITORIAL REVISTA CIENCIA DIGITAL



 Efraín Velasteguí López¹

Contacto: Ciencia Digital, Jardín Ambateño, Ambato- Ecuador

Teléfono: 0998235485 - 032511262

Publicación:

w: www.cienciadigitaleditorial.com

e: luisefrainvelastegui@cienciadigital.org

Editora Ejecutiva

Dr. Tatiana Carrasco R.

Director General

Dr.C. Efraín Velasteguí PhD.

¹ **Efraín Velasteguí López: Magister** en Tecnología de la Información y Multimedia Educativa, **Magister** en Docencia y Currículo para la Educación Superior, **Doctor (PhD)** en Ciencia Pedagógicas por la Universidad de Matanza Camilo Cien Fuegos Cuba, **Postdoctorado** en Estrategias Didácticas para la Investigación y **Postdoctorado** en Currículo y Formación de Investigadores por la Universidad Nacional Sur del Lago Jesús María Semprum (UNESUR) Venezuela, **Doctor Hispano Honoris Causa** por la Universidad Anglo Hispano Mexicana, cuenta con más de **200 publicaciones** en revista indexadas en Latindex Scielo y Scopus, 21 ponencias a nivel nacional e internacional, mas **30 libros** con ISBN, registrada en la cámara ecuatoriano del libro, **tres patente** de la marca Ciencia Digital, Acreditación en la categorización de **investigadores** nacionales y extranjeros Registro REG-INV-18-02074, **Director, editor** de las revistas indexadas en Latindex Catalogo 2.0, Ciencia Digital, Visionario Digital, Explorador Digital, Conciencia Digital, Anatomía Digital, Alfa Publicaciones, Ciencia & Turismo y la **editorial Ciencia Digital** Registro editorial N.º 663. Cámara Ecuatoriana del Libro, director de la **Red de Investigación Ciencia Digital**, emitido mediante Acuerdo N.º.SENESCYT-2018-040, con número de registro REG-RED-18-0063.

EJEMPLAR GRATUITO
PROHIBIDA SU VENTA

El “copyright” y todos los derechos de propiedad intelectual y/o industrial sobre el contenido de esta edición son propiedad de CDE. No está permitida la reproducción total y/o parcial de esta obra, ni su tratamiento informático, ni la transmisión de ninguna forma o por cualquier medio, electrónico, mecánico, por fotocopia o por registro u otros medios, salvo cuando se realice con fines académicos o científicos y estrictamente no comerciales y gratuitos, debiendo citar en todo caso a la editorial.

TABLA DE CONTENIDO

CAPÍTULO 1. EXPERIMENTACIÓN AGROPECUARIA Y EL DISEÑO EXPERIMENTAL	15
1.1. Experimentación agropecuaria y el diseño experimental.....	19
1.2. Diferencia entre investigación y experimento	22
1.3. Tipos de estudios.....	26
1.3.1. Estudio observacional.....	26
1.3.2. Estudio experimental.....	26
1.3.3. Estudios Mixtos	27
1.3.4. Experimentos comparativos.....	28
1.4. Otros estudios	28
1.4.1. Según el propósito de estudio	28
1.4.2. Según la profundización del estudio.....	29
1.4.3. Según el tipo de “data”.....	29
1.4.4. Según su temporalidad	30
1.4.5. Según el tipo de inferencia.....	30
1.5. El método científico para experimentos agropecuarios.....	30
1.5.1. Orígenes y su desarrollo histórico.....	31
1.5.2. Las etapas clásicas del método científico	31
1.5.3. Etapas en la experimentación agropecuaria	32
1.6. Especificación de problema	34
1.6.1. Objetivos y metas	37
1.6.2. Hipótesis.....	39
1.6.3. Hipótesis nula y alternativa	39
1.6.4. Importancia de la Prueba de Hipótesis	41
1.6.5. Muestra o unidades experimentales.....	42
1.6.6. Diseño del experimento	47
1.6.7. Modelo lineal aditivo.....	54
1.6.8. Recolección de la data.....	63
1.6.9. Interpretación y comparación.....	64

1.6.10.	Conclusiones y toma de decisiones	67
1.7.	Variables respuesta y su clasificación.....	68
1.7.1.	Variables cualitativas.....	69
1.7.2.	Variables cuantitativas	70
1.8.	Estadística descriptiva e importancia del coeficiente de variación	70
1.8.1.	Sorteo en un Diseño Completamente al Azar	72
1.8.2.	Sorteo en un Diseño de Bloques Completamente al Azar.....	81
	Referencias bibliográficas del capítulo	89
CAPITULO II. EXPLORACIÓN Y TRATAMIENTO DE DATOS		
EXPERIMENTALES.....		92
2.1.	Obtención y organización de la “data”	95
2.2.	Análisis descriptivo de la “data”	99
2.3.	Exploración de Supuestos en Análisis Estadístico: Normalidad, Homogeneidad de la Varianza e Independencia.....	109
2.3.1.	Ejemplos de Violaciones de Normalidad	109
2.3.2.	. Homogeneidad de la Varianza	110
2.3.3.	Consecuencias de la Violación de Homogeneidad	110
2.4.	Independencia.....	110
2.4.1.	Ejemplos de Violaciones de Independencia.....	110
2.4.2.	Evaluación de la Normalidad	110
2.4.3.	Evaluación de la Homogeneidad de la Varianza.....	111
2.4.4.	Evaluación de la Independencia	111
2.5.	Implicaciones de la Violación de Supuestos	112
2.5.1.	Estrategias para Manejar Violaciones de Supuestos	112
2.5.2.	Transformaciones de Datos	112
2.5.2.1.	Métodos No Paramétricos	112
2.5.2.2.	Modelos Mixtos	113
2.6.	Ejercicios sobre Modelos Lineales Aditivos Simples en la Experimentación Animal	113
2.6.1.	Soluciones a los Ejercicios sobre Modelos Lineales Aditivos Simples en la Experimentación Animal	118
2.6.2.	Cálculo de Indicadores Bio productivos para la Producción de Leche y Carne	121
2.6.2.1.	Indicadores Bio productivos en la Producción de Leche.....	121

2.6.2.2.	Indicadores Bioroductivos en la Producción de Carne	122
2.6.3.	Comprobación de supuestos mediante gráficas.....	123
2.6.4.	Consideraciones Generales.....	130
2.7.	Impacto de la Genética	130
2.8.	Prácticas de Manejo	130
2.9.	Sostenibilidad	130
	Referencias bibliográficas del capítulo	131
	CAPITULO III: MODELOS LINEALES ADITIVOS DE ORDENACIÓN SIMPLE	
	(ONEWAY).....	134
3.1.	Características del Diseño Completamente al Azar (DCA)	136
3.1.1.	Definición y Principios Básicos del DCA.....	137
3.1.2.	Aspectos del Diseño Completamente al Azar (DCA).....	140
3.1.3.	Limitaciones y Consideraciones del DCA.....	144
3.1.4.	Ventajas y desventajas del DCA	147
3.1.5.	Ventajas.....	149
3.1.6.	Desventajas	150
3.2.	Modelo lineal aditivo para el DCA	151
3.2.1.	Fundamentos del Modelo Lineal Aditivo en el DCA.....	152
3.2.2.	Aplicaciones del Modelo Lineal Aditivo en el DCA.....	156
3.2.3.	Fórmulas de cálculo para el DCA con igual repetición	159
3.3.	Importancia para el Análisis de Varianza (ANOVA).....	163
3.4.	Aplicación del DCA con igual repetición	163
3.5.	Análisis de Varianza (ANOVA) en DCA con desigual repetición	167
3.6.	Fórmulas de cálculo para el DCA con desigual repetición	167
3.7.	Aplicación del DCA con desigual repetición	170
	Referencias bibliográficas del capítulo	174
	CAPITULO IV: MODELOS LINEALES ADITIVOS DE ORDENACIÓN	
	MÚLTIPLE.....	176
4.1.	Características del Diseño de Bloques Completamente al Azar (DBCA)	177
4.2.	Ventajas y desventajas del DBCA	180
4.3.	Modelo lineal aditivo para el DBCA	181
4.4.	Fórmulas de cálculo para el DBCA.....	182

4.5.	Aplicación del DBCA	184
4.6.	Características del Diseño Cuadrado Latino (DCL)	189
4.7.	Ventajas y desventajas del DCL	193
4.7.1.	Ventajas del Diseño Cuadrado Latino	194
4.7.2.	Desventajas del Diseño Cuadrado Latino	194
4.7.3.	Modelo lineal aditivo para el DCL	196
4.7.4.	Fundamentos del Modelo Lineal Aditivo	196
4.7.5.	Interpretación de los Componentes del Modelo	196
4.8.	Fórmulas de cálculo para el DCL	197
4.9.	Aplicación del DCL	200
	Referencias bibliográficas del capítulo	210
	CAPÍTULO V: PRUEBAS DE SIGNIFICACIÓN O COMPARACIÓN DE MEDIAS DE LOS TRATAMIENTOS.....	212
5.1.	Características de la comparación de medias de los tratamientos.....	215
5.2.	Diferencia mínima significativa.....	217
5.4.	Prueba de Tukey	229
5.5.	Prueba de Scheffé.....	237
	Referencias bibliográficas del capítulo	243
	CAPITULO VI: APLICACIÓN DE LOS MODELOS LINEALES ADITIVOS SIMPLES MEDIANTE EL INFOSTAT	245
6.1.	Fundamentos Teóricos de los Modelos Lineales	246
6.2.	Consideraciones Prácticas y Limitaciones	246
6.3.	Selección de Tamaño Muestral y Poder Estadístico	247
6.4.	Aplicación de los modelos lineales aditivos simples mediante el InfoStat 249	
6.5.	Entorno del InfoStat.....	251
6.5.1.	Interfaz General y Navegación	253
6.5.2.	Gestión y Manejo de Datos	253
6.5.3.	Visualización y presentación de Resultados	254
6.5.4.	Soporte y Actualización	254
6.5.5.	Integración con R.....	255
6.6.	Protocolo para la exploración y tratamiento de “data” con el uso de las librerías del “R”	255

6.7.3.	Ejecución y Resultados.....	263
6.7.4.	Comparación de Medias.....	264
6.8.	Uso del InfoStat en un modelo DBCA	265
6.9.	Uso del InfoStat en un modelo DCL	271
6.10.	Aplicación del InfoStat en los modelos lineales aditivos simples	278
	Referencias bibliográficas.....	295

AGRADECIMIENTOS

A todas las personas que apoyaron el desarrollo de este libro

INTRODUCCIÓN

El manejo de las herramientas para establecer el análisis y diseño de experimentos, son imponderablemente necesarias para la formación del profesional, puesto que se complementa con la metodología de la investigación y permite que con el uso del método científico, se puedan tomar decisiones certeras en el campo agropecuario. Este conjunto de conocimientos desarrolla la experticia para poder elegir los diferentes modelos lineales aditivos, en los experimentos simples, así como; en los factoriales además, permiten establecerse pronósticos futuros en base a datos históricos en las distintas variables zootécnicas y veterinarias, que se exteriorizan en la dinámica de la producción animal.

Con el análisis y diseño de experimentos anclados a los modelos matemáticos, el profesional aplicará las estrategias estadísticas requeridas, según el caso particular que debe tratar para resolver problemas de la producción pecuaria basado en una metodología científica; describiendo y comparando diferentes realidades para comprobar si los datos zootécnicos corresponden o no a los estándares internacionales, nacionales o locales; con el fin de diseñar su propio plan de manejo productivo eficiente, utilizando los tratamientos específicos para el efecto.

Con estas competencias el profesional pecuario logrará tomar decisiones científicas para resolver los problemas de la dinámica productiva en los sectores sociales y productivos del país.



CAPÍTULO I

EXPERIMENTACIÓN
AGROPECUARIA Y EL
DISEÑO
EXPERIMENTAL

La experimentación agropecuaria es un componente fundamental en la investigación agrícola que busca optimizar la producción de cultivos y la cría de animales. A medida que la población mundial crece y se enfrenta a desafíos como el cambio climático, la escasez de recursos y la demanda de alimentos de calidad, la necesidad de desarrollar y aplicar métodos científicos en la agricultura se vuelve cada vez más crítica. La experimentación agropecuaria permite a los investigadores y agricultores evaluar nuevas técnicas, variedades de cultivos, prácticas de manejo y tecnologías, contribuyendo así a la sostenibilidad y la eficiencia del sector agropecuario.

La experimentación en el ámbito agropecuario no solo se limita a la mejora de la productividad, sino que también abarca aspectos como la conservación del medio ambiente, la salud del suelo y la biodiversidad. A través de ensayos controlados, los investigadores pueden obtener datos precisos que permiten tomar decisiones informadas sobre el manejo de recursos naturales. Por ejemplo, el uso de fertilizantes y pesticidas puede ser evaluado para determinar su impacto en el rendimiento de los cultivos y en la salud del ecosistema.

Además, la experimentación agropecuaria es esencial para el desarrollo de variedades de cultivos que sean más resistentes a enfermedades, plagas y condiciones climáticas adversas. La biotecnología y la genética son herramientas cruciales en este proceso, permitiendo la creación de cultivos que no solo son más productivos, sino también más nutritivos y adaptados a las necesidades de las comunidades locales.

El diseño experimental es el proceso de planificar un experimento de manera que se obtengan resultados válidos y confiables. En la experimentación agropecuaria, un diseño adecuado es crucial para minimizar el sesgo y la variabilidad, lo que permite obtener conclusiones precisas y aplicables. Un buen diseño experimental incluye la selección adecuada de tratamientos, la asignación aleatoria de tratamientos a las unidades experimentales y la repetición de ensayos para garantizar la robustez de los resultados.

Existen varios tipos de diseños experimentales que se pueden aplicar en la investigación agropecuaria:

1. Diseños Completamente Aleatorizados: En este tipo de diseño, los tratamientos se asignan aleatoriamente a las unidades experimentales. Es sencillo de implementar y es adecuado cuando las unidades experimentales son homogéneas.
2. Diseños en Bloques Aleatorizados: Este diseño se utiliza cuando hay variabilidad en las unidades experimentales. Se agrupan en bloques homogéneos y dentro de cada bloque se asignan aleatoriamente los tratamientos. Esto ayuda a controlar la variabilidad y mejora la precisión de las estimaciones.
3. Diseños Factoriales: Permiten estudiar el efecto de dos o más factores simultáneamente. Por ejemplo, se podría investigar cómo diferentes dosis de fertilizantes y tipos de riego afectan el rendimiento de un cultivo. Este enfoque es eficiente y proporciona información sobre interacciones entre factores.
4. Diseños de Ensayo de Campo: En el contexto agropecuario, los ensayos de campo son esenciales para evaluar el rendimiento de cultivos en condiciones reales. Estos diseños deben considerar factores como la topografía, el tipo de suelo y las condiciones climáticas.

Al planificar un experimento agropecuario, es importante considerar varios aspectos:

- Definir claramente qué se quiere investigar y cuáles son las hipótesis a probar. Esto guiará todo el proceso de diseño y ejecución.
- Identificar las variables dependientes e independientes que se evaluarán. Las variables dependientes son aquellas que se medirán como resultado de los tratamientos aplicados, mientras que las independientes son los factores que se manipulan en el experimento.

- Determinar el número adecuado de repeticiones y el tamaño de la muestra es esencial para asegurar que los resultados sean representativos y estadísticamente significativos.
- Planificar cómo se analizarán los datos recolectados. Esto incluye la elección de pruebas estadísticas adecuadas que se alineen con los objetivos del estudio.

La experimentación agropecuaria tiene numerosas aplicaciones prácticas que impactan directamente en la vida de los agricultores y en la seguridad alimentaria. Por ejemplo, los ensayos de variedades de cultivos pueden ayudar a identificar las más adecuadas para una región específica, considerando factores como el clima, el tipo de suelo y las prácticas de manejo.

Asimismo, la investigación sobre el manejo integrado de plagas (MIP) permite a los agricultores adoptar enfoques más sostenibles, combinando métodos biológicos, culturales y químicos para controlar plagas de manera efectiva, reduciendo así la dependencia de pesticidas químicos.

Los estudios sobre prácticas de conservación del suelo, como la rotación de cultivos y el uso de cubiertas vegetales, también son ejemplos de cómo la experimentación puede contribuir a la sostenibilidad agropecuaria. Estas prácticas ayudan a mejorar la estructura del suelo, aumentar su fertilidad y reducir la erosión, lo que a su vez beneficia la producción a largo plazo.

A pesar de sus beneficios, la experimentación agropecuaria enfrenta varios desafíos. Uno de los principales es la variabilidad inherente de los sistemas agroecosistémicos, que puede dificultar la obtención de resultados consistentes. Además, la falta de recursos, tanto financieros como humanos, puede limitar la capacidad de llevar a cabo investigaciones exhaustivas.

Otro desafío importante es la necesidad de traducir los resultados de la investigación en prácticas aplicables y accesibles para los agricultores. La comunicación efectiva de los hallazgos y la capacitación en nuevas tecnologías son aspectos cruciales para asegurar que los beneficios de la investigación lleguen a quienes más los necesitan.

La experimentación agropecuaria y el diseño experimental son pilares fundamentales en la búsqueda de soluciones innovadoras y sostenibles para los desafíos del sector agropecuario. A través de un enfoque científico riguroso, es posible mejorar la producción de alimentos, conservar los recursos naturales y contribuir al bienestar de las comunidades rurales. A medida que se avanza hacia un futuro en el que la seguridad alimentaria y la sostenibilidad son cada vez más prioritarias, la investigación agropecuaria seguirá desempeñando un papel crucial en el desarrollo de prácticas agrícolas más eficientes y responsables.

1.1. Experimentación agropecuaria y el diseño experimental

La investigación científica es considerada como la búsqueda permanente de la verdad por medio de métodos, objetivos adecuados y precisos; uno de los cuales es la experimentación, que consiste en realizar prácticas destinadas a demostrar, comprobar o descubrir fenómenos o principios básicos.

Se cuenta con un experimento cuando en la práctica se va a probar una hipótesis; y se tiene una investigación cuando se estudia la causa y el efecto. Es decir, en un experimento se observa únicamente los efectos y es de aplicación práctica inmediata, ya sea para el científico o para la comunidad; en tanto que la investigación es de aplicación mediata y puede ser evolucionista, conduciendo entonces a idear nuevas técnicas o a modificar las existentes.

La experimentación agropecuaria puede ser tratada como arte y como ciencia. Como arte porque se requiere de habilidad para ingeniar, planear o aplicar con conjunto de técnicas con el propósito de que sean eliminadas causas extrañas a la realización correcta del experimento de campo, laboratorio o invernadero; y, como ciencia, porque hay que aplicar el método científico y un conjunto de conocimientos científicos para el desarrollo de tecnologías que conduzcan a la formación de nuevos tipos de plantas, animales o nuevas prácticas agropecuarias que conlleven hacia el incremento de la producción y/o productividad, (Reyes, 2003).

Con este antecedente se puede indicar que en la experimentación agropecuaria su objetivo primordial es comprobar en la práctica del campo (cultivar o manejo

animal), una hipótesis formulada sobre alguna estrategia productiva o sanitaria que pueden intervenir en su productividad de una o más cosechas o sacas. Esto indicaría por ejemplo que si en alguna región específica una línea genética vacuna lechera es más productiva que las razas o variedades de nuestra zona, y que al introducirlas podría elevar considerablemente el rendimiento lechero del sector, basado en las suposiciones de lo acontecido en otro territorio o lo que sugieren los conocimientos técnicos pecuarios, y se debe adquirir la certidumbre de que tal suposiciones son bien fundamentadas, se debe acudir a la experimentación agropecuaria, con el fin de corroborar en el territorio nuestro con las condiciones establecidas (ambiente), (De La Loma, 1996).

La importancia que tiene la experimentación agropecuaria es que sobre ella descansa el progreso de la agricultura y ganadería mundial, con nuevos retos como la maximización de la productividad por unidad de superficie, puesto que quizá en el mismo territorio o peor aún en menor cantidad de superficie cultivada debido al avance demográfico, hay que producir para alimentar a una población creciente. Para cumplir esta misión, las instituciones de educación superior deben sintonizarse con esta realidad, aplicando estrategias como la agricultura de precisión y inclusive simulaciones con el uso de la inteligencia artificial; todos estos conceptos deben pasar por el crisol de experimentación para que pueda ser aceptada, divulgada y aplicada en las unidades productivas agropecuarias.

Gran parte del desarrollo de los países depende de la generación de alimentos, tomando en cuenta las técnicas amigables al medio ambiente, así como la seguridad alimentaria; para el efecto acudimos al método científico y el aporte de algunas ciencias como la agroecología, zootecnia, veterinaria, biología, etc. Ciertas estrategias que pudieron crearse en otras localidades con resultados muy satisfactorios, pero que requieren de protocolos científicos para que sean probadas, adaptadas y posteriormente aplicadas con técnicas propias respetando el sincretismo que tienen los productores agropecuarios; dentro de estas metodologías científicas, el diseño de experimentos mediante sus modelos lineales aditivos juega un rol de gran importancia (Moscoso Gómez, 2019)

Precisamente el diseño de experimentos nació con la agricultura en la década de 1920 y fue desarrollado por el británico Ronald Fisher quien prestaba sus

servicios en la Estación Experimental de Rothamsted en Inglaterra, sus estudios agrícolas tuvieron como objetivo mejorar la productividad de los cultivos. A partir de estas investigaciones, Fisher comenzó a formalizar métodos que permitieran diseñar experimentos de manera rigurosa, con el fin de obtener conclusiones válidas y confiables (Fisher, 1937).

Fisher introdujo conceptos fundamentales en el diseño de experimentos, tales como la aleatorización, la replicación y el bloqueo, los cuales se utilizan para controlar la variabilidad y asegurar que los resultados no se deban al azar. También es reconocido por introducir el análisis de varianza (ANOVA o ADEVA) como herramienta para analizar los resultados experimentales.

De este modo, el diseño de experimentos se originó en el contexto de la agricultura, pero rápidamente se expandió a otras áreas, como la biología, la industria y las ciencias sociales, debido a su utilidad para evaluar y mejorar procesos mediante experimentación controlada.

El estadístico de Fisher, establece que la “data” debe cumplir con todos los supuestos para que puedan aplicarse los modelos lineales aditivos convencionales, sin embargo hoy se ha cambiado el paradigma del enfoque estadístico, con técnicas mucho más modernas debido a que en la actualidad ya los datos no deben necesariamente adaptarse a los modelos fijos y cerrados, por el contrario hoy mediante las técnicas de modelación “el modelo se puede ajustar a nuestros datos”, es una suerte de elasticidad en el tratamiento y procesamiento para encontrar resultados más cercanos a nuestras realidades.

En la historia de la investigación, se tiene como premisa su ejecución directamente en los campus académicos, o en los centros de investigación; haciendo que sus resultados no lleguen directamente a los beneficiarios (productores), rompiendo la misión de la universidad “resolver problemas de la sociedad mediando la investigación”. Sin embargo, en la última década las investigaciones son adaptativas, donde la mayoría de las observaciones y experimentos se realizan en los campos de los productores agropecuarios; aunque se pueden inclusive combinar estudios observacionales con experimentales, llamados mixtos; por ejemplo, en los bosques o páramos donde se investiga secuestro de carbono, biomasa, carga animal, etc. Por ejemplo, si

en los sistemas de páramo se desea investigar los diferentes tipos de ecosistemas con la finalidad de aprovecharlos, se podría establecer su condición y tendencia (observacional) y dentro de los mismos analizar la carga animal soportable sin destruir su condición (experimental), y en el caso que se pronostique lo contrario, las medidas que se podrían tomar para evitar su deterioro.

Un forestal compara distintos tipos de bosques y quiere aprovecharlos.

1.2. Diferencia entre investigación y experimento

La producción agropecuaria, ha sido el pilar de la civilización humana desde los albores de la historia, y en la actualidad, el rol que juega en el desarrollo económico y la seguridad alimentaria sigue siendo vital. Dentro de este ámbito, la investigación y la experimentación agropecuaria han evolucionado para enfrentar los retos de producir alimentos de manera eficiente y sostenible; en este capítulo se debe analizar a estos dos conceptos, su evolución y su impacto en la ciencia agrícola contemporánea, destacando los aportes metodológicos más relevantes.

La investigación agropecuaria puede definirse como el conjunto de estudios destinados a comprender, mejorar y optimizar la producción agrícola y ganadera. La investigación en este campo tiene sus raíces en la observación empírica. Durante siglos, los agricultores aprendieron de la experiencia, transmitiendo sus conocimientos de generación en generación. Sin embargo, este enfoque empírico, si bien valioso, carecía del rigor científico necesario para generalizar los resultados a diferentes contextos agroecológicos.

Con la llegada de la Revolución Científica y la consolidación del método científico en el siglo XVII, se abrió un nuevo camino para la investigación agrícola. Este proceso culminó en el siglo XX con la institucionalización de la investigación agropecuaria a través de centros especializados, universidades y estaciones experimentales. En este sentido, la investigación dejó de ser una serie de observaciones informales para convertirse en un proceso estructurado, basado en el análisis de datos y la validación científica de los resultados.

Dentro del ámbito de la investigación agropecuaria, la experimentación ha sido la herramienta principal para probar hipótesis y generar nuevo conocimiento. La experimentación agropecuaria consiste en la aplicación controlada de técnicas y métodos de cultivo y cría, para observar los efectos de diferentes factores (como fertilizantes, riego, tipos de suelo, plagas, razas selectas para producción de leche y carne, etc.) en el rendimiento agrícola o la productividad animal. Este enfoque permite a los científicos aislar variables y determinar relaciones de causa y efecto de manera objetiva.

Un avance clave en la experimentación agropecuaria fue el concepto que introdujo Fisher sobre el diseño de experimentos (DOE), lo que permitió optimizar la investigación agrícola al estructurar los ensayos de manera sistemática, incorporando principios como la aleatorización, la replicación y el bloqueo para minimizar el error experimental y aumentar la precisión de los resultados (Fisher, 1937). Estas técnicas siguen siendo la base de la experimentación en la agricultura moderna.

La investigación y la experimentación están profundamente interconectadas en el campo agropecuario. La investigación proporciona el marco teórico y las preguntas que la experimentación pretende responder. Por ejemplo, una investigación puede plantear la hipótesis de que una nueva variedad de maíz forrajero es más resistente a las plagas. Para confirmar esta hipótesis, se diseña un experimento en el que la variedad de alfalfa es plantada en diferentes condiciones controladas, comparando los resultados con variedades convencionales; este tipo de interacciones permite que la investigación teórica y la práctica experimental se complementen mutuamente.

El valor o la utilidad de la experimentación es que proporciona datos concretos que pueden ser replicados, algo que no es posible únicamente con observaciones teóricas o investigaciones de campo sin control experimental. Esto es especialmente importante en la producción agropecuaria, donde las variables o razas zootécnicas que afectan la producción son múltiples y complejas, como el clima, la calidad del suelo, las técnicas de cultivo, y las plagas o enfermedades. La experimentación rigurosa es la única manera de obtener datos fiables y válidos para mejorar la toma de decisiones (Montgomery, 2019).

Por otro lado, a lo largo del siglo XX y XXI, la tecnología ha jugado un papel crucial en el avance de la experimentación agropecuaria. Los avances en genética, biotecnología, sensores remotos y análisis de datos han permitido a los investigadores diseñar experimentos mucho más precisos y obtener un mayor control sobre las variables que influyen en los resultados.

Por ejemplo, el uso de tecnologías de sensores de suelo permite monitorizar las condiciones de humedad y nutrientes, condición y tendencia de pastizales con un nivel de precisión sin precedentes, mejorando la planificación de los experimentos y optimizando además la toma de decisiones sobre riego, fertilización, alimentación, capacidad de carga. De manera similar, la tecnología genética ha transformado la experimentación agrícola y pecuaria al permitir la modificación de plantas y animales para mejorar su rendimiento y resistencia a factores externos como plagas, enfermedades y condiciones ambientales extremas. Además, el uso de modelos estadísticos avanzados y algoritmos de aprendizaje automático ha revolucionado la manera en que se analizan los datos experimentales. Hoy en día, los investigadores pueden procesar grandes cantidades de datos obtenidos de ensayos de campo, lo que permite un análisis más profundo y la optimización de los recursos (Borlaug, 2007).

Uno de los aspectos más relevantes de la investigación y experimentación agropecuaria es su contribución a la sostenibilidad. En un mundo en el que la población sigue creciendo, la demanda de alimentos aumenta mientras que los recursos como el agua y la tierra agrícola son limitados. La investigación agropecuaria busca constantemente nuevas formas de aumentar la productividad agrícola sin comprometer los recursos naturales a largo plazo.

La agricultura de precisión, basada en la experimentación científica, ha permitido optimizar el uso de insumos como fertilizantes y agua, reduciendo el impacto ambiental. Además, el desarrollo de nuevas variedades de cultivos más resistentes al estrés climático y a las plagas ha sido clave para mejorar la seguridad alimentaria, especialmente en las regiones más vulnerables a los efectos del cambio climático.

Un claro ejemplo es la Revolución Verde, liderada por Norman Borlaug, quien, mediante experimentación e investigación genética, desarrolló variedades de trigo de alto rendimiento, lo que contribuyó significativamente a reducir el hambre en muchas partes del mundo. Este avance fue el resultado directo de décadas de investigación científica combinada con experimentación en diferentes contextos agrícolas (Borlaug, 2007).

La investigación y la experimentación agropecuaria han sido y continúan siendo fundamentales para el progreso agrícola. Mientras que la investigación proporciona el conocimiento teórico necesario para entender los complejos fenómenos biológicos y ambientales que afectan la agricultura, la experimentación permite poner a prueba ese conocimiento en condiciones controladas, obteniendo datos válidos y replicables. Juntas, estas disciplinas han permitido a la humanidad enfrentar con éxito desafíos críticos como la productividad agrícola, la sostenibilidad y la seguridad alimentaria.

A medida que las tecnologías avanzan y las demandas sobre los sistemas agrícolas aumentan, la importancia de continuar invirtiendo en investigación y experimentación será más crucial que nunca. La agricultura debe adaptarse rápidamente a un entorno cambiante, y solo mediante el conocimiento científico podremos garantizar un futuro sostenible y seguro para las próximas generaciones.

En la práctica podemos manifestar que la investigación es una búsqueda permanente de la verdad usando métodos, objetivos adecuados y precisos; por su parte la experimentación es un método de la investigación que permite realizar operaciones y prácticas destinadas a demostrar, comprobar o descubrir fenómenos puntuales. La experimentación agropecuaria en particular usa pruebas, ensayos, observaciones, análisis o estudios prácticos de lo que le interesa al productor para resolver sus problemas más sentidos dentro de sus sistemas de producción. Se considera un experimento cuando se comprueba con la práctica una hipótesis formulada, por ejemplo, un ensayo destinado a medir el rendimiento de seis variedades de alfalfa. En cambio, se considera una investigación cuando se estudia la causa y el efecto con varios experimentos en muchas locaciones y en época seca y húmeda, por ejemplo. En un experimento

solo se aprecia únicamente los efectos y es de aplicación inmediata en la zona donde fue establecido, en contraste la investigación puede ser de aplicación mediata y tiende a ser evolucionista pudiendo conducir a idear nuevas técnicas o a modificar las existentes; usualmente los dos términos se involucran y son inseparables (Reyes, 2003).

1.3. Tipos de estudios

1.3.1. Estudio observacional

Tiene como característica que sobre un proceso existente se observan dentro de los registros de información una o más variables aleatorias; por ejemplo, al analizar el carbono presente en el suelo se observarían distintos tipos de ecosistemas, diferentes cambios de uso de la tierra, y se mira el efecto que tendría esos cambios de uso o el tipo de ecosistema prístino sobre la acumulación o pérdida de carbono. Entonces la finalidad de este tipo de estudios es la exploración, descripción y confirmar hipótesis (Casanoves, 2019).

1.3.2. Estudio experimental

Un experimento está constituido por una prueba o serie de pruebas en las cuales se inducen cambios deliberados en las variables de entrada (fuentes de variación) de un proceso o sistema de manera que sea posible observar e identificar las causas de los cambios en la respuesta de salida (variables respuesta). Esta aseveración demuestra demasiada ingeniería, es más aplicativo para la industria, control de calidad o de procesos (Montgomery, 2019).

Además, un estudio experimental agropecuario en particular comprende pruebas, ensayos, observaciones análisis o estudios prácticos de todo lo que interesa a la agricultura y ganadería. Se considera un experimento, probar con la práctica una hipótesis formulada (por ejemplo, un ensayo sobre el rendimiento de seis variedades de alfalfa); por lo tanto, es considerado como un **arte** por la habilidad para ingeniar, planear y aplicar técnicas; y, **ciencia** por la aplicación del método científico (Reyes, 2003)

Un estudio experimental es una reproducción restringida de la realidad con el fin de observar los efectos de su manipulación planificada. Si nos fijamos en la

última expresión se construyen las unidades experimentales para simular a la realidad (parcelas o grupos pequeños de animales), ya que no se puede trabajar con la población sino con una muestra de ella procurando que sea representativa. Con este antecedente se pueden mirar los efectos de esa manipulación planificada (situación observada), que se convertirían en nuestros tratamientos los mismos que son decididos estratégica y técnicamente por el investigador (Di Rienzo & Casanoves, 1999)

Por ejemplo, si se analiza la voluntad (cuantía) de la población para pagar por un servicio ecosistémico, como lo hace Costa Rica que son pioneros por pago de este tipo de propuesta donde a los productores reciben un pago (definido por estudio experimental) para mantener los servicios ecosistémicos como la regulación de la calidad del aire o de la fertilidad de los suelos, programas de recreación y ecoturismo, etc. (Casanoves, 2019).

1.3.3. Estudios Mixtos

Combinan los dos conceptos, por ejemplo si se pretende estudiar 3 tipos de dietas para mejorar la producción de leche de vacas Holstein, Jersey y Brown Swiss, cada parámetro técnico productivo (producción vaca día) es observacional propiamente dicho porque no podemos manipular la raza individual que es propia de sus “genes”, en cambio si se puede establecer el comportamiento de las mismas cuando difiere el porcentaje de proteína entregado a través del balanceado, dentro de cada raza y compararlo con un testigo.

En esencia no se puede manipular a un estudio observacional, la raza no puedo cambiarla, sus estándares productivos se expresan en sus genes siempre que les proporcionemos un buen medio ambiente, por lo tanto, solo se observa su comportamiento; en cambio si a dos razas de porcinos le induzco a que consuman 3 suplementos alimenticios diferentes para medir su comportamiento productivo, entonces estamos en un experimento ya que he manipulado su realidad incluyendo los tratamientos (suplementos).

La finalidad de usar los tipos de estudio es para comprobar hipótesis, modelar y predecir usando los modelos lineales aditivos.

1.3.4. Experimentos comparativos

Son los típicos del “Diseño de Experimentos”, ya que consisten en la aplicación de los llamados tratamientos a un conjunto de unidades experimentales, siendo éstas de acuerdo a la realidad: cajas Petri, parcelas, animales u una parte de ellos, una planta, una hoja, etc. Es esta práctica se valora y compara las respuestas obtenidas desde diferentes tratamientos. Son utilizados cuando se requiere incrementar la precisión de las diferencias realizadas y/o ampliar el espacio de las “inferencias” (Casanoves, 2019)

1.4. Otros estudios

En la dinámica investigativa se pueden encontrar algunas metodologías que suelen emplearse en base a la realidad que se presente al momento de comprobar una hipótesis y producir información a través de la extracción de una “data” apegada a una realidad en el momento oportuno, es así como se cuenta con los siguientes casos:

1.4.1. Según el propósito de estudio

- **Estudio teórico:** Cuando se genera conocimiento sin colocarlo en práctica; recopila datos y genera conceptos; puede considerarse una revisión bibliográfica científica en las plataformas científica sobre un problema específico como la ascitis en aves de carne, por ejemplo.
- **Estudio Aplicado:** Su finalidad es resolver un problema, busca y consolidar el conocimiento para su aplicación, se divide en: Aplicación tecnológica, cuando se aplica sistemáticamente los conocimientos científicos o habilidades técnicas para desarrollar soluciones innovadoras (Ponce, Mario;, 2014), por ejemplo, para aplicar el ordeño mecánico en las granjas de producción lechera, debieron realizar pruebas específicas antes de implementar esta tecnología. Aplicada científica que se caracteriza por buscar convertir el conocimiento teórico en práctico y útil (Comunicación institucional, 2020), teóricamente se conoce que en los 4 estómagos de los rumiantes se produce una fermentación anaeróbica del alimento que consumen, pero si se entrega grano partido como el maíz y fermentado en fundas herméticas, se

facilitaría el metabolismo con la consecuencia en el mejoramiento de la conversión alimenticia del animal en cuestión.

1.4.2. Según la profundización del estudio

- **Estudio exploratorio:** Se caracteriza por analizar información específica NO estudiada, permite un primer acercamiento, para luego hacer más detallada; por ejemplo, un sondeo para explorar la predisposición de una comunidad para la adopción de tecnologías nuevas en la crianza de animales.
- **Estudio descriptivo:** Realiza un informe detallado sobre un determinado fenómeno, buscando información clara sobre el objeto, cuando se ha adoptado una tecnología en varias unidades productivas y se requiere describir el grado de adopción auscultando variables sociales y productivas, por ejemplo; en este tipo de estudios es muy frecuente el uso de la estadística descriptiva.
- **Estudio explicativo:** Establece relación causa – efecto, e intenta establecer las consecuencias de un fenómeno; por ejemplo, si se desea mejorar un pastizal que no presta las condiciones óptimas para la alimentación ganadera y se pretende introducir especies nuevas que podrían aportar mayor rendimiento y calidad que se refleja en el incremento los índices productivos de los animales, estos últimos son consecuencia de las estrategias que se han aplicado.

1.4.3. Según el tipo de “data”.

- **Estudio cualitativo:** Rescata las individualidades, las particularidades del campo social con base en conocimientos y saberes de los individuos, puede considerarse dentro de este tipo de estudio a un análisis de las preferencias de un determinado grupo humano a los productos que se realizan con valor agregado de la producción agropecuaria como lácteos, mermeladas, etc., y como se comercializaría en función de su estrato social.
- **Estudio cuantitativo:** Emplea un método estadístico y matemático. Utiliza una gran cantidad de variables que provienen de distintos datos, en la praxis es más relacionado con los experimentos que se plantean

con modelo lineales aditivos en nuestra área como establecer el comportamiento biológico de cerdos frente a 4 dietas alimenticias; donde se miden variables como consumo de alimento, peso vivo, conversión alimenticia, etc.

1.4.4. Según su temporalidad

- **Estudio longitudinal:** Observa a un individuo o evento durante un tiempo para identificar los cambios en las variables analizadas. En este caso podríamos citar una caracterización estática de agroecosistemas de producción agropecuaria de pequeña escala, donde se conoce una línea base de las variables productivas, sociales y ecológicas de la zona de estudio.
- **Estudio transversal:** Compara característica de varios individuos o eventos en forma secuencial y en momentos específicos. En el mismo caso de la caracterización de agroecosistemas, cuando se ha realizado algunas de ellas en la dinámica del tiempo, donde inclusive se pueden generar modelos lineales de comportamiento de las diferentes variables analizadas y predecir su relación futura.

1.4.5. Según el tipo de inferencia

- **Estudio deductivo:** Estudia la realidad y verifica o refuta una hipótesis, saca conclusiones en base a una premisa planteada (verdades), verifica proposiciones.
- **Estudio inductivo:** Genera conocimiento con información específicos para crear nuevas teorías. Analiza elementos y fenómenos a partir de la observación.

1.5. El método científico para experimentos agropecuarios

El método científico es la base sobre la cual se construyen los conocimientos y descubrimientos de la ciencia moderna. Consiste en una serie de etapas sistemáticas diseñadas para investigar fenómenos, formular hipótesis, experimentar, y alcanzar conclusiones objetivas y replicables. Podemos explorar los orígenes y evolución del método científico, sus principales etapas, y su relevancia en la investigación actual, destacando cómo esta metodología ha

permitido el avance del conocimiento de forma fiable y rigurosa para el desarrollo de los pueblos.

1.5.1. Orígenes y su desarrollo histórico

El método científico tal como lo conocemos hoy en día es el resultado de una evolución gradual de pensamiento crítico y sistematización que comenzó con los filósofos de la antigua Grecia. Aunque personajes como Aristóteles (384-322 A.C.) establecieron principios de observación y lógica en la indagación científica, su enfoque estaba limitado por la falta de experimentación empírica (Popper, 2002).

El verdadero precursor del método científico fue Roger Bacon en el siglo XIII, quien en su obra *Opus Majus* (1267) defendió el uso de la experimentación para probar ideas, combinando la observación con el razonamiento deductivo. Sin embargo, fue en el siglo XVII cuando figuras como Francis Bacon y Galileo Galilei revolucionaron el proceso investigativo. Francis Bacon, en su libro *Novum Organum*, proponía la inducción sistemática y el uso de la experimentación como una manera de superar los errores de la percepción humana, dando un fuerte impulso al empirismo en la ciencia. Galileo, por otro lado, aplicó estos principios al estudio de la física, estableciendo la importancia de la observación cuantitativa y el experimento como pruebas necesarias para validar hipótesis (Bacon, 1620).

Ya en nuestros tiempos el método científico consiste en un protocolo sistemático y ordenado; en donde se debe conocer previamente, los materiales, instrumentos, así como las técnicas necesarias, que permitan alcanzar un objetivo o probar una hipótesis, procurando flexibilidad, objetivismo y crítica; partiendo de una relativa realidad para resolver problemas (Moscoso G., Moreno, Moscoso M., & Armijos, 2022).

1.5.2. Las etapas clásicas del método científico

El método científico, desde su consolidación, ha seguido un conjunto de etapas básicas que han sido formalizadas a lo largo de los siglos.

- **Observación:** En esta etapa, el investigador observa un fenómeno o conjunto de fenómenos, estableciendo una base empírica para formular preguntas.
- **Formulación de la Hipótesis:** Basándose en las observaciones, se plantea una hipótesis, que es una proposición que intenta explicar el fenómeno observado.
- **Experimentación:** La hipótesis es puesta a prueba mediante experimentos controlados que permitan verificar o refutar su validez.
- **Análisis de Resultados:** Los datos obtenidos se analizan estadística y lógicamente, evaluando si respaldan o contradicen la hipótesis.
- **Conclusión:** Se establece si la hipótesis inicial es válida o si debe rechazarse. Esta etapa también puede llevar a nuevas preguntas y, por ende, a nuevas investigaciones.
- **Comunicación y Revisión:** Los resultados y conclusiones se comparten con la comunidad científica, lo cual permite la replicación del estudio y una revisión crítica de los hallazgos (Chalmers, 2013).

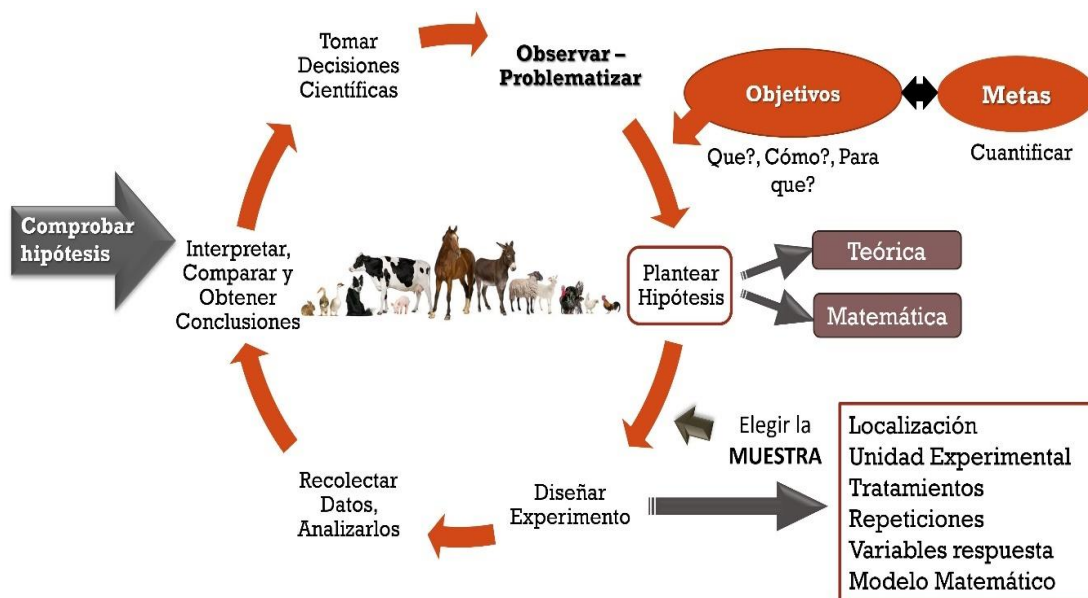
La importancia de estas etapas radica en que permiten una estructura replicable y verificable que facilita la construcción del conocimiento científico. La publicación y revisión de los resultados son fundamentales para el progreso científico, ya que otros investigadores pueden validar, refutar o extender los descubrimientos.

El valor del método científico reside en su capacidad de adaptación a nuevas tecnologías y descubrimientos, lo que le permite seguir siendo relevante en un mundo en constante cambio. Gracias a este enfoque sistemático, podemos avanzar en la comprensión del universo, resolver problemas complejos y construir un conocimiento fiable y duradero que beneficie a la humanidad.

1.5.3. Etapas en la experimentación agropecuaria

En el caso específico de las ciencias agropecuarias tradicionalmente se cuenta con un protocolo muy detallado que se ilustra en la figura 1.1.

Figura 1.1. Esquema del método científico para la comprobación de hipótesis en la Ciencias Agropecuarias



Elaborado y adaptado por: Moscoso, 2025

El protocolo para realizar un experimento en el área agropecuaria con uso del método científico pasa por varias etapas fundamentales que deben formalizarse en una planificación o lo que comúnmente denominamos “anteproyecto” con características definidas que se pueden sintetizar en las siguientes:

- La simplicidad, en la selección de tratamientos, el modelo lineal aditivo más apegado a la realidad, y en estricta relación con los objetivos y las hipótesis planteados.
- Un cierto grado de precisión, con la finalidad que el experimento pueda ser capaz de medir las reales diferencias entre los tratamientos que reflejen las respuestas deseadas por el investigador. Esto indica un diseño apropiado y un número suficiente de repeticiones por tratamiento.
- Anulación del error sistémico mediante la planeación óptima de tal manera que las unidades experimentales que reciban un tratamiento no difieran sistemáticamente de aquellas que reciben un diferente tratamiento, esto permite que se obtenga una estimación imparcial de cada resultado generado, sin establecer un sesgo a favor o en contra de ningún tratamiento.
- Las conclusiones deben presentar un rango de validez lo más amplio como sea posible. Un experimento que se replica en tiempo y el espacio puede

subir sustancialmente el rango de validez de sus conclusiones, aprovechando debidamente el ejercicio experimental. La utilización de un conjunto factorial de tratamientos puede ser otro medio para incrementar el indicador. En un experimento factorial, los efectos de un factor son evaluados por interacción con los niveles de un segundo factor de estudio.

- Todo experimento tiene un grado de incertidumbre en referencia a la validez de las conclusiones, por esta razón debe ser concebido de tal modo que resulte posible calcular la probabilidad de obtener los resultados de sus variables medidas y observadas, solamente debido únicamente al azar (Little & Hills, 1991).

Tomando como premisa que, en este tipo de procedimientos es importante indicar que la acción experimental se desarrolla dentro de un agroecosistema donde confluyen componentes de producción agrícola, pecuario e inclusive agroindustrial; y el método científico es aplicado para auscultar y solucionar sus problemas, planteando estrategias técnicas que permiten comprobar una hipótesis planteada.

1.6. Especificación de problema

Es quizá la génesis de un experimento, debemos procurar ser claros, precisos y concisos con lo que estamos tratando. Es un paso decisivo de todo el protocolo, debido a que, si se considera que el método científico ayuda a resolver problemas de la dinámica productiva, es en esta tarea que debemos utilizar todas las estrategias, así como el conocimiento de la realidad circundante para ubicar al “problema”.

Puede considerarse una interrogante o situación no resuelta que se busca comprender, entender, describir, explicar o transformar mediante un proceso sistemático. (Hernández Sampieri, R., Fernández Collado, C., & Baptista Lucio, P, 2014) ; puede construirse desde una interrogante que sea clara, precisa y delimitada que guíe a la indagación científica; con el fin de garantizar que los resultados tengan impacto teórico, práctico o social. Para que un problema sea considerado adecuado en un contexto de investigación, debe cumplir con las siguientes características:

- Tener relevancia, que aborde una necesidad importante o un vacío de conocimiento. Según (Creswell, 2014) "los problemas bien definidos permiten generar nuevas aportaciones significativas al campo de estudio".
- Clara delimitación clara, especificando el alcance del problema para evitar ambigüedades.
- Viabilidad, que pueda ser abordado con los recursos disponibles, incluyendo tiempo, presupuesto y acceso a datos.

En la identificación del problema, el investigador debe reunir ciertas características para observar una determinada realidad; de hecho, *la observación* juega un papel decisivo en todo el proceso, puesto que, si se cumple debidamente con esta actividad, el problema se define con eficiencia y se puede cumplir con todo el protocolo para resolverlo.

La formulación adecuada del problema es la base de toda investigación, ya que orienta los objetivos, el diseño metodológico, la toma de la data y su procesamiento, así como la interpretación de resultados; la fundamentación lógica de la investigación está en esta fase considerándose como la mitad del trabajo investigativo resuelto (Kerlinger, F. N., & Lee, H. B., 2002).

Como se analizó anteriormente, se logra identificar el problema siempre y cuando el investigador sea observador, talento que permite acercarse a la realidad tal cual como se manifiesta, problematizando, definiendo los puntos críticos (Anguera, M. T., Blanco-Villaseñor, A., Losada, J. L., & Sánchez-Algarra, P., 2014).

Comúnmente confundimos los términos entre “ver” y “observar”, el observador profundiza y analiza en base a sus conocimientos del área de competencia, en cambio quienes solo ven superficialmente un fenómeno no establecen la diferencia entre lo normal y lo anormal. Para que se desarrolle este instintivo concepto, el observador una vez que se encuentra en el entorno por ejemplo dentro de una granja productiva, analiza registros, parámetros técnicos, compara con indicadores que la comunidad científica los ha popularizado y se fija en las fortalezas de debilidades del sistema.

Un ejemplo práctico es cuando visitamos una granja de cobayos y analizamos su composición poblacional, que debe reunir características como pozas de 2 m² donde se albergan 1 macho y 10 hembras, cuenta con crías destetadas, levante y engorde, etc., analiza indicadores como número de crías al nacimiento, edad y pesos al destete, período y peso de saca (desde el nacimiento hasta la venta de engordados), todo registrado en la data que obligatoriamente deben tener las granjas, caso contrario deberá evaluar haciendo un muestreo aleatorio, tomando datos y analizándolos.

Luego de esta actividad, entonces con juicio de valor podrá indicar técnicamente si existe o no un problema que puede ser productivo o sanitario, el mismo indudablemente se reflejará en los indicadores económicos. La importancia de la observación en la identificación del problema se debe generalmente a los siguientes aspectos:

- Genera preguntas relevantes: Al observar fenómenos o situaciones, se pueden identificar vacíos de conocimiento, debilidades o inconsistencias que justifican la necesidad de investigar.
- Conecta la teoría con la práctica: La observación ayuda a contextualizar teorías abstractas, adaptándolas a problemas concretos.
- Permite comprender fenómenos dinámicos: Es útil para explorar fenómenos complejos, como interacciones humanas, cambios sociales o procesos biológicos.
- Facilita la identificación de patrones: Al observar repetidamente, se pueden identificar regularidades que den lugar a interrogantes importantes.

La observación permite al investigador acercarse a la realidad tal como se manifiesta, obteniendo datos iniciales para identificar áreas problemáticas; sin embargo aunque es una técnica poderosa, puede estar influida por la subjetividad del investigador o por sesgos derivados del contexto; por esta razón, es esencial complementar la observación con métodos rigurosos para confirmar los hallazgos (Hernández Sampieri, R., Fernández Collado, C., & Baptista Lucio, P, 2014); en el caso de la producción agropecuaria, de productos tradicionales y

animales de granja existen estándares para su manejo, que facilitan la comprensión del entorno.

Entre los elementos que se toman en cuenta pueden incluir la contextualización tanto teórica (indicadores de manejo) como práctica (experiencia de expertos); de este elemento se debe construir posteriormente la justificación, hipótesis y objetivos.

Supongamos que es evidente el descenso de la productividad en una finca de vacas lecheras, para plantear el problema debemos considerar la siguiente interrogante:

"¿Qué aspectos están afectando producción vaca / día en hembras Holstein dentro de la finca? Para llegar a esta interrogante es obvio que el investigador debió fijarse en los registros de producción, así como en los estándares, por ejemplo, se conoce que una vaca Holstein alta cruza en la sierra ecuatoriana produce aproximadamente 25 litros / día; y puede afectarse por estrés, alimentación no balanceada (contenido proteico), morbilidad, longitud de la lactancia, etc.; por lo tanto nos obligamos a ir descartando cada uno de estos aspectos para llegar a establecer una verdadera hipótesis que facilitará los objetivos mediante la aplicación de los tratamientos o estrategias que podrían permitir resolver ese problema.

1.6.1. Objetivos y metas

En la mayoría de las investigaciones, antes de plantear los objetivos se precisa especificar primero el problema a resolver, aunque en algunas se gestan desde los objetivos directamente (Moscoso G., Moreno, Moscoso M., & Armijos, 2022). En el caso de la dinámica dentro de los agroecosistemas agropecuarios, siempre se utiliza el método científico para resolver los problemas que se presentan continuamente, por lo tanto, identificado el "problema" nos conduce a resolverlo con el uso de estrategias que puedan desarrollarse sin contratiempos dentro del ambiente en donde se está produciendo.

Los objetivos de investigación son enunciados específicos que describen lo que se pretende lograr con el estudio y delimitan su alcance. Pretenden guiar el proceso de la investigación proporcionando un protocolo para recolectar datos

de las variables respuesta y su posterior análisis e interpretación (Hernández Sampieri, R., Fernández Collado, C., & Baptista Lucio, P, 2014). Los objetivos establecen lo que pretendemos conocer o analizar y se formulan en forma clara, precisa y alcanzable.

En el caso de la producción ganadera, si se establece una debilidad en algún proceso, hay que probar diferentes estrategias para poder resolverlo, dando origen al planteamiento de los **objetivos** que deben estructurarse con 3 componentes a manera de cuestionamiento:

- ¿Qué?
- ¿Cómo?
- ¿Para qué?

Supongamos a un técnico zootecnista, el mismo que en base a los registros productivos de una granja en donde se cría ganado porcino de raza Duroc, encuentra un peso promedio de 67,5 kilos a la edad de 6 meses de crecimiento o ceba, cuando el indicador es de 90 kilos; entonces se establece un problema con un desfase de 22,5 kilos menor al peso ideal. Inmediatamente se hace un análisis de su salud, condiciones de crianza (densidad por metro cuadrado, por ejemplo), alimentación, etc. La problemática se centra en el alimento que debe contener entre 15 a 18% de proteína bruta (Murcia, N. Savio, M., Cora, F., & Beneitez, A., 2021) y apenas están entregando un “pienso” con 13%; por lo tanto se pretende resolver el problema con la elaboración de 3 dietas alimenticias a base de 16, 17 y 18% de proteína, para medir el comportamiento del engorde de estos cerdos en los parámetros productivos como: pesos, ganancias de peso, consumo de alimento, conversión alimenticia, etc.

Planteamos el objetivo usando inicialmente un verbo, el mismo que puede identificarse metodológicamente con la ayuda de la taxonomía de Bloom; en este caso elegiremos el verbo “**evaluar**”. Si tomamos como base el *¿qué?*: se indicaría la acción es decir evaluaremos el comportamiento biológico de los cerdos; *¿cómo?* con el suministro de las 3 dietas proteicas; y, *¿para qué?* para cumplir o acercarnos al estándar (90 kilos de peso) en la misión de resolver el problema que se indicó anteriormente.

La redacción del objetivo puede quedar de la siguiente manera: *“Evaluar el comportamiento biológico de cerdos Duroc mediante 3 dietas proteicas (16, 17 y 18%), para alcanzar el indicador de la raza”*.

En el caso que el protocolo exija una **meta**, el mismo objetivo estructurado podría convertirse en una meta cuando su equivalente es cuantificado, por ejemplo: *“Valorar los parámetros productivos en cerdos Duroc mediante 3 dietas proteicas (16, 17 y 18%), para procurar un peso de 85 kilos a los 180 días de ceba”*. La redacción expresa *qué se pretende realizar*, con la alimentación de cerdos valorando las 3 dietas hasta en lo posible alcanzar los estándares de la raza.

1.6.2. Hipótesis

Una vez que en secuencia se define el problema, se han planteado los objetivos y las metas, debe construirse la hipótesis que consiste en una respuesta tentativa del investigador para solucionar el problema determinado con anterioridad, el mismo que es sujeto del experimento (Moscoso G., Moreno, Moscoso M., & Armijos, 2022). En nuestro caso ¿Cómo se debe solucionar la merma del peso en los cerdos Duroc?, entonces, gracias al conocimiento del técnico se plantea la elaboración de 3 diferentes balanceados alimenticios que contengan 16, 17 y 18% de proteína en la dieta; que a nuestro juicio pudiera corregir la merma de peso y probablemente alcanzar el estándar requerido (85 kilos en 100 días); para el efecto se pueden generar 2 tipos de hipótesis: una de carácter teórico y otra desde el punto de vista matemático, tomando en consideración la modelación matemática.

1.6.3. Hipótesis nula y alternativa

La hipótesis nula (H_0) y la hipótesis alternativa (H_1) son conceptos fundamentales en la estadística inferencial y en la metodología de la investigación. Las dos forman son parte fundamental del proceso de contrastación de hipótesis, donde se busca determinar si existe suficiente evidencia para rechazar la hipótesis nula en favor de la alternativa.

- *Hipótesis Nula (H_0)*

La hipótesis nula representa la afirmación que se somete a prueba y que generalmente expresa la ausencia de efecto (no efectividad de los tratamientos o estrategias usadas para resolver el problema), es decir no hay relación o diferencia significativa en un estudio. En términos estadísticos, la hipótesis nula asume que cualquier variación observada en los datos se debe al azar o a fluctuaciones naturales en la muestra (Montgomery, 2019).

Si expresamos en términos de hipótesis nula nuestro ejemplo, la **hipótesis teórica** puede expresarse de la siguiente manera: “*Las 3 dietas alimenticias suministradas a los cerdos Duroc en ceba, no mejoraran el peso de saca*”. Esta mención hace entender que los 3 niveles de proteína no solucionarían el problema ya que el indicador se mantendría alrededor de 67,5 kilos como al inicio. Desde el punto de vista estadístico, la **hipótesis matemática** suele expresarse en términos de igualdad de la siguiente manera:

$$H_0: \mu_1 = \mu_2 = \mu_3$$

Esta expresión indica que las 3 medias de la variable respuesta (peso de saca de los cerdos) encontradas son iguales desde el punto de vista estadístico, que se comprobaría científicamente luego con los modelos matemáticos que elijamos para el cálculo de análisis de varianza y la separación de medias, como se describirán en los capítulos posteriores.

- **Hipótesis Alternativa (H_1 o H_a)**

La hipótesis alternativa es la contraparte de la hipótesis nula y representa la afirmación que el investigador espera demostrar; es decir previamente se afirma que las acciones, estrategias o tratamientos aplicadas para resolver el problema, van a reflejar un efecto positivo en las variables respuesta. Indica la existencia de una diferencia significativa o un efecto real (Walpole, R. Myers, R. & Myers, S., 2012).

En nuestro ejemplo, la **hipótesis teórica** podría indicar lo siguiente: “*La utilización de 3 dietas alimenticias (16, 17 y 18% de proteína) influirá en el mejoramiento del proceso biológico de ceba para crianza de cerdos Duroc*”

elevando el peso de saca"; es una afirmación (hipótesis alternativa) que igualmente debe ser comprobada científicamente.

Si se refiere a una **hipótesis matemática** debemos concentrarnos en la dinámica de un modelo matemático, es decir que:

$$H_1: \mu_1 \neq \mu_2 \neq \mu_3$$

Indicaría que existe una diferencia significativa sin establecer una dirección; ya que se la conocería con otra prueba conocida como separación de medias o prueba de significancia.

1.6.4. Importancia de la Prueba de Hipótesis

El procedimiento estadístico busca recolectar evidencia en la muestra o las unidades experimentales, para decidir si se rechaza la hipótesis nula en favor de la alternativa. Sin embargo, no se "acepta" a la hipótesis nula (H_0); solo se concluye si hay suficiente evidencia para rechazarla o no (Fisher, 1937).

En la tabla 1.1. se explica la definición esquemática de las hipótesis que se pueden utilizar para usarlas en el método científico que explican previamente lo que podría ocurrir con las variables respuesta, cuando se asignan los tratamientos que son conducidos a resolver un problema.

Tabla 1.1. Características principales de las hipótesis nula y alternativa para comprobarlas con los diferentes modelos lineales aditivos.

Característica	Hipótesis Nula (H_0)	Hipótesis Alternativa (H_1)
Definición	No hay efecto o diferencia significativa	Hay efecto o diferencia significativa
Expresión matemática	$H_0: \mu_1 = \mu_2 = \mu_3$	$H_1: \mu_1 \neq \mu_2 \neq \mu_3$
Conclusión posible	No se rechaza (se asume como válida) o se rechaza	Se apoya si se rechaza H_0
Relación con el error	Se comete error tipo I si se rechaza erróneamente	Se comete error tipo II si no se rechaza H_0 cuando es falsa

Fuente: Adaptado de (Walpole, R. Myers, R. & Myers, S., 2012).

1.6.5. Muestra o unidades experimentales

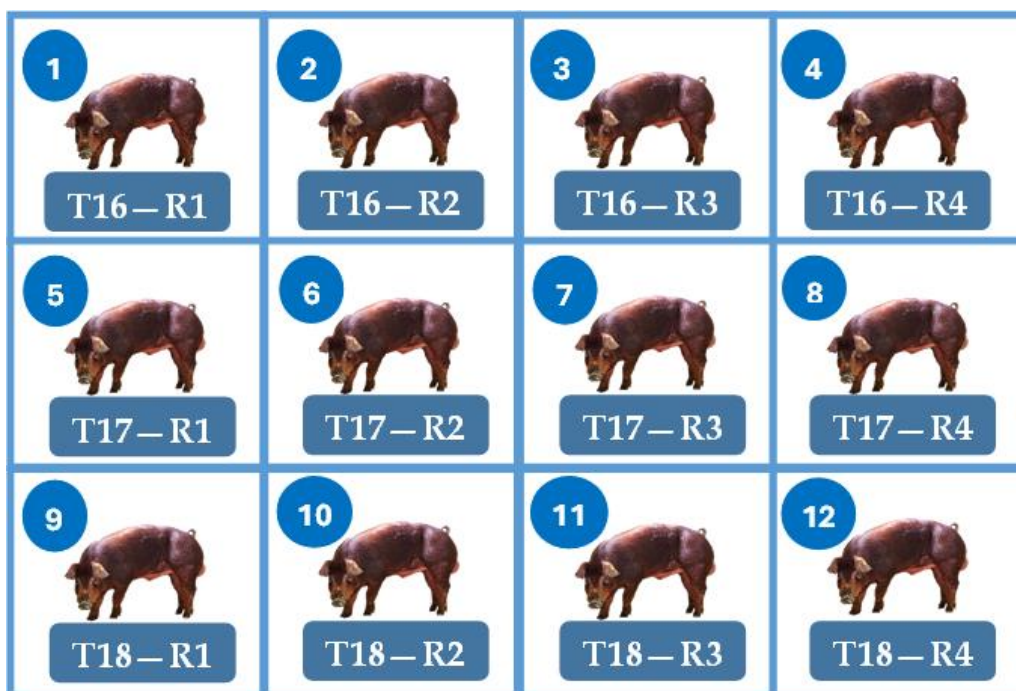
Para concretar la investigación una vez definido el problema, los objetivos, las hipótesis, y los tratamientos, a estos últimos se les debe evaluar en las unidades experimentales que son en esencia donde se expresan los mismos y sale la “data” midiendo las variables respuesta a cada estrategia que pensamos pueden aportar para mejorar los rendimientos o solucionar la problemática encontrada.

Unidad Experimental, es la mínima unidad que recibe un único tratamiento, en el caso del campo pecuario por ejemplo si se ensayan diferentes estrategias de fertilización en el pasto azul, sobre las parcelas de 4 por 8 (32 m²) entonces cada parcela es una unidad experimental (Casanoves, 2019)

En nuestro caso retomando el ejemplo de los cerdos Duroc, cada uno de ellos se pueden convertir en una unidad experimental, la misma que para poder sujetarse al modelo lineal aditivo requiere de varias de ellas que recibirían el mismo tratamiento, es decir que si se consideran las 3 dietas de ceba cada una de ellas debe contar con al menos 4 ejemplares consideradas como repeticiones, entonces si la dieta del 16% de proteína requiere 4 cerdos, igual con las demás (17% y 18%) es decir se contarán con un total mínimo de 12 cerdos recién destetados (12 unidades experimentales).

En los casos que se necesite mejorar la precisión del experimento se puede utilizar un mayor número de animales, jugando con el tamaño de la unidad experimental (TUE), es decir con el número de animales por tratamiento y por repetición. Esto se explicaría que si se pretende instalar un ensayo se requieren formar 12 unidades experimentales y si en cada una de ellas solo se elige 1 animal, se contarían con 12 cerdos; pero si por algún motivo se muere uno de ellos automáticamente se pierde una unidad experimental, en cambio si el TUE sube a 2 se usarán 24 cerdos lo que corresponde a 2 en cada unidad experimental, y si un animal se pierde, el dato del segundo validaría el diseño experimental y se podría calcular el análisis de varianza (ADEVA), para comprobar la hipótesis planteada con el modelo matemático que corresponda.

Figura 1.2. Esquema de la distribución de las 12 unidades experimentales para un experimento en cerdos Duroc con 3 tratamientos alimenticios, 4 repeticiones y un tamaño de la unidad experimental de 1 (total 12 animales)

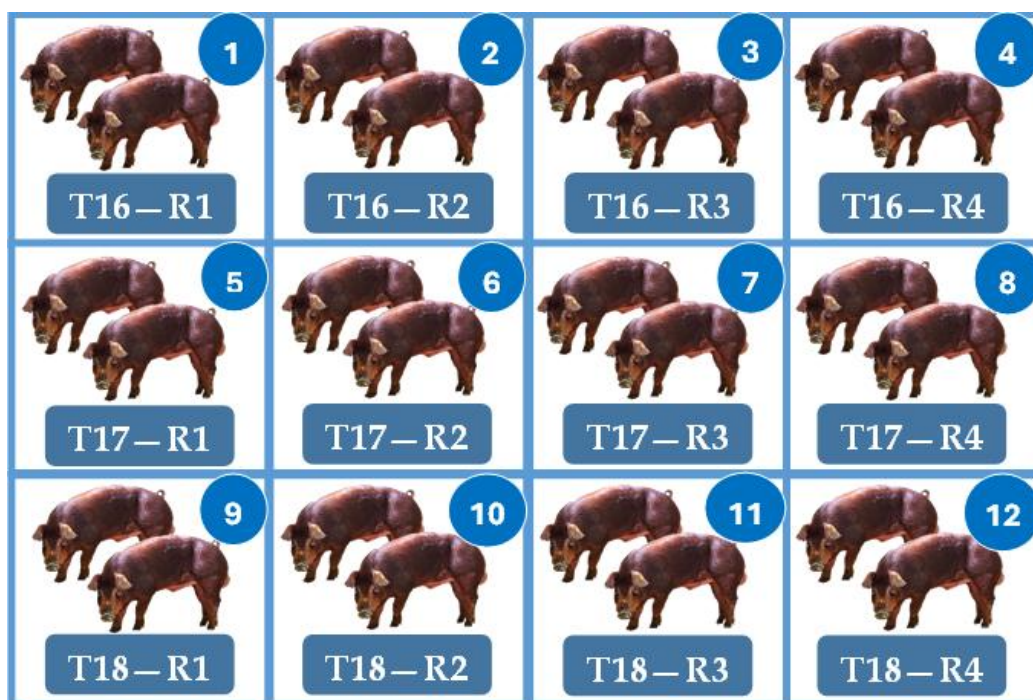


Elaborado por: Moscoso, 2025

Si se analiza la figura 1.2., podemos indicar que son 12 unidades experimentales que son representadas en cada cuadrado en donde se ubica un cerdo, cada una de ellas está codificada de la siguiente manera: T es el tratamiento 16 es la dieta con 16% de proteína, y R representa a la repetición; es decir que T18R3 representaría al cerdo que recibirá una dieta con 18% de proteína y es la tercera repetición. Al tratarse de una TUE de 1 solo lechón, son 12 animales en 12 unidades experimentales.

En cambio, en la figura 1.3., se esquematizan las mismas 12 unidades experimentales, pero con un TUE de 2 lechones, lo que representaría un total de 24 animales distribuidos uniformemente 2 en cada encierro (cuadrado).

Figura 1.3. Esquema de la distribución de las 12 unidades experimentales para un experimento en cerdos Duroc con 3 tratamientos alimenticios, 4 repeticiones y un tamaño de la unidad experimental de 2 (total 24 animales)

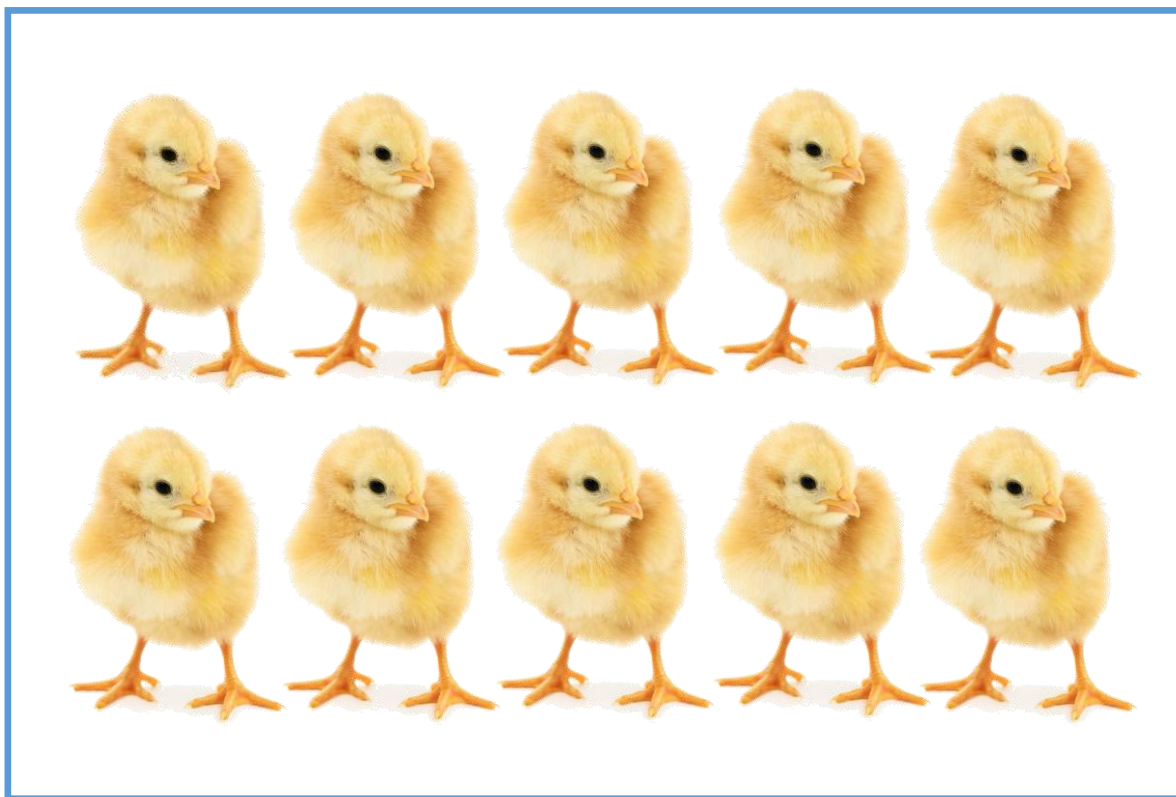


Elaborado por: Moscoso, 2025

Para que se cumpla con la extracción de información, para evaluar cada tratamiento se requiere de un dato de cada unidad experimental; por ejemplo, si se trata del último caso (Figura 1.3.) de cada unidad que forman 2 cerdos se toman los pesos de los mismos y se promedia para que se conformen los resultados experimentales de las variables como: peso, ganancia de peso, consumo de alimento, conversión alimenticia, etc.

En el caso de animales menores se suelen realizar unidades experimentales con un mayor tamaño, 20 pollos bb, 5 cuyes destetados, en este caso se procedería así a formarlos de la siguiente manera: Supongamos que la unidad experimental de un ensayo en aves de carne (pollos broiler), se conforme por celdas de 22 (4 m²), en su interior se ubican 10 entonces el tamaño de la unidad experimental es 10 (TUE = 10), probablemente cada encierro o jaula podría ser:

Figura 1.4. Conformación de una unidad experimental con un tamaño de 10 pollos broiler bb, en jaulas de evaluación.



Elaborado por: Moscoso, 2025

Esta práctica nos permite asegurar la data cuando se evalúen las variables respuesta; ya que si se pierden (muerte o pérdida) dentro de cada unidad experimental por ejemplo los valores semanales se evalúan con el peso de los animales que restan, reportando un promedio de cada jaula.

La técnica del diseño experimental puede exigir además el uso de las submuestras que es una pseudo repetición, y esto puede convertirse en el principal motivo para rechazar un trabajo de investigación; por lo tanto, de cada unidad experimental se debe tomar y registrar solo un dato que se debe proceder de forma técnica (promedio). Por ejemplo, si se evalúa el ataque de hongos en cajas de piña y se construye la unidad experimental con 6 piñas en cada caja y un total de 3 cajas, entonces estará conformada por un grupo de 18 piñas, por lo tanto, se evalúa una variable como la incidencia de hongos y el dato que saldrá de cada unidad experimental sería 0 de 18, 2 de 18, o 5 de 18, 8 de 18 o 18 de

18 como máximo, transformándose en una distribución binomial, es decir el número de éxitos en 18 ensayos o 18 piñas (Casanoves, 2019).

El problema de la utilización de pseudo réplicas es que los datos obtenidos de estas no son “independientes”; por el contrario, el protocolo para el procesamiento estadístico, uno de los supuestos es la independencia de ellos, consecuentemente los datos deben provenir de unidades experimentales totalmente independientes.

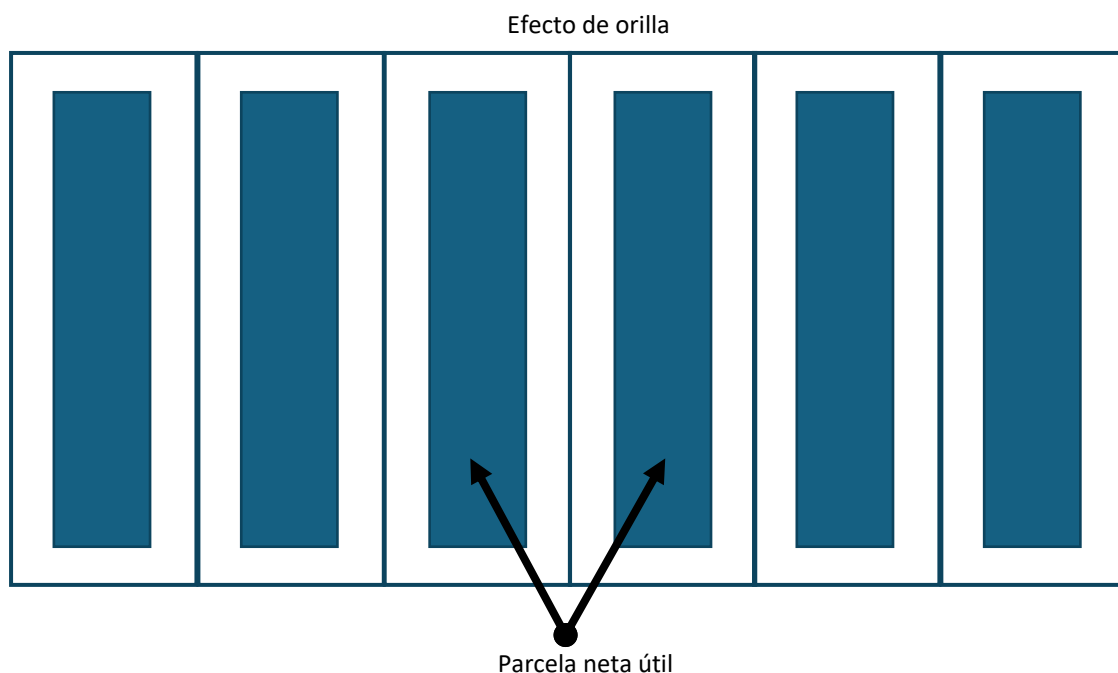
Se evita estos problemas si se usa un tamaño de unidad experimental manejable y sobre todo que sean lo más homogéneas posibles, animales de la misma edad, raza o línea genética, peso uniforme, etc.

En el caso de los experimentos agrícolas o sobre pasturas y forrajes; la unidad experimental es la parcela que debe ser delineada según los requerimientos del investigador. El tamaño está determinado por el área de terreno, los recursos económicos disponibles, calidad del suelo (heterogeneidad), el objetivo del experimento, los métodos de manejo (control de malezas, riegos, etc), así como la especie cultivada (Little & Hills, 1991). Es mejor sacrificar el tamaño de las parcelas y subir el número de repeticiones por tratamiento, prefiriendo las de forma rectangular (largas y angostas), ya que las cuadradas usan mayor perímetro y se dificulta manejarlas; si se conoce la variabilidad del suelo (análisis físico – químico) se orienta la mayor longitud de la parcela o unidad experimental en sentido de dicha variabilidad. El tamaño adecuado influye en el grado de precisión del experimento, así como en la mejor estimación de los efectos de los tratamientos (Reyes, 2003).

Para evitar que exista traslape de los tratamientos en las parcelas de experimentos con pasturas, si se trabaja con niveles de fertilización, control de plagas y enfermedades, métodos de riego etc., se procura que exista una distancia razonable entre parcelas contiguas para evitar el “efecto de orilla y competencia mutua de tratamientos” como se aprecia en la figura 1.5., se destina una distancia de 0,8 a 1 m que es razonable para evitar este fenómeno (traslape de tratamientos) y se conforma cada parcela neta útil, dentro de la cual se

tomarán las respectivas variables respuesta como altura de planta, tallos por planta, peso de campo, etc.

Figura 1.5. Esquema para conformación de unidades experimentales de investigaciones en pastizales para evitar la competencia mutua de tratamientos.



Elaborado por: Moscoso, 2025

En resumen, la unidad experimental es el material o lugar sobre el cual se aplican los tratamientos en estudio. Por ejemplo: una parcela, una maceta o grupo de macetas, un marrano o grupo de marranos, un conjunto de semillas, etc. Es característica de las unidades experimentales que muestran variación aun cuando se les aplique igual tratamiento (Reyes, 2003); por esta razón se debe procurar usar unidades que sean lo más uniformes posibles (suelo o animales homogéneos)

1.6.6. Diseño del experimento

Quizá este apartado es el más importante del método científico, ya que nos permite establecer las pautas para poder comprobar las hipótesis por métodos numéricos, en este caso definir el modelo lineal aditivo que mejor se apegue a las condiciones de las unidades experimentales y los tratamientos designados.

Para el efecto es necesario realizar algunas puntualizaciones en los conceptos necesarios para cumplir con el diseño de un experimento.

- Tratamiento: Son las estrategias, procedimientos o variables cuyos efectos se van a medir y que se han construido en base a la experticia del investigador con el fin de resolver el problema encontrado; por ejemplo: 5 dosis de alimento, 4 tipos de cruzamiento, 3 niveles de fertilización para pastos.
- Testigo: Existe una ligera divergencia al considerarlo como tratamiento; sin embargo, es parte de estas estrategias que sirven como fuente de comparación, testeo, además se encarga de medir el resultado del experimento, por lo que en ciertos casos su elección es importante. Cuando realizamos la observación de una determinada realidad “in situ”, y al comparar ciertas variables con estándares hasta encontrar el problema, entonces esta realidad se convierte en el “testigo” que puede conducir a una futura comparación entre los otros tratamientos elegidos para tomar las mejores decisiones.
- Repetición: Para formar las unidades experimentales se considera su tamaño que depende de las condiciones donde se realizará el experimento; una vez definidas estas sirven para albergar a los tratamientos (al azar); sin embargo, no se puede depender de una sola unidad para cada tratamiento; entonces se deben utilizar varias unidades experimentales del mismo tratamiento lo que se conoce como repetición; como se explicó en la figura 1.3. El número de repeticiones al diseñar un ensayo debe ser tal que los grados de libertad para el error experimental (en el modelo) sea mayor de 10 y nunca menor a 4 (Casanoves, 2019); se recomienda en la producción agropecuaria que al menos sean 4.

No debemos confundir a las repeticiones con “*submuestras*”, las mismas que suelen integrarse dentro de las unidades experimentales según su tamaño por ejemplo pueden ser 4 animales (submuestras) por tratamiento y por repetición. Para el modelo matemático pueden funcionar con un promedio de cada unidad experimental (convirtiéndose en pseudo réplicas) como indicamos anteriormente; o individualmente, en este caso el modelo

debe contener el error de las unidades experimentales y el error de las submuestras por separado. Todos los modelos lineales aditivos pueden diseñarse con submuestras si el caso lo amerita y las condiciones lo permiten (sobre todo las económicas).

- **Error experimental:** En la dinámica del manejo de una investigación, al aplicar los tratamientos podemos cometer ciertos errores que se las identifica como:
 - ✓ Variaciones pertinentes: Producidas por la acción propia en la elección de los tratamientos y de las unidades experimentales que pueden animales, suelo, plata; que por lo general tienen su propia dinámica individual en su desarrollo o proceso biológico; sin embargo, esto puede controlarse con el modelo lineal aditivo.
 - ✓ Variaciones no pertinentes: Se producen por causas extrañas que suelen disfrazar el efecto de los tratamientos, pero también podrían controlarse parcialmente con los métodos numéricos estadísticos (Reyes, 2003).

El error experimental debe suponerse es igual para todas las unidades experimentales, puesto que su distribución es normal con la media cero ($\mu = 0$) y su variación se estima con la varianza (σ^2). No puede eliminarse el error, solo se disminuye con estrategias para mejorar la estimación de los efectos de los tratamientos al medir cada variable respuesta, para mermar este efecto se puede recomendar:

- ✓ Usar unidades experimentales lo más uniformes posibles, como animales de la misma raza, sexo, edad, peso; un suelo homogéneo.
- ✓ El tamaño de la unidad experimental adecuado (TUE), recordemos que si elevamos el número, existe más fidelidad en los datos obtenidos, así como se evita la pérdida de unidades experimentales por daño o muerte.
- ✓ En experimentos con plantas se debe eliminar el efecto de orilla o competencia mutua de tratamientos usando la “parcela neta útil”.

- ✓ Utilizar un número eficiente de repeticiones para cada tratamiento (tomar en cuenta los grados de libertad del error cuando se apique el modelo matemático)
 - ✓ Manejo uniforme de las unidades experimentales, si se evalúa el peso debe ser el un solo día con la misma balanza, si vacunamos será en todas las unidades experimentales (animales), cambios de cama a todos por igual; en el caso de pastizales los riegos, densidades de siembra, fertilización, control de plagas, deshierbas será realizado conforme a un calendario, en mismo día para todas las parcelas experimentales.
 - ✓ Designar los tratamientos en las unidades experimentales procurando igualdad de condiciones, con la finalidad que, si alguno resultó superior al resto en las variables respuesta, tenga la oportunidad de mostrarlo.
 - ✓ En el sorteo de los tratamientos en base al coeficiente de variación, elegir el modelo lineal aditivo que mejor se adapte a las condiciones del experimento, y será al azar, aleatoria o estocásticamente pudiendo utilizar algunos métodos.
 - ✓ Aplicar los estadísticos que permitan separar las causas de variación (fuente de variación), para aprovechar eficientemente los resultados (De La Loma, 1996).
- Comparación: Para comprobar científicamente una hipótesis se requiere obligatoriamente utilizar métodos numéricos en este caso un análisis de varianza que me permite indicar si existe o no diferencia entre los tratamientos aplicados; sin embargo, este estadístico (Fisher) no necesariamente me reporta ¿qué protocolo o que tratamiento fue el más o menos favorecido al analizar una variable respuesta?; por lo tanto, debemos usar la *comparación* de medias de los tratamientos.
- Contraste: En cambio los contrastes son un conjunto de coeficientes que permiten formular y comprobar hipótesis específicas de nuestro interés (Casanoves, 2019); por ejemplo, si en un experimento queremos testear a un testigo (T_0) con el resto de los tratamientos ($T_1+T_2+T_3$); entonces su hipótesis nula será:

$$H_0: \mu T_0 = \frac{\mu T_1 + \mu T_2 + \mu T_3}{3}$$

Para construir el contraste ortogonal, asignamos coeficientes de manera que la suma sea cero. Para el tratamiento testigo (T_0): Coeficiente = 3; para los otros tratamientos (T_1, T_2, T_3): Coeficiente = -1 cada uno. Esto asegura que la suma de los coeficientes sea cero: $3 + (-1) + (-1) + (-1) = 0$. El contraste ortogonal sería: $C = 3\mu T_0 - \mu T_1 - \mu T_2 - \mu T_3$. Entonces si la diferencia es significativa se indicaría que en el testeo del testigo respecto a resto de tratamientos, hay una distancia y no existe evidencia para aceptar la hipótesis nula planteada.

Los contrastes pueden ser: polinómicos, ortogonales y no ortogonales (Casanoves, 2019).

En general para poder diseñar el experimento podemos plantearnos las siguientes interrogantes:

- ¿Cuántos tratamientos se desean estudiar?
- ¿Cuáles son las unidades experimentales?
- ¿Cuáles son las unidades de observación (TUE)?
- ¿Cuántas veces necesita observar la respuesta (repeticiones)?
- ¿Es necesario observar la evolución de la respuesta en el tiempo (clima o época seca y húmeda)?
- ¿Son las unidades experimentales homogéneas (coeficiente de variación)?
- ¿Son las unidades experimentales suficientes para realizar todos los tratamientos?
- ¿Cómo se asignan los tratamientos a las unidades experimentales?
- El objetivo del experimento, ¿es la comparación y/o la estimación? Si se trata de comparar la producción de pasto avena con fertilización y sin fertilización puedo usar 6 macetas (serían 3 repeticiones cada tratamiento), pero esos resultados ¿puedo estimar en una hectárea?, probablemente NO ya que para eso requeriría varias parcelas no macetas como unidades experimentales.
- ¿Los tratamientos tienen alguna estructura? Si son mono factoriales o poli factoriales (2 o más factores)

- ¿Puede el diseño resultante ser analizado y/o las comparaciones deseadas llevadas a cabo?
- ¿Cómo se modelan los datos? Hace referencia a los modelos lineales aditivos fijos, los mixtos, etc.

Con estos conceptos previos y tomando en cuenta el experimento que se pretende diseñar sobre alimentación en cerdos Duroc, quedaría de la siguiente manera:

Problema

Bajo peso promedio de cerdos Duroc (67,5 kg) a los 6 meses de edad, en comparación con el indicador de la raza (90 kg).

Hipótesis

Hipótesis teórica:

La utilización de 3 dietas alimenticias (16, 17 y 18% de proteína) influirá en el mejoramiento del proceso biológico de ceba para crianza de cerdos Duroc elevando el peso de saca

Hipótesis matemática:

$$H_1: \mu_1 \neq \mu_2 \neq \mu_3$$

Objetivo:

Evaluar el comportamiento biológico de cerdos Duroc mediante 3 dietas proteicas (16, 17 y 18%), para alcanzar el indicador de la raza.

Meta:

Valoración de los parámetros productivos en cerdos Duroc mediante 3 dietas proteicas (16, 17 y 18%), para procurar un peso de 85 kilos a los 180 días de ceba.

Unidad experimental:

Cada unidad experimental está integrada por 1 cerdo macho Duroc destetado de 15 kilos en su peso corporal y 21 días de edad.

Tratamientos:

- Dieta con 16% de proteína, código T1
- Dieta con 17% de proteína, código T2
- Dieta con 18% de proteína, código T3

Repeticiones:

4 repeticiones por tratamiento, es decir si nos referimos a las unidades experimentales entonces:

Total, de Unidades Experimentales = N° de tratamientos N° de repeticiones

$$UE = 3 * 4 = 12$$

Por lo tanto, en el experimento en mención se requerirán de 12 lechones destetados

Tabla 1.2. Esquema de un experimento para alimentación de cerdos con 3 tratamientos y 4 repeticiones

Tratamientos	Código	T.U.E.	Repeticiones	Total U.E./Trat.
Dieta con 16% de proteína	T1	1	4	4
Dieta con 16% de proteína	T2	1	4	4
Dieta con 16% de proteína	T3	1	4	4
Total, de unidades experimentales				12

Elaborado y adaptado por: Moscoso, 2025

En el diseño del experimento además debe tomarse en cuenta que aún si se pretende resolver el problema propio de la finca en nuestro caso de una piara, es imponderable meditar que puede servir para aplicarlo en otras granjas aledañas al sector, aporte científico que influye en la productividad y utilidad neta del mismo. Según el problema establecido, puede ocurrir una variante de una vía como es el caso del incipiente contenido de proteína en los cerdos entonces

se genera el experimento simple (unifactoriales); pero puede ocurrir que se pretenda evaluar las 3 dietas no solo en machos sino además en las hembras; en este caso se trataría de un experimento factorial con 2 factores de estudio (bifactorial): las 3 dietas (factor A) con los 2 sexos macho y hembra (factor B), y si se siguen usando 4 repeticiones entonces en el diseño tendríamos lo siguiente:

$$\text{Total de Unidades Experimentales} = \text{N}^\circ \text{ de } \mathbf{tratamientos} * \text{N}^\circ \text{ de repeticiones}$$

$$\text{Total de Unidades Experimentales} = (\mathbf{factor\ A} * \mathbf{factor\ B}) * \text{N}^\circ \text{ de repeticiones}$$

$$UE = (3 * 2) * 4 = 24 \text{ unidades experimentales o 24 cerdos destetados}$$

Al establecer los factores de estudio siempre primero será el que inicialmente nos condujo al experimento (dietas) y luego los otros que se presenten en el desarrollo del trabajo (sexo, localidad); estos pueden ser cuantitativos como el caso de los niveles de proteína en la dieta, o cualitativos como es el sexo, la raza (Moscoso G., Moreno, Moscoso M., & Armijos, 2022).

En ensayos sobre fertilización, plagas y enfermedades, dosis de siembra, métodos de riego de pastos, se deben ubicar en zonas representativas del sector ganadero o agropecuario. Al montar el experimento se debe cuidar los errores humanos u operacionales, se repetirá el proceso en otras localidades con las mismas condiciones climatológicas, etc. Hay que conocer la tecnología del cultivo y la dinámica de la producción animal.

El investigador debe ser sensible a los problemas de los productores pequeños, medianos e industriales, así como conocer de los canales de comercialización y los consumidores (para poder segmentar el mercado y trabajar con compradores cautivos).

1.6.7. Modelo lineal aditivo

La finalidad de un proceso investigativo es probar las hipótesis planteadas con anterioridad sobre las estrategias (tratamientos) que se han elegido técnicamente pretendiendo resolver un problema que puede ser productivo, económico, etc. Para que se pueda concretar el protocolo se acude a los algoritmos matemáticos que existen, pudiendo encontrarse metodologías como el “*Análisis de varianza*” (ADEVA o ANOVA por sus siglas en inglés), que se logra

mediante un estadístico denominado “*Fisher*” (por su creador Ronald Fisher 1921); utilizándolo para determinar la significación estadística de la diferencia en las medias de dos o más tratamientos experimentales; evidentemente que se refiere a los promedios de las unidades experimentales que sorteadas al azar les asignaron el mismo tratamiento mediante sus repeticiones; las respuestas pueden ser:

- Los tratamientos tienen diferencias significativas ($P \leq 0,05$)
- Los tratamientos tienen diferencias altamente significativas ($P \leq 0,01$)
- Los tratamientos no tienen diferencias significativas ($P > 0,05$)

En los dos primeros casos se podría indicar que no existe evidencia para aceptar la hipótesis nula, es decir que la verdadera sería la hipótesis alternativa (H_1). En el tercer caso en cambio la hipótesis nula (H_0) se ubicaría como verdadera.

Como se indicó para elegir el modelo lineal óptimo es preciso analizar mediante la estadística descriptiva a las unidades experimentales (muestras) que se van a usar para el experimento, las mismas deben reunir condiciones como:

- Representar a la población (en nuestro caso a los semovientes de la granja)
- Elegirlas al azar (randomización, estocásticamente)
- Contar con una distribución normal o estándar
- Su coeficiente de variación debe ser relativamente bajo

En el caso de las unidades de experimentación animal, tienen variación en su capacidad fisiológica individual (unos ganan más peso que otros, aunque se aplique igual tratamiento). Para medir la variabilidad del material experimental es necesario aplicar un ENSAYO EN BLANCO, que consiste en someter a todos los animales a un mismo régimen alimenticio y manejo uniforme un tiempo determinado y al final se determina cuáles semovientes son los más homogéneos en la ganancia de peso (usar el coeficiente de variación), si resulta bajo, el modelo lineal más recomendable sería un Diseño Completamente al Azar; en cambio sí debemos usar las 2 categorías (de buena ganancia y de ganancia media) puedo distribuir los tratamientos con un diseño experimental

que corrija esos estratos (más conveniente elegir al Diseño de Bloques Completos al Azar). Lo ideal sería usar gemelos univitelinos; aunque podemos ayudarnos con animales de la misma camada, misma edad, igual peso y sexo.

En los experimentos sobre pasturas donde se usa parcelas como unidades experimentales, el suelo es muy heterogéneo y se manifiesta con diferente intensidad de humedad, fertilización, pendiente, distribución de plagas, etc. La fertilidad no sigue una secuencia o patrón uniforme, aunque puede también reflejar alguna tendencia (mayor fertilidad en una sola región); igualmente se recomienda aplicar en el área a usar de suelo un ENSAYO EN BLANCO, que se caracteriza porque en todas las parcelas ya trazadas se siembra un solo cultivo de ciclo corto (rábano por ejemplo) y en la cosecha se evalúa la producción (kg/parcela), luego se identifica el coeficiente de variación para asignarle el modelo que mejor se adapte a su realidad; puede también construirse curvas de fertilidad, para unir puntos de igual producción para trazar las parcelas y sortear los tratamientos.

Con el antecedente indicado para designar el modelo respectivo de un experimento es preciso calcular el coeficiente de variación para en base a su valor, establecer el más apto para las condiciones de las unidades experimentales.

Un modelo lineal aditivo es una herramienta estadística que se utiliza para capturar relaciones no lineales entre una variable respuesta (peso, convención alimenticia, altura de planta, etc) y una o más variables predictoras (tratamientos). Difiere de los modelos tradicionales que asumen una relación lineal directa, este modelo aditivo utiliza funciones de suavizado para representar la relación entre las variables indicadas (Morales, 2023).

Básicamente un modelo lineal aditivo con una variable predictora X (tratamientos) y una variable respuesta (productivas o económicas en nuestro caso); se expresa como:

$$Y = f(X) + \epsilon$$

Donde $f(X)$ es una función de suavizado que captura la relación entre X e Y ; y ϵ representa el error o la variabilidad no explicada por el modelo (Morales, 2023). Esta orientación es bastante flexible permitiendo modelar las relaciones complejas que normalmente no se podrían capturar con los modelos simples; son especialmente útiles en situaciones donde la relación entre las variables no es lineal y puede variar de manera impredecible (ShallBD, 2023), en nuestro caso específico por las condiciones ambientales y la individualidad de las unidades experimentales (animales, plantas, suelo).

Este modelo inicial cuando se trata de experimentos agropecuarios podría expresarse matemáticamente como se indica a continuación:

$$Y_{ij} = \mu + \tau_i + \epsilon_{ij}$$

Donde:

Y_{ij} = Variable respuesta que corresponde a la observación de la i -ésima repetición del i -ésimo tratamiento

μ = Media general

τ_i = Efecto del i -ésimo tratamiento

ϵ_{ij} = Efecto aleatorio de la ij -ésima unidad experimental (Casanoves, 2019)

El modelo corresponde a un Diseño Completamente al Azar (DCA), más simple y seguro como se tratará en los próximos capítulos cuando se calcule el análisis de varianza (ADEVA). Antes de aplicar un análisis de varianza, se puede comparar directamente las medias de los tratamientos mediante la técnica estadística denominada intervalos de confianza, que fue el génesis tanto del análisis de varianza como de la separación de medias, para el efecto se usaba la prueba “t-student” que comparaba en pares, y cuando existían muchos tratamientos, se producía un sobre uso de información y falta de independencia en la prueba de hipótesis; pero si el experimento solo se desarrolla con 2 tratamientos con la prueba “t” se puede realizar un ADEVA y separación de medias al mismo tiempo, ya que compara solo un par. En cambio, si se dispone desde 3 tratamientos en adelante, el análisis de varianza solo prueba la hipótesis que entre las 3 o más medias existen diferencias, pero no indica que tratamiento

es el más o menos beneficiado, por lo tanto, si se conoce que existen estas diferencias significativas (hipótesis alternativa verdadera), obligatoriamente hay que utilizar la prueba de separación de medias de los tratamientos (Duncan, Tukey, DMS, etc); ya que solo tratándose de 3 en la “t-student” nos obligaría a realizar comparaciones de T_1 vs T_2 , T_1 vs T_3 ; y T_2 vs T_3 ; tres veces la prueba. Si se cuenta con más medias se deben calcular todas las combinaciones posibles con la expresión:

$$c = \frac{t*(t-1)}{2}; \text{ donde:}$$

- c = Número de combinaciones pares para la prueba “t-student)
- t = Número de medias de los tratamientos en consideración

Si por ejemplo se estudiaría 10 tratamientos, aplicando la ecuación tendríamos 45 pruebas “t” o comparaciones, que dificulta la toma de decisión respecto al tratamiento más favorecido y en cada variable respuesta. Por esta razón Ronald Fisher propone el modelo del análisis de varianza en donde cada observación de la variable respuesta (Y_{ij}), que proviene de los i -esimos tratamientos y j -esimas repeticiones, se debe a una media general (μ) más el efecto de los i -esimos tratamientos (τ_i) y aditivamente el efecto aleatorio del error experimental asociado a los i -esimos tratamientos con j -esimas repeticiones (ϵ_{ij}) (Casanoves, 2019).

Hay que tomar en cuenta las suposiciones que se presentan en el término del error experimental que se expresa en:

$$e_{ij} \sim N(0, \sigma^2)$$

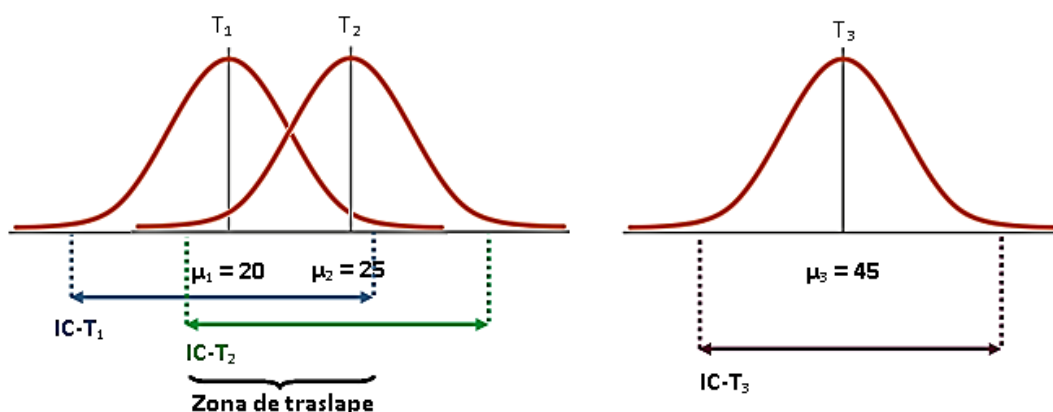
Debe existir una distribución normal, con media cero (0) y varianzas homogéneas

Además, es preciso que existe independencia de las observaciones:

$$\text{cov}(e_{ij}, e_{i'j'}) = 0 \quad \forall ij \neq i'j'$$

En términos de independencia, la covarianza entre el error proveniente de una observación y el error de la otra es igual a cero (0) para todas las combinaciones posibles; si en todas ellas la covarianza es cero, entonces los datos no están relacionados (independencia); en el caso que no se cumpliera esta característica, la prueba de varianza puede fracasar y llevarnos a conclusiones erróneas. En los ensayos con pastos, si retomamos el efecto de borde que se indicó anteriormente, cuando no se considera este particular es decir no trabajamos con la parcela neta útil, se produce una variación no pertinente y el efecto de borde se hace presente por lo tanto la “data” que sale no muestra independencia; en el caso de los animales, cuando se evalúa raciones alimenticias por ejemplo y no disponen de comederos individuales por cada tratamiento y repetición, pueden enmascarse los resultados y evitar la independencia de los mismos.

Figura 1.6. Esquema para la dinámica y funcionamiento de un modelo lineal en una hipótesis alternativa.



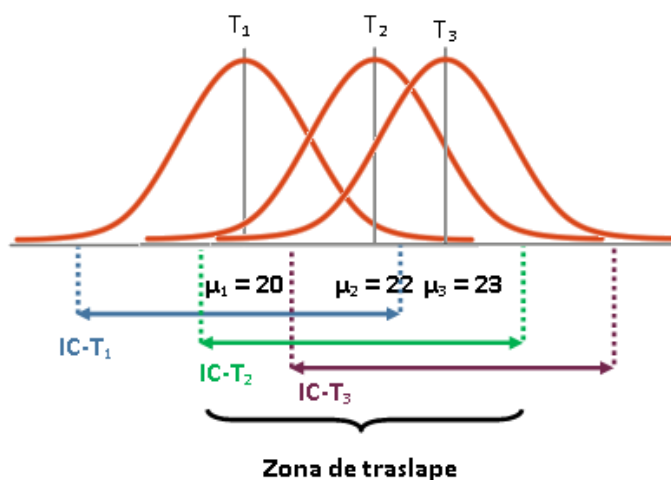
Elaborado por: Moscoso, 2025

Si contamos con 3 muestras o en el caso 3 tratamientos que tienen sus respectivas medias ($T_1 = 20$, $T_2 = 25$ y $T_3 = 45$), con solo observar la figura 1.6, nos daríamos cuenta que las dos primeras campanas se traslapan, en cambio la tercera se separa de ellas considerablemente, podríamos inclusive mirar los intervalos de confianza (IC) de cada una de las distribuciones y efectivamente decidir sobre las hipótesis. Sin realizar un análisis de varianza podemos concluir que la hipótesis verdadera sería la alternativa:

$$H_1: \mu_1 \neq \mu_2 \neq \mu_3$$

En otras palabras, no hay evidencia estadística para aceptar la hipótesis nula; sin embargo, al comparar T_1 vs T_2 no se podría indicar diferencias (por el traslape), pero al contrario en la comparación de los pares T_1 vs T_3 o T_2 vs T_3 si se considera una distancia significativa. En otras palabras, para la primera comparación podría ser verdadera la hipótesis nula, y lo contrario para las dos últimas.

Figura 1.7. Esquema para la dinámica y funcionamiento de un modelo lineal en una hipótesis alternativa.



Elaborado por: Moscoso, 2025

Si por el contrario las medias fueran similares como se aprecia en la figura 1.7., al fijarnos en los tres intervalos de confianza (IC) todos se traslapan por lo tanto estadísticamente no tienen diferencias significativas lo que indicaría que la hipótesis verdadera es la nula:

$$H_0: \mu_1 = \mu_2 = \mu_3$$

Para entender los componentes del modelo lineal aditivo, tomemos como referencia la figura 1.6., en las 3 medias muestrales (20, 25 y 45) y suponiendo que cada muestra tenga 5 datos, tendríamos un total de 15; entonces la media global será el promedio de las medias muestrales es decir: $(20 + 25 + 45)/3 = 30$; no obstante si en uno de los tratamientos existe una unidad experimental perdida, no se da esta tendencia por desbalance. En el ejemplo en mención con muestras balanceadas se pueden calcular los efectos de cada una de las medias de los tratamientos con referencia a la media global $\mu = 30$, tendríamos que:

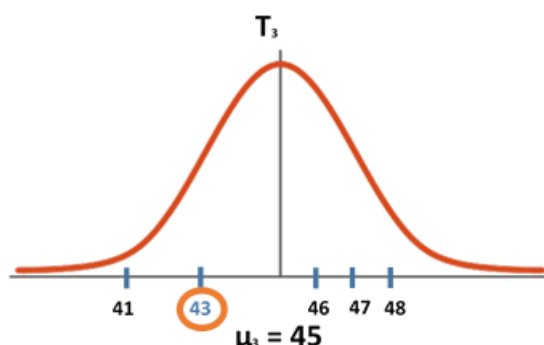
Efecto 1: $\bar{x}_1 - \mu$; $\tau_1 = 20 - 30 = -10$

Efecto 2: $\bar{x}_2 - \mu$; $\tau_2 = 25 - 30 = -5$

Efecto 3: $\bar{x}_3 - \mu$; $\tau_3 = 45 - 30 = +15$

Entonces la suma de los efectos en las medias respecto a la media general es cero (0), que explica los componentes del modelo.

Figura 1.8. Análisis de los componentes de un modelo aditivo



Elaborado por: Moscoso, 2025

Recordemos nuestro modelo $Y_{ij} = \mu + \tau_i + \epsilon_{ij}$ y según la figura 1.8., el *tratamiento 3* está conformado por 5 unidades experimentales (repeticiones), y los datos obtenidos son: 41, 43, 46, 47 y 48; que tendría una media de $\mu_3 = 45$; tomamos en consideración el segundo dato que es 43 el mismo que se encuentra a 2 puntos debajo del promedio ($45 - 43$); entonces aplicando el modelo tendríamos que:

$$Y_{ij} = 43$$

$$\mu = 30$$

$$\tau_i = +15 \text{ (corresponde al efecto del tratamiento } \tau_3 \text{)}$$

$$\epsilon_{ij} = (\bar{x}_3 - \mu_3) = 43 - 45 = -2 \text{ (es el término del error respecto a su media)}$$

Entonces:

$Y_{ij} = \mu + \tau_i + \epsilon_{ij}$ para este caso sería: $43 = 30 + 15 - 2$; este ensayo corresponde a la aplicación del modelo del análisis de varianza, el mismo que

verifica los efectos de los tratamientos; entonces las varianzas son más o menos homogéneas y los errores tendrían similar variabilidad, por esta razón la prueba de la homogeneidad de la varianza, así como la prueba de normalidad se realiza sobre los residuos y no en base a la variable respuesta. Los errores del modelo es la distancia que tienen con respecto a la media cada uno de los datos dentro del tratamiento en mención (Casanoves, 2019).

Si al comprobar una hipótesis, *la nula* (H_0) resulta *verdadera*; lo que indicaría es que todas las medias de los tratamientos son iguales por lo tanto τ_i sería cero (0), dando lugar a un *modelo reducido*; que se expresaría:

$$Y_{ij} = \mu + \epsilon_{ij}$$

Según estos criterios cuando comprobamos las hipótesis de trabajo tenemos solo 2 escenarios:

- Que la *hipótesis nula* (H_0) sea verdadera (todas las medias de los tratamientos son iguales) y se use el modelo reducido $Y_{ij} = \mu + \epsilon_{ij}$; o
- Que la hipótesis alternativa (H_1) sea verdadera (al menos un par de medias de los tratamientos no son iguales) y se use el modelo completo $Y_{ij} = \mu + \tau_i + \epsilon_{ij}$

Uno de los protocolos para la experimentación es decidir sobre el uso del modelo matemático; en el caso de la producción agropecuaria se usa como indicador comúnmente el coeficiente de variación (CV), de una variable representativa; luego de someter a las unidades experimentales a un “ensayo en blanco”, si se encuentra homogeneidad entonces se puede asumir que la hipótesis nula es verdadera y se optaría por el reducido en el **sorteo**; aunque acudamos luego al completo finalizando la evaluación, y la pregunta que nos formulamos será: ¿Cómo se establece si el modelo más simple es adecuado para los datos observados?, al responder esta interrogante nos planteamos lo siguiente:

- Se asume que la hipótesis nula es cierta, todas las medias de los tratamientos son iguales, hasta que se demuestre lo contrario.
- Se calcula una medida de credibilidad de la hipótesis nula, y se logra con el número de repeticiones que se hayan planteado al inicio (sorteo).

- La medida de credibilidad es una probabilidad que se conoce como p-valor; que evidentemente se consigue con un análisis de varianza dentro de un diseño de experimentos que se distancia cuando se trata de un conjunto de datos sobre un estudio observacional. En el caso de p-valor (p-value) se toma en cuenta el mismo para poder decidir:
 - ✓ $P\text{-value} \leq 0,05 > 0,01$ entonces se rechaza la hipótesis nula y la diferencia de esos tratamientos es significativa.
 - ✓ $P\text{-value} \leq 0,01$ igualmente se rechaza la hipótesis nula y su diferencia es altamente significativa.
 - ✓ $P\text{-value} > 0,05$ la hipótesis nula será verdadera, indicando que no existe diferencia en los tratamientos aplicados, según la variable respuesta considerada.

En los dos primeros casos al escribir científicamente ya no se estila indicar que la hipótesis nula es falsa, en cambio se usa la terminología señalando que “*no existió evidencia estadística para aceptar la hipótesis nula*”.

- ✓ Cuanto menor es el p-value, menos verosímil es la hipótesis nula, como se explicó en el párrafo anterior.
- ✓ Se fija un umbral por debajo del cual la hipótesis nula se rechaza
- ✓ El umbral se conoce como nivel de significación y se simboliza con α , en nuestro caso se acostumbra a trabajar con $\alpha 0,05$ y $\alpha 0,01$.

1.6.8. Recolección de la data

Superados los anteriores protocolos del método científico, una vez que se sortee los tratamientos en las unidades experimentales y se determine el modelo lineal aditivo que se utilizará para comprobar las hipótesis, entonces inicia el trabajo de campo, período en el que se recolecta los resultados experimentales de cada una de las variables respuesta.

Durante esta etapa, suelen cometerse errores de manejo, o variables no pertinentes, por lo tanto, deben evitarle y se recomienda:

- En experimentos con pastos: evitar hacer tomas en los bordes de las parcelas en variables como: altura de planta, número de tallos por planta,

hojas por tallo, índice de área foliar; las mismas que se evaluará preferentemente concéntricamente y al azar. Si se riega una parcela deben proceder con todas, la toma de datos se programa en un determinado día para todo el campo experimental evaluando las unidades experimentales completas.

- En experimentos con animales: las evaluaciones se realizarán en mismo día programado para todas las unidades, usar la misma balanza cuando se trate del peso, disponer de comederos y bebederos individuales en cada unidad experimental, evitar el estrés, utilizar debidamente las tablas de alimentación para proporcionar la cantidad de la dieta que deben consumir evitando sobrealimentación, desperdicios o falta de alimento; vacunaciones, desinfecciones, limpiezas, etc, realizar periódicamente según el protocolo en su totalidad y el mismo día (no podemos vacunar por ejemplo parcialmente a unos animales y luego a otros).

Toda la data debe ser anotada en una libreta de campo y luego tabulada en una hoja de Excel por ejemplo para luego hacer el tratamiento de los datos y finalmente su procesamiento para el análisis de varianza y la separación de medias.

1.6.9. Interpretación y comparación

Cuando se ha culminado con el trabajo de campo y generado la “data” de las variables respuesta, así como el tratamiento de los datos para superar los 3 supuestos necesarios previo al procesamiento estadístico, se procede a la comprobación de las hipótesis planteadas mediante el análisis de varianza y de ser necesaria la separación de medias; debemos interpretar los resultados experimentales obtenidos.

Cuando se analiza un ADEVA se cuenta con un conjunto de resultados que se expresan dentro de una tabla como la que continúa:

Tabla 1.3. Resultados de un análisis de varianza en un experimento donde se midió el efecto de la zanahoria en el consumo de alimento de cerdos mestizos.

Fuente de variación	Suma de cuadrados	Grados de libertad	Cuadrado Medio	Fisher Calculado	F. Tab 0,05	F. Tab 0,01	P-value
Total	183,54	19					
Tratamientos	132,10	4	33,03	9,63	3,06	4,89	0,0005
Error	51,44	15	3,43				

Elaborado y adaptado por: Moscoso, 2025

Si consideramos la tabla 1.3., una ADEVA está conformado por la fuente de variación que contiene la información en base al modelo matemático elegido en el sorteo de tratamientos al inicio del experimento; además la suma de cuadrados, los grados de libertad, el cuadrado medio o varianza, y el estadístico Fisher calculado, que en este caso se obtiene dividiendo la varianza de los tratamientos para la varianza del error; luego contamos con los valores de Fisher tabular, que se identifican en la respectiva tabla en base a los grados de libertad de los tratamientos y el error, y finalmente contamos con el P-valor o P-value.

Existen dos métodos para declarar que una determinada hipótesis resultó verdadera o falsa:

- En base al Fisher calculado comparándolo con el tabular, en el ejemplo mencionado el valor F.Cal. es de 9,63; que supera a los valores tabulares (3,06 a 0,05; y 4,89 a 0,01) por lo tanto se podría declarar que existen diferencias altamente significativas en el consumo de alimento de dietas en base a zanahoria, en cerdos mestizos. En otras palabras, no hay evidencia estadística para aceptar la hipótesis nula; sin embargo, el ADEVA no nos indicaría que dieta fue la que mayor consumo tuvo por lo que convendría realizar un análisis adicional que se conoce como separación de medias de los tratamientos.
- El segundo método es fijarse en el P-value, que en la tabla se describe con el valor de 0,005; como se indicó si el P-value $\leq$ 0,01 igualmente se rechaza la hipótesis nula ya que la diferencia de los tratamientos es altamente significativa.

Si se analizan los dos estadísticos; aparentemente se pensaría que son antagónicos, ya que en el un caso cuando se supera el valor calculado con el tabulado, se habla de una hipótesis alternativa como verdadera; en cambio con el P-valor o P-value si el indicador calculado es menor que los niveles de significación (α) se opta por rechazar la hipótesis nula, para entender mejor se esquematizan los 2 procedimientos en la table 1.4.

Tabla 1.4. Identificación de la hipótesis verdadera en base a los 2 estadísticos (Fisher y P-value).

Hipótesis verdadera	Fisher			P-value		Resultado
	Calculado	F-Tab. α 0,05	F-Tab. α 0,01	α 0,01	α 0,05	
Nula (H_0)	F.Cal	Menor	Menor	Mayor	Mayor	No significativo
Alternativa (H_1)	F.Cal	Mayor	Menor	Mayor	Menor	Significativo
	F.Cal	Mayor	Mayor	Menor	Menor	Altamente Significativo

Elaborado y adaptado por: Moscoso, 2025

Luego de la comprobación de la hipótesis mediante el análisis de varianza, si se trata de una hipótesis nula verdadera, no se procede al cálculo de la separación de medias; en cambio si es la alternativa, entonces hace falta identificar a los tratamientos menos y más favorecidos que nos faciliten una toma de decisión más óptima; acción que se denomina comparación de medias de los tratamientos, existiendo pruebas como: Duncan, Tukey, Scheffe, ect., las que se analizarán detenidamente en el capítulo V.

Al identificar los mejores tratamientos en cada variable respuesta, además debemos compararlos con los estándares nacionales o mundiales, así como con los valores que se identificaron al inicio del diagnóstico en donde se detectó el problema.

Este procedimiento tiene su explicación en la realidad de la dinámica dentro de un agroecosistema, entendiéndose que el uso del método científico se ha justificado para resolver la problemática y debe analizarse el “antes y después” del experimento, si valió o no le pena establecerlo o si el mismo sirve para

replantear nuevos experimentos dirigidos a solucionar el problema en su totalidad.

1.6.10. Conclusiones y toma de decisiones

Es la última actividad del protocolo científico que se utiliza para establecer el paquete tecnológico o las acciones técnicas que se requieren implementar para poder resolver el problema presentado inicialmente y mejorar la productividad, así como los estándares económicos de la finca.

Para desarrollar un juicio de valor antes de la decisión, es importante listar cada una de las variables productivas y económicas que fueron sometidas al modelo matemático e identificar dentro de los tratamientos el conjunto de ellas que resultaron con mejores promedios y así mejorar la comprensión holísticamente; en la tabla 1.5., se puede apreciar el procedimiento

Tabla 1.5. Análisis de los tratamientos favorecidos por las técnicas estadísticas en cada una de las variables respuesta de una investigación.

Variable	Tratamientos favorecidos			
	T ₁	T ₂	T ₃	T ₄
Peso final, kg			x	
Ganancia de peso, kg			x	
Consumo de alimento, kg		x		
Conversión alimenticia			x	
Costo por kg de ganancia de peso				x
Relación beneficio / costo			x	

Elaborado y adaptado por: Moscoso, 2025

Según el ejemplo se puede apreciar que en la mayoría de las variables el tratamiento 3 fue el más beneficiado, aunque en otras hubo diferente respuesta. Para tomar la decisión final debemos fijarnos que las variables respuesta se clasifican en dos grandes grupos que son: *productivas* y *económicas*; que pueden llevarnos a tomar decisiones equivocadas por ejemplo si pueden existir ventajas productivas con un determinado tratamiento, no debemos precipitarnos a elegirlo directamente, ya que podría resultar muy costoso o el beneficio costo por ejemplo puede reflejar un valor menor que 1, entonces no sería

recomendable sugerirlo debido a que en el ejercicio económico estamos trabajando a pérdida así se demuestre ventajas productivas; hay que hacer un balance entre estos dos grupos y sabiamente recomendar lo debido.

1.7. Variables respuesta y su clasificación

En la investigación científica las variables juegan un importante papel, ya que permiten medir y analizar fenómenos de manera cuantitativa o cualitativa. Dentro de este contexto, la variable respuesta es fundamental, debido a que la misma refleja el resultado o efecto que se espera observar o explicar a partir de otros factores llamados variables independientes que son los tratamientos.

Según (Montgomery, 2019), la variable respuesta, también conocida como variable dependiente (Y), es *"la medición o el resultado de interés que se intenta modelar y explicar en un experimento o análisis estadístico"*. Es decir, es la variable que se ve afectada por cambios en las variables independientes (X) o de control (Hernández, R., Fernández, C., & Baptista, P, 2014), que generalmente es manejada por el investigador; por ejemplo, los niveles de proteína que se utilizarán en la preparación de balanceado animal que pretende evaluar el comportamiento nutricional.

Por su parte, (Hernández Sampieri, R., Fernández Collado, C., & Baptista Lucio, P, 2014) definen la variable respuesta como: *"el fenómeno que se observa, mide y analiza para determinar si existe alguna relación o efecto significativo ante la manipulación o variación de otras variables"*.

En otras palabras, la variable respuesta es el objeto final de observación que permitirá extraer conclusiones sobre el fenómeno de estudio. Una variable respuesta es aquella que nos permite evaluar a las estrategias o los tratamientos que se escogieron para resolver un problema conocido al inicio del método científico; en ellas se genera una "data" que permite usar el modelo matemático con el fin de comprobar hipótesis y seleccionar la mejor propuesta para cambiar la condición productiva o económica de una granja agropecuaria.

De acuerdo con (Triola, 2018), la variable respuesta debe reunir las siguientes características:

- Depende directamente de las condiciones o valores de las variables independientes (tratamientos aplicados).
- Puede ser cuantitativa (cuando se mide numéricamente) o cualitativa (cuando se clasifica en categorías).
- Es la variable que se desea *predecir, estimar o controlar* en un proceso de modelado o análisis.

El análisis de la variable respuesta es esencial en todo diseño de experimentos, ya que es a partir de sus resultados que se comprueba o refuta una hipótesis; además se destaca que un buen diseño experimental debe asegurar que la variable respuesta sea sensible a las variaciones de las variables controladas (tratamientos), para garantizar la validez estadística de las conclusiones (Montgomery, 2019).

Como se indicó la “data” evidentemente está conformada por algunos datos considerados como “información” y que puede ser escrita o electrónica como cifras en bruto, imágenes, audios, videos, que representa a las variables (Toledo, 2018), las mismas que pueden dividirse en:

1.7.1. Variables cualitativas

Caracterizadas porque reflejan cualidades o también llamados atributos, estos valores o modalidades no se pueden asociar normalmente a un número, por lo tanto, se limita a no realizar operaciones aritméticas con las mismas; por ejemplo, el sexo, color de pelaje, color de ojos, la raza, diagnóstico del paciente, etc. Pueden subdividirse en:

- Nominal: si los valores o modalidades no se pueden ordenar, como el sexo, la raza, forma de la hoja de un pasto.
- Ordinal: cuando los valores o modalidades pueden ser sujetos de ordenamiento como la intensidad del color de la lana de oveja o de los ojos del cuy macabeo, la mejoría de un tratamiento (Gómez, D. et. al, 2013)

1.7.2. Variables cuantitativas

Sus valores son numéricos, pueden sujetarse a operaciones algebraicas y medibles como el número de abortos en madres porcinas, número de gazapos por parto en conejas, peso de ovinos al nacimiento. Estas se caracterizan por distribuirse en 2 grupos:

- Discretas: cuando solamente toman valores enteros, sin fraccionamiento como el número de corderos destetados, número de partos por año, total de vacas gestantes.
- Continuas: si presentan valores enteros, pero también en fracciones, es decir que, entre dos valores, son posibles infinitos valores intermedios (Gómez, D. et. al, 2013) como puede ser la altura a la cruz, el peso vivo de bovinos de carne, la conversión alimenticia

Según (Casanoves, 2019), siempre que una variable cuantitativa se mida es continua, en cambio cuando se pueda contar es discreta, comúnmente discretizamos muchas variables como la temperatura, la edad. Hay que tomar en consideración estos aspectos.

1.8. Estadística descriptiva e importancia del coeficiente de variación

En el área agropecuaria se acostumbra investigar utilizando a seres vivos como unidades experimentales que tienen su propio individualismo y dinámica, por esta razón se hace preciso contar con un “ensayo en blanco”, que debe proyectar resultados en una “data” de la variable más lógica como puede ser una ganancia de peso (en animales) o un peso de campo por parcela (en vegetales o pastos). Con estos valores obtenidos se puede usar la estadística descriptiva para analizar su tendencia y sobre todo el coeficiente de variación que responde a la ecuación:

$$CV = \sigma / \bar{x} * 100$$

Donde:

CV = Coeficiente de variación

σ = Desviación estándar

$\bar{x}$ = Media general

El coeficiente de variación es el estadístico que nos advierte sobre la homogeneidad o heterogeneidad del material experimental (unidades experimentales), se expresa en porcentaje (%). Mientras más bajo sea este coeficiente, se indica mayor homogeneidad, caso contrario si se eleva representa lo contrario (heterogeneidad).

Si se califica como homogénea a la muestra; entonces se utiliza el modelo lineal aditivo más simple denominado Diseño Completamente al Azar (DCA); por el contrario, si es alto el coeficiente de variación (CV) se podría seleccionar al Diseño de Bloques Completos al Azar (DBCA) o al Diseño Cuadrado Latino (DCL) según el sentido de la fuente de variación.

Una vez establecido el modelo, puede procederse al sorteo de los tratamientos en cada unidad experimental pero estocásticamente (al azar), técnica de Fisher que pretende estimar en forma insesgada el error experimental; además el sorteo destruye la correlación de los errores y hacen más válidas las pruebas de significación o separación de medias, se randomiza con tablas de números al azar, la tómbola, papeles marcados con las unidades experimentales, etc (Casanoves, 2019).

Como se indicó el modelo más simple es el DCA, que se asigna dependiendo de la realidad del experimento, así como de las condiciones logísticas que se pueda encontrar al momento de la instalación del ensayo o experimento.

Tomando en consideración las experiencias que durante más de 50 años se ha contado en la Escuela Superior Politécnica de Chimborazo sobre experimentos agropecuarios, dependerá la decisión de la unidad experimental y el sitio donde se realice la prueba, tomando parte de protocolo de (González, 1976) podemos recomendar que se aplique el modelo lineal aditivo conocido como Diseño Completamente al Azar en unidades experimentales homogéneas que son analizadas su **coeficiente de variación**, luego del ensayo en blanco y que sus valores pueden ubicarse considerando lo siguiente:

- Experimentos de laboratorio (cajas Petri) menos de 1%

- En invernaderos de plantaciones o galpones de crianza animal de 5 a máximo 8%
- En el aire libre como pastizales o rejos de ganado puede aceptarse hasta el 20% de variación, aunque esto causa un grado de polémica; sin embargo, debemos apegarnos a las condiciones reales del campo.

Si los valores sobrepasan esta política, como se explicó anteriormente se optará por utilizar otros modelos como el Diseño de Bloques Completos al Azar o el Cuadrado Latino.

Procederemos a analizar algunos casos prácticos que pueden considerarse en la experimentación pecuaria.

1.8.1. Sorteo en un Diseño Completamente al Azar

Se pretende realizar un ensayo de cuyes macabeos para evaluar el comportamiento productivo cuando se utilizan 2 sistemas de crianza: en jaula y en el piso; se entiende que el problema encontrado fue la morbilidad de los ejemplares por coccidiosis adquirido por las camas húmedas del piso en época de invierno.

Para el efecto se cuentan con un total de 24 animales de una misma línea genética, de edad similar y seleccionados de madres con un segundo parto, con pesos al destete que van entre 300 y 330 gramos aproximadamente. En la tabla 1.6., se describen los datos de los 24 ejemplares que se utilizarían para sortear los 2 tratamientos en cuestión: Manejo en jaula y en piso.

Tabla 1.6. Data inicial de la ganancia de peso de cuyes luego del ensayo en blanco, que serán convertidos en unidades experimentales para un ensayo.

N°	Peso de cuyes, gr		Ganancia de Peso, gr.
	Inicial	Final	
1	302,39	310,53	8,14
2	391,22	399,36	8,14
3	283,46	291,60	8,14
4	301,80	309,97	8,17

5	387,02	395,32	8,30
6	336,83	345,35	8,52
7	352,35	360,88	8,53
8	347,33	356,00	8,67
9	304,04	312,84	8,80
10	256,11	265,06	8,95
11	361,82	370,79	8,97
12	395,09	404,09	9,00
13	326,93	336,07	9,14
14	337,86	347,02	9,16
15	390,96	400,18	9,22
16	356,68	365,98	9,30
17	257,08	266,47	9,39
18	393,16	402,60	9,44
19	311,51	320,98	9,47
20	376,62	386,15	9,53
21	396,10	405,77	9,67
22	294,25	303,96	9,71
23	360,29	370,01	9,72
24	333,44	343,24	9,80

Elaborado por: Moscoso, 2025

Las unidades observacionales (24 cuyes) se han sometido al ensayo en blanco por 15 días donde recibieron todas la misma alimentación, manejo uniforme, manejados en pozas de piso como se acostumbra tradicionalmente, al inicio se tomó un peso y luego de este período se evaluó esta misma variable a los 15 días, para obtener la ganancia con la expresión siguiente: $GP = PF - PI$, donde

GP es la ganancia de peso, PI el peso inicial y PF peso final; por lo tanto por ejemplo para el cobayo número 11 inició con 361,82 gramos, al término de los 15 días subió a 370,79 gramos; por lo tanto tendríamos una ganancia de peso de 8,97 gramos (tabla 1.6.)

Para establecer el experimento se toma en cuenta los 2 tratamientos con las respectivas repeticiones, para construir las unidades experimentales por lo tanto si son 2 sistemas (jaula y piso) y 6 repeticiones, hacen un total de 12 unidades experimentales. Si se cuentan con 24 cobayos destetados, entonces el tamaño de cada unidad experimental será 2; lo que se entendería en la tabla 1.7.

Tabla 1.7. Esquema de un experimento para dos sistemas de crianza de cuyes (Jaula y Piso) con 6 repeticiones

Tratamientos	Código	T.U.E.	Repeticiones	Unid. Exp./Trat.	Total Animales (unid. Observacional)
Sistema en Jaula	T1	2	6	6	12
Sistema en Piso	T2	2	6	6	12
Total				12	24

Elaborado por: Moscoso, 2025

Figura 1.9. Esquema de las unidades experimentales conformadas con 2 animales cada una previo al sorteo de tratamientos.

Poza 1	Poza 2	Poza 3	Poza 4	Poza 5	Poza 6
 8,14	 8,14	 8,30	 8,53	 8,80	 8,97
 8,14	 8,17	 8,52	 8,67	 8,95	 9,00
$\bar{x}$ 8,14	$\bar{x}$ 8,16	$\bar{x}$ 8,41	$\bar{x}$ 8,60	$\bar{x}$ 8,87	$\bar{x}$ 8,99
Poza 7	Poza 8	Poza 9	Poza 10	Poza 11	Poza 12
 9,14	 9,22	 9,39	 9,47	 9,67	 9,72
 9,16	 9,30	 9,44	 9,53	 9,71	 9,80
$\bar{x}$ 9,15	$\bar{x}$ 9,26	$\bar{x}$ 9,42	$\bar{x}$ 9,50	$\bar{x}$ 9,69	$\bar{x}$ 9,76

Elaborado por: Moscoso, 2025

Entonces se configuran las unidades experimentales con 2 animales de similar peso como se aprecia en la figura 1.9., con su respectiva data de la ganancia de peso, para efectos del sorteo debemos hacer un promedio de cada poza (luego será denominada unidad experimental). Por ejemplo, la poza 2 integrada por 2 cuyes el 3 y 4 con pesos de 8,14 y 8,17 gramos en su orden, lo que reflejaría un promedio de 8,16; y así en todas las demás hasta llegar al casillero número 12 que contiene a 2 cuyes de 9,72 y 9,80 gramos que representa una media de 9,76 gramos (ver figura 1.9.).

Con estos promedios de las ganancias de peso es decir 12 datos se procede a calcular los estadísticos necesarios para poder encontrar el coeficiente de variación y determinar si es bajo para poder sortear dentro de cada poza a los tratamientos con las repeticiones que les corresponda y con un modelo matemático llamado DCA.

Tabla 1.8. Distribución de la ganancia de peso en gramos, de cuyes luego del ensayo en blanco, y promediados por cada unidad experimental, previo al sorteo

Poza 1	Poza 2	Poza 3	Poza 4	Poza 5	Poza 6
8,14	8,16	8,41	8,60	8,87	8,99
Poza 7	Poza 8	Poza 9	Poza 10	Poza 11	Poza 12
9,15	9,26	9,42	9,50	9,69	9,76

Elaborado por: Moscoso, 2025

Se procede al cálculo de la media aritmética con los datos de la tabla 1.8.

$$\bar{x} = \frac{\sum X_i}{n}$$

Donde:

$\bar{x}$ = Media aritmética

$\sum X_i$ = Sumatoria de todas y cada una de las observaciones

n = Número de observaciones

Entonces:

$$\bar{x} = \frac{8,14 + 8,16 + 8,41 + 8,60 + 8,87 + 8,99 + 9,15 + 9,26 + 9,42 + 9,50 + 9,69 + 9,76}{12}$$

$$\bar{x} = \frac{107,04}{12} = 9,00$$

Continuamos con el cálculo de la desviación estándar

$$\sigma = \sqrt{\frac{\Sigma(X_i - \bar{x})^2}{(n - 1)}}$$

Donde:

σ = Desviación estándar

$\Sigma(X_i - \bar{x})^2$ = Sumatoria del cuadrado de las desviaciones de cada observación con referencia a su media

$n - 1$ = Grados de libertad de las observaciones

Tabla 1.9. Esquema del procedimiento para el cálculo de la desviación estándar

Pozas	Ganancia de Peso, gr. ($\bar{x}$)	$(X_i - \bar{x})$	$(X_i - \bar{x})^2$
1	8,14	-0,86	0,73
2	8,16	-0,84	0,71
3	8,41	-0,58	0,34
4	8,60	-0,40	0,16
5	8,87	-0,12	0,01
6	8,99	-0,01	0,00
7	9,15	0,15	0,02
8	9,26	0,27	0,07
9	9,42	0,42	0,18
10	9,50	0,51	0,26
11	9,69	0,69	0,48
12	9,76	0,76	0,59
Suma (Σ)	107,94	0,00	3,54
			$\Sigma(X_i - \bar{x})^2$

Elaborado por: Moscoso, 2025

Con el fin de reemplazar la ecuación debemos configurar la tabla 1.9., para el cálculo de cada componente.

En la columna 3 se calcula cada una de las desviaciones ($X_i - \bar{x}$) lo que indicaría que cada observación debe restarse de la media aritmética, por ejemplo, en la poza 1 ($8,14 - 9,00$) nos arroja como resultado $-0,86$; y se procede con los demás datos; la suma de las desviaciones debe ser cero (0).

La columna 4 en cambio encierra los valores de dichas desviaciones, pero elevadas al cuadrado, en el mismo ejemplo $(-0,86)^2$ sería $0,73$ la suma total es el valor para reemplazar en la ecuación.

$$\sigma = \sqrt{\frac{3,54}{(12 - 1)}} = 0,567$$

Ahora procedemos al cálculo del coeficiente de variación

$$CV = \frac{\sigma}{\bar{x}} \times 100 = \frac{0,567}{9,00} = 0,0631 \times 100 = 6,31\%$$

Recordemos que se trata de un experimento dentro de un galpón que para calificar a las unidades experimentales (pozas) como “homogéneas” se requiere de un coeficiente de variación hasta el 8%, por lo tanto el 6,31%, está dentro de los rangos recomendados para sortear los tratamientos con un modelo lineal aditivo Diseño Completamente al Azar, el mismo que se expresa $Y_{ij} = \mu + \tau_i + \epsilon_{ij}$; para lo cual se procede con el protocolo que continúa.

Sorteo al azar o estocásticamente, en cada poza se le designa con una tómbola o papeles identificados con los respectivos códigos a los tratamientos y sus repeticiones. Supongamos que se sorteó y las unidades experimentales quedaron como se expresa en la tabla 1.10.

Tabla 1.10. Sorteo de tratamientos con un DCA para un ensayo de cuyes criados en Jaula y Piso

Pozas	Ganancia de peso, gr	Tratamiento	Repetición
1	8,14	Jaula	6
2	8,16	Piso	3
3	8,41	Piso	5
4	8,60	Jaula	3
5	8,87	Jaula	1
6	8,99	Piso	1

7	9,15	Jaula	2
8	9,26	Jaula	4
9	9,42	Piso	4
10	9,50	Piso	2
11	9,69	Jaula	5
12	9,76	Piso	6

Elaborado por: Moscoso, 2025

Entonces el esquema de las unidades experimentales sería como se presenta en la figura 1.10.

Figura 1.10. Sorteo de tratamientos con un DCA para un ensayo de cuyes criados en Jaula y Piso una vez asignadas las repeticiones

J6	P3	P5	J3	J1	P1
 8,14	 8,14	 8,30	 8,53	 8,80	 8,97
 8,14	 8,17	 8,52	 8,67	 8,95	 9,00
$\bar{x}$ 8,14	$\bar{x}$ 8,16	$\bar{x}$ 8,41	$\bar{x}$ 8,60	$\bar{x}$ 8,87	$\bar{x}$ 8,99
J2	J4	P4	P2	J5	P6
 9,14	 9,22	 9,39	 9,47	 9,67	 9,72
 9,16	 9,30	 9,44	 9,53	 9,71	 9,80
$\bar{x}$ 9,15	$\bar{x}$ 9,26	$\bar{x}$ 9,42	$\bar{x}$ 9,50	$\bar{x}$ 9,69	$\bar{x}$ 9,76

Elaborado por: Moscoso, 2025

Como ejemplo indicaremos la poza 9 que luego del sorteo se identificó como P4, el código representa a el tratamiento de crianza en Piso y la repetición 4; a la poza 11 le correspondió el tratamiento en Jaula y la repetición 5; así sucesivamente.

Tabla 1.11. Consolidado de la data en el sorteo de tratamientos con un DCA para un ensayo de cuyes criados en Jaula y Piso

Tratamientos	Repeticiones						Suma de Tratamientos $\Sigma \tau_i$
	I	II	III	IV	V	VI	
Jaula	8,87	9,15	8,60	9,26	9,69	8,14	53,71
Piso	8,99	9,50	8,16	9,42	8,41	9,76	54,24
Tot al $\Sigma X_{..}$							107,95
Me día $\bar{X}$							9,00

Elaborado por: Moscoso, 2025

Posteriormente, se construye una tabla en donde se identifican los datos de los tratamientos y repeticiones sorteadas y que corresponde a cada unidad experimental (ver tabla (1.11.)

Antes de proceder a calcular los estadísticos en función del modelo matemático, es preciso indicar que para comprobar que el sorteo esté bien realizado, debemos superar los siguientes supuestos:

- Que el coeficiente de variación del primer ADEVA debe mantenerse o descender de primer calculado (datos sueltos sin sorteo)
- El en ese ADEVA (análisis de varianza) la hipótesis nula debe ser verdadera, en otras palabras, el Fisher calculado debe ser menor al Fisher tabular; o el P-value debe superar al nivel de significancia α 0,05.

Para el cálculo del ADEVA en un modelo DCA se procede de la siguiente manera:

Cálculo del factor de corrección FC, tomando en cuenta el número de tratamientos "t" es 2 y con 6 repeticiones (r).

$$F.C. = \frac{(\sum X.)^2}{(t * r)} = \frac{(107,95)^2}{(2 * 6)} = \frac{11653,20}{12} = 971,1$$

Suma de cuadrados totales

$$SC_{TOT} = \{(8,87)^2 + (9,15)^2 + (8,60)^2 + \dots + (9,76)^2\} - 971,10$$

$$SC_{TOT} = 974,64 - 971,10 = 3,54$$

Suma de cuadrados de los tratamientos

$$SC_{TRAT} = \frac{\sum \tau_{ij}^2}{r} - FC$$

$$SC_{TRAT} = \left\{ \frac{[(53,71)^2 + (54,24)^2]}{6} \right\} - 971,10$$

$$SC_{TRAT} = \left\{ \frac{5826,74}{6} \right\} - 971,10 = (971,12 - 971,10) = 0,02$$

Suma de cuadrados del error

$$SC_{ERROR} = SC_{TOTAL} - SC_{TRAT} = (3,54 - 0,02) = 3,52$$

Grados de libertad (GL)

$$GL_{\text{total}} = (n - 1)$$

$$n = tr = (26) = 12$$

$$GL_{\text{totales}} = (12-1) = 11$$

$$GL_{\text{tratamientos}} = (t - 1) = (2 - 1) = 1$$

$$GL_{\text{error}} = GL_{\text{total}} - GL_{\text{tratamientos}} = (11 - 1) = 10$$

Cuadrado Medio (CM)

$$CM_{\text{TRAT}} = \frac{SC_{\text{TRAT}}}{GL_{\text{TRAT}}} = \frac{0,02}{1} = 0,02$$

$$CM_{\text{ERROR}} = \frac{SC_{\text{ERROR}}}{GL_{\text{ERROR}}} = \frac{3,52}{10} = 0,352$$

Fisher calculado

$$F_{\text{CAL}} = \frac{CM_{\text{TRAT}}}{CM_{\text{ERROR}}} = \frac{0,02}{0,352} = 0,06656$$

Coeficiente de variación del ADEVA

$$CV = \frac{\sqrt{CM_{\text{ERROR}}}}{\bar{X}} * 100$$

$$CV = \frac{\sqrt{0,352}}{9,00} * 100 = 6,59\%$$

Tabla 1.12. Resultados del Análisis de Varianza mediante un sorteo de tratamientos con un DCA para un ensayo de cuyes criados en Jaula y Piso

Fuente de variación	Suma de cuadrados	Grados de libertad	Cuadrado medio	Fcal	F.Tab 0,05	F.Tab 0,01	Signif.
Total	3,54	11					
Tratamientos	0,02	1	0,02	0,067	4,96	10,04	NS
Error	3,52	10,00	0,35				

Elaborado por: Moscoso, 2025

Comprobación del sorteo

Como se indicó anteriormente, luego de calcular el primer ADEVA realizado el sorteo al azar mediante un modelo lineal aditivo DCA, se debe comprobar y superar los 2 supuestos:

- El coeficiente de variación de datos sueltos fue de 6,31%; mientras que en el ADEVA se logró un 6,59% muy similar lo que indicaría un procedimiento efectivo.
- Además, el F. Cal 0,067 es menor que el tabular (4,96) por lo tanto no existen diferencias en los tratamientos recién asignados, lo que demostraría a la hipótesis nula como verdadera, ratificando que la asignación de los tratamientos fue positiva.

Jamás deberíamos tener una hipótesis alternativa como verdadera en un sorteo de, ya que aún no se prueban los tratamientos y no contamos con data de variables respuesta. Si no se cumple con este concepto, cometeríamos una variación no pertinente, y desde el inicio el experimento estaría sesgado.

Si luego del ensayo en blanco y el sorteo se obtiene una respuesta antagónica (p-value menor que 0,05), se debería sortear nuevamente a los tratamientos con un modelo diferente, que puede ser un Diseño de Bloques al Azar (DBCA) o el Cuadrado Latino (DCL), en dependencia del coeficiente de variación.

1.8.2. Sorteo en un Diseño de Bloques Completamente al Azar

Se requiere mejorar la ovulación de ovejas Rambouillet mediante 4 dietas de *flushing* (suplementación nutricional realizada en el periodo de pre - empadre y empadre para mejorar la eficiencia reproductiva de hembras ovinas), ya que un problema descubierto por el técnico es la baja concepción de hembras debido a la ausencia de una alimentación que cubra los requisitos proteicos.

Se disponen de 16 ovejas de 3 años y segundo parto, las misma que se convertirán en las unidades experimentales; se las sometió a un ensayo en blanco y los resultados de la ganancia de peso de describen en la tabla 1.13.

Tabla 1.13. Data inicial de la ganancia de peso de ovejas Rambouillet luego del ensayo en blanco, que serán convertidos en unidades experimentales para un ensayo de flushing.

N°	Peso de ovejas, Kg		Ganancia de Peso, Kg.	$(X_i - \bar{x})$	$(X_i - \bar{x})^2$
	Inicial	Final			
1	71,12	74,22	3,10	0,70	0,50
2	70,01	72,90	2,89	0,49	0,24
3	69,02	71,78	2,76	0,36	0,13
4	68,56	71,40	2,84	0,44	0,20
5	68,23	70,76	2,53	0,13	0,02
6	67,23	69,69	2,46	0,06	0,00
7	66,22	68,77	2,55	0,15	0,02
8	65,11	67,60	2,49	0,09	0,01
9	64,56	67,00	2,44	0,04	0,00
10	63,85	66,23	2,38	-0,02	0,00
11	62,48	64,97	2,49	0,09	0,01
12	61,12	63,55	2,43	0,03	0,00
13	60,34	62,01	1,67	-0,73	0,53
14	60,27	62,05	1,78	-0,62	0,38
15	59,89	61,72	1,83	-0,57	0,32
16	58,93	60,62	1,69	-0,71	0,50
Total			38,33	0,00	2,86
Promedio $\bar{x}$			2,44		$\Sigma(X_i - \bar{x})^2$

Elaborado por: Moscoso, 2025

El procedimiento es similar al anterior

Cálculo del promedio

$$\bar{x} = \frac{38,33}{16} = 2,44$$

Desviación estándar en base a la tabla 1.13.

$$\sigma = \sqrt{\frac{2,86}{(16 - 1)}} = 0,4367$$

Procedemos con la obtención del coeficiente de variación que es el indicador de la homogeneidad o heterogeneidad de las que serán unidades experimentales (ovejas Rambouillet).

$$CV = \frac{\sigma}{\bar{x}} \times 100 = \frac{0,4367}{2,44} = 18,23\%$$

Podemos apreciar que el coeficiente de variación es alto, en estos estudios se podría aceptar hasta 15% probablemente para sortear con un DCA. En este caso debemos utilizar el Diseño de Bloques Completos al Azar, que consiste en ordenar de mayor a menor o viceversa los valores de la ganancia de peso (tabla 1.13.) y bloquear en sentido contrario a la pendiente de variación.

Para el diseño del experimento debemos considerar las 4 dietas de flushing como tratamientos, cada uno de ellos con 4 repeticiones lo que indicaría un total de 16 unidades experimentales con un tamaño de 1 oveja madre cada una; al bloquear entonces cada bloque que es sinónimo de “repetición” contará con 1 solamente; lo que indicaría que en un mismo bloque no debemos contar con más de 1 mismo tratamiento. En nuestro ejercicio entonces los 16 datos debemos organizarlos en 4 grupos homogéneos, por ejemplo, un bloque (color verde) puede estar conformado por 4 ovejas con los primeros pesos (3,10; 2,89; 2,76; y 2,84 kilos), como se explica en la figura 1.11.

Figura 1.11. Sorteo de tratamientos con un DBCA para un ensayo de ovejas que pretende probar 4 dietas de flushing con 4 repeticiones cada una.

N°	Ganancia de Peso, gr.	Color del bloque	Pendiente ganancia de peso (de mayor a menor) 
1	3,10	Bloque verde	
2	2,89		
3	2,76		
4	2,84		
5	2,53	Bloque naranja	
6	2,46		
7	2,55		
8	2,49		
9	2,44	Bloque turqueza	
10	2,38		
11	2,49		
12	2,43		
13	1,67	Bloque ciruela	
14	1,78		
15	1,83		
16	1,69		

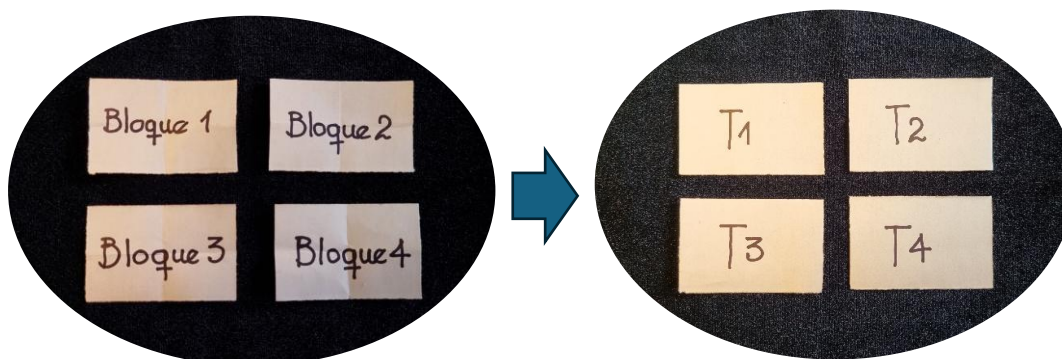


Bloqueo en sentido contrario a la
pendiente (colores)

Elaborado por: Moscoso, 2025

Una vez constituidos los grupos de datos ordenados, se sortea primero el bloque y dentro de cada uno los 4 tratamientos. En el campo pueden utilizarse papeles membretados como se aprecia en la figura 1.11., a los mismos se los dobla y se somete al azar en un recipiente para seleccionar el color del bloque con su respectiva numeración y posteriormente con la misma metodología a los tratamientos codificados desde T1 a T4 estocásticamente en cada bloque como se expresa en la tabla 1.14.

Figura 1.11. Codificación para el sorteo de tratamientos mediante un Diseño de Bloques Completamente al Azar, tanto en bloques como en tratamientos



Elaborado por: Moscoso, 2025

Tabla 1.14. Sorteo con modelo DBCA sobre ganancia de peso de ovejas Rambouillet luego del ensayo en blanco, asignando bloques y tratamientos al azar.

N°	Ganancia de Peso, gr.	Repetición (Bloque)	Tratamiento
1	3,10	3	T3
2	2,89	3	T4
3	2,76	3	T1
4	2,84	3	T2
5	2,53	1	T2
6	2,46	1	T1
7	2,55	1	T4
8	2,49	1	T3
9	2,44	4	T1
10	2,38	4	T3
11	2,49	4	T4
12	2,43	4	T2
13	1,67	2	T4
14	1,78	2	T1
15	1,83	2	T3
16	1,69	2	T2

Elaborado por: Moscoso, 2025

En este caso el bloque verde por efectos del azar le correspondió la repetición 3 y dentro del mismo a la oveja 1 con 3,10 kg de ganancia de peso se le asigna el tratamiento T3; a la 2 (2,89 kg) el tratamiento T4; a la 3 (2,76 Kg) le tocó la dieta T1, mientras que la cuarta oveja que tuvo una ganancia de 2,84 kg en el sorteo alcanzó el tratamiento T2.

Es importante indicar que en el modelo lineal aditivo del DBCA, se ha sumado una fuente de variación adicional que son las repeticiones o bloques, quedando nuestra ecuación de la siguiente manera:

$$Y_{ij} = \mu + \tau_i + \beta_j + \epsilon_{ij}$$

Donde:

- Y_{ij} = Variable ganancia de peso que corresponde a la observación de la i-ésima repetición del i-ésimo tratamiento
- μ = Media general
- τ_i = Efecto del i-ésimo tratamiento (dietas de flushing)
- β_j = Efecto de las j-ésimas repeticiones (bloques)
- ϵ_{ij} = Efecto aleatorio de la ij-ésima unidad experimental

Entonces la data generada en el ensayo en blanco; antes del inicio del experimento propiamente dicho quedaría como se expresa en la tabla 1.15.

Tabla 1.15. Data consolidada de la ganancia de peso de ovejas Rambouillet luego del ensayo en blanco, asignando al modelo DBCA, para el sorteo de tratamientos (dietas de flusing)

Tratamientos	Repeticiones o bloques				Suma de tratamientos
	I	II	III	IV	
T1 (Dieta flushing 1)	2,46	1,78	2,76	2,44	9,44
T2 (Dieta flushing 2)	2,53	1,69	2,84	2,43	9,49
T3 (Dieta flushing 3)	2,49	1,83	3,1	2,38	9,80
T4 (Dieta flushing 4)	2,55	1,67	2,89	2,49	9,60
Suma de repeticiones	10,03	6,97	11,59	9,74	$\Sigma X..$ 38,33
Media $\bar{X}$					2,40

Elaborado por: Moscoso, 2025

Para el tratamiento de la información del ADEVA en un modelo DBCA se procederá con el siguiente protocolo:

Cálculo del factor de corrección FC, considerando el número de tratamientos “t” es 4 y 4 repeticiones (r).

$$F.C. = \frac{(\Sigma X..)^2}{(t * r)} = \frac{(38,33)^2}{(4 * 4)} = \frac{1469,19}{16} = 91,82$$

Suma de cuadrados totales

$$SC_{TOT} = \{(2,46)^2 + (1,78)^2 + (2,76)^2 + \dots + (2,89)^2\} - 2,49$$

$$SC_{TOT} = 94,69 - 91,82 = 2,86$$

Suma de cuadrados de los tratamientos

$$SC_{TRAT} = \frac{\sum \tau_i^2}{r} - FC$$

$$SC_{TRAT} = \left\{ \frac{[(9,44)^2 + (9,49)^2 + (9,80)^2 + (9,60)^2]}{4} \right\} - 91,82$$

$$SC_{TRAT} = \left\{ \frac{367,37}{4} \right\} - 91,82 = (91,84 - 91,82) = 0,0191$$

Suma de cuadrados de las repeticiones

$$SC_{REP} = \frac{\sum \beta_j^2}{t} - FC$$

$$SC_{REP} = \left\{ \frac{[(10,03)^2 + (6,97)^2 + (11,59)^2 + (9,74)^2]}{4} \right\} - 91,82$$

$$SC_{REP} = \left\{ \frac{378,38}{4} \right\} - 91,82 = (94,59 - 91,82) = 2,7701$$

Suma de cuadrados del error

$$SC_{ERROR} = SC_{TOTAL} - SC_{TRAT} - SC_{REP} = (2,86 - 0,02 - 2,77) = 0,0722$$

Grados de libertad (GL)

$$GL_{total} = (n - 1)$$

$$n = tr = (44) = 16$$

$$GL_{totales} = (16-1) = 15$$

$$GL_{tratamientos} = (t - 1) = (4 - 1) = 3$$

$$GL_{repeticiones} = (r - 1) = (4 - 1) = 3$$

$$GL_{error} = GL_{total} - GL_{tratamientos} - GL_{repeticiones} = (15 - 3 - 3) = 9$$

Cuadrado Medio (CM)

$$CM_{TRAT} = \frac{SC_{TRAT}}{GL_{TRAT}} = \frac{0,0191}{3} = 0,0064$$

$$CM_{REP} = \frac{SC_{REP}}{GL_{REP}} = \frac{2,7701}{3} = 0,9234$$

$$CM_{ERROR} = \frac{SC_{ERROR}}{GL_{ERROR}} = \frac{0,07222}{9} = 0,0080$$

Fisher calculado

$$F_{CAL-TRAT} = \frac{CM_{TRAT}}{CM_{ERROR}} = \frac{0,0064}{0,0080} = 0,07993$$

Coefficiente de variación del ADEVA

$$CV = \frac{\sqrt{CM_{ERROR}}}{\bar{X}} * 100 = \frac{\sqrt{0,0080}}{2,3956} * 100 = 3,74\%$$

Tabla 1.16. Resultados del Análisis de Varianza mediante un sorteo de tratamientos con un DBCA para un ensayo de flushing en ovejas Rambouillet.

Fuente de variación	Suma de cuadrados	Grados de libertad	Cuadrado medio	Fcal	F.Tab 0,05	F.Tab 0,01	Signif .
Total	2,861	15					
Tratamientos	0,019	3	0,006	0,79	3,86	6,99	NS
Repeticiones	2,770	3	0,923	115,09	3,86	6,99	
Error	0,072	9	0,008				

Elaborado por: Moscoso, 2025

Comprobación del sorteo

Luego del “bloqueo” de la fuente de variación (repeticiones o bloques) y el cálculo del primer ADEVA realizado el sorteo al azar mediante un modelo lineal aditivo DBCA, se comprobarán los 2 supuestos:

- El coeficiente de variación inicial de datos sueltos fue de 18,23%; considerado muy alto por lo que se decidió cambiar al modelo DBCA. Si comparamos con el estadístico del ADEVA en el que se obtuvo un 3,74% podemos con certeza indicar que se procedió debidamente.
- Por otro lado, el F. Cal de los tratamientos 0,794 es menor que el tabular (3,86), es decir no existen diferencias significativas en los tratamientos recién sorteados, lo que señalaría a la hipótesis nula como verdadera, confirmando que la asignación de los tratamientos fue positiva con el modelo indicado.

Por otro lado, si se aprecia en la tabla 1.16., al Fisher de las repeticiones evidentemente es extremadamente alto superando a los valores tabulares, esto indica que la variación ha sido corregida con la incorporación de los “bloques o repeticiones” en el modelo lineal aditivo.

Referencias bibliográficas del capítulo

- Anguera, M. T., Blanco-Villaseñor, A., Losada, J. L., & Sánchez-Algarra, P. (2014). Diseños observacionales: ajuste y aplicación en psicología del deporte. Cuadernos de psicología del deporte, 63-76.
- Bacon, F. (1620). Novum Organum Scientiarum. Francifcum Moiardum.
- Borlaug, N. (2007). Feeding a Hungry World. Science.
- Casanoves, F. (2019). Diplomado Internacional en Bioestadística. Análisis estadístico en estudios comparativos. Turrialba, Costa Rica: CATIE.
- Chalmers, A. (2013). What is This Thing Called Science? Open University Press. Indianapolis, Cambridge: University of Queensland Press.
- Comunicación institucional. (2020). Ibero Tijuana. Obtenido de <https://blogposgrados.tijuana.iberomx/investigacion-aplicada/>

- Creswell, J. W. (2014). Research Design: Qualitative, Quantitative, and Mixed Methods Approaches. Los Angeles, USA: SAGE Publications.
- De La Loma, J. L. (1996). Experimentación agrícola. México: Hispano América.
- Di Rienzo, J., & Casanoves, F. (1999). Estadística para las ciencias agropecuarias. Córdoba, Argentina: Brujas.
- Fisher, R. (1937). Design of experiments. Edinburgo: Oliver & Boyd.
- Gómez, D. et. al. (2013). Introducción a la Estadística. Alicante, España: Universidad de Alicante.
- González, G. (1976). Métodos Estadísticos y principios de Diseño Experimental. Quito, Ecuador: Universidad Central del Ecuador.
- Hernández Sampieri, R., Fernández Collado, C., & Baptista Lucio, P. (2014). Metodología de la investigación. México, México: McGraw-Hill.
- Hernández, R., Fernández, C., & Baptista, P. (2014). Metodología de la investigación. Mexico: McGraw-Hill.
- Kerlinger, F. N., & Lee, H. B. (2002). Foundations of Behavioral Research. Harcourt Brace. Méico, México: McGraw-Hill.
- Little, T., & Hills, J. (1991). Métodos estadísticos para la investigación en la agricultura. California: Trillas.
- Montgomery, D. (2019). Design and analysis of experiments. Nueva Delhi: John Wiley & Sons.
- Morales, J. (2023). Modelos estadísticos con R. Obtenido de Modelos aditivos lineales: https://bookdown.org/j_morales/weblinmod/05ML-SMOOTH.html
- Moscoso G., M., Moreno, M., Moscoso M., N., & Armijos, R. (2022). Metodología de la investigación científica y su aplicación en las ciencias agropecuarias. Riobamba: ESPOCH.

- Moscoso Gómez, M. E. (2019). Modelo de uso racional para el desarrollo sustentable de agroecosistemas de "El Guzo" cantón Penipe, Ecuador. Lima: Universidad Agraria La Molina.
- Murcia, N. Savio, M., Cora, F., & Beneitez, A. (2021). Principios Básicos de Nutrición Porcina. La Pampa, Argentina: INTA.
- Ponce, M. (2014). Caos y ciencia. Obtenido de <https://www.caosyciencia.com/investigacion-tecnologica/>
- Popper, K. (2002). The Logic of Scientific Discovery. Londres: Taylor & Francis Group.
- Reyes, P. (2003). Diseño de epeprimentos aplicados. México: Trillas.
- ShallBD. (31 de marzo de 2023). Local regressiom: choises. Obtenido de shallbd.com: <https://shallbd.com/es/comprender-los-modelos-aditivos-definicion-y-aplicaciones/>
- Toledo, G. (31 de enero de 2018). Hugepdf. Obtenido de Bioestadística: Tema 1: El método Estadístico, partes de la Estadística: https://hugepdf.com/download/tema-1-el-metodo-estadistico-partes-de-la_pdf
- Triola, M. F. (2018). Estadística. Madrid, España: Pearson Educación.
- Walpole, R. Myers, R. & Myers, S. (2012). Probability and statistics for engineers and scientistis. Boston, USA: Global edition.



CAPÍTULO II

EXPLORACIÓN Y
TRATAMIENTO DE
DATOS
EXPERIMENTALES

La exploración y tratamiento de datos experimentales es un componente esencial en el proceso de investigación científica. En un mundo donde la cantidad de datos generados es cada vez mayor, la capacidad para manejar, interpretar y extraer conclusiones significativas de estos datos se ha convertido en una habilidad crucial para investigadores en diversas disciplinas. Desde la biología hasta la ingeniería, la correcta manipulación de datos experimentales no solo permite validar hipótesis, sino que también impulsa el avance del conocimiento y la innovación.

Los datos experimentales son el resultado de observaciones y mediciones realizadas durante un estudio que se expresan en las variables respuesta (ganancia de peso, conversión alimenticia, costo por kilogramo de ganancia de peso, etc.). Estos pueden abarcar una amplia gama de formatos, desde números y estadísticas hasta imágenes y textos. La calidad y la precisión de estos datos son fundamentales, ya que cualquier error en la recolección o el tratamiento puede llevar a conclusiones incorrectas y, en última instancia, a decisiones erróneas. Por lo tanto, la exploración y el tratamiento de datos son pasos críticos en el ciclo de vida de la investigación.

La exploración de datos implica el análisis preliminar de los datos recolectados para comprender su estructura, patrones y características. Este proceso es vital para identificar tendencias, anomalías y relaciones entre variables. A través de técnicas de visualización, como gráficos y diagramas, los investigadores pueden obtener una representación visual de los datos, facilitando la identificación de patrones que podrían no ser evidentes a simple vista. Algunas herramientas y técnicas comunes en la exploración de datos incluyen:

1. **Estadísticas Descriptivas:** Estas proporcionan un resumen de las características principales de los datos, como la media, la mediana, la moda, la varianza y los percentiles. Estas métricas ayudan a comprender la distribución y la dispersión de los datos.
2. **Visualización de Datos:** Gráficos de dispersión, histogramas, diagramas de caja y gráficos de líneas son ejemplos de visualizaciones que permiten a los investigadores observar relaciones y tendencias en los datos.

3. Análisis de Correlación: Este análisis ayuda a determinar si existe una relación entre dos o más variables, lo que puede ser crucial para formular hipótesis y entender el fenómeno estudiado.
4. Detección de Anomalías: Identificar valores atípicos o anomalías en los datos es esencial, ya que estos pueden influir significativamente en los resultados y conclusiones del estudio.

Una vez que se ha explorado y comprendido la estructura de los datos, el siguiente paso es el tratamiento de los mismos. Este proceso implica la limpieza, transformación y análisis de los datos con el fin de prepararlos para un análisis más profundo. El tratamiento adecuado de los datos es fundamental para asegurar la validez y la confiabilidad de los resultados. Algunos pasos clave en el tratamiento de datos incluyen:

1. Limpieza de Datos: Este paso implica la identificación y corrección de errores en los datos, como valores faltantes, duplicados o inconsistencias. La limpieza es esencial para garantizar su precisión y representatividad.
2. Transformación de Datos: A menudo, los datos necesitan ser transformados para ser utilizados en análisis estadísticos o modelos. Esto puede incluir la normalización, estandarización o la creación de nuevas variables a partir de las existentes.
3. Análisis Estadístico: Dependiendo de los objetivos de la investigación, se pueden aplicar diferentes métodos estadísticos para analizar los datos. Esto puede incluir pruebas de hipótesis, análisis de varianza (ANOVA), regresión y análisis multivariante, entre otros.
4. Modelado Predictivo: En muchos casos, los investigadores utilizan técnicas de modelado para predecir resultados basados en los datos experimentales. Esto puede incluir la creación de modelos de regresión, árboles de decisión o redes neuronales, dependiendo de la complejidad del problema.

A pesar de la importancia de la exploración y tratamiento de datos, existen varios desafíos que los investigadores deben enfrentar. Uno de los principales desafíos es la calidad de los datos. Los datos incompletos o inexactos pueden llevar a

conclusiones erróneas, lo que subraya la necesidad de un enfoque riguroso en la recolección y limpieza de datos.

Otro desafío es la complejidad de los métodos estadísticos y analíticos. Los investigadores deben tener una comprensión sólida de las técnicas que utilizan para asegurarse de que están aplicando los métodos adecuados para sus datos y preguntas de investigación.

Además, la creciente cantidad de datos generados en diversas disciplinas plantea un desafío adicional. La capacidad para manejar grandes volúmenes de datos, conocidos como "big data", requiere herramientas y técnicas avanzadas que pueden no estar al alcance de todos los investigadores.

La exploración y tratamiento de datos experimentales son procesos críticos en la investigación científica que permiten a los investigadores extraer conclusiones significativas y validar hipótesis. A medida que la cantidad de datos disponibles sigue creciendo, la capacidad para manejar y analizar estos datos se vuelve cada vez más crucial. A través de técnicas adecuadas de exploración y tratamiento, los investigadores pueden contribuir al avance del conocimiento y la innovación en sus respectivos campos, asegurando que los datos se utilicen de manera efectiva para abordar los desafíos del mundo real.

2.1. Obtención y organización de la “data”

La obtención y organización de datos es un proceso fundamental en la investigación científica y en la toma de decisiones informadas en diversos campos, desde la salud hasta la ingeniería y las ciencias sociales. En un mundo donde la información se genera a un ritmo acelerado, la capacidad para recolectar, gestionar y estructurar datos de manera efectiva se ha vuelto crucial. Este epígrafe explora las metodologías y herramientas actuales para la obtención y organización de datos, así como su importancia en el contexto de la investigación moderna.

La obtención de datos puede realizarse a través de diversas técnicas y metodologías, que se pueden clasificar en dos categorías principales: datos primarios y datos secundarios. Los datos primarios son aquellos que se

recolectan directamente de la fuente a través de métodos como encuestas, entrevistas, experimentos y observaciones.

La recolección de datos primarios permite a los investigadores obtener información específica y relevante para sus preguntas de investigación.

Las encuestas son una herramienta común para la recolección de datos primarios. Se diseñan para recopilar información sobre actitudes, comportamientos y características de una población específica. Según Galesic y Bosnjak (2019), la calidad de los datos obtenidos a través de encuestas depende en gran medida del diseño del cuestionario y de la forma en que se administran las encuestas.

Los experimentos controlados son otra forma de obtener datos primarios. En un entorno experimental, los investigadores manipulan variables independientes para observar su efecto en variables dependientes. Este enfoque permite establecer relaciones causales y es fundamental en disciplinas como la psicología y la medicina (Field, 2020).

Los datos secundarios son aquellos que han sido recolectados por otros investigadores o entidades. Estos datos pueden provenir de bases de datos, informes estadísticos, artículos académicos y registros administrativos. La utilización de datos secundarios puede ser ventajosa, ya que permite a los investigadores acceder a información que de otro modo sería costosa o difícil de obtener.

Existen numerosas bases de datos disponibles que contienen datos relevantes para diversas investigaciones. Por ejemplo, la base de datos PubMed proporciona acceso a una amplia gama de investigaciones en el campo de la salud, mientras que el Banco Mundial ofrece datos económicos y sociales a nivel global. Según Chen et al. (2021), el uso de datos secundarios puede acelerar el proceso de investigación y proporcionar una base sólida para nuevos estudios.

La revisión de literatura también es una forma de obtener datos secundarios. A través de un análisis exhaustivo de estudios previos, los investigadores pueden

identificar tendencias, vacíos en la investigación y áreas que requieren más atención (Hernández et al., 2020).

Una vez que se han obtenido los datos, es crucial organizarlos de manera efectiva para facilitar su análisis y uso posterior. La organización de datos implica la estructuración, almacenamiento y gestión de la información recolectada.

La estructuración de datos se refiere a la forma en que se organizan y presentan los datos. Esto puede incluir la creación de bases de datos, tablas y hojas de cálculo que permitan un fácil acceso y análisis. Según Kitchin (2019), una buena estructuración de datos no solo mejora la eficiencia del análisis, sino que también minimiza el riesgo de errores.

Las bases de datos relacionales son una forma común de organizar datos. Utilizan tablas para almacenar información y permiten la creación de relaciones entre diferentes conjuntos de datos. Esto facilita la consulta y el análisis de datos de manera eficiente (Elmasri & Navathe, 2016).

Las hojas de cálculo son otra herramienta popular para la organización de datos. Permiten a los investigadores ingresar y manipular datos de manera sencilla. Sin embargo, es importante tener en cuenta que, aunque las hojas de cálculo son útiles para análisis simples, pueden volverse ineficaces para conjuntos de datos grandes o complejos (Pérez, 2020).

El almacenamiento de datos es un aspecto crítico que debe considerarse durante la organización de la información. Los investigadores deben asegurarse de que sus datos estén almacenados de manera segura y accesible.

El almacenamiento en la nube ha ganado popularidad en los últimos años, ya que permite a los investigadores acceder a sus datos desde cualquier lugar y en cualquier momento. Plataformas como Google Drive y Dropbox ofrecen soluciones de almacenamiento escalables y seguras (Liu et al., 2021).

Aunque el almacenamiento en la nube es conveniente, algunas organizaciones prefieren almacenar datos en bases de datos locales por razones de seguridad y control. Esto es especialmente relevante en campos como la salud, donde la privacidad de los datos es de suma importancia (Smith et al., 2022).

La gestión de datos implica el mantenimiento y la actualización de los datos a lo largo del tiempo. Esto incluye la implementación de protocolos para la limpieza de datos, la verificación de la calidad de los datos y la documentación de los procesos de recolección y organización.

La limpieza de datos es un paso esencial en la organización de datos. Implica la identificación y corrección de errores, valores faltantes y duplicados. Según Rahm y Do (2019), la limpieza de datos puede mejorar significativamente la calidad de los análisis y resultados.

La documentación adecuada de los procesos de obtención y organización de datos es crucial para la reproducibilidad de la investigación. Esto incluye la creación de metadatos que describan el contenido, la estructura y la calidad de los datos (Higgins et al., 2020).

La obtención y organización de datos son procesos interrelacionados que impactan directamente en la calidad y la validez de la investigación. Una adecuada obtención de datos permite a los investigadores abordar preguntas de investigación relevantes y significativas, mientras que una organización efectiva facilita el análisis y la interpretación de los resultados.

Además, la transparencia en los procesos de obtención y organización de datos es fundamental para la credibilidad de la investigación. La capacidad de otros investigadores para replicar estudios y validar resultados depende en gran medida de la claridad y la precisión en la documentación de los datos (Stodden et al., 2016).

La obtención y organización de datos son componentes críticos en el proceso de investigación que determinan la calidad y la validez de los resultados. A medida que la cantidad de datos disponibles sigue creciendo, la capacidad para manejar y estructurar esta información de manera efectiva se vuelve cada vez más crucial. La implementación de metodologías adecuadas y el uso de herramientas tecnológicas modernas son esenciales para optimizar estos procesos y garantizar que los datos se utilicen de manera efectiva en la toma de decisiones informadas.

Para el caso específico de la investigación agropecuaria, la data se obtiene en el período conocido como “trabajo de campo” en donde se obtendrá información de las variables respuesta cuali y cuantitativas según el caso, evitando caer en los llamados errores no pertinentes, como pueden ser uso de balanzas sin calibración para evaluar el peso de animales, no completar la toma de datos en todas las unidades experimentales en mismo día, cambio frecuente de operarios para el manejo de los animales que puede producir estrés, deterioro de las hojas de campo con datos anteriores que no fueron trasladados oportunamente a una base, etc.

En el campo se utiliza comúnmente una libreta para anotaciones donde se escriben los datos en base a un calendario establecido (cada semana por ejemplo); sin embargo, los investigadores suelen acumular información y luego digitalizar en una hoja de cálculo; se recomienda que tras cada evento de evaluación la información sea trasladada a la base informática y no descartar las libretas de campo para futuras verificaciones si es que las hubiere.

Los equipos digitales que en la actualidad se usan para toma de datos en experimentos agrícolas deben ser verificados, calibrados y/o encerados para evitar sesgos en la información, y se recomienda que el mismo equipo se use en toda el trabajo de campo.

2.2. Análisis descriptivo de la “data”

El análisis descriptivo de datos es un paso crucial en la investigación científica, especialmente en la experimentación animal. Este tipo de análisis permite resumir y visualizar las características principales de los datos recolectados, facilitando la comprensión de patrones y tendencias antes de realizar análisis más complejos.

En el contexto de la experimentación animal, los modelos lineales aditivos (GLM) son herramientas poderosas que permiten analizar la relación entre variables y explorar cómo diferentes factores influyen en los resultados observados.

El análisis descriptivo implica la utilización de estadísticas y visualizaciones para resumir y presentar datos de manera efectiva. Esto incluye el cálculo de medidas de tendencia central, como la media, la mediana y la moda, así como medidas

de dispersión, como la varianza y el rango intercuartílico. Estas métricas proporcionan un primer vistazo a la distribución de los datos y ayudan a identificar patrones significativos.

Media: Es el promedio aritmético de un conjunto de datos. Por ejemplo, si se mide el peso de un grupo de conejos en un experimento, la media proporcionará una idea general del peso promedio de los animales en el estudio (Cohen et al., 2021).

Mediana: Es el valor que divide el conjunto de datos en dos mitades. En estudios donde los datos pueden estar sesgados por valores extremos, la mediana puede ser una mejor representación del centro de los datos (Hernández et al., 2020).

Moda: Es el valor que aparece con mayor frecuencia en el conjunto de datos. Por ejemplo, si en un experimento de comportamiento animal se observa que un determinado comportamiento se repite con más frecuencia, la moda puede indicar este hallazgo (Smith & Jones, 2022).

Además, las medidas de dispersión como la varianza, mide la variabilidad de los datos respecto a la media. Una alta varianza indica que los datos están muy dispersos, lo que puede ser relevante en la experimentación animal al evaluar la respuesta a un tratamiento (López et al., 2021).

Entre tanto la desviación estándar, proporciona una medida de la dispersión en las mismas unidades que los datos originales. Por otra parte, el Rango intercuartílico (IQR), mide la diferencia entre el tercer cuartil (Q3) y el primer cuartil (Q1), proporcionando una indicación de la dispersión de los datos en el medio 50%.

Para ilustrar cómo se aplica el análisis descriptivo en la experimentación animal utilizando modelos lineales aditivos, consideremos algunos ejemplos prácticos.

Ejemplo 1

- Efecto de un Tratamiento Farmacológico en el Comportamiento

Se supone que un investigador está analizando el efecto de un nuevo fármaco en el comportamiento de conejos en un laberinto. Se recolectan datos sobre el

tiempo que tardan los conejos en completar el laberinto, así como otras variables como la cantidad de errores cometidos.

Resolución

1. Análisis Descriptivo:

- **Media:** Se calcula la media del tiempo de finalización del laberinto para cada grupo de tratamiento (fármaco vs. placebo).
- **Mediana:** Se determina la mediana para identificar si hay algún sesgo en los tiempos de finalización.
- **Varianza y Desviación Estándar:** Se evalúa la varianza para entender la consistencia de los tiempos de finalización entre los sujetos.

2. Visualización:

- **Gráfico de Caja:** Un gráfico de caja puede mostrar la mediana, los cuartiles y los posibles extremos en los tiempos de finalización, permitiendo una comparación visual clara entre los grupos.

3. Interpretación:

- Los resultados descriptivos pueden indicar si el fármaco ha tenido un efecto significativo en la reducción del tiempo de finalización del laberinto.

Ejemplo 2

Estudio de la eficiencia alimentaria en animales de granja

En un estudio sobre la eficiencia alimentaria en pollos, se recolectan datos sobre la cantidad de alimento consumido y el peso ganado durante un período específico.

Resolución

1. Análisis Descriptivo:

- **Media:** Se calcula la media del peso ganado por pollo y la media del alimento consumido.
- **Desviación Estándar:** Se evalúa la desviación estándar del peso ganado para determinar la variabilidad entre los pollos.

2. Visualización:

- **Histogramas:** Se utilizan histogramas para visualizar la distribución del peso ganado y la cantidad de alimento consumido. Esto puede revelar si hay un sesgo en los datos o si siguen una distribución normal.

3. Interpretación:

- Un análisis descriptivo puede ayudar a identificar patrones en la eficiencia alimentaria, lo que puede ser útil para ajustar las dietas y mejorar la producción.

1.1.1. Ejercicios y Soluciones

Ejercicio 1

Efecto de un Tratamiento Farmacológico en el Comportamiento

Supongamos que se han recolectado los siguientes datos sobre el tiempo (en segundos) que tardan conejos en completar un laberinto, divididas en dos grupos: tratamiento y control.

- **Grupo Tratamiento:** 30, 28, 32, 29, 31
- **Grupo Control:** 35, 34, 36, 37, 38

Preguntas:

1. Calcular la media, mediana y desviación estándar para ambos grupos.
2. Crear un gráfico de caja para visualizar los datos.

Soluciones:**1. Cálculos:**

- **Grupo Tratamiento:**

- Media: $30+28+32+29+31=150$
 $150/5=30$
- Mediana: 30 (valor medio de los datos ordenados: 28, 29, 30, 31, 32)
- Desviación estándar:
 - Varianza: $(30-30)^2+(28-30)^2+(32-30)^2+(29-30)^2+(31-30)^2=25$
 $25/5=5$
 - Desviación estándar: $2\approx 1.412\approx 1.41$

- **Grupo Control:**

- Media: $35+34+36+37+38=170$
 $170/5=34$
- Mediana: 36 (valor medio de los datos ordenados: 34, 35, 36, 37, 38)
- Desviación estándar:
 - Varianza: $(35-36)^2+(34-36)^2+(36-36)^2+(37-36)^2+(38-36)^2=25$
 $25/5=5$
 - Desviación estándar: $2\approx 1.412\approx 1.41$

2. Gráfico de Caja:

- Se crea un gráfico de caja utilizando software como R, Python o Excel, donde se representen los tiempos de finalización para ambos grupos, mostrando la mediana y los cuartiles.

Ejercicio 2

Estudio de la eficiencia alimentaria en animales de granja

Se han recolectado los siguientes datos sobre el peso ganado (en kg) por un grupo de pollos durante un período de alimentación:

- **Pesos Ganados:** 1.5, 1.6, 1.4, 1.7, 1.8, 1.6, 1.5, 1.7

Preguntas:

1. Calcular la media, mediana y rango intercuartílico (IQR).
2. Crear un histograma para visualizar la distribución del peso ganado.

Soluciones:

1. Cálculos:

- **Media:** $1.5+1.6+1.4+1.7+1.8+1.6+1.5+1.7=15.75$
 $15.75/8=1.96875$
- **Mediana:** (1.5, 1.5, 1.6, 1.6, 1.7, 1.7, 1.8)
 $\rightarrow 1.6+1.6=3.2$
 $3.2/2=1.6$
- **Rango Intercuartílico (IQR):**
 - Q1 (1.5) y Q3 (1.7) $\rightarrow IQR = 1.7-1.5=0.2$

2. Histograma:

- Se crea un histograma utilizando software como R, Python o Excel, donde se representen los pesos ganados, mostrando la frecuencia de cada rango de peso.

Ejemplo 3

- Aplicación de la estadística descriptiva en resultados experimentales de ganancias de peso de cuyes peruanos sometidos a 4 dietas alimenticias.

En el establecimiento del experimento se pretendió resolver el problema de falta proteica en la dieta de los cobayos, por lo que se prepararon 4 dietas (14%, 16%, 18% y 20%) con 4 repeticiones cada una y los resultados fueron:

Tratamiento 1 (14% de proteína): 70, 61, 49, 25

Tratamiento 2 (16% de proteína): 66, 59, 45, 69

Tratamiento 3 (18% de proteína): 72, 51, 77, 82

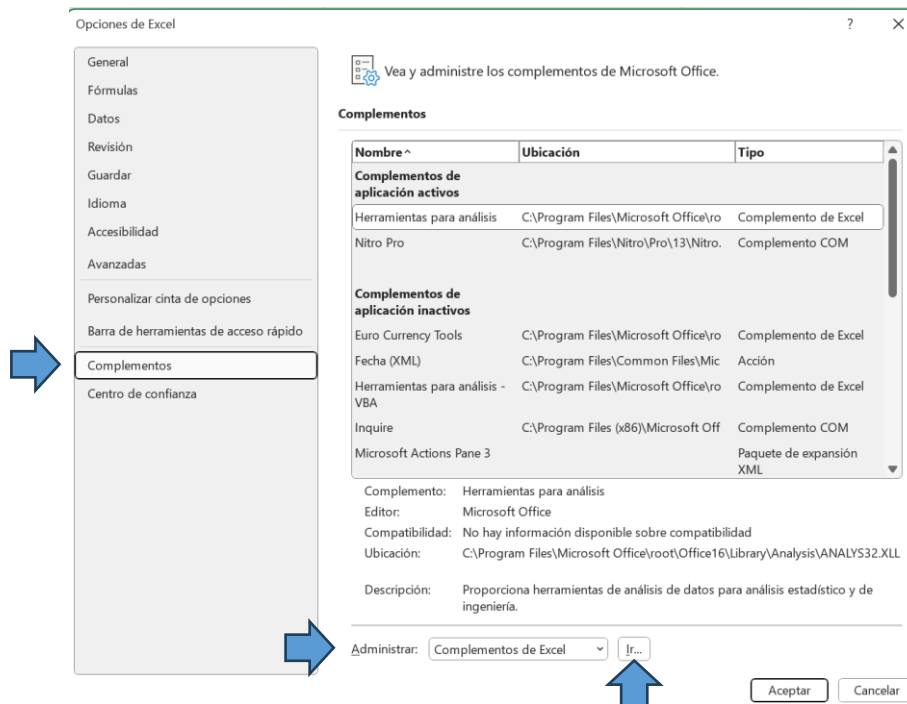
Tratamiento 4 (20% de proteína): 95, 92, 20, 98

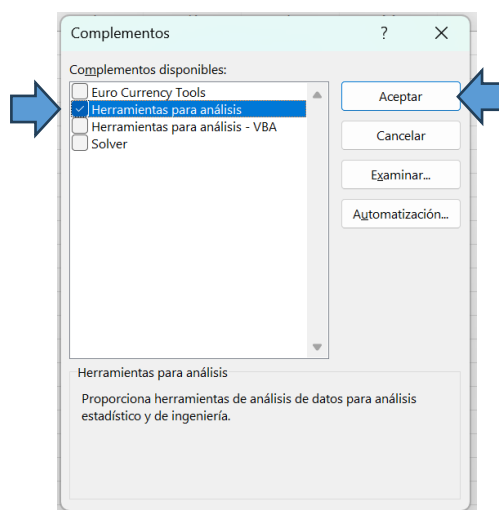
Resolución

1. Análisis Descriptivo:

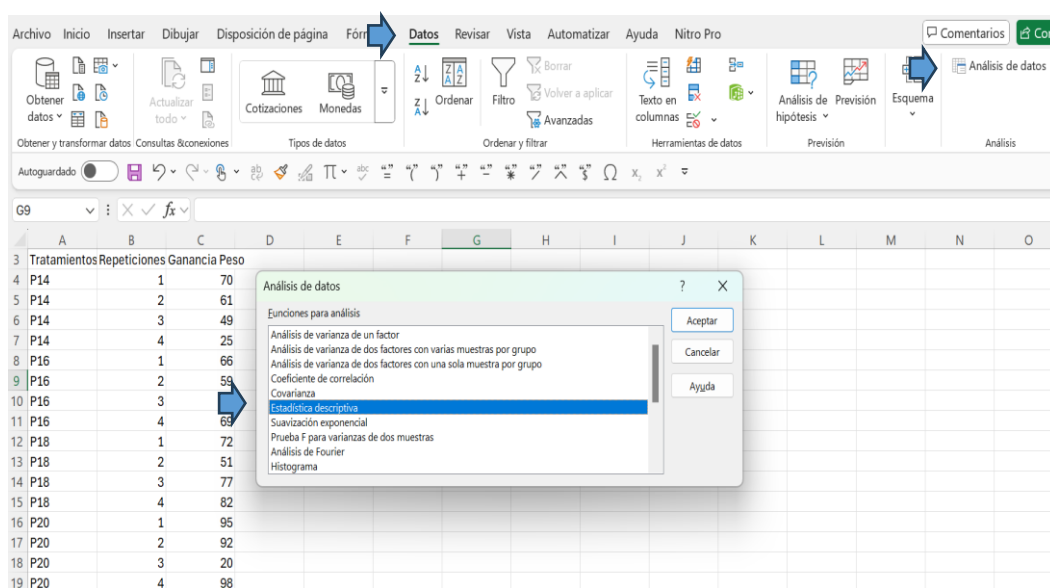
Para este análisis se tomarán en cuenta todos los datos obtenidos discriminando los tratamientos, para establecer una descripción mediante las medidas de tendencia central, de dispersión y de simetría.

Se puede hacer uso del Excel utilizando la función “análisis de datos”, la ventana se activa desde la pestaña: archivo, opciones, complementos, ir a complementos y activar herramientas para análisis.

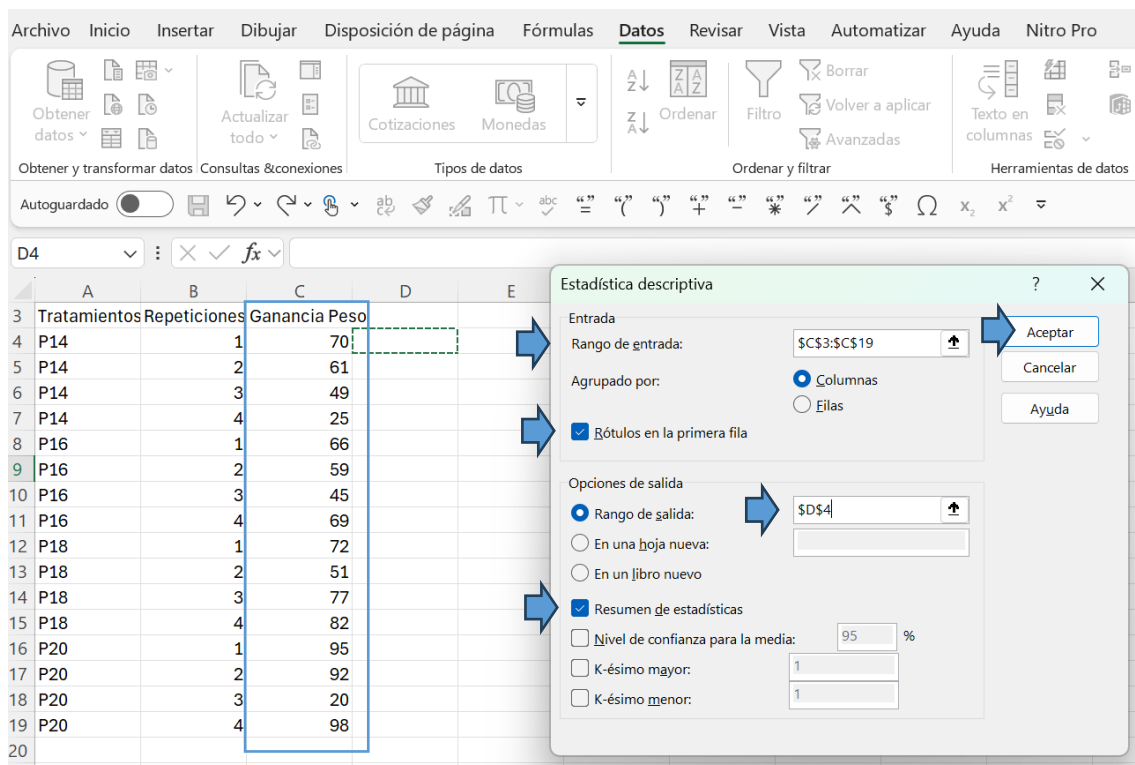




Se procede desde la pestaña datos, análisis de datos, estadística descriptiva.



Luego se define el rango de entrada que en el caso que amerita es la variable ganancia de peso, luego se activa rótulos en la primera fila puesto que la selección es desde el nombre de la variable; posteriormente se establece el rango de salida en una celda vacía y se activa resumen de estadísticas.

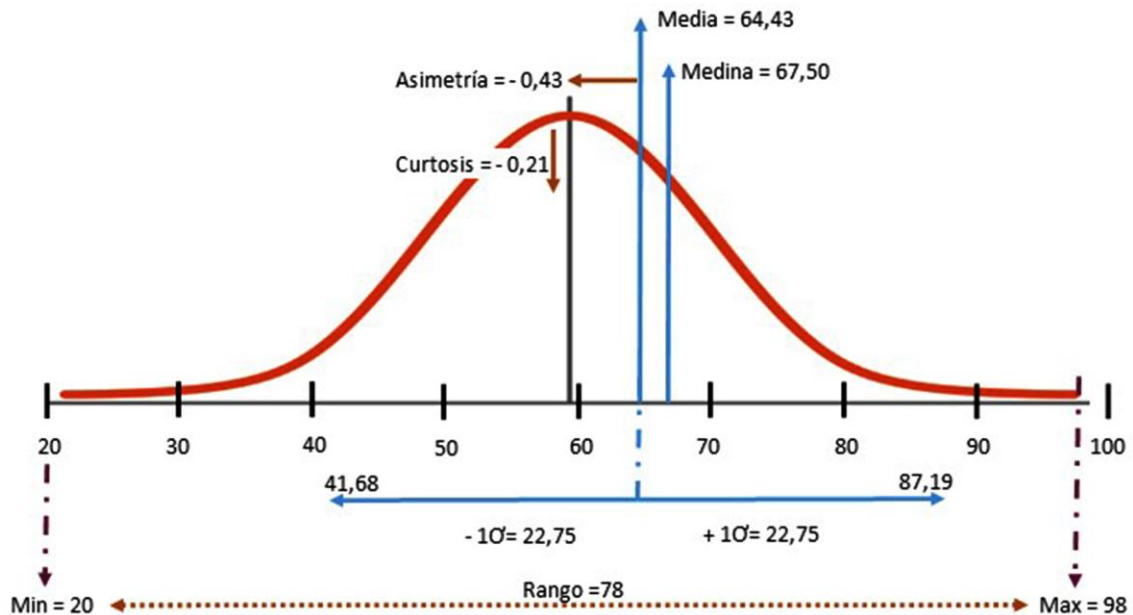


Se visualizarán los resultados y se calcula el coeficiente de variación dividiendo la celda de la desviación estándar para la media.

Ganancia Peso	
Media	64,4375
Error típico	5,68841568
Mediana	67,5
Moda	#N/D
Desviación estándar	22,7536627
Varianza de la muestra	517,729167
Curtosis	-0,2166896
Coeficiente de asimetría	-0,435552
Rango	78
Mínimo	20
Máximo	98
Suma	1031
Cuenta	16
CV	0,35311213

2. Visualización:

- **Gráfico de la distribución:** Un gráfico puede darnos una idea de los estadísticos y una primera aproximación de normalidad, por ejemplo.



3. Interpretación:

- En la gráfica se puede comprobar que no hay total normalidad, por un lado, las medidas de tendencia central no coinciden, la distribución es amodal, la desviación es grande, su coeficiente de variación pasa del 35%, además no hay simetría, puesto que hay una desviación hacia la izquierda y una ligera platicurtosis. Este primer análisis puede dirigirnos a pensar que probablemente se deberán suavizar los datos antes de proceder a aplicar cualquier modelo lineal aditivo que verifique el efecto de los tratamientos; para cumplir con este procedimientos existen varias metodologías para con certeza alcanzar los 3 supuestos que requerimos.

2.3.Exploración de Supuestos en Análisis Estadístico: Normalidad, Homogeneidad de la Varianza e Independencia

La validación de supuestos es un aspecto fundamental en el análisis estadístico, especialmente en los contextos de investigación científica. La correcta interpretación de los resultados de un estudio depende en gran medida de la veracidad de los supuestos bajo los cuales se realizan los análisis.

En este epígrafe, se explorarán tres supuestos clave: normalidad, homogeneidad de la varianza y la independencia. Se discutirá su importancia, métodos para evaluarlos y su implicación en el análisis de datos, con un enfoque en investigaciones recientes.

Los supuestos estadísticos son condiciones que deben cumplirse para que los resultados de un análisis sean válidos y generalizables. Cuando estos supuestos son violados, los resultados pueden ser engañosos, lo que lleva a conclusiones incorrectas. Por lo tanto, es esencial que los investigadores realicen una evaluación exhaustiva de estos supuestos antes de proceder con el análisis de datos.

El supuesto de normalidad establece que los datos deben seguir una distribución normal. Este supuesto es crítico para muchos métodos estadísticos, incluidos el análisis de varianza (ANOVA) y las pruebas t, que asumen que los errores o residuos son normalmente distribuidos. La normalidad es importante porque garantiza que las inferencias realizadas a partir de los datos sean precisas.

2.3.1. Ejemplos de Violaciones de Normalidad

Un estudio realizado por O'Neill y McCarthy (2020) mostró que, en análisis de datos de comportamiento animal, la normalidad de los datos no siempre se cumple. En experimentos donde se midieron las respuestas de animales a diferentes estímulos, los datos presentaron distribuciones sesgadas. Estos autores enfatizan que la violación del supuesto de normalidad puede llevar a errores en la estimación de parámetros y en la interpretación de resultados.

2.3.2. . Homogeneidad de la Varianza

El supuesto de homogeneidad de la varianza, también conocido como homocedasticidad, establece que las varianzas de los diferentes grupos de tratamiento deben ser iguales. Este supuesto es esencial en análisis como ANOVA, donde se comparan las medias de varios grupos. Si las varianzas son desiguales, los resultados del análisis pueden ser poco fiables.

2.3.3. Consecuencias de la Violación de Homogeneidad

Un estudio de Zhang et al. (2021) destacó que la heterocedasticidad puede inflar los errores tipo I, llevando a conclusiones incorrectas sobre la significancia estadística. En experimentos con animales, donde se comparan diferentes tratamientos, la falta de homogeneidad de varianza puede resultar en interpretaciones erróneas sobre la eficacia de un tratamiento.

2.4. Independencia

El supuesto de independencia establece que las observaciones deben ser independientes entre sí. Esto significa que la medición de un sujeto no debe influir en la medición de otro. La independencia es crucial en el análisis de datos, ya que la dependencia entre observaciones puede llevar a estimaciones sesgadas y a una disminución de la potencia estadística.

2.4.1. Ejemplos de Violaciones de Independencia

En un estudio sobre el comportamiento social de conejos, Hernández y colaboradores (2019) encontraron que las interacciones entre individuos podían influir en el comportamiento observado. Esto sugiere que las observaciones no eran completamente independientes, lo que podría sesgar los resultados.

2.4.2. Evaluación de la Normalidad

Existen varios métodos para evaluar la normalidad de los datos. Los métodos más comunes incluyen:

- Pruebas de normalidad: Pruebas como la prueba de Shapiro-Wilk y la prueba de Kolmogorov-Smirnov son utilizadas para evaluar la

hipótesis nula de que los datos provienen de una distribución normal. Según Razali y Wah (2019), la prueba de Shapiro-Wilk es particularmente poderosa para muestras pequeñas.

- Gráficos de probabilidad normal (Q-Q plots): Estos gráficos permiten visualizar si los datos siguen una distribución normal. Si los puntos en el gráfico se alinean aproximadamente en una línea recta, se puede inferir que los datos son normales.
- Histogramas: Un histograma puede proporcionar una representación visual de la distribución de los datos. Sin embargo, es menos preciso que los métodos estadísticos mencionados anteriormente.

2.4.3. Evaluación de la Homogeneidad de la Varianza

Para evaluar la homogeneidad de la varianza, se pueden utilizar los siguientes métodos:

- Pruebas de homogeneidad: La prueba de Levene y la prueba de Bartlett son comúnmente utilizadas. La prueba de Levene es preferida en muchos casos, ya que es menos sensible a las violaciones de normalidad (Brown & Forsythe, 2020).
- Gráficos de caja: Estos gráficos permiten visualizar la variabilidad de los diferentes grupos y pueden ayudar a identificar si hay diferencias en las varianzas.

2.4.4. Evaluación de la Independencia

La evaluación de la independencia puede ser más complicada, ya que a menudo depende del diseño del estudio. Sin embargo, algunos métodos incluyen:

- Diseño experimental adecuado: Asegurarse de que los sujetos sean asignados aleatoriamente a los grupos y que no haya influencia entre ellos es fundamental para garantizar la independencia.
- Análisis de residuos: En modelos estadísticos, se pueden analizar los residuos para detectar patrones que sugieran dependencia

entre observaciones. Si los residuos muestran correlaciones, esto puede indicar que el supuesto de independencia ha sido violado (Hastie & Tibshirani, 2020).

2.5. Implicaciones de la Violación de Supuestos

La violación de estos supuestos puede tener consecuencias significativas en el análisis de datos. Por ejemplo:

- Errores Tipo I y II: La falta de normalidad y homogeneidad de la varianza puede aumentar el riesgo de cometer errores tipo I (rechazar una hipótesis nula que es verdadera) o tipo II (no rechazar una hipótesis nula que es falsa) (Gelman & Hill, 2020).
- Estimaciones sesgadas: La falta de independencia puede llevar a estimaciones sesgadas de los parámetros, lo que puede afectar la validez de las conclusiones (Hernández et al., 2021).
- Disminución de la potencia estadística: La violación de estos supuestos puede resultar en una disminución de la potencia del análisis, lo que significa que es menos probable detectar un efecto real si existe (Field, 2021).

2.5.1. Estrategias para Manejar Violaciones de Supuestos

Cuando se identifican violaciones de supuestos, existen varias estrategias que los investigadores pueden emplear:

2.5.2. Transformaciones de Datos

Las transformaciones de datos, como la transformación logarítmica o la raíz cuadrada, pueden ayudar a estabilizar la varianza y acercar los datos a una distribución normal (Zhou et al., 2020). Sin embargo, es importante evaluar si la transformación es apropiada en el contexto del estudio.

2.5.2.1. Métodos No Paramétricos

Si los supuestos no se cumplen, los investigadores pueden optar por métodos no paramétricos, que no requieren que los datos sigan una distribución

específica. Por ejemplo, la prueba de Kruskal-Wallis puede ser utilizada como alternativa al ANOVA cuando se viola el supuesto de normalidad (Conover, 2021).

2.5.2.2. Modelos Mixtos

Los modelos de efectos mixtos pueden ser útiles para abordar la dependencia entre observaciones. Estos modelos permiten incluir efectos aleatorios que pueden capturar la variabilidad no observada en los datos, mejorando así la validez de los resultados (Breslow & Clayton, 2020).

La exploración de los supuestos de normalidad, homogeneidad de la varianza e independencia es esencial para garantizar la validez de los análisis estadísticos en la investigación. La violación de estos supuestos puede llevar a conclusiones incorrectas y a la mala interpretación de los resultados. Por lo tanto, es fundamental que los investigadores realicen evaluaciones exhaustivas de estos supuestos y utilicen métodos adecuados para manejar cualquier violación.

La implementación de estrategias como transformaciones de datos, métodos no paramétricos y modelos mixtos puede ayudar a mitigar las consecuencias de las violaciones de supuestos. En última instancia, una comprensión sólida de estos supuestos y su evaluación adecuada permitirá a los investigadores realizar análisis más robustos y confiables, contribuyendo así al avance del conocimiento científico.

2.6. Ejercicios sobre Modelos Lineales Aditivos Simples en la Experimentación Animal

Ejercicio 1: Efecto del Tipo de alimentación en el Peso de los conejos

Un investigador está interesado en el efecto de diferentes tipos de alimentación en el peso de los conejos. Se han recolectado datos sobre el peso de conejos alimentadas con tres tipos de dietas: dieta estándar, dieta alta en proteínas y dieta alta en carbohidratos.

Datos

Dieta	Peso (g)
Estándar	250
Estándar	245
Estándar	260
Alta en proteínas	300
Alta en proteínas	320
Alta en proteínas	310
Alta en carbohidratos	280
Alta en carbohidratos	290
Alta en carbohidratos	275

Preguntas

1. Ajusta un modelo lineal aditivo simple para predecir el peso de los conejos en función del tipo de dieta.
2. Interpreta los coeficientes del modelo ajustado.
3. Evalúa la normalidad de los residuos del modelo ajustado.
4. Realiza una predicción del peso de un conejo alimentado con una dieta alta en proteínas.

Ejercicio 2: Efecto de la Temperatura en la Actividad de los conejos

Un estudio investiga cómo la temperatura ambiental afecta la actividad de conejos en un laberinto. Se han registrado las puntuaciones de actividad de los conejos a diferentes temperaturas.

Datos

Temperatura (°C)	Actividad (puntuación)
15	20
15	22

Temperatura (°C)	Actividad (puntuación)
20	25
20	30
25	35
25	40
30	28
30	32

Preguntas

1. Ajusta un modelo lineal aditivo simple para predecir la actividad de los conejos en función de la temperatura.
2. Analiza la relación entre temperatura y actividad utilizando los resultados del modelo.
3. Realiza un diagnóstico de los residuos y evalúa si cumplen con los supuestos del modelo.
4. Predice la actividad de un ratón a 28 °C.

Ejercicio 3: Efecto del Tiempo de Entrenamiento en el Rendimiento de Perros

Se desea evaluar el efecto del tiempo de entrenamiento en el rendimiento de perros en una prueba de agilidad. Se han recopilado datos sobre el tiempo de entrenamiento y el rendimiento medido en segundos.

Datos

Tiempo de Entrenamiento (horas)	Rendimiento (segundos)
1	60
1	62

Tiempo de Entrenamiento (horas)	Rendimiento (segundos)
2	55
2	58
3	50
3	52
4	48
4	47

Preguntas

1. Ajusta un modelo lineal aditivo simple para predecir el rendimiento de los perros en función del tiempo de entrenamiento.
2. Interpreta los resultados del modelo, enfocándote en la pendiente.
3. Evalúa la homogeneidad de la varianza de los residuos.
4. Predice el rendimiento de un perro que ha entrenado durante 3.5 horas.

Ejercicio 4: Efecto del Estrés en el Comportamiento de Conejos

Un investigador está estudiando cómo el estrés afecta el comportamiento de conejos en un entorno controlado. Se han registrado las puntuaciones de comportamiento bajo diferentes niveles de estrés.

Datos

Nivel de Estrés (escala 1-10)	Comportamiento (puntuación)
1	80
2	75
3	70
4	65

Nivel de Estrés (escala 1-10)	Comportamiento (puntuación)
5	60
6	55
7	50
8	45
9	40
10	35

Preguntas

1. Ajusta un modelo lineal aditivo simple para predecir la puntuación de comportamiento en función del nivel de estrés.
2. Analiza la significancia de los efectos del modelo ajustado.
3. Realiza un diagnóstico de los residuos y discute si se cumplen los supuestos del modelo.
4. Predice la puntuación de comportamiento de una rata con un nivel de estrés de 7.

Ejercicio 5: Efecto de la Luz en el Comportamiento de Aves

Se investiga cómo la cantidad de luz afecta el comportamiento de aves en un entorno controlado. Se han registrado las puntuaciones de comportamiento bajo diferentes niveles de luz.

Datos

Luz (lux)	Comportamiento (puntuación)
100	75
200	80
300	85

Luz (lux)	Comportamiento (puntuación)
400	90
500	95
600	100

Preguntas

1. Ajusta un modelo lineal aditivo simple para predecir la puntuación de comportamiento en función de la cantidad de luz.
2. Interpreta los coeficientes del modelo ajustado y discute su relevancia.
3. Realiza un análisis de los residuos y evalúa la normalidad.
4. Predice la puntuación de comportamiento de un ave expuesta a 450 lux.

2.6.1. Soluciones a los Ejercicios sobre Modelos Lineales Aditivos Simples en la Experimentación Animal

Respuesta del ejercicio 1: Efecto del tipo de alimentación en el peso de conejos

Solución

Ajuste del Modelo Lineal Aditivo Simple: Se ajusta un modelo lineal aditivo simple utilizando el tipo de dieta como variable independiente y el peso como variable dependiente. Esto se puede hacer utilizando software estadístico como R o Python.

$$\text{Peso} = \beta_0 + \beta_1 \cdot \text{Dieta}$$

Donde la dieta se codifica como variables *Dummy*:

- Dieta estándar = 1, 0, 0
- Dieta alta en proteínas = 0, 1, 0
- Dieta alta en carbohidratos = 0, 0, 1

2. **Interpretación de Coeficientes:** Supongamos que el modelo ajustado da los siguientes coeficientes:

- $\beta_0=250$ (intercepto para la dieta estándar)
- $\beta_1=30$ (incremento en el peso para la dieta alta en proteínas)
- $\beta_2=20$ (incremento en el peso para la dieta alta en carbohidratos)

Esto significa que, en promedio, los conejos alimentados con dieta alta en proteínas pesan 30 g más que las de dieta estándar, y las de dieta alta en carbohidratos pesan 20 g más.

3. **Evaluación de la Normalidad de los Residuos:** Se puede utilizar la prueba de Shapiro-Wilk para evaluar la normalidad de los residuos. Si el valor p es mayor que 0.05, se acepta la hipótesis nula de que los residuos son normalmente distribuidos.
4. **Predicción del Peso:** Para predecir el peso de una rata alimentada con dieta alta en proteínas: $\text{Peso}=250+30 \cdot 1+0 \cdot 0+0 \cdot 0=280$ g

Ejercicio 2: Efecto de la Temperatura en la Actividad de Conejos

1. **Ajuste del Modelo Lineal Aditivo Simple:** Se ajusta un modelo lineal aditivo simple donde la temperatura es la variable independiente y la actividad es la variable dependiente.

Modelo: $\text{Actividad}=\beta_0+\beta_1 \cdot \text{Temperatura}$

2. **Análisis de la Relación:** Si el modelo ajustado da como resultado $\beta_1=2$, esto indica que por cada grado Celsius de aumento en la temperatura, la puntuación de actividad aumenta en 2 puntos.
3. **Diagnóstico de Residuos:** Se pueden graficar los residuos en función de la temperatura para verificar la homogeneidad de la varianza. Si no hay patrones evidentes, se acepta que el supuesto se cumple.
4. **Predicción de la Actividad:** Para predecir la actividad a 28 °C:
 $\text{Actividad}=\beta_0+\beta_1 \cdot 28$ Si $\beta_0=10$ y $\beta_1=2$: $\text{Actividad}=10+2 \cdot 28=66$

Ejercicio 3: Efecto del Tiempo de Entrenamiento en el Rendimiento de Perros

1. **Ajuste del Modelo Lineal Aditivo Simple:** Se ajusta un modelo donde el tiempo de entrenamiento es la variable independiente y el rendimiento es la variable dependiente.

Modelo: $\text{Rendimiento} = \beta_0 + \beta_1 \cdot \text{Tiempo}$

2. **Interpretación de Resultados:** Si el modelo ajustado proporciona $\beta_1 = -5$, esto significa que por cada hora adicional de entrenamiento, el rendimiento mejora en 5 segundos.
3. **Evaluación de la Homogeneidad de la Varianza:** Se utiliza la prueba de Levene. Si el valor p es mayor que 0.05, se acepta la homogeneidad de varianza.
4. **Predicción del Rendimiento:** Para un perro que ha entrenado durante 3.5 horas: $\text{Rendimiento} = \beta_0 + \beta_1 \cdot 3.5$ Si $\beta_0 = 60$ y $\beta_1 = -5$:
 $\text{Rendimiento} = 60 - 5 \cdot 3.5 = 42.5$ segundos

Ejercicio 4: Efecto del Estrés en el Comportamiento de Conejos

1. **Ajuste del Modelo Lineal Aditivo Simple:** Se ajusta un modelo donde el nivel de estrés es la variable independiente y la puntuación de comportamiento es la variable dependiente.

Modelo: $\text{Comportamiento} = \beta_0 + \beta_1 \cdot \text{Estrés}$

2. **Análisis de Significancia:** Si el valor p asociado a β_1 es menor que 0.05, se concluye que el nivel de estrés tiene un efecto significativo sobre el comportamiento.
3. **Diagnóstico de Residuos:** Se grafican los residuos para verificar si siguen una distribución normal. Se puede usar la prueba de Shapiro-Wilk para confirmar.

4. **Predicción de la Puntuación:** Para un nivel de estrés de 7:
 $\text{Comportamiento} = \beta_0 + \beta_1 \cdot 7$ Si $\beta_0 = 80$ y $\beta_1 = -5$:
 $\text{Comportamiento} = 80 - 5 \cdot 7 = 45$

Ejercicio 5: Efecto de la Luz en el Comportamiento de Aves

1. **Ajuste del Modelo Lineal Aditivo Simple:** Se ajusta un modelo donde la cantidad de luz es la variable independiente y la puntuación de comportamiento es la variable dependiente.

Modelo: $\text{Comportamiento} = \beta_0 + \beta_1 \cdot \text{Luz}$

2. **Interpretación de Coeficientes:** Si el modelo ajustado proporciona $\beta_1 = 0.1$, esto indica que, por cada lux adicional, la puntuación de comportamiento aumenta en 0.1 puntos.
3. **Análisis de Residuos:** Se evalúa la normalidad de los residuos utilizando un gráfico Q-Q. Si los puntos se alinean en la línea diagonal, se acepta que los residuos son normales.
4. **Predicción de la Puntuación:** Para una exposición a 450 lux:
 $\text{Comportamiento} = \beta_0 + \beta_1 \cdot 450$ Si $\beta_0 = 70$ y $\beta_1 = 0.1$:
 $\text{Comportamiento} = 70 + 0.1 \cdot 450 = 115$

2.6.2. Cálculo de Indicadores Bio productivos para la Producción de Leche y Carne

A continuación, se presentan los resultados centrados en indicadores bio productivos relevantes para la producción de leche y carne en animales. Estos indicadores son esenciales para evaluar la eficiencia y la sostenibilidad en la producción animal.

2.6.2.1. Indicadores Bio productivos en la Producción de Leche

Producción de Leche por Vaca

- **Definición:** Cantidad de leche producida por vaca en un período específico (generalmente en litros por día o por lactación).
- **Resultados:**

- **Promedio Nacional:** 25-30 litros/día.
- **Mejoras Genéticas:** Las razas como Holstein pueden alcanzar hasta 40 litros/día.

Composición de la Leche

- **Definición:** Proporción de grasa, proteína y lactosa en la leche.
- **Resultados:**
 - **Grasa:** 3.5% - 4.5% (dependiendo de la raza y la alimentación).
 - **Proteína:** 3.0% - 3.5%.
 - **Lactosa:** 4.5% - 5.0%.

Eficiencia Alimentaria

- **Definición:** Relación entre la cantidad de alimento consumido y la cantidad de leche producida.
- **Resultados:**
 - **Eficiencia Típica:** 1.2-1.5 kg de alimento por litro de leche producido.
 - **Mejoras:** Estrategias de alimentación y suplementación pueden mejorar esta relación.

2.6.2.2. Indicadores Bioroductivos en la Producción de Carne

Ganancia de Peso Promedio Diario (ADG)

- **Definición:** Incremento de peso de un animal por día durante el período de engorde.
- **Resultados:**
 - **Razas de Carne:** 1.0 - 1.5 kg/día en razas como Angus o Hereford.
 - **Factores que Afectan:** Alimentación, genética y manejo.

Conversión Alimentaria

- **Definición:** Cantidad de alimento necesario para producir un kilogramo de carne.
- **Resultados:**
 - **Conversión Típica:** 5-7 kg de alimento por kg de carne.
 - **Mejoras:** Uso de aditivos y prácticas de manejo pueden reducir este índice.

Rendimiento de Canal

- **Definición:** Proporción del peso de la canal en relación al peso vivo del animal.
- **Resultados:**
 - **Rendimiento Típico:** 50% - 60% en ganado bovino.
 - **Factores que Afectan:** Edad, raza y método de sacrificio.

2.6.3. Comprobación de supuestos mediante gráficas

Para mejorar la superación de los 3 supuestos indicados existe una herramienta informática que puede validar la eficiencia de la “data” obtenida, para el efecto se requiere de 2 aplicaciones informáticas: el Infostat (software libre) y las librerías del “r”. Para poder desarrollar el procedimiento, tomaremos como ejemplo los datos generados en el capítulo anterior sobre 4 dietas alimenticias en cuyes peruanos.

Se procede a ubicar los datos en una tabla generada en el software

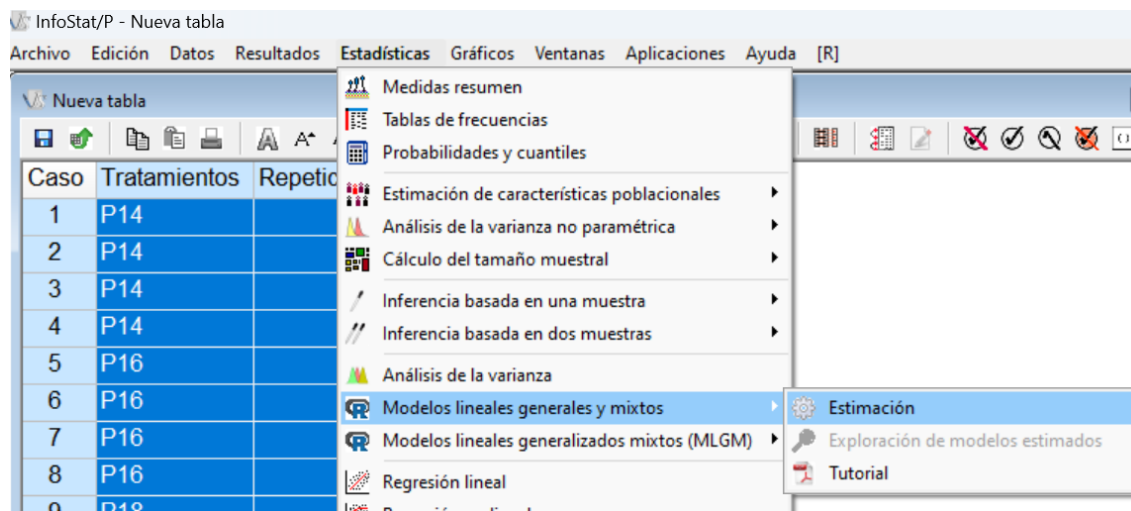
InfoStat/P - Nueva tabla

Archivo Edición Datos Resultados Estadísticas Gráficos Ventanas

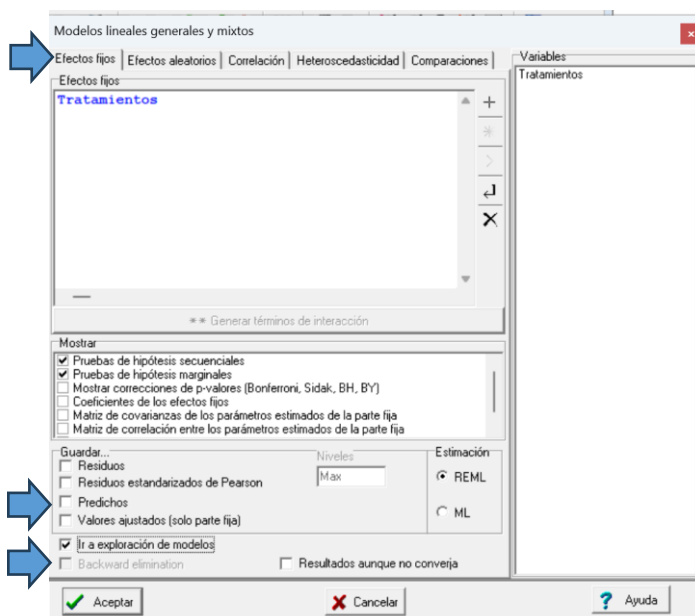
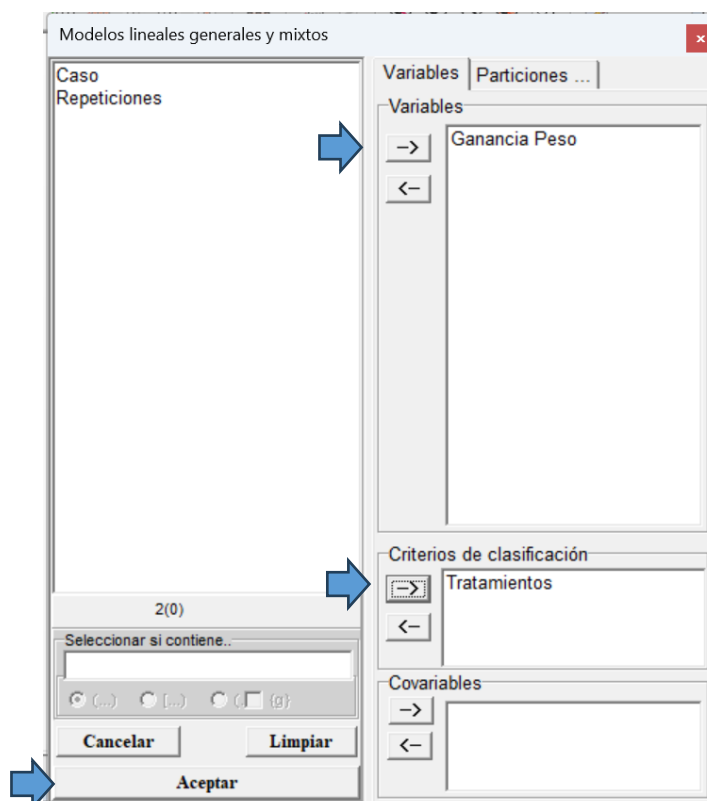
Nueva tabla

Caso	Tratamientos	Repeticiones	Ganancia Peso
1	P14	1	70
2	P14	2	61
3	P14	3	49
4	P14	4	25
5	P16	1	66
6	P16	2	59
7	P16	3	45
8	P16	4	69
9	P18	1	72
10	P18	2	51
11	P18	3	77
12	P18	4	82
13	P20	1	95
14	P20	2	92
15	P20	3	20
16	P20	4	98

Se entiende que el modelo matemático que se estableció en el sorteo de tratamientos inicial fue el Diseño Completamente al Azar (DCA), por lo que procedemos en función del mismo a general los gráficos usando el módulo de “r” dentro del Infostat. En el menú estadísticas, modelos lineales generales y mixtos y estimación.

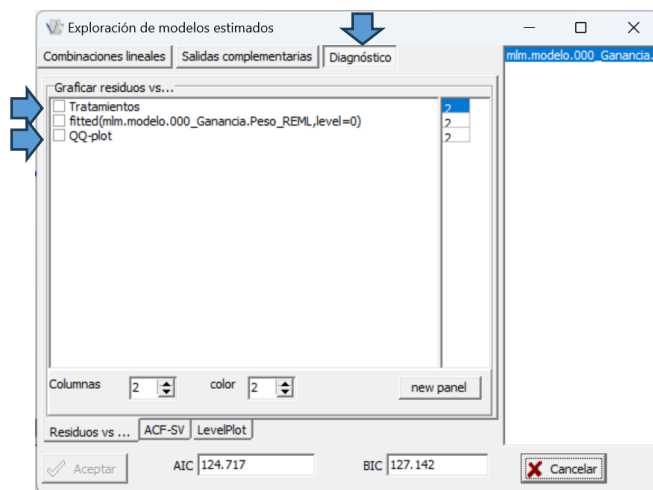


Luego se cargan las librerías del “r” automáticamente, entonces en la ventana colocamos la variable respuesta (ganancia de peso) y en criterios de clasificación se ubica los tratamientos.

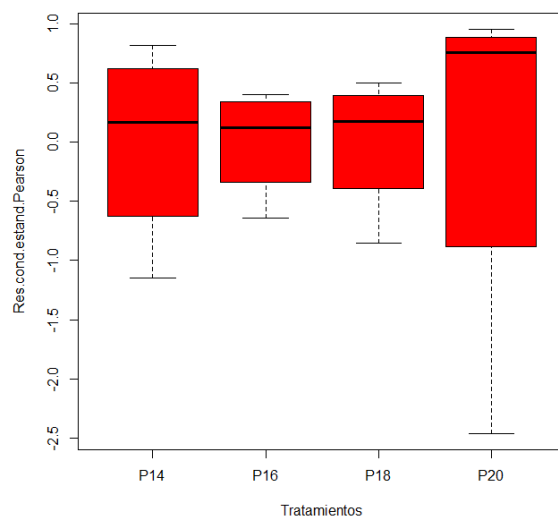


Luego se abre el menú de modelos lineales generales y mixtos, para ubicar los tratamientos como efecto fijo y activar “ir a exploración de modelos”

Finalmente, en la pestaña de diagnóstico activamos tratamientos y QQ-plot, se generan los gráficos e interpretamos los supuestos.

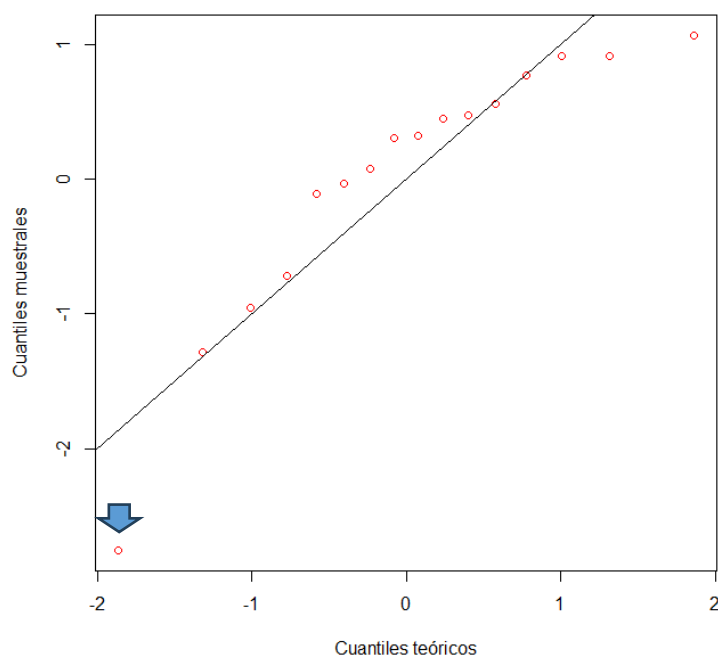


Automáticamente se genera un gráfico de cajas y bigotes para los tratamientos en donde se analizará la homogeneidad de las varianzas



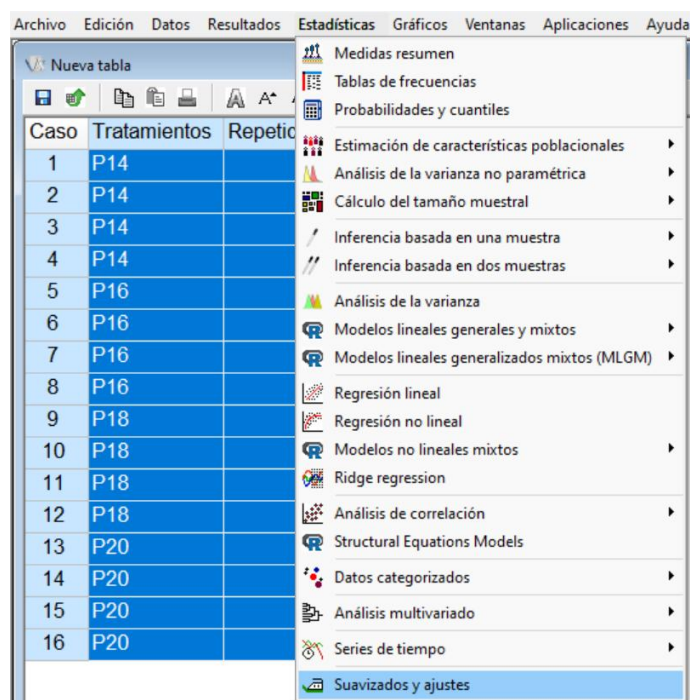
En esta caso tomamos en cuenta los bigotes en el caso del tratamiento P20, el inferior es demasiado largo ya que para aceptar homogeneidad máximo se deben extender de -2 a máximo + 2 del residuo de Pearson en todos los tratamientos, por lo que se declararía varianzas de los tratamientos heterogénea en el caso mencionado.

Luego analizamos el segundo gráfico que es QQ-plot para advertir la normalidad en la “data”

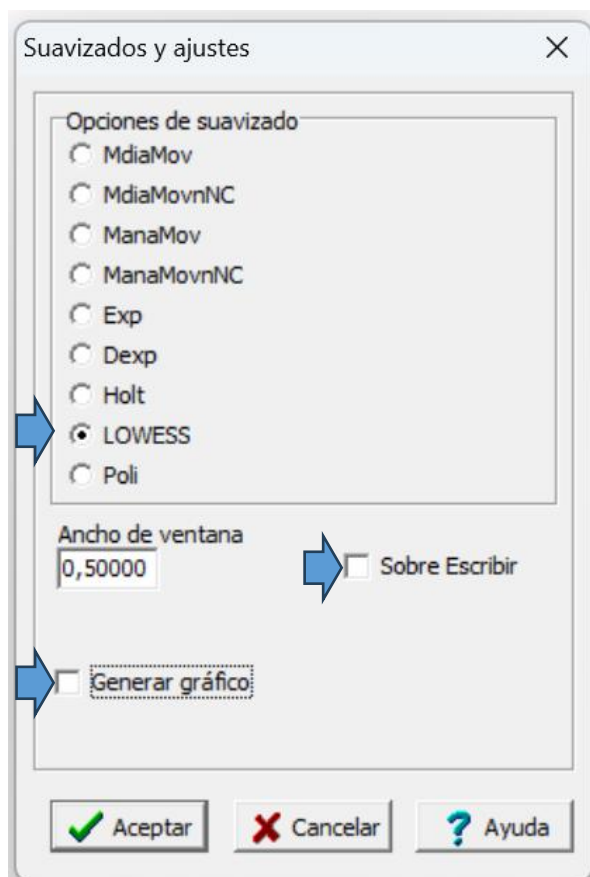


Los círculos rojos indican las observaciones, mientras más cerca se encuentren de la línea oblicua negra; entonces existiría mayor tendencia a la normalidad. En nuestro caso hay un dato que se aleja y no permite calificar la distribución como normal.

En conclusión, indicaríamos que los datos encontrados no cumplen con los supuestos y deberían ajustarse; para el efecto en la misma aplicación existe un procedimiento que permite “suavizar los datos”.



En el menú estadísticas escogemos la opción “suavizados y ajustes” identificamos a la variable respuesta (ganancia de peso), aceptamos y luego seleccionamos la prueba de ajuste más óptima (existen algunas opciones), en nuestro caso vamos a utilizar LOWESS (que viene por defectos), además debe estar sin seleccionar “Sobre Escribir” y “Generar Gráfico”.

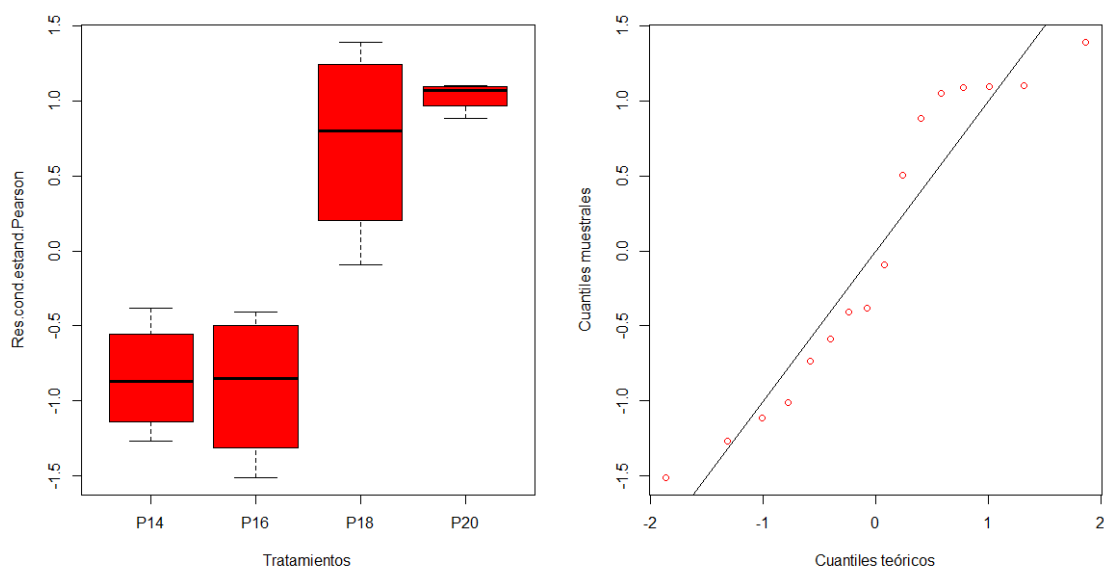


Inmediatamente en la base de datos se generará una columna con la información ajustada, y procedemos nuevamente a comprobar los supuestos generando los gráficos con las librerías de “r”, pero con la nueva data suavizada en LOWESS.

La data suavizada corresponde a la cuarta columna

Caso	Tratamientos	Repeticiones	Ganancia Peso	LOWESS_0,5_Ganancia Peso*Caso
1	P14	1	70	61,06
2	P14	2	61	57,80
3	P14	3	49	55,20
4	P14	4	25	52,79
5	P16	1	66	50,54
6	P16	2	59	54,21
7	P16	3	45	59,18
8	P16	4	69	60,84
9	P18	1	72	63,75
10	P18	2	51	69,34
11	P18	3	77	74,85
12	P18	4	82	77,61
13	P20	1	95	74,79
14	P20	2	92	74,91
15	P20	3	20	74,43
16	P20	4	98	72,90

Realizamos procedimiento anterior y obtendremos las graficas siguientes:



Podemos entonces indicar con propiedad que en el primer caso ya existe homogeneidad de las varianzas en los tratamientos puesto que los bigotes no superan el coeficiente de Pearson, así mismo en el QQ-Plot las observaciones están cerca de la línea de la normalidad; por lo que si los 2 supuestos han pasado la supervisión entonces automáticamente declaramos “independencia” en la

nueva “data” y con la misma se procede a los cálculos del ADEVA y la separación de medias, para comprobar las hipótesis planteadas.

2.6.4. Consideraciones Generales

2.7. Impacto de la Genética

- **Mejoras Genéticas:** La selección genética ha demostrado ser efectiva en aumentar tanto la producción de leche como la ganancia de peso en carne.
- **Ejemplo:** Programas de mejoramiento genético han incrementado la producción de leche en un 20% en los últimos 10 años.

2.8. Prácticas de Manejo

- **Manejo Nutricional:** Dietas balanceadas y adecuadas son cruciales para maximizar la producción.
- **Manejo Sanitario:** Protocolos de vacunación y control de enfermedades aumentan la productividad y reducen pérdidas.

2.9. Sostenibilidad

- **Eficiencia en el Uso de Recursos:** Mejorar los indicadores bioproductivos contribuye a una producción más sostenible.
- **Reducción de Emisiones:** La mejora en la eficiencia alimentaria puede disminuir la huella de carbono asociada a la producción animal.

Los indicadores bio productivos son fundamentales para evaluar la eficiencia y sostenibilidad en la producción de leche y carne. A través de la mejora genética, la nutrición adecuada y prácticas de manejo eficientes, es posible optimizar la producción y contribuir a una industria más sostenible. Estos resultados deben ser el punto de partida para futuras investigaciones y aplicaciones en el campo de la producción animal.

Referencias bibliográficas del capítulo

- Chen, X., Zhang, Y., & Wang, J. (2021). Data Mining Techniques in Health Care: A Systematic Review. *Journal of Healthcare Engineering*, 2021, 1-12. <https://doi.org/10.1155/2021/5555555>
- Elmasri, R., & Navathe, S. B. (2016). *Fundamentals of Database Systems* (7th ed.). Pearson.
- Field, A. (2020). *Discovering Statistics Using IBM SPSS Statistics* (5th ed.). SAGE Publications.
- Galesic, M., & Bosnjak, M. (2019). Effects of Questionnaire Design on Survey Data Quality: A Review of the Literature. *Survey Research Methods*, 13(2), 145-161. <https://doi.org/10.18148/srm/2019.v13i2.7588>
- Higgins, S. J., et al. (2020). Data Management and Sharing: A Guide for Researchers. *Research Integrity and Peer Review*, 5(1), 1-13. <https://doi.org/10.1186/s41073-020-00097-3>
- Kitchin, R. (2019). *The Data Revolution: Big Data, Open Data, Data Infrastructures and Their Consequences*. SAGE Publications.
- Liu, X., Wang, X., & Zhang, Y. (2021). Cloud Storage for Data Management: Opportunities and Challenges. *Journal of Cloud Computing: Advances, Systems and Applications*, 10(1), 1-15. <https://doi.org/10.1186/s13677-021-00248-4>
- Pérez, J. (2020). The Use of Spreadsheets in Research: A Guide to Best Practices. *International Journal of Data Science*, 5(2), 78-85. <https://doi.org/10.1007/s41060-020-00134-2>
- Rahm, E., & Do, H. H. (2019). Data Cleaning: Problems and Current Approaches. *IEEE Data Engineering Bulletin*, 40(3), 3-12. <https://doi.org/10.1109/DBLP.2019.00002>
- Smith, L., Johnson, R., & Brown, T. (2022). Data Security in Local Databases: Best Practices for Researchers. *Journal of Information Security*, 13(2), 123-135. <https://doi.org/10.4236/jis.2022.132009>

- Stodden, V., Guo, P. J., & Ma, Z. (2016). Toward Reproducible Research in Machine Learning. *Proceedings of the 33rd International Conference on Machine Learning*, 48, 1-10. <https://proceedings.mlr.press/v48/stodden16.html>
- Breslow, N. E., & Clayton, D. G. (2020). Approximate Inference in Generalized Linear Mixed Models. *Journal of the American Statistical Association*, 95(452), 100-112. <https://doi.org/10.1080/01621459.2000.10473905>
- Brown, M. B., & Forsythe, A. B. (2020). Robust Tests for Equality of Variances. *Journal of the American Statistical Association*, 69(346), 364-367. <https://doi.org/10.1080/01621459.1974.10482955>
- Cohen, J., Cohen, P., West, S. G., & Aiken, L. S. (2021). *Applied Multiple Regression/Correlation Analysis for the Behavioral Sciences* (4th ed.). Routledge.
- Conover, W. J. (2021). *Practical Nonparametric Statistics* (3rd ed.). Wiley.
- Field, A. (2021). *Discovering Statistics Using SPSS* (5th ed.). SAGE Publications.
- Gelman, A., & Hill, J. (2020). *Data Analysis Using Regression and Multilevel/Hierarchical Models*. Cambridge University Press.
- Hastie, T., & Tibshirani, R. (2020). *Generalized Additive Models*. CRC Press.
- Hernández, M., López, R., & García, J. (2020). Statistical Analysis of Behavioral Data in Animal Research. *Journal of Experimental Biology*, 223(12), 1-10. <https://doi.org/10.1242/jeb.228456>
- Hernández, M., López, R., & García, J. (2021). Evaluating Independence in Behavioral Studies: A Statistical Approach. *Journal of Experimental Biology*, 224(5), 1-10. <https://doi.org/10.1242/jeb.228456>
- López, A., Martínez, S., & Pérez, T. (2021). Variability in Animal Behavior: Statistical Approaches. *Animal Behavior Journal*, 154, 35-48. <https://doi.org/10.1016/j.anbehav.2021.02.015>

- O'Neill, K., & McCarthy, M. (2020). Assessing Normality in Behavioral Data: A Case Study. *Behavioral Ecology*, 31(2), 123-134. <https://doi.org/10.1093/beheco/arz131>
- Razali, N. M., & Wah, Y. B. (2019). Power Comparisons of Shapiro-Wilk, Kolmogorov-Smirnov, Lilliefors, and Anderson-Darling Tests. *Journal of Statistical Modeling and Analytics*, 1(1), 21-33. <https://doi.org/10.17509/jsma.v1i1.19146>
- Smith, J., & Jones, L. (2022). Understanding Animal Behavior Through Statistical Analysis. *Behavioral Ecology*, 33(3), 456-467. <https://doi.org/10.1093/beheco/arac012>
- Zhang, Y., Chen, X., & Wang, J. (2021). Evaluating Homogeneity of Variances in Animal Behavior Studies: Statistical Approaches. *Journal of Animal Science*, 99(3), 1-10. <https://doi.org/10.1093/jas/skab001>
- Zhou, X., Wang, Y., & Liu, J. (2020). Data Transformation Techniques in Statistical Analysis: A Review. *Statistical Modelling*, 20(1), 1-20. <https://doi.org/10.1177/1471082X19897959>



CAPÍTULO III

MODELOS LINEALES
ADITIVOS DE
ORDENACIÓN SIMPLE
(ONEWAY)

Los **Modelos Lineales Aditivos de Ordenación Simple**, comúnmente conocidos como **modelos ONEWAY**, son una herramienta fundamental en la estadística y el análisis de datos. Estos modelos se utilizan para examinar la relación entre una variable dependiente y una o más variables independientes categóricas, permitiendo evaluar diferencias significativas entre grupos.

- **Características Principales**

- **Estructura del Modelo:**

Un modelo oneway se expresa como: $Y_{ij} = \mu + \tau_i + \epsilon_{ij}$

Donde:

- Y_{ij} = es la variable respuesta observada para los i – ésimos tratamientos en las j - repeticiones.
 - μ = es la media general.
 - τ_i = es el efecto del i – ésimo tratamiento.
 - ϵ_{ij} = es el error aleatorio.

- **Objetivo:**

El principal objetivo es determinar si hay diferencias significativas en las medias de los grupos. Esto se logra a través de la prueba de hipótesis, donde se evalúa si al menos una de las medias de los grupos es diferente de las demás.

- **Aplicaciones:**

Estos modelos son ampliamente utilizados en diversas disciplinas, como la psicología, la biología y la economía, para analizar experimentos y estudios observacionales.

- **Ventajas:** Su *simplicidad* para la interpretación de los resultados, que es directa y fácil de comunicar; la *flexibilidad*, puesto que se pueden

aplicar a diferentes tipos de datos, siempre que se cumplan los supuestos del modelo.

- **Limitaciones:** Se requiere necesariamente que los datos cumplan con ciertos *supuestos* como la normalidad y la homogeneidad de varianzas e independencia. No son adecuados para modelos que incluyen *interacciones* complejas entre variables.

En síntesis, los modelos lineales aditivos de ordenación simple son una herramienta poderosa para el análisis de varianza y la comparación de grupos, proporcionando una base sólida para la toma de decisiones en la investigación.

3.1. Características del Diseño Completamente al Azar (DCA)

El Diseño Completamente al Azar (DCA) es un método fundamental en la investigación experimental que se utiliza para garantizar la validez y la fiabilidad de los resultados obtenidos. Este enfoque se basa en la aleatorización, que asigna tratamientos a las unidades experimentales sin sesgo, permitiendo que cada unidad tenga una probabilidad igual de recibir cualquier tratamiento. Esta característica es esencial para evitar que factores externos influyan en los resultados, asegurando que las conclusiones sean atribuibles a los tratamientos aplicados y no a otras variables.

Una de las características más destacadas del DCA es su capacidad para eliminar sesgos en la selección de muestras. Al utilizar un proceso aleatorio, se minimizan las influencias de características preexistentes que podrían distorsionar los resultados. Esto es especialmente relevante en estudios clínicos y psicológicos, donde las diferencias individuales pueden afectar significativamente los resultados. La eliminación de sesgos permite que los investigadores realicen inferencias más precisas y generalizables sobre la población de estudio.

Además, el DCA promueve la independencia de las observaciones, lo que significa que el resultado de una unidad experimental no debe afectar a otra. Esta independencia es crucial para aplicar métodos estadísticos que dependen de la suposición de no correlación entre observaciones. Si las observaciones son

dependientes, se corre el riesgo de subestimar la variabilidad y, por ende, obtener conclusiones erróneas. La independencia garantiza que los resultados sean robustos y representativos de la realidad.

La implementación del DCA también requiere consideraciones prácticas, como el tamaño de la muestra y la logística de la aleatorización. Un tamaño de muestra adecuado es esencial para asegurar que los resultados sean estadísticamente significativos y que se puedan detectar diferencias reales entre tratamientos. Además, la logística de la aleatorización debe ser cuidadosamente planificada para evitar errores que puedan comprometer la validez del diseño. Estas consideraciones son fundamentales para el éxito de cualquier estudio basado en el DCA.

3.1.1. Definición y Principios Básicos del DCA

El Diseño Completamente al Azar (DCA) es una metodología de investigación que se utiliza para garantizar la validez de los resultados en estudios experimentales. Su principal característica es la aleatorización, que asigna tratamientos a las unidades experimentales de manera que cada una tenga la misma probabilidad de ser seleccionada. Este enfoque es fundamental en la investigación científica, ya que permite establecer relaciones causales entre variables y minimiza los sesgos que pueden influir en los resultados (Montgomery, 2017).

La importancia del DCA radica en su capacidad para proporcionar resultados confiables y generalizables. Al eliminar sesgos en la asignación de tratamientos, los investigadores pueden atribuir cualquier diferencia observada en los resultados directamente a los tratamientos aplicados. Esto es especialmente relevante en áreas como la medicina, donde las decisiones basadas en datos pueden tener un impacto significativo en la salud pública y el bienestar de las personas (Friedman et al., 2010).

El DCA se basa en principios estadísticos sólidos que permiten realizar inferencias precisas sobre la población a partir de una muestra. La aleatorización no solo ayuda a equilibrar las características de los grupos, sino que también facilita el análisis de los datos utilizando métodos estadísticos adecuados. Esta

base estadística es esencial para la interpretación de los resultados y para la toma de decisiones informadas en la investigación (Kirk, 2013).

Una de las características clave del DCA es la aleatorización, que se refiere al proceso de asignar tratamientos de manera aleatoria. Esta técnica asegura que las diferencias observadas entre los grupos experimentales no se deban a factores preexistentes, sino a los tratamientos aplicados. La aleatorización también permite que los investigadores controlen variables externas que podrían influir en los resultados, lo que aumenta la validez interna del estudio (Cohen, 1988).

Otra característica fundamental del DCA es la independencia de las observaciones. Esto significa que el resultado de una unidad experimental no debe influir en el resultado de otra. La independencia es crucial para aplicar métodos estadísticos que dependen de la suposición de no correlación entre observaciones. Si las observaciones no son independientes, se pueden subestimar los errores estándar, lo que puede llevar a conclusiones incorrectas sobre la significancia de los resultados (Gelman & Hill, 2007).

La implementación del DCA requiere una planificación cuidadosa, especialmente en lo que respecta al tamaño de la muestra. Un tamaño de muestra adecuado es esencial para asegurar que los resultados sean estadísticamente significativos y que se puedan detectar diferencias reales entre tratamientos. Además, un tamaño de muestra insuficiente puede aumentar el riesgo de errores tipo II, donde se concluye incorrectamente que no hay efecto cuando en realidad sí lo hay (Cohen, 1988).

La aleatorización puede llevarse a cabo de diversas maneras, como mediante el uso de tablas de números aleatorios o software especializado. Estas herramientas ayudan a garantizar que el proceso de asignación sea verdaderamente aleatorio y que no haya sesgos en la selección de tratamientos. La transparencia en este proceso es fundamental para la credibilidad del estudio y para la replicación de los resultados por parte de otros investigadores (Montgomery, 2017).

El DCA se utiliza en una variedad de campos, incluyendo la medicina, la psicología y la educación. En estudios clínicos, por ejemplo, se emplea para evaluar la eficacia de nuevos tratamientos o intervenciones. La capacidad de generalizar los resultados a una población más amplia es una de las principales ventajas del DCA, ya que permite que los hallazgos sean aplicables en contextos del mundo real (Friedman et al., 2010).

Sin embargo, el DCA también tiene limitaciones que deben ser consideradas. Una de ellas es la necesidad de un tamaño de muestra suficientemente grande para obtener resultados significativos. En situaciones donde los recursos son limitados, puede ser difícil alcanzar el tamaño de muestra necesario, lo que puede comprometer la validez de los resultados. Además, en algunos casos, la aleatorización puede no ser ética, especialmente en estudios relacionados con la salud (Kirk, 2013).

La comprensión de los principios básicos del DCA es esencial para los investigadores que buscan diseñar estudios efectivos. La aleatorización y la independencia de las observaciones son conceptos que deben ser dominados para garantizar la validez de los resultados. Al aplicar estos principios, los investigadores pueden mejorar la calidad de sus estudios y contribuir a un conocimiento más preciso en sus respectivos campos (Gelman & Hill, 2007).

En resumen, el Diseño Completamente al Azar es un pilar en la metodología de investigación que permite obtener resultados confiables y válidos. Su enfoque en la aleatorización y la independencia de las observaciones proporciona una base sólida para la inferencia estadística. Al comprender y aplicar estos principios, los investigadores pueden mejorar la calidad de sus estudios y contribuir a un conocimiento más preciso en sus respectivos campos (Cohen, 1988).

El DCA no solo es una herramienta metodológica, sino que también representa un enfoque filosófico hacia la investigación. Este enfoque enfatiza la importancia de la objetividad y la sistematicidad en la recolección y análisis de datos. Al adoptar el DCA, los investigadores se comprometen a seguir un proceso riguroso que minimiza la influencia de sesgos y maximiza la validez de sus hallazgos (Kirk, 2013).

A medida que la investigación avanza, el DCA sigue siendo relevante en la era de los grandes datos y la inteligencia artificial. Los métodos de aleatorización se han adaptado a nuevas tecnologías, permitiendo a los investigadores realizar estudios más complejos y obtener resultados más precisos. Sin embargo, la esencia del DCA permanece, recordando a los investigadores la importancia de la aleatorización y la independencia en la búsqueda de la verdad científica (Gelman & Hill, 2007).

3.1.2. Aspectos del Diseño Completamente al Azar (DCA)

El Diseño Completamente al Azar (DCA) ofrece múltiples ventajas que lo convierten en una opción preferida en la investigación experimental. Una de las principales ventajas es su capacidad para eliminar sesgos en la asignación de tratamientos. Al asignar tratamientos de manera aleatoria, se asegura que cada unidad experimental tenga la misma probabilidad de recibir cualquier tratamiento, lo que minimiza la influencia de variables confusoras (Montgomery, 2017). Esta característica es fundamental para garantizar la validez interna del estudio.

Otra ventaja significativa del DCA es la generalización de resultados. Al eliminar sesgos y asegurar la aleatorización, los resultados obtenidos en el estudio pueden ser aplicables a una población más amplia. Esto es especialmente relevante en estudios clínicos, donde los hallazgos pueden influir en las decisiones de tratamiento para diferentes grupos de pacientes (Friedman et al., 2010). La capacidad de generalizar los resultados aumenta la relevancia y utilidad práctica de la investigación.

El DCA también permite una mayor simplicidad en el análisis estadístico. La aleatorización facilita la aplicación de métodos estadísticos estándar, como ANOVA y regresión, que requieren la suposición de independencia entre las observaciones (Kirk, 2013). Esta simplicidad en el análisis ayuda a los investigadores a interpretar los resultados de manera más clara y a comunicar sus hallazgos de forma efectiva.

Además, el DCA promueve la reproducibilidad de los estudios. La aleatorización y la transparencia en el proceso de asignación permiten que otros investigadores

reproduzcan el estudio de manera más sencilla. Esto es crucial en la ciencia, donde la reproducibilidad es un principio fundamental para validar los hallazgos (Gelman & Hill, 2007). La capacidad de replicar resultados aumenta la confianza en la validez de las conclusiones.

Una ventaja adicional del DCA es su flexibilidad en el diseño experimental. Los investigadores pueden adaptar el DCA a diferentes contextos y tipos de estudios, desde ensayos clínicos hasta estudios de mercado. Esta flexibilidad permite que el DCA sea utilizado en una variedad de disciplinas, lo que lo convierte en una herramienta valiosa para investigadores de diversas áreas (Cohen, 1988). La adaptabilidad del DCA contribuye a su popularidad en la investigación.

El DCA también facilita la identificación de relaciones causales. Al eliminar sesgos y controlar variables confusoras, los investigadores pueden establecer relaciones más claras entre las variables independientes y dependientes. Esto es esencial en campos como la medicina, donde es crucial determinar la eficacia de un tratamiento (Friedman et al., 2010). La capacidad de identificar relaciones causales contribuye a la base de conocimientos en diversas disciplinas.

Otra ventaja del DCA es su capacidad para manejar la variabilidad. La aleatorización ayuda a equilibrar las características de las unidades experimentales, lo que reduce la variabilidad no deseada en los resultados. Esto es especialmente importante en estudios donde la variabilidad puede afectar significativamente las conclusiones (Montgomery, 2017). Al manejar la variabilidad de manera efectiva, el DCA mejora la precisión de los resultados.

El DCA también permite la evaluación de múltiples tratamientos de manera simultánea. Al asignar tratamientos aleatoriamente, los investigadores pueden comparar diferentes intervenciones en un mismo estudio, lo que optimiza el uso de recursos y tiempo. Esta capacidad para evaluar múltiples tratamientos es particularmente útil en ensayos clínicos donde se están probando nuevas terapias (Kirk, 2013). La evaluación simultánea de tratamientos proporciona información valiosa para la toma de decisiones.

Una ventaja importante del DCA es su potencial para reducir el tamaño de la muestra necesario para obtener resultados significativos. Al eliminar sesgos y

controlar la variabilidad, es posible detectar diferencias reales con un tamaño de muestra más pequeño en comparación con otros diseños experimentales (Cohen, 1988). Esto es particularmente ventajoso en situaciones donde los recursos son limitados o el reclutamiento de participantes es difícil.

El DCA también contribuye a la mejora de la calidad de la investigación. La implementación de un diseño riguroso y aleatorio ayuda a los investigadores a evitar errores sistemáticos que pueden comprometer la validez de los resultados. Esta mejora en la calidad de la investigación es esencial para avanzar en el conocimiento científico y para informar políticas basadas en evidencia (Gelman & Hill, 2007). La calidad de la investigación es un factor clave en la confianza pública en los hallazgos científicos.

Otra ventaja del DCA es su aplicabilidad en estudios longitudinales. Aunque el DCA se utiliza comúnmente en estudios transversales, también puede adaptarse a diseños longitudinales donde se evalúan tratamientos a lo largo del tiempo. La aleatorización sigue siendo efectiva en estos contextos, lo que permite a los investigadores obtener información valiosa sobre la evolución de los efectos de un tratamiento (Friedman et al., 2010). La flexibilidad del DCA lo hace útil en diversas configuraciones de investigación.

El DCA también permite una evaluación más precisa de los efectos de interacción. Al tener un diseño aleatorio, los investigadores pueden explorar cómo diferentes tratamientos interactúan entre sí y afectan los resultados. Esto es particularmente importante en estudios donde se sospecha que las interacciones pueden influir en la efectividad de un tratamiento (Montgomery, 2017). La capacidad de evaluar efectos de interacción contribuye a una comprensión más completa de los fenómenos estudiados.

Además, el DCA fomenta la transparencia en la investigación. Al documentar el proceso de aleatorización y la asignación de tratamientos, los investigadores pueden proporcionar una justificación clara de sus métodos. Esta transparencia es esencial para la revisión por pares y para la aceptación de los hallazgos en la comunidad científica (Kirk, 2013). La transparencia en la investigación es fundamental para mantener la integridad científica.

El DCA también tiene la ventaja de mejorar la comunicación de resultados. Los resultados obtenidos a partir de un diseño aleatorio suelen ser más fáciles de entender y comunicar a audiencias no técnicas. Esto es importante para la divulgación científica y para la toma de decisiones informadas en políticas públicas (Gelman & Hill, 2007). La capacidad de comunicar resultados de manera efectiva es esencial para el impacto de la investigación.

Otra ventaja del DCA es su apoyo en el desarrollo de hipótesis. La aleatorización permite a los investigadores explorar diferentes tratamientos y sus efectos, lo que puede llevar a la formulación de nuevas hipótesis para investigaciones futuras. Este proceso de descubrimiento es fundamental para el avance del conocimiento científico (Friedman et al., 2010). La generación de nuevas hipótesis es un componente clave en la investigación.

El DCA también permite una evaluación más clara de los efectos a largo plazo. Al realizar un seguimiento de las unidades experimentales a lo largo del tiempo, los investigadores pueden observar cómo los efectos de un tratamiento pueden cambiar. Esta capacidad para evaluar efectos a largo plazo es crucial en campos como la medicina, donde los efectos de un tratamiento pueden no ser evidentes de inmediato (Montgomery, 2017). La evaluación de efectos a largo plazo contribuye a una comprensión más completa de los tratamientos.

Además, el DCA ayuda a identificar subgrupos dentro de la población. Al asignar tratamientos aleatoriamente, los investigadores pueden analizar cómo diferentes subgrupos responden a un tratamiento específico. Esta información es valiosa para personalizar tratamientos y mejorar la eficacia de las intervenciones (Kirk, 2013). La identificación de subgrupos es esencial para la medicina personalizada y para mejorar la atención al paciente.

El DCA también fomenta la colaboración entre investigadores. La claridad y la transparencia en el diseño experimental facilitan la colaboración entre diferentes grupos de investigación. Esta colaboración es fundamental para abordar preguntas complejas y para avanzar en el conocimiento científico (Gelman & Hill, 2007). La colaboración entre investigadores puede llevar a descubrimientos significativos y a una mejor comprensión de los fenómenos estudiados.

Otra ventaja del DCA es su contribución a la ética en la investigación. Al utilizar la aleatorización, los investigadores pueden asegurar que todos los participantes tengan la misma oportunidad de recibir un tratamiento, lo que es fundamental para la equidad en la investigación. Esta equidad es especialmente importante en estudios clínicos donde se evalúan tratamientos que pueden afectar la salud de los participantes (Friedman et al., 2010). La ética en la investigación es un principio fundamental que debe ser respetado.

El DCA también permite un análisis más robusto de los datos. La aleatorización ayuda a controlar variables confusoras, lo que permite a los investigadores aplicar análisis estadísticos más sofisticados. Esto es esencial para obtener conclusiones precisas y para informar decisiones basadas en evidencia (Montgomery, 2017). La robustez en el análisis de datos es crucial para la validez de los hallazgos.

Además, el DCA facilita la identificación de efectos adversos. Al tener un diseño aleatorio, los investigadores pueden observar y documentar efectos adversos de manera más efectiva. Esta identificación es crucial en estudios clínicos, donde la seguridad del paciente es una prioridad (Kirk, 2013). La identificación de efectos adversos contribuye a la mejora de los tratamientos y a la protección de los pacientes.

3.1.3. Limitaciones y Consideraciones del DCA

El Diseño Completamente al Azar (DCA), aunque es una metodología poderosa, presenta varias limitaciones y consideraciones que los investigadores deben tener en cuenta. Una de las principales limitaciones es la dependencia del tamaño de la muestra. Para que el DCA sea efectivo, se requiere un tamaño de muestra suficientemente grande para asegurar que los grupos sean comparables y que los resultados sean estadísticamente significativos (Cohen, 1988). Si el tamaño de la muestra es pequeño, puede haber un mayor riesgo de errores tipo I y tipo II, comprometiendo la validez de los resultados.

Otra limitación del DCA es la dificultad en la implementación práctica. En algunos contextos, puede ser complicado asignar tratamientos de manera completamente aleatoria, especialmente en estudios donde los participantes no

pueden ser seleccionados al azar debido a consideraciones éticas o logísticas (Friedman et al., 2010). Esto puede llevar a un diseño experimental que no cumple con los principios del DCA, afectando la validez de los hallazgos.

Además, el DCA asume que todos los participantes tienen características homogéneas. Sin embargo, en la práctica, las poblaciones pueden ser heterogéneas, lo que puede influir en la respuesta a los tratamientos. Esta variabilidad no controlada puede introducir sesgos y afectar la generalización de los resultados (Montgomery, 2017). Por lo tanto, es crucial considerar las características de la población al diseñar el estudio.

El DCA también puede ser sensible a la falta de adherencia de los participantes. Si algunos participantes no siguen el tratamiento asignado, esto puede introducir un sesgo en los resultados. La falta de adherencia puede ser un problema común en estudios clínicos, donde los pacientes pueden no seguir las recomendaciones del tratamiento (Kirk, 2013). Esta falta de adherencia puede afectar la validez interna del estudio y complicar la interpretación de los resultados.

Otra consideración importante es el costo y los recursos necesarios para implementar un DCA. La aleatorización puede requerir más tiempo y recursos en comparación con otros diseños experimentales, especialmente en estudios a gran escala. Los investigadores deben evaluar si los beneficios del DCA justifican los costos adicionales (Gelman & Hill, 2007). En situaciones donde los recursos son limitados, puede ser necesario considerar diseños alternativos.

El DCA también puede presentar desafíos en la recolección de datos. La necesidad de seguir un protocolo de aleatorización puede complicar la logística del estudio, especialmente en ensayos clínicos donde se requiere un seguimiento riguroso de los participantes (Friedman et al., 2010). La complejidad en la recolección de datos puede aumentar el riesgo de errores y afectar la calidad de los resultados.

Otra limitación del DCA es la imposibilidad de controlar todas las variables confusoras. Aunque el diseño aleatorio ayuda a equilibrar las características de los grupos, no siempre es posible controlar todas las variables que pueden influir en los resultados (Cohen, 1988). Esto puede ser especialmente problemático en

estudios en los que hay múltiples factores que interactúan entre sí. La falta de control sobre estas variables puede limitar la capacidad de hacer inferencias causales.

Además, el DCA puede no ser adecuado para todos los tipos de estudios. En investigaciones observacionales o en estudios donde la aleatorización no es ética, el DCA puede no ser la mejor opción. Los investigadores deben considerar el contexto y los objetivos del estudio al elegir el diseño más apropiado (Montgomery, 2017). La elección del diseño debe basarse en una evaluación cuidadosa de las circunstancias específicas del estudio.

La duración del estudio también puede ser una limitación del DCA. En algunos casos, los efectos de un tratamiento pueden no ser evidentes de inmediato, lo que puede requerir un seguimiento prolongado. La aleatorización puede complicar la planificación de estudios a largo plazo, especialmente si los participantes se retiran o si hay cambios en las condiciones externas que pueden afectar los resultados (Kirk, 2013). La duración del estudio debe ser considerada al diseñar experimentos.

Otra consideración es la ética de la aleatorización. En algunos casos, puede no ser ético asignar tratamientos de manera aleatoria, especialmente en estudios relacionados con la salud. Los investigadores deben equilibrar la necesidad de obtener resultados válidos con la responsabilidad ética de proteger a los participantes (Friedman et al., 2010). La ética en la investigación es un aspecto fundamental que no debe ser ignorado.

El DCA también puede ser limitado por la falta de control sobre el entorno en el que se llevan a cabo los experimentos. Las condiciones externas pueden influir en los resultados, y la aleatorización no puede controlar todos estos factores. Esto puede ser particularmente relevante en estudios de campo, donde las condiciones pueden ser difíciles de controlar (Gelman & Hill, 2007). La falta de control sobre el entorno puede afectar la validez externa del estudio.

Además, la interpretación de los resultados puede ser complicada en un DCA. Aunque la aleatorización ayuda a establecer relaciones causales, los investigadores deben tener cuidado al interpretar los resultados, especialmente

si hay factores no controlados que pueden influir en los hallazgos (Montgomery, 2017). La interpretación cuidadosa de los resultados es crucial para evitar conclusiones erróneas.

El DCA también puede ser afectado por la variabilidad en la respuesta al tratamiento. Las diferencias individuales en la respuesta a los tratamientos pueden complicar la interpretación de los resultados y hacer que sea difícil detectar efectos significativos (Kirk, 2013). Esta variabilidad en la respuesta puede ser un desafío en estudios clínicos, donde las diferencias en la biología de los pacientes pueden influir en los resultados.

Otra limitación es la posibilidad de sesgos en la selección de participantes. Aunque el DCA busca eliminar sesgos mediante la aleatorización, la forma en que se seleccionan los participantes puede introducir sesgos en el estudio. Los investigadores deben asegurarse de que la selección de participantes sea representativa de la población objetivo (Cohen, 1988). La selección adecuada de participantes es esencial para la validez del estudio.

Además, el DCA puede no ser efectivo en situaciones donde los efectos del tratamiento son pequeños. En tales casos, puede ser difícil detectar diferencias significativas entre grupos, incluso con un diseño aleatorio (Friedman et al., 2010). Esto puede ser un desafío en estudios clínicos donde se están evaluando tratamientos que tienen efectos modestos.

La complejidad del análisis de datos también puede ser una limitación del DCA. Aunque el diseño aleatorio facilita ciertos análisis estadísticos, los investigadores deben estar preparados para manejar la complejidad que puede surgir al analizar datos de estudios aleatorios (Gelman & Hill, 2007). La complejidad en el análisis de datos requiere habilidades estadísticas avanzadas y puede representar un desafío para algunos investigadores.

3.1.4. Ventajas y desventajas del DCA

El Diseño Completamente al Azar (DCA) es una metodología ampliamente utilizada en la investigación experimental que presenta tanto ventajas como desventajas. Entre las principales ventajas se encuentra su capacidad para eliminar sesgos en la asignación de tratamientos. Al asignar tratamientos de

manera aleatoria, se asegura que cada participante tenga la misma probabilidad de recibir cualquier intervención, lo que minimiza la influencia de variables confusoras y mejora la validez interna del estudio (Montgomery, 2017). Esta característica es fundamental para obtener resultados confiables y aplicables a la población general.

Otra ventaja significativa del DCA es su flexibilidad. Este diseño puede adaptarse a diversos contextos y tipos de estudios, desde ensayos clínicos hasta investigaciones de mercado. Esta versatilidad permite a los investigadores aplicar el DCA en una variedad de disciplinas, lo que aumenta su relevancia y utilidad práctica (Kirk, 2013). La capacidad de adaptarse a diferentes situaciones es un factor clave que contribuye a la popularidad del DCA en la investigación.

Sin embargo, el DCA también presenta desventajas que deben ser consideradas. Una de las principales limitaciones es la dependencia del tamaño de la muestra. Para que el DCA sea efectivo, se requiere un tamaño de muestra suficientemente grande que garantice la comparabilidad de los grupos y la significancia estadística de los resultados (Cohen, 1988). Si el tamaño de la muestra es pequeño, puede haber un mayor riesgo de errores en las conclusiones, lo que compromete la validez del estudio.

Además, la implementación práctica del DCA puede ser complicada en algunos contextos. Situaciones en las que los participantes no pueden ser asignados al azar por razones éticas o logísticas pueden limitar la aplicabilidad del DCA (Friedman et al., 2010). Esta dificultad en la implementación puede llevar a la necesidad de considerar otros diseños experimentales que se adapten mejor a las circunstancias específicas del estudio.

Otra desventaja del DCA es la dificultad para controlar variables confusoras. Aunque la aleatorización ayuda a equilibrar las características de los grupos, no siempre es posible controlar todas las variables que pueden influir en los resultados (Montgomery, 2017). Esto puede ser especialmente problemático en estudios donde múltiples factores interactúan, lo que limita la capacidad de hacer inferencias causales precisas.

3.1.5. Ventajas

El Diseño Completamente al Azar (DCA) es una metodología fundamental en la investigación experimental que ofrece múltiples beneficios. Al permitir la asignación aleatoria de tratamientos a los participantes, el DCA minimiza sesgos y proporciona una base sólida para inferencias causales. A continuación, se presentan diez ventajas clave de este enfoque:

1. Eliminación de sesgos: La aleatorización asegura que los grupos sean comparables, reduciendo la influencia de variables confusoras.
2. Validez interna: Al controlar mejor las variables, se incrementa la validez interna del estudio, lo que permite conclusiones más precisas sobre las relaciones causales.
3. Flexibilidad: El DCA puede aplicarse en diversas disciplinas, desde la medicina hasta la psicología, adaptándose a diferentes contextos de investigación.
4. Simplicidad en el análisis: Los datos obtenidos en un DCA suelen ser más fáciles de analizar estadísticamente, lo que facilita la interpretación de los resultados.
5. Generalización de resultados: La aleatorización mejora la capacidad de generalizar los hallazgos a la población más amplia, aumentando la aplicabilidad de los resultados.
6. Control de variabilidad: Al distribuir aleatoriamente los participantes, se controla la variabilidad entre grupos, lo que mejora la precisión de las estimaciones.
7. Facilitación de meta-análisis: Los estudios que utilizan DCA son más fáciles de incluir en meta-análisis, lo que puede aumentar el poder estadístico y la robustez de las conclusiones.
8. Reducción de la influencia de factores externos: La aleatorización ayuda a mitigar el impacto de factores externos que pueden afectar los resultados del estudio.
9. Mejora en la calidad de los datos: La asignación aleatoria puede llevar a una mejor adherencia al tratamiento, lo que resulta en datos más confiables.

10. Ética en la investigación: Cuando se justifica, el DCA puede ser considerado ético, ya que todos los participantes tienen la misma oportunidad de recibir el tratamiento.

3.1.6. Desventajas

A pesar de sus numerosas ventajas, el Diseño Completamente al Azar (DCA) también presenta ciertas desventajas que los investigadores deben considerar cuidadosamente. Estas limitaciones pueden afectar la implementación y la interpretación de los resultados. A continuación, se enumeran diez desventajas importantes:

1. Dependencia del tamaño de la muestra: Un tamaño de muestra pequeño puede comprometer la validez de los resultados y aumentar el riesgo de errores estadísticos.
2. Dificultades en la implementación: En algunos contextos, la asignación aleatoria puede ser poco práctica o poco ética, limitando la aplicabilidad del DCA.
3. Control limitado de variables confusoras: Aunque se busca equilibrar los grupos, no siempre se pueden controlar todas las variables que pueden influir en los resultados.
4. Problemas de adherencia: La falta de adherencia al tratamiento puede sesgar los resultados, afectando la validez interna del estudio.
5. Costos y recursos: La implementación del DCA puede requerir más tiempo y recursos en comparación con otros diseños experimentales.
6. Complejidad en el análisis: Aunque el análisis de datos puede ser más sencillo, la variabilidad en las respuestas puede complicar la interpretación de los resultados.
7. Ética de la aleatorización: En ciertos estudios, puede no ser ético asignar tratamientos de manera aleatoria, lo que limita el uso del DCA.
8. Dificultades en estudios a largo plazo: La duración del estudio puede ser un desafío, especialmente si los efectos del tratamiento no son inmediatos.

9. Sesgos en la selección de participantes: La forma en que se seleccionan los participantes puede introducir sesgos, afectando la representatividad del estudio.
10. Menor idoneidad para estudios exploratorios: Para investigaciones que buscan explorar nuevas ideas, el DCA puede ser demasiado rígido y limitar la creatividad en el diseño experimental.

El Diseño Completamente al Azar (DCA) es una herramienta valiosa en la investigación que ofrece una serie de ventajas significativas, como la eliminación de sesgos y la mejora de la validez interna. Sin embargo, también presenta desventajas que pueden afectar su implementación y la interpretación de los resultados.

Es fundamental que los investigadores evalúen cuidadosamente tanto los beneficios como las limitaciones del DCA al diseñar sus estudios. Al hacerlo, pueden maximizar la calidad y la aplicabilidad de sus hallazgos, contribuyendo así al avance del conocimiento en sus respectivas disciplinas. La elección del diseño adecuado depende de las circunstancias específicas del estudio y de los objetivos de investigación.

3.2. Modelo lineal aditivo para el DCA

El modelo lineal aditivo es una de las metodologías más utilizadas en el análisis de datos experimentales, especialmente en el contexto del Diseño Completamente al Azar (DCA). Este enfoque estadístico permite a los investigadores evaluar la relación entre una variable dependiente y múltiples variables independientes, facilitando la identificación de efectos significativos de los tratamientos aplicados. En un entorno donde se busca minimizar sesgos y maximizar la validez interna, el modelo lineal aditivo se convierte en una herramienta esencial para el análisis de datos.

Una de las características distintivas del modelo lineal aditivo es su simplicidad. Al asumir que los efectos de los factores son aditivos, permite a los investigadores construir modelos que son fáciles de interpretar y aplicar. Esta simplicidad es particularmente valiosa en estudios donde se desea comunicar resultados a audiencias no técnicas, asegurando que los hallazgos sean

accesibles y comprensibles. Además, la estructura lineal del modelo facilita la estimación de efectos y la realización de inferencias estadísticas.

El modelo lineal aditivo también es flexible, lo que permite su aplicación en una variedad de contextos y disciplinas. Desde la medicina hasta la psicología y la economía, este modelo puede adaptarse a diferentes tipos de datos y diseños experimentales. Esta versatilidad lo convierte en una opción atractiva para investigadores que buscan un enfoque robusto y confiable para el análisis de datos obtenidos a través del DCA. La capacidad de manejar múltiples factores y sus interacciones es una de las razones por las que este modelo se ha vuelto tan popular en la investigación.

Además, el uso del modelo lineal aditivo en el DCA permite realizar análisis de varianza (ANOVA), una técnica que ayuda a determinar si existen diferencias significativas entre los grupos tratados. Esta capacidad para evaluar la variabilidad entre tratamientos es crucial para validar los resultados experimentales y hacer inferencias sobre la población general. Al proporcionar una base estadística sólida, el modelo lineal aditivo contribuye a la credibilidad y la validez de las conclusiones extraídas del estudio.

Sin embargo, es importante reconocer que, aunque el modelo lineal aditivo tiene muchas ventajas, también presenta ciertas limitaciones. Por ejemplo, el supuesto de linealidad puede no ser válido en todos los contextos, lo que podría afectar la precisión de las estimaciones. Además, la presencia de interacciones complejas entre factores puede requerir modelos más sofisticados que vayan más allá de un enfoque puramente aditivo. Por lo tanto, los investigadores deben evaluar cuidadosamente la adecuación del modelo lineal aditivo para sus datos y objetivos específicos.

3.2.1. Fundamentos del Modelo Lineal Aditivo en el DCA

El modelo lineal aditivo es un enfoque estadístico que se utiliza para analizar la relación entre una variable dependiente y una o más variables independientes. En el contexto del Diseño Completamente al Azar (DCA), este modelo es fundamental para evaluar los efectos de diferentes tratamientos asignados a los participantes (Montgomery, 2017). La premisa básica del modelo es que la

variable dependiente puede ser expresada como una combinación lineal de los efectos de los factores involucrados, lo que permite interpretar los resultados de manera clara y concisa.

Una de las características más importantes del modelo lineal aditivo es su capacidad para manejar múltiples variables. Esto significa que los investigadores pueden incluir varios factores en el análisis, lo que les permite explorar interacciones y efectos combinados (Gelman & Hill, 2007). Al hacerlo, se obtiene una comprensión más completa de cómo diferentes tratamientos afectan la variable de interés. Esta flexibilidad es especialmente útil en estudios donde se desea evaluar la eficacia de múltiples intervenciones simultáneamente.

El modelo lineal aditivo también se basa en ciertos supuestos estadísticos que deben cumplirse para que los resultados sean válidos. Estos supuestos incluyen la linealidad, la independencia de los errores, la homocedasticidad y la normalidad de los residuos (Kutner, Nachtsheim, & Neter, 2004). Si bien estos supuestos son fundamentales para la validez del modelo, es importante que los investigadores realicen pruebas adecuadas para asegurarse de que se cumplen. De no ser así, podrían considerar transformaciones de datos o el uso de modelos alternativos.

La interpretación de los coeficientes en un modelo lineal aditivo es relativamente directa. Cada coeficiente representa el cambio esperado en la variable dependiente por cada unidad de cambio en la variable independiente correspondiente, manteniendo constantes las demás variables (Field, 2013). Esto permite a los investigadores cuantificar el impacto de cada tratamiento de manera clara, facilitando la comunicación de los resultados a audiencias tanto técnicas como no técnicas.

En el contexto del DCA, el modelo lineal aditivo permite realizar análisis de varianza (ANOVA), una técnica que ayuda a determinar si existen diferencias significativas entre los grupos tratados (Cohen, 1988). ANOVA descompone la variabilidad total en componentes atribuibles a diferentes fuentes, lo que permite evaluar la efectividad de los tratamientos en comparación con un grupo de control. Este enfoque es esencial para validar los resultados experimentales y hacer inferencias sobre la población general.

Una de las ventajas del modelo lineal aditivo es su robustez ante violaciones de supuestos. Aunque es ideal que se cumplan los supuestos, en la práctica, el modelo puede seguir proporcionando resultados útiles incluso cuando algunos de ellos son violados (Huitema & McKean, 2000). Esto lo convierte en una herramienta valiosa en situaciones donde los datos no se ajustan perfectamente a las condiciones ideales, permitiendo a los investigadores obtener conclusiones significativas.

El modelo lineal aditivo también es esencial para la modelización de datos longitudinales. En estudios donde se recopilan datos en múltiples puntos en el tiempo, este enfoque permite analizar cómo los tratamientos afectan a los participantes a lo largo del tiempo (McCullagh & Nelder, 1989). Esto es particularmente relevante en investigaciones médicas y psicológicas, donde los efectos de un tratamiento pueden no ser inmediatos y pueden variar con el tiempo.

Además, el modelo lineal aditivo puede ser extendido para incluir interacciones entre variables. Esto significa que los investigadores pueden explorar cómo el efecto de un tratamiento puede depender del nivel de otra variable (Ritchie, Lewis, & Elam, 2003). Por ejemplo, el impacto de un fármaco puede variar según la edad o el género del paciente. Incluir estas interacciones en el modelo puede proporcionar una comprensión más matizada de los resultados y ayudar a identificar subgrupos que podrían beneficiarse más de un tratamiento específico.

La validación del modelo es un aspecto crucial en el análisis de datos. Los investigadores deben asegurarse de que el modelo se ajuste bien a los datos observados y que sea capaz de predecir resultados en nuevos conjuntos de datos (Weisberg, 2005). Esto se puede lograr mediante técnicas como la validación cruzada, que permite evaluar la capacidad predictiva del modelo en datos no utilizados durante el ajuste inicial. Un modelo bien validado es esencial para garantizar la credibilidad de los hallazgos de investigación.

El uso del modelo lineal aditivo en el DCA también permite la evaluación de la significancia estadística de los efectos observados. A través de pruebas de hipótesis, los investigadores pueden determinar si las diferencias en los

resultados entre grupos son lo suficientemente grandes como para ser consideradas significativas (Cohen, 1988). Esto ayuda a evitar conclusiones erróneas basadas en variaciones aleatorias y proporciona un marco sólido para la toma de decisiones informadas.

La interpretación de los resultados en el contexto del modelo lineal aditivo requiere una comprensión clara de las métricas estadísticas involucradas. Los investigadores deben estar familiarizados con conceptos como el valor p , los intervalos de confianza y los tamaños del efecto (Field, 2013). Estos elementos son fundamentales para comunicar los hallazgos de manera efectiva y para garantizar que los resultados sean interpretados correctamente por otros investigadores y profesionales en el campo.

El modelo lineal aditivo también se puede utilizar en combinación con otras técnicas estadísticas, como la regresión múltiple. Esta combinación permite a los investigadores explorar relaciones más complejas y ajustar sus modelos para tener en cuenta múltiples factores simultáneamente (Gelman & Hill, 2007). Al integrar diferentes enfoques, los investigadores pueden obtener una visión más completa y precisa de los efectos de los tratamientos en sus estudios.

La visualización de datos es otro aspecto importante al trabajar con el modelo lineal aditivo. Utilizar gráficos y diagramas para representar los resultados puede ayudar a los investigadores a comunicar sus hallazgos de manera más efectiva (Montgomery, 2017). Gráficos de dispersión, diagramas de caja y gráficos de líneas son herramientas útiles para ilustrar las relaciones entre variables y destacar diferencias significativas entre grupos.

Además, el modelo lineal aditivo es particularmente útil en el contexto de experimentos controlados. En estos estudios, los investigadores pueden manipular variables independientes de manera precisa y observar sus efectos en la variable dependiente (Kutner, Nachtsheim, & Neter, 2004). Esto permite establecer relaciones causales más claras y proporciona un marco sólido para la inferencia estadística.

La revisión de la literatura existente sobre el modelo lineal aditivo y su aplicación en el DCA es esencial para comprender su evolución y relevancia en la

investigación actual. A medida que los métodos estadísticos continúan desarrollándose, es importante que los investigadores se mantengan actualizados sobre las mejores prácticas y enfoques emergentes en el análisis de datos (Weisberg, 2005). Esto incluye estar al tanto de nuevas técnicas y software que pueden facilitar el uso del modelo lineal aditivo.

Como se explicó en capítulos anteriores este primer modelo solo considera como su principal fuente de variación a los tratamientos, intuyendo que las unidades experimentales son totalmente homogéneas, de hecho es el más simple y efectivo, y se sintetiza en que cada variable respuesta esta en función de la media general, el efecto de los tratamientos y del error experimental.

$$Y_{ij} = \mu + \tau_i + \epsilon_{ij}$$

3.2.2. Aplicaciones del Modelo Lineal Aditivo en el DCA

El modelo lineal aditivo se ha consolidado como una herramienta fundamental en el análisis de datos experimentales, especialmente en el contexto del Diseño Completamente al Azar (DCA). Su principal aplicación radica en la evaluación de los efectos de diferentes tratamientos asignados a los sujetos de estudio. Este enfoque permite a los investigadores descomponer la variabilidad total en componentes atribuibles a los tratamientos, facilitando la identificación de diferencias significativas entre grupos (Montgomery, 2017).

Una de las ventajas más destacadas del modelo lineal aditivo es su simplicidad en la interpretación de los resultados. Cada coeficiente del modelo representa el efecto de una variable independiente sobre la variable dependiente, lo que permite a los investigadores cuantificar de manera clara el impacto de cada tratamiento. Esta claridad es especialmente valiosa cuando se comunican los hallazgos a audiencias no técnicas, asegurando que los resultados sean accesibles y comprensibles (Field, 2013).

El modelo lineal aditivo también es flexible, permitiendo la inclusión de múltiples factores en el análisis. Esto significa que los investigadores pueden evaluar simultáneamente los efectos de diferentes tratamientos y sus interacciones, lo que proporciona una comprensión más completa de cómo estas variables

influyen en la variable de interés (Gelman & Hill, 2007). Esta capacidad para manejar múltiples variables es esencial en estudios donde se busca evaluar la eficacia de diversas intervenciones.

Además, el modelo permite realizar análisis de varianza (ANOVA), lo que ayuda a determinar si las diferencias observadas entre los grupos son estadísticamente significativas (Cohen, 1988). ANOVA descompone la variabilidad total en componentes atribuibles a los tratamientos y a los errores, proporcionando un marco robusto para la evaluación de la efectividad de los tratamientos. Esto es crucial para validar los resultados experimentales y hacer inferencias sobre la población general.

Una de las aplicaciones más relevantes del modelo lineal aditivo es en el ámbito de la investigación médica. Los ensayos clínicos, que a menudo utilizan diseños completamente al azar, se benefician enormemente de este enfoque. Al aplicar el modelo, los investigadores pueden evaluar la eficacia de nuevos tratamientos, comparando sus efectos con un grupo de control, lo que es fundamental para la aprobación de medicamentos (Huitema & McKean, 2000).

En el ámbito de la psicología, el modelo lineal aditivo se utiliza para analizar la efectividad de diferentes intervenciones terapéuticas. Por ejemplo, al evaluar programas de tratamiento para trastornos de ansiedad, los investigadores pueden aplicar el modelo para determinar qué enfoques son más efectivos y en qué condiciones. Esto permite personalizar las intervenciones y mejorar los resultados para los pacientes (Weisberg, 2005).

El modelo lineal aditivo también es útil en investigaciones educativas, donde se evalúan diferentes métodos de enseñanza. Al aplicar este enfoque, los investigadores pueden comparar el rendimiento académico de estudiantes que han recibido diferentes tipos de instrucción, identificando qué métodos resultan más efectivos. Esta información es valiosa para diseñar currículos y mejorar la calidad de la educación (Ritchie, Lewis, & Elam, 2003).

Otra ventaja significativa del modelo es su robustez ante violaciones de supuestos. Aunque es ideal que se cumplan los supuestos de linealidad y homocedasticidad, el modelo puede seguir proporcionando resultados útiles

incluso cuando estos supuestos no se cumplen estrictamente (Kutner, Nachtsheim, & Neter, 2004). Esto lo convierte en una herramienta valiosa en situaciones donde los datos no se ajustan perfectamente a las condiciones ideales.

El modelo lineal aditivo también se puede aplicar en el análisis de datos longitudinales. En estudios donde se recopilan datos en múltiples momentos, este enfoque permite a los investigadores analizar cómo los tratamientos afectan a los participantes a lo largo del tiempo. Esto es especialmente relevante en investigaciones médicas y psicológicas, donde los efectos de un tratamiento pueden no ser inmediatos y pueden variar con el tiempo (McCullagh & Nelder, 1989).

La visualización de datos es otro aspecto importante que se beneficia del uso del modelo lineal aditivo. Utilizar gráficos y diagramas para representar los resultados puede ayudar a los investigadores a comunicar sus hallazgos de manera más efectiva. Gráficos de dispersión, diagramas de caja y gráficos de líneas son herramientas útiles para ilustrar las relaciones entre variables y destacar diferencias significativas entre grupos (Gelman & Hill, 2007).

En el contexto de experimentos controlados, el modelo lineal aditivo permite a los investigadores manipular variables independientes de manera precisa y observar sus efectos en la variable dependiente. Esto facilita el establecimiento de relaciones causales más claras y proporciona un marco sólido para la inferencia estadística. Esta capacidad es fundamental para la validez de los hallazgos experimentales (Montgomery, 2017).

La validación del modelo es un aspecto crucial en el análisis de datos. Los investigadores deben asegurarse de que el modelo se ajuste bien a los datos observados y que sea capaz de predecir resultados en nuevos conjuntos de datos. Esto se puede lograr mediante técnicas como la validación cruzada, que permite evaluar la capacidad predictiva del modelo en datos no utilizados durante el ajuste inicial (Weisberg, 2005).

El modelo lineal aditivo también permite la evaluación de la significancia estadística de los efectos observados. A través de pruebas de hipótesis, los

investigadores pueden determinar si las diferencias en los resultados entre grupos son lo suficientemente grandes como para ser consideradas significativas. Esto ayuda a evitar conclusiones erróneas basadas en variaciones aleatorias y proporciona un marco sólido para la toma de decisiones informadas (Cohen, 1988).

En el ámbito de la economía, el modelo lineal aditivo se utiliza para analizar el impacto de políticas económicas. Al evaluar cómo diferentes políticas afectan indicadores económicos, los investigadores pueden identificar qué enfoques son más efectivos y en qué condiciones. Esto permite a los responsables de políticas tomar decisiones informadas basadas en evidencia sólida (Ritchie, Lewis, & Elam, 2003).

El uso del modelo lineal aditivo también fomenta la transparencia en la investigación. Al proporcionar un marco claro y estructurado para el análisis de datos, los investigadores pueden documentar sus métodos y resultados de manera más efectiva. Esto facilita la replicación de estudios y el avance del conocimiento en diversas disciplinas (Gelman & Hill, 2007).

La formación en el uso del modelo lineal aditivo es crucial para garantizar que los investigadores puedan aplicar esta metodología de manera efectiva. La educación en estadística y análisis de datos debe ser parte integral de la formación de los investigadores, permitiéndoles utilizar herramientas estadísticas de manera competente y ética. Esto no solo mejora la calidad de la investigación, sino que también contribuye al avance del conocimiento en diversas disciplinas (Field, 2013).

3.2.3. Fórmulas de cálculo para el DCA con igual repetición

El Diseño Completamente al Azar (DCA) es una de las metodologías más utilizadas en la investigación experimental, especialmente en campos como la biología, la medicina y la psicología; así como en nuestro caso la agropecuaria. Este enfoque permite a los investigadores evaluar el efecto de diferentes tratamientos o condiciones sobre una variable dependiente, asegurando que la asignación de tratamientos se realice de manera aleatoria. La aleatorización es

crucial, ya que minimiza sesgos y asegura que los resultados sean generalizables a la población más amplia.

En el contexto del DCA, las fórmulas de cálculo son herramientas esenciales que permiten a los investigadores analizar y interpretar los datos obtenidos de sus experimentos. Estas fórmulas ayudan a calcular medidas estadísticas clave, como medias, varianzas y sumas de cuadrados, que son fundamentales para llevar a cabo un análisis de varianza (ADEVA o ANOVA). Este último es un método que permite determinar si existen diferencias significativas entre los tratamientos aplicados, proporcionando así una base sólida para la toma de decisiones.

El DCA con igual repetición se refiere a la situación en la que cada tratamiento se aplica el mismo número de veces a diferentes sujetos. Esta característica simplifica el análisis estadístico, ya que permite una comparación directa entre los tratamientos sin complicaciones adicionales. La igual repetición también ayuda a aumentar la precisión de las estimaciones, ya que se reduce la variabilidad dentro de cada grupo de tratamiento, lo que a su vez mejora la capacidad del estudio para detectar efectos reales.

La correcta aplicación de las fórmulas de cálculo en el DCA es fundamental para garantizar la validez de los resultados. Los investigadores deben familiarizarse con las diferentes métricas y métodos de análisis para asegurar que sus conclusiones sean robustas y confiables. Esto incluye no solo el cálculo de medias y varianzas, sino también la interpretación de los resultados del ANOVA, que proporciona información sobre la significancia estadística de los efectos observados.

A lo largo de este epígrafe, se explorarán en detalle las fórmulas de cálculo específicas utilizadas en el DCA con igual repetición. Se presentarán ejemplos prácticos y se discutirán las implicaciones de estos cálculos en la interpretación de los resultados experimentales. Al final, se espera que los lectores comprendan la importancia de estas fórmulas en el contexto del DCA y cómo pueden aplicarse para mejorar la calidad de la investigación científica.

El procedimiento es el siguiente:

- Se organiza la información en un cuadro de resultados
- Calculamos las medias de los tratamientos, $\bar{X}_{\tau_i} = \frac{\sum \tau_i}{r_i}$ que en nuestro caso sería la suma de las observaciones de cada tratamiento, dividido para el total de repeticiones en cuestión.
- Sumatoria total $\Sigma X_{..} = X_1 + X_2 + X_3 + \dots + X_n$ que corresponde a la sumatoria del universo de las observaciones encontradas en la “data” de cada variable respuesta.
- Media general $\bar{X} = \frac{\Sigma X_{..}}{n}$ obtenida de la sumatoria de total dividido para las observaciones.
- Sacamos un factor de corrección $F.C. = \frac{(\Sigma X_{..})^2}{(t \cdot r)}$ expresión que nos apoyará para la suma de cuadrados de los componentes en las fuentes de variación; y se obtiene del cuadrado de la sumatoria total dividido para el universo de datos es decir en un DCA de igual repetición sería el número de tratamientos por las repeticiones.
- Procedemos al cálculo de las sumas de cuadrados. En la suma de cuadrados totales $SC_{TOT} = \Sigma X_{ij}^2 - FC$ la expresión indica la suma de todos y cada uno de los datos encontrados y elevados al cuadrado, resultado que luego se resta del factor de corrección
- Suma de cuadrados los tratamientos $SC_{TRAT} = \frac{\Sigma \tau_i^2}{r} - FC$ es la suma de cada grupo de observaciones por cada tratamiento al cuadrado y dividido para las repeticiones (DCA con igual repetición) y extraído el factor de corrección
- Suma de cuadrados del error $SC_{ERROR} = SC_{TOTAL} - SC_{TRAT}$ encontrado por diferencia de las sumas de cuadrados, del total menos SC de tratamientos.
- Una vez calculadas las sumas de cuadrados, el siguiente paso es determinar los grados de libertad (gl). Los grados de libertad totales se

obtienen del total de observaciones menos uno $GL_{TOT} = (n - 1)$; luego los GL de los tratamientos se calculan con el número de tratamientos menos uno $GL_{TRAT} = (t - 1)$; y los grados de libertad del error se obtienen por diferencia de los grados de libertad totales menos los grados de libertad de los tratamientos $GL_{ERROR} = GL_{TOT} - GL_{TRAT}$. Estos grados de libertad son fundamentales para calcular los estadísticos F.

- Los cuadrados medios o varianza de los tratamientos se procede a obtenerlos dividiendo la suma de cuadrados de los tratamientos para sus grados de libertad $CM_{TRAT} = \frac{SC_{TRAT}}{GL_{TRAT}}$; en cambio en el cuadrado medio del error se continúa con proceso similar, dividiendo su correspondiente suma de cuadrados para los grados de libertad $CM_{ERROR} = \frac{SC_{ERROR}}{GL_{ERROR}}$
- El estadístico Fisher calculado, así mismo se obtiene dividiendo el cuadrado medio de los tratamientos por el cuadrado medio del error $F_{CAL} = \frac{CM_{TRAT}}{CM_{ERROR}}$. Un valor alto de F indica que la variabilidad entre tratamientos es mayor que la variabilidad dentro de los tratamientos, lo que sugiere que existen diferencias significativas entre las medias de los tratamientos.
- Encontramos el coeficiente de variación $CV = \frac{\sqrt{CM_{ERROR}}}{\bar{X}} * 100$
- Para interpretar los resultados del análisis de varianza, se compara el valor de F calculado con un valor crítico de F obtenido de tablas estadísticas, según el nivel de significancia deseado (por lo general, 0.05). Si el valor de F calculado es mayor que el valor crítico, se rechaza la hipótesis nula, que establece que no hay diferencias significativas entre los tratamientos. Esto implica que al menos uno de los tratamientos tiene un efecto distinto sobre la variable dependiente.
- Es importante recordar que, aunque ANOVA puede indicar que hay diferencias significativas, no especifica cuáles tratamientos son diferentes. Para determinar qué tratamientos difieren entre sí, se pueden realizar análisis post hoc, como la prueba de Tukey o la prueba de

Bonferroni. Estas pruebas permiten realizar comparaciones múltiples de las medias de los tratamientos y ayudan a identificar cuáles son significativamente diferentes.

3.3. Importancia para el Análisis de Varianza (ANOVA)

Las métricas de medias y varianzas son cruciales para el análisis de varianza (ANOVA), ya que permiten descomponer la variabilidad total en componentes atribuibles a los tratamientos y a la variabilidad interna. En ANOVA, se comparan las varianzas entre tratamientos con la varianza dentro de los tratamientos para determinar si hay diferencias significativas entre las medias de los tratamientos. Esta comparación es fundamental para validar los resultados del experimento.

Un valor alto del estadístico F, que se calcula como la razón entre la varianza entre tratamientos y la varianza dentro de los tratamientos, sugiere que las diferencias observadas entre tratamientos son estadísticamente significativas. Esto permite a los investigadores hacer inferencias sobre el efecto de los tratamientos aplicados, facilitando la toma de decisiones informadas en base a los resultados del análisis.

3.4. Aplicación del DCA con igual repetición

Ejercicio 3.1.

Se estableció un experimento para mejorar el rendimiento del pasto festuca que alimentará a ovinos en la sierra ecuatoriana, se prueban 3 niveles de fertilización nitrogenada (75, 150, 225 kg/ha), con 5 repeticiones por tratamiento.

Problema:

Se encontró una producción de apenas 7,5 toneladas por hectárea de festuca en propiedades altas de la provincia de Chimborazo, cantón Guamote; considerado deficiente ya que según los estándares en ese tipo de ecosistemas es de 10 a 15 t/ha/corte.

Objetivo:

Mejorar el rendimiento de *Festuca arundinacea* mediante 3 niveles de fertilización nitrogenada (75, 150, 225 kg/ha), en granjas del cantón Guamote.

Hipótesis:

- *Teórica*

Los tres niveles de fertilización nitrogenada no presentaron un rendimiento aceptable en la productividad de *Festuca arundinacea* en las granjas del cantón Guamote.

- *Matemática*

$$H_0 = \mu_{n57} = \mu_{n150} = \mu_{n225}$$

Resultados experimentales

Tratamientos (t)	Repeticiones (r)					Suma tratam. $\sum \tau_i$	Media tratam. $\bar{X}_{\tau_i}$
	I	II	III	IV	V		
N (75 kg/ha)	14,8	14,6	14,7	14,5	15,0	73,6	14,7
N (150 kg/ha)	25,1	25,4	25,1	25,0	25,2	125,8	25,2
N (225 kg/ha)	32,6	32,4	32,2	32,6	32,1	161,9	32,4
						361,3	24,1
						Total $\sum X_{..}$	Media $\bar{X}$

Factor de corrección

$$F.C. = \frac{(\sum X_{..})^2}{(t * r)}$$

$$F.C. = \frac{(361,3)^2}{(3 * 5)} = \frac{130537,69}{15} = 8702,51$$

Suma de cuadrados

- Suma de cuadrados totales

$$SC_{TOT} = \sum X_{ij}^2 - FC$$

$$SC_{TOT} = [(14,8)^2 + (14,6)^2 + (14,7)^2 + \dots + (32,1)^2] - 8702,51$$

$$SC_{TOT} = 9491,29 - 8702,51 = 788,78$$

- Suma de cuadrados de los tratamientos

$$SC_{\text{TRAT}} = \frac{\sum \tau_i^2}{r} - FC$$

$$SC_{\text{TRAT}} = \frac{[(73,6)^2 + (125,8)^2 + (161,9)^2]}{5} - 8702,51$$

$$SC_{\text{TRAT}} = \frac{47454,21}{5} - 8702,51 = 9490,84 - 8702,51 = 788,33$$

- Suma de cuadrados del error

$$SC_{\text{ERROR}} = SC_{\text{TOTAL}} - SC_{\text{TRAT}}$$

$$SC_{\text{ERROR}} = 788,78 - 788,33 = 0,45$$

Grados de libertad

$$GL_{\text{TOT}} = (n - 1) = 15 - 1 = 14$$

$$GL_{\text{TRAT}} = (t - 1) = 3 - 1 = 2$$

$$GL_{\text{ERROR}} = GL_{\text{TOT}} - GL_{\text{TRAT}} = 14 - 2 = 12$$

Cuadrados medios

- Cuadrados medios de los tratamientos

$$CM_{\text{TRAT}} = \frac{SC_{\text{TRAT}}}{GL_{\text{TRAT}}}$$

$$CM_{\text{TRAT}} = \frac{788,33}{2} = 394,165$$

- Cuadrados medios del error

$$CM_{\text{ERROR}} = \frac{SC_{\text{ERROR}}}{GL_{\text{ERROR}}}$$

$$CM_{\text{ERROR}} = \frac{0,45}{12} = 0,037$$

- Estadístico Fisher calculado

$$F_{\text{CAL}} = \frac{CM_{\text{TRAT}}}{CM_{\text{ERROR}}}$$

$$F_{\text{CAL}} = \frac{394,165}{0,037} = 10557,98$$

- Coeficiente de variación del ADEVA

$$CV = \frac{\sqrt{CM_{\text{ERROR}}}}{\bar{X}} * 100$$

$$CV = \frac{\sqrt{0,037}}{24,09} * 100 = \frac{0,193}{24,09} * 100 = 0,80\%$$

ADEVA

Fuente de variación	Suma de cuadrados	Grados de libertad	Cuadrado medio	Fcal	F.Tab 0,05	F.Tab 0,01	Signif.
Total	788,777	14					
Tratamientos	788,329	2	394,165	10557,98	3,89	6,93	**
Error	0,448	12	0,037				
CV	0,80%						

Comparar el F. Cal con el F. Tab en la respectiva tabla con 2 grados de libertad para el numerador (tratamientos) y 12 grados de libertad para el denominador (error experimental), con los 2 niveles de significancia. Sería $\alpha 0,05 = 3,89$; y, $\alpha 0,01 = 6,93$.

Análisis del estadístico Fisher

Por cuanto el Fisher calculado (10557,98) supera a los valores del Fisher tabular (3,89 de $\alpha 0,05$ y 6,93 de $\alpha 0,01$), entonces la respuesta es altamente significativa.

Interpretación

Según el análisis de varianza, la respuesta del rendimiento de la *Festuca arundinacea*, frente a la aplicación de 3 dosis de fertilizante nitrogenado, tuvo diferencias altamente significativas ($P \leq 0,01$); por lo tanto, no existe evidencia para aceptar la hipótesis nula.

3.5. Análisis de Varianza (ANOVA) en DCA con desigual repetición

El análisis de varianza (DEVA) es una técnica estadística utilizada para comparar las medias de tres o más grupos y determinar si existen diferencias significativas entre ellos. En el contexto de un diseño completamente al azar (DCA) con desigual repetición, permite evaluar el efecto de diferentes tratamientos con distintas repeticiones sobre una variable dependiente. Este método es esencial para validar hipótesis en experimentos donde se busca entender la variabilidad entre tratamientos. Esta es una gran ventaja en referencia a los otros modelos lineales aditivos, ya que requieren de igual número de repeticiones por tratamiento.

3.6. Fórmulas de cálculo para el DCA con desigual repetición

El Diseño Completamente al Azar (DCA) es una de las metodologías más utilizadas en la investigación experimental debido a su flexibilidad y simplicidad. Sin embargo, cuando se trabaja con grupos de tratamientos que no tienen el mismo número de repeticiones, es necesario aplicar un enfoque específico que contemple esta desigualdad. Este tipo de diseño permite a los investigadores explorar el efecto de diferentes tratamientos sobre una variable dependiente, adaptándose a situaciones donde la asignación de sujetos a tratamientos no es uniforme.

En un DCA con desigual repetición, las fórmulas de cálculo son fundamentales para analizar los datos de manera efectiva. A diferencia del DCA con igual repetición, donde las sumas de cuadrados y los grados de libertad se calculan de manera más sencilla, en el caso de desigual repetición se requiere un análisis más detallado. Esto incluye la consideración de las variaciones en el número de observaciones por tratamiento, lo que puede influir en la interpretación de los resultados.

El procedimiento para calcular el ADEVA en un DCA con desigual número de repeticiones es el siguiente:

- Se organiza la información en un cuadro de resultados

- Calculamos Las medias de los tratamientos, $\bar{X}_{n_i} = \frac{\sum \tau_i}{r_i}$ que en esta ocasión individualmente se suma cada observación en cada tratamiento y se divide por el número de repeticiones que tenga la misma.
- Sumatoria total $\sum X_{..} = X_1 + X_2 + X_3 + \dots + X_n$ obtenida de la sumatoria de todas y cada una de las observaciones encontradas en la “data”
- Media general $\bar{X} = \frac{\sum X_{..}}{n}$ obtenida de la sumatoria de total dividido para las observaciones en cuestión.
- Sacamos un factor de corrección $F.C. = \frac{(\sum X_{..})^2}{(n)}$, siendo “n” el conteo de todos los datos de la distribución. Igual este valor nos servirá para el cálculo de suma de cuadrados de los componentes de las fuentes de variación.
- Posteriormente calculamos las sumas de cuadrados. El procedimiento es similar al anterior en la suma de cuadrados totales $SC_{TOT} = \sum X_{ij}^2 - FC$ en consecuencia se suma de todos y cada uno de los valores encontrados y son elevados al cuadrado, su respuesta se resta del factor de corrección.
- Para la suma de cuadrados los tratamientos, cambia ligeramente su procedimiento anterior, en función de la siguiente expresión:

$$SC_{TRAT} = \left[\frac{\sum \tau_1^2}{r_1} + \frac{\sum \tau_2^2}{r_2} + \frac{\sum \tau_3^2}{r_3} + \dots + \frac{\sum \tau_n^2}{r_n} \right] - FC$$

Es decir, se va sumando los cuadrados de cada tratamiento individualmente y se divide por el número de repeticiones que correspondo a cada uno; luego ese resultado se resta del factor de corrección.

- Igualmente, para la suma de cuadrados del error se usa la ecuación $SC_{ERROR} = SC_{TOTAL} - SC_{TRAT}$ por diferencia de las sumas de cuadrados, del total menos SC de tratamientos.

- Los grados de libertad (gl) son obtenidos de similar acción que en el caso anterior: grados de libertad totales es igual a las observaciones menos uno $GL_{TOT} = (n - 1)$; para los de tratamientos, número de tratamientos menos uno $GL_{TRAT} = (t - 1)$; mientras que para el error, se procede por diferencia de los grados de libertad totales menos los grados de libertad de los tratamientos $GL_{ERROR} = GL_{TOT} - GL_{TRAT}$.
- La varianza de los tratamientos o sus cuadrados medios serán obtenidos dividiendo la suma de cuadrados de los tratamientos para sus grados de libertad $CM_{TRAT} = \frac{SC_{TRAT}}{GL_{TRAT}}$. En el caso del cuadrado medio del error se divide su suma de cuadrados del error para sus grados de libertad $CM_{ERROR} = \frac{SC_{ERROR}}{GL_{ERROR}}$
- El estadístico Fisher calculado se encuentra dividiendo el cuadrado medio de los tratamientos entre el cuadrado medio del error $F_{CAL} = \frac{CM_{TRAT}}{CM_{ERROR}}$.
- Procedemos a encontrar el coeficiente de variación con la expresión matemática: $CV = \frac{\sqrt{CM_{ERROR}}}{\bar{X}} * 100$
- Una vez que se ha calculado el valor de F, es crucial interpretarlo en el contexto del análisis. Para ello, se compara el valor de F calculado con un valor crítico obtenido de la tabla de distribuciones F, que depende de los grados de libertad correspondientes y del nivel de significancia deseado. Si el valor de F calculado es mayor que el valor crítico, se puede concluir que existen diferencias significativas entre al menos dos de los tratamientos. Esta comparación es un paso clave en el proceso de toma de decisiones basado en los resultados del ANOVA.

Es importante también considerar el nivel de significancia que se ha establecido antes de realizar el análisis. Comúnmente, se utiliza un nivel de significancia del cinco por ciento, lo que indica que se acepta un riesgo del cinco por ciento de rechazar la hipótesis nula cuando en realidad es verdadera. Este nivel de significancia ayuda a determinar el umbral para la comparación del valor de F.

Además, en contextos de investigación, se pueden establecer niveles más estrictos o más laxos según la naturaleza del estudio.

Un ejemplo práctico puede ilustrar mejor este proceso. Supongamos que se están comparando tres tratamientos diferentes en un experimento agrícola. Después de realizar el ANOVA, se obtiene un valor de F calculado de 4.5. Al consultar la tabla de distribuciones F con los grados de libertad correspondientes, se encuentra un valor crítico de 3.2. Dado que 4.5 es mayor que 3.2, se puede concluir que hay diferencias significativas entre los tratamientos, sugiriendo que al menos uno de ellos tiene un efecto diferente en el rendimiento de los cultivos.

Además de determinar la significancia, el ANOVA también permite realizar comparaciones post-hoc para identificar cuáles tratamientos son significativamente diferentes entre sí. Estas pruebas adicionales son esenciales para comprender mejor los resultados y proporcionar recomendaciones prácticas basadas en los hallazgos. Algunas de las pruebas más comunes incluyen la prueba de Tukey y la prueba de Bonferroni, que ayudan a controlar el error tipo I al realizar múltiples comparaciones.

3.7. Aplicación del DCA con desigual repetición

Ejercicio 3.2.

En un experimento se evaluó la ganancia de peso (kg) de cerdos mestizos alimentados con dietas en base varios niveles de sorgo en reemplazo del maíz (0, 25, 50, 75 y 100%)

Problema:

El alto costo del maíz en temporada de escasez compromete el valor del balanceado comercial, por lo que se pretende solucionar reemplazándolo con el sorgo, un energético también pero más económico.

Objetivo:

Evaluar la ganancia de peso de los cerdos mestizos mediante dietas preparadas con diferentes niveles de sorgo en reemplazo del maíz, con el fin de abaratar costos del alimento balanceado.

Hipótesis:

- *Teórica*

El reemplazo de varios niveles de sorgo en la dieta de cerdos sobre el maíz, no presenta diferencia en la ganancia de peso con referencia al testigo.

- *Matemática*

$$H_0 = \mu_{s0} = \mu_{s25} = \mu_{s50} = \mu_{s75} = \mu_{s100}$$

Resultados experimentales

Tratamientos (Niv.sorgo)	Repeticiones				Suma tratam.	r	Media tratam.
	I	II	III	IV			
0%	40,63	40,63	47,28		128,54	3	42,85
25%	44,12	45,78	48,27	47,5	185,67	4	46,42
50%	46,8	48,46			95,26	2	47,63
75%	41,71	38,02	37,06	42,88	159,67	4	39,92
100%	41,62	37,47	49,93	39,96	168,98	4	42,25
$\Sigma X_{..}$					738,12	17	43,42
Suma total					Total	n	Media

$$\Sigma X_{..} = X_1 + X_2 + X_3 + \dots + X_n$$

$$\Sigma X_{..} = 40,63 + 40,63 + 47,28 + \dots + 49,93 + 39,96 = 738,12$$

Media de tratamientos

$$\bar{X}_1 = \frac{\Sigma \tau_1}{r_1} = \left[\left(\frac{128,54}{3} \right) \right] = 42,85$$

$$\bar{X}_2 = \frac{\Sigma \tau_2}{r_2} = \left[\left(\frac{185,67}{4} \right) \right] = 46,42$$

$$\bar{X}_3 = \frac{\Sigma \tau_3}{r_3} = \left[\left(\frac{95,26}{2} \right) \right] = 47,63$$

$$\bar{X}_4 = \frac{\Sigma \tau_4}{r_4} = \left[\left(\frac{159,67}{4} \right) \right] = 39,92$$

$$\bar{X}_5 = \frac{\sum \tau_5}{r_5} = \left[\left(\frac{168,98}{4} \right) \right] = 42,25$$

Media total o general

$$\bar{X} = \frac{\sum X..}{n} = \frac{738,12}{17} = 43,42$$

Factor de corrección

$$F.C. = \frac{(\sum X..)^2}{n}$$

$$F.C. = \frac{(738,12)^2}{17} = \frac{544821,13}{17} = 32048,30$$

Suma de cuadrados

- Suma de cuadrados totales

$$SC_{TOT} = \sum X_{ij}^2 - FC$$

$$SC_{TOT} = [(40,63)^2 + (40,63)^2 + (47,28)^2 + \dots + (49,93)^2 + (39,96)^2] - 32048,30$$

$$SC_{TOT} = 32327,64 - 32048,30 = 279,34$$

- Suma de cuadrados de los tratamientos

$$SC_{TRAT} = \left[\frac{\sum \tau_1^2}{r_1} + \frac{\sum \tau_2^2}{r_2} + \frac{\sum \tau_3^2}{r_3} + \dots + \frac{\sum \tau_n^2}{r_n} \right] - FC$$

$$SC_{TRAT} = \left[\frac{(128,54)^2}{3} + \frac{(185,67)^2}{4} + \frac{(95,26)^2}{2} + \frac{(159,67)^2}{4} + \frac{(168,98)^2}{4} \right] - 32048,30$$

$$SC_{TRAT} = [5507,51 + 8618,34 + 4537,23 + 6373,63 + 7138,56] - 32048,30$$

$$SC_{TRAT} = 32175,27 - 32048,30 = 126,97$$

- Suma de cuadrados del error

$$SC_{ERROR} = SC_{TOTAL} - SC_{TRAT}$$

$$SC_{ERROR} = 279,34 - 126,97 = 152,375$$

Grados de libertad

$$GL_{TOT} = (n - 1) = 17 - 1 = 16$$

$$GL_{TRAT} = (t - 1) = 5 - 1 = 4$$

$$GL_{ERROR} = GL_{TOT} - GL_{TRAT} = 16 - 4 = 12$$

Cuadrados medios

- Cuadrados medios de los tratamientos

$$CM_{TRAT} = \frac{SC_{TRAT}}{GL_{TRAT}}$$

$$CM_{TRAT} = \frac{126,97}{4} = 31,742$$

- Cuadrados medios del error

$$CM_{ERROR} = \frac{SC_{ERROR}}{GL_{ERROR}}$$

$$CM_{ERROR} = \frac{152,375}{12} = 12,698$$

- Estadístico Fisher calculado

$$F_{CAL} = \frac{CM_{TRAT}}{CM_{ERROR}}$$

$$F_{CAL} = \frac{31,742}{12,698} = 2,98$$

- Coeficiente de variación del ADEVA

$$CV = \frac{\sqrt{CM_{ERROR}}}{\bar{X}} * 100$$

$$CV = \frac{\sqrt{12,698}}{43,42} * 100 = \frac{3,563}{43,42} * 100 = 8,21\%$$

ADEVA

Fuente de variación	Suma de cuadrados	Grados de libertad	Cuadrado medio	Fcal	F.Tab 0,05	F.Tab 0,01	Signif.
Total	279,342	16					
Tratamientos	126,967	4	31,742	2,50	3,26	5,41	*
Error	152,375	12	12,698				
CV	8,21%						

Finalmente se compara el estadístico F. Cal con el F. Tab en la tabla correspondiente con 4 grados de libertad para el numerador (tratamientos) y 12 grados de libertad para el denominador (error experimental), utilizando los 2 niveles de significancia. Sería $\alpha 0,05 = 3,26$; y, $\alpha 0,01 = 5,41$.

Análisis del estadístico Fisher

El Fisher calculado (2,50) supera a los valores del Fisher tabular al nivel de significancia $\alpha 0,05$ (3,26) pero es menor que $\alpha 0,01$; la diferencia es solo significativa.

Interpretación

En base al resultado alcanzado en el análisis de varianza (ADEVA), la ganancia de peso de cerdos mestizos que se alimentaron con diferentes niveles de sorgo en reemplazo del maíz tuvo diferencias significativas ($P \leq 0,05$), esto indicaría que no hay evidencia estadística para aceptar la hipótesis nula.

Referencias bibliográficas del capítulo

- Montgomery, D. C. (2017). Design and Analysis of Experiments (9th ed.). Wiley.
- Kirk, R. E. (2013). Experimental Design: Procedures for the Behavioral Sciences (4th ed.). Sage Publications.
- Field, A. (2018). Discovering Statistics Using IBM SPSS Statistics (5th ed.). Sage Publications.
- Huitema, B. E., & McKean, J. W. (2000). Designing For Analysis of Variance. Wiley.

- Cohen, J. (1988). *Statistical Power Analysis for the Behavioral Sciences* (2nd ed.). Lawrence Erlbaum Associates.
- Hinkle, D. E., Wiersma, W., & Jurs, S. G. (2003). *Applied Statistics for the Behavioral Sciences* (5th ed.). Houghton Mifflin.
- Tabachnick, B. G., & Fidell, L. S. (2013). *Using Multivariate Statistics* (6th ed.). Pearson.
- Friedman, L. M., Furberg, C., & DeMets, D. L. (2010). *Fundamentals of Clinical Trials* (4th ed.). Springer.
- Kirk, R. E. (2013). *Experimental Design: Procedures for the Behavioral Sciences* (4th ed.). Sage Publications.
- Cohen, J. (1988). *Statistical Power Analysis for the Behavioral Sciences* (2nd ed.). Lawrence Erlbaum Associates.
- Gelman, A., & Hill, J. (2007). *Data Analysis Using Regression and Multilevel/Hierarchical Models*. Cambridge University Press.
- Kutner, M. H., Nachtsheim, C. J., & Neter, J. (2004). *Applied Linear Statistical Models* (5th ed.). McGraw-Hill.
- Field, A. (2013). *Discovering Statistics Using IBM SPSS Statistics* (5th ed.). Sage Publications.
- Huitema, B. E., & McKean, J. W. (2000). *Designing For Analysis of Variance*. Wiley.
- Ritchie, J., Lewis, J., & Elam, G. (2003). Designing and Selecting Samples. In J. Ritchie & J. Lewis (Eds.), *Qualitative Research Practice: A Guide for Social Science Students and Researchers*. Sage Publications.
- McCullagh, P., & Nelder, J. A. (1989). *Generalized Linear Models* (2nd ed.). Chapman and Hall.
- Weisberg, S. (2005). *Applied Linear Regression* (3rd ed.). Wiley.



CAPÍTULO IV

MODELOS LINEALES
ADITIVOS DE
ORDENACIÓN
MÚLTIPLE

En el ámbito de la estadística y el diseño experimental, los modelos lineales aditivos de ordenación múltiple se presentan como herramientas fundamentales para el análisis de datos. Este enfoque permite evaluar y comparar distintos tratamientos o condiciones en experimentos, facilitando la interpretación de resultados y la toma de decisiones informadas.

En este contexto, exploraremos dos de los diseños más utilizados: el Diseño de Bloques Completamente al Azar (DBCA) y el Diseño Cuadrado Latino (DCL). A lo largo de este epígrafe, se abordarán las características, ventajas y desventajas de cada uno de estos diseños, proporcionando una comprensión profunda de su aplicación en la investigación.

En la sección se describirán las características del DBCA, seguido por un análisis de sus ventajas y desventajas. Posteriormente, se presentará un modelo lineal aditivo específico para el DBCA, junto con las fórmulas de cálculo pertinentes en la sección. Para consolidar el aprendizaje, se incluirán ejercicios prácticos.

Luego, se realizará un análisis similar del DCL, comenzando con sus características, seguido de las ventajas y desventajas. Posteriormente, se presentará el modelo lineal aditivo correspondiente al DCL, junto con las fórmulas de cálculo. Finalmente, se ofrecerán ejercicios prácticos relacionados con el DCL.

A través de este recorrido, se busca proporcionar una base sólida en el uso de modelos lineales aditivos en diseños experimentales, permitiendo a los investigadores aplicar estos conceptos de manera efectiva en sus estudios.

4.1. Características del Diseño de Bloques Completamente al Azar (DBCA)

El Diseño de Bloques Completamente al Azar (DBCA) es un método estadístico ampliamente utilizado en la investigación experimental que permite controlar la variabilidad entre tratamientos. Este diseño se basa en la asignación aleatoria de tratamientos a unidades experimentales dentro de bloques homogéneos, lo que ayuda a reducir el error experimental y a mejorar la precisión de las estimaciones. Según Gómez y López (2021), el DBCA es especialmente útil en

situaciones donde las unidades experimentales difieren en características que no son de interés para el estudio, pero que influyen en la respuesta medida.

Una de las principales características del DBCA es su capacidad para agrupar unidades experimentales en bloques que son similares en ciertos aspectos. Este agrupamiento permite que las variaciones dentro de cada bloque sean menores que las variaciones entre bloques. De acuerdo con Martínez et al. (2022), la identificación de factores de bloqueo es crucial para el éxito del diseño, ya que un mal agrupamiento conlleva a conclusiones erróneas sobre los efectos de los tratamientos.

El DBCA también se distingue por su simplicidad y flexibilidad. Según Rojas y Pérez (2023), este diseño es fácil de implementar y se adapta a una variedad de contextos experimentales. La asignación aleatoria de tratamientos dentro de cada bloque asegura que los efectos de los tratamientos se puedan estimar de manera precisa, lo que es fundamental para la validez de los resultados. Esta aleatorización es un principio clave en la investigación científica, ya que minimiza el sesgo y permite generalizar los hallazgos a una población más amplia.

Otra característica importante del DBCA es su capacidad para manejar múltiples tratamientos. En estudios donde se evalúan varios tratamientos simultáneamente, el DBCA permite una comparación directa entre ellos dentro de cada bloque. Según González et al. (2024), esto no solo facilita la evaluación de la eficacia de los tratamientos, sino que también optimiza el uso de recursos al permitir que se realicen múltiples comparaciones en un solo experimento. Además, el diseño es adecuado para experimentos con un número limitado de replicaciones, lo que lo hace accesible para investigadores con recursos limitados.

Sin embargo, a pesar de sus ventajas, el DBCA también presenta limitaciones. Una de las principales críticas es que requiere un conocimiento previo de las características de las unidades experimentales para establecer bloques efectivos. Si los bloques no se definen correctamente, el diseño pierde su efectividad. Según López y Morales (2021), es fundamental realizar un análisis

previo para identificar las variables que influyen en la respuesta y, a partir de ahí, definir los bloques adecuados.

Además, el DBCA es menos eficiente en situaciones donde hay interacciones significativas entre los tratamientos y los factores de bloqueo. En tales casos, el Diseño de Bloques Incompletos o el Diseño Factorial podrían ser más apropiados. Por ejemplo, Aredo (2021) argumentó que, si los efectos de los tratamientos dependen de las condiciones específicas de cada bloque, el DBCA podría no ser la mejor opción, ya que no permite evaluar estas interacciones de manera efectiva.

La implementación del DBCA también implica consideraciones prácticas. La asignación aleatoria de tratamientos dentro de cada bloque debe ser rigurosa para asegurar la validez del diseño. Según Vargas et al. (2022), es recomendable utilizar software estadístico para realizar esta asignación, lo que ayuda a evitar errores humanos y garantizar que el proceso sea completamente aleatorio. Además, la planificación cuidadosa del experimento, incluida la determinación del tamaño de muestra adecuado, es esencial para maximizar la potencia estadística y la capacidad de detectar efectos significativos.

En términos de análisis de datos, el DBCA permite el uso de modelos lineales aditivos para evaluar los efectos de los tratamientos. Estos modelos incluyen términos para los efectos de los bloques y los tratamientos, lo que permite una estimación precisa de las diferencias entre tratamientos. Según Ramírez y Soto (2023), el análisis de varianza (ANOVA) es una técnica comúnmente utilizada para evaluar los datos obtenidos de experimentos DBCA, proporcionando una forma efectiva de determinar si hay diferencias significativas entre los tratamientos.

Finalmente, es importante señalar que el DBCA se utiliza en una amplia gama de disciplinas, desde la agricultura hasta la medicina. Su versatilidad lo convierte en una herramienta valiosa para investigadores que buscan evaluar tratamientos en condiciones controladas. Según Martínez et al. (2021), el uso del DBCA en estudios agrícolas ha permitido a los investigadores evaluar el impacto de diferentes fertilizantes y prácticas de cultivo en el rendimiento de los cultivos, contribuyendo a prácticas agrícolas más sostenibles y eficientes.

En conclusión, el Diseño de Bloques Completamente al Azar es un enfoque poderoso en la investigación experimental que ofrece numerosas ventajas en términos de control de la variabilidad y la comparación de tratamientos. Sin embargo, su efectividad depende en gran medida de la correcta identificación de los bloques y la rigurosidad en la implementación del diseño. A medida que la investigación continúa evolucionando, el DBCA seguirá siendo una herramienta fundamental para los científicos que buscan obtener resultados precisos y confiables en sus estudios.

4.2. Ventajas y desventajas del DBCA

El DBCA es un método ampliamente utilizado en experimentación animal y en estudios agrícolas, ya que permite controlar la variabilidad entre tratamientos al agrupar experimentalmente las unidades en bloques homogéneos. Una de las principales ventajas del DBCA es su simplicidad, lo que facilita su implementación y análisis. Al asignar tratamientos al azar dentro de cada bloque, se minimiza el sesgo y se asegura que cada tratamiento tenga la misma probabilidad de ser asignado a cualquier unidad experimental, lo que mejora la validez interna del estudio (Montgomery, 2017). Además, el DBCA es efectivo cuando se espera que la variabilidad entre bloques sea mayor que la variabilidad dentro de los bloques, lo que permite detectar diferencias significativas entre tratamientos con mayor eficacia.

Sin embargo, el DBCA también presenta desventajas. Una de las principales limitaciones es que requiere que los bloques sean homogéneos y que se identifique adecuadamente la fuente de variabilidad que se desea controlar. Si los bloques no son bien definidos o si hay factores no controlados que afectan la respuesta, los resultados resultan engañosos (Cochran & Cox, 1957). Además, el DBCA es menos eficiente en comparación con otros diseños experimentales, como el Diseño de Bloques Incompletos o el Diseño Factorial, especialmente cuando se tienen múltiples factores a considerar. Esto resulta en una menor precisión y un mayor error experimental, lo que dificulta la interpretación de los resultados (Patterson & Thompson, 1971).

4.3. Modelo lineal aditivo para el DBCA

El Modelo Lineal Aditivo es fundamental en el análisis de datos provenientes de un Diseño de Bloques Completamente al Azar (DBCA). Este modelo permite descomponer la variabilidad observada en los datos en componentes atribuibles a los tratamientos y a los bloques, facilitando así la evaluación de la efectividad de los tratamientos aplicados. En un contexto de DBCA, el modelo se expresa como:

$$Y_{ij} = \mu + \tau_i + \beta_j + \epsilon_{ij}$$

Donde:

- Y_{ij} = es la variable respuesta observada para los i – ésimos tratamientos en los j - esimos bloques o repeticiones,
- μ = es la media general,
- τ_i = representa el efecto los i -ésimos tratamientos,
- β_j = es el efecto de los j – ésimos bloques, y
- ϵ_{ij} = es el error aleatorio asociado a la observación.

Este enfoque permite identificar cómo cada tratamiento influye en la respuesta, controlando simultáneamente la variabilidad debida a los bloques (Montgomery, 2017).

Una de las ventajas del Modelo Lineal Aditivo en el contexto del DBCA es su capacidad para proporcionar estimaciones precisas de los efectos de los tratamientos, incluso en presencia de variabilidad no controlada entre bloques. Esto se traduce en una mayor potencia estadística para detectar diferencias significativas entre tratamientos.

Sin embargo, es crucial que los supuestos del modelo, como la normalidad de los errores y la homogeneidad de varianzas, se cumplan para asegurar la validez de los resultados (Cochran & Cox, 1957). Si estos supuestos no se satisfacen, las inferencias realizadas a partir del modelo pueden ser erróneas, lo que subraya la importancia de un análisis cuidadoso de los datos y la verificación de los supuestos del modelo.

El Modelo Lineal Aditivo aplicado en el Diseño de Bloques Completamente al Azar (DBCA) tiene diversas aplicaciones en el campo de agropecuaria. Se utiliza para evaluar el efecto de diferentes tratamientos de fertilización, variedades de cultivos o prácticas de manejo sobre el rendimiento de las cosechas, controlando la variabilidad debida a factores ambientales como el tipo de suelo o condiciones climáticas.

4.4. Fórmulas de cálculo para el DBCA

El Diseño de Bloques Completamente al Azar (DBCA) se basa en varias fórmulas clave que permiten analizar los datos y extraer conclusiones significativas sobre los efectos de los tratamientos. A continuación, se presentan las fórmulas más relevantes utilizadas en el análisis de varianza de un DBCA.

- Organización de la información en una tabla de resultados para cada variable respuesta
- Se obtiene las medias de los tratamientos, $\bar{X}_{\tau_i} = \frac{\sum \tau_i}{r_i}$ que hace referencia a la suma de las observaciones de cada tratamiento, dividido para el total de sus repeticiones.
- Procedemos con el cálculo de la suma de las repeticiones o bloques utilizando $\sum \beta_j = \frac{b_j}{r_i}$, esta expresión indica que se sumarían las observaciones en cada uno de los bloques o repeticiones b_j , y se divide para el número de datos correspondiente r_i .
- Luego se extrae la sumatoria total $\sum X_{..} = X_1 + X_2 + X_3 + \dots + X_n$ que es la suma de cada dato encontrado y tratado, según su variable respuesta.
- Para la media general se utiliza $\bar{X} = \frac{\sum X_{..}}{n}$ que representa a la sumatoria de total dividido para el número de observaciones.
- Procedemos con el factor de corrección F. C. $= \frac{(\sum X_{..})^2}{(t \cdot r)}$ que servirá para extraer la suma de cuadrados de los componentes en las fuentes de

variación (total, tratamientos, bloques); y se obtiene extrayendo el cuadrado de la sumatoria total $\sum X_{..}$ dividido para el general de los datos o también la multiplicación entre el número de tratamientos y las repeticiones o bloques ($t \cdot r$).

- Inmediatamente se calculan las sumas de cuadrados; primeramente de los totales $SC_{TOT} = \sum X_{ij}^2 - FC$ la ecuación establece la suma de todos y cada uno de los datos encontrados de la variable respuesta y son elevados al cuadrado, resultado que se restará del factor de corrección
- En la suma de cuadrados de los tratamientos, se usa $SC_{TRAT} = \frac{\sum \tau_{ij}^2}{r} - FC$ que se refiere a la suma de cada grupo de observaciones individuales y por tratamiento al cuadrado y dividido para el número de repeticiones; cuyo resultado se restaría del factor de corrección
- Luego se extrae la suma de cuadrados de las repeticiones, con la formula $SC_{REP} = \frac{\sum \beta_j^2}{t} - FC$ que representaría a la suma de cada observación que corresponde a un determinado bloque elevado al cuadrado dividiéndolo para el número de tratamientos; luego su resultado se resta igual del factor de corrección.
- Como se ha indicado anteriormente la suma de cuadrados del error se obtiene por sustracción $SC_{ERROR} = SC_{TOTAL} - SC_{TRAT} - SC_{REP}$ desde la suma total menos los tratamientos y repeticiones.
- Seguidamente se encuentran los grados de libertad (gl); en los totales se realiza la resta de todas las observaciones menos uno $GL_{TOT} = (n - 1)$; los de tratamientos en cambio será obtenido del número de tratamientos menos uno $GL_{TRAT} = (t - 1)$ así como en los bloques el total de repeticiones menos uno ($r - 1$); mientras que, para el error, encontramos por diferencia de los grados de libertad totales menos los grados de libertad de los tratamientos, y grados de libertad de las repeticiones $GL_{ERROR} = GL_{TOT} - GL_{TRAT} - GL_{REP}$.

- La varianza de los tratamientos o también denominados cuadrados medios serán alcanzados dividiendo la suma de cuadrados de los tratamientos para sus grados de libertad $CM_{TRAT} = \frac{SC_{TRAT}}{GL_{TRAT}}$. Luego los cuadrados medios de las repeticiones serán derivados de la suma de cuadrados de los bloques para sus grados de libertad $CM_{REP} = \frac{SC_{TRAT}}{GL_{TRAT}}$. En cambio, el cuadrado medio del error se obtendrá dividiendo la suma de cuadrados del error para los grados de libertad $CM_{ERROR} = \frac{SC_{ERROR}}{GL_{ERROR}}$
- El estadístico Fisher calculado que corresponde a los tratamientos será encontrado dividiendo el cuadrado medio de los tratamientos y el cuadrado medio del error $F_{CAL.TRAT} = \frac{CM_{TRAT}}{CM_{ERROR}}$. En cambio, el Fisher calculado para las repeticiones o bloques se extrae del cuadrado medio de los bloques entre el cuadrado medio del error $F_{CAL.REP} = \frac{CM_{REP}}{CM_{ERROR}}$
- Luego se procede con el coeficiente de variación en base a la expresión matemática: $CV = \frac{\sqrt{CM_{ERROR}}}{\bar{X}} * 100$
- Finalmente se analiza el valor del Fisher calculado en referencia a la respectiva tabla para declarar verdadera o no a la hipótesis nula.

Estas fórmulas son fundamentales para realizar un análisis estadístico adecuado en un diseño de bloques completamente al azar. Permiten descomponer la variabilidad total en componentes atribuibles a los tratamientos y bloques, facilitando la interpretación de los resultados y la toma de decisiones informadas en la investigación experimental.

4.5. Aplicación del DBCA

Ejercicio 4.1.

En una hacienda del oriente ecuatoriano se evaluó el efecto de diferentes niveles de pollinaza como raciones suplementarias en la etapa de levante y engorde en toretes Brown Swiss sobre su ganancia de peso (kg); los tratamientos fueron distribuidos con un modelo DBCA.

Problema:

El engorde de bovinos en el oriente tiene debilidad en la alimentación por sus praderas que no tienen complementos vegetales proteicos en sus pastizales, por lo tanto, debe suministrarse raciones suplementarias que resultan muy costosas, por lo que se recomienda prepararlas con residuos como la pollinaza que gracias a la baja digestibilidad de las aves, sus excretas tratadas pueden entregar un nivel interesante.

Objetivo:

Evaluar la ganancia de peso diaria de toretes Brown Swiss, levantados con varios niveles de pollinaza en la dieta complementaria, con el fin de corregir su déficit alimenticio.

Hipótesis:

- *Teórica*

Las 4 dietas complementaria a base de pollinaza, utilizadas para el levante de toretes Brown Swiss respondieron igual que el testigo (0%) en la ganancia de peso.

- *Matemática*

$$H_0 = \mu_{p0} = \mu_{p20} = \mu_{p40} = \mu_{p60} = \mu_{p80}$$

Resultados experimentales (Ajustados con Lowess)

Tratamientos Pollinaza (t)	Repeticiones (r)				Suma Tratam $\sum \tau_i$	Media Tratam $\bar{X}_{\tau_i}$
	I	II	III	IV		
0%	0,428	0,430	0,433	0,439	1,730	0,433
20%	0,450	0,471	0,532	0,602	2,055	0,514
40%	0,687	0,799	0,915	1,003	3,404	0,851
60%	1,058	1,077	1,058	0,973	4,166	1,042
80%	0,889	0,800	0,705	0,607	3,001	0,750
Suma Repeticiones $\bar{X}_{\beta_j}$	3,512	3,577	3,643	3,624	14,356	0,718
					Total $\sum X_{..}$	Media $\bar{X}$

Factor de corrección

$$F.C. = \frac{(\sum X..)^2}{(t * r)}$$

$$F.C. = \frac{(14,356)^2}{(5 * 4)} = \frac{206,095}{20} = 10,305$$

Suma de cuadrados

- Suma de cuadrados totales

$$SC_{TOT} = \sum X_{ij}^2 - FC$$

$$SC_{TOT} = [(0,428)^2 + (0,430)^2 + \dots + (0,705)^2 + (0,607)^2] - 10,305$$

$$SC_{TOT} = 11,413 - 10,305 = 1,108$$

- Suma de cuadrados de los tratamientos

$$SC_{TRAT} = \frac{\sum \tau_i^2}{r} - FC$$

$$SC_{TRAT} = \frac{[(1,730)^2 + (2,055)^2 + (3,404)^2 + (4,166)^2 + (3,001)^2]}{4} - 10,305$$

$$SC_{TRAT} = \frac{45,165}{4} - 10,305 = 11,291 - 10,305 = 0,986$$

- Suma de cuadrados de los bloques o repeticiones

$$SC_{REP} = \frac{\sum \beta_j^2}{t} - FC$$

$$SC_{REP} = \frac{[(3,512)^2 + (3,577)^2 + (3,643)^2 + (3,624)^2]}{5} - 10,305$$

$$SC_{REP} = \frac{51,534}{5} - 10,305 = 10,307 - 10,305 = 0,002$$

- Suma de cuadrados del error

$$SC_{ERROR} = SC_{TOTAL} - SC_{TRAT} - SC_{REP}$$

$$SC_{\text{ERROR}} = 1,108 - 0,986 - 0,002 = 0,120$$

Grados de libertad

$$GL_{\text{TOT}} = (n - 1) = 20 - 1 = 19$$

$$GL_{\text{TRAT}} = (t - 1) = 5 - 1 = 4$$

$$GL_{\text{REP}} = (r - 1) = 4 - 1 = 3$$

$$GL_{\text{ERROR}} = GL_{\text{TOT}} - GL_{\text{TRAT}} - GL_{\text{REP}} = 19 - 4 - 3 = 12$$

Cuadrados medios

- Cuadrados medios de los tratamientos

$$CM_{\text{TRAT}} = \frac{SC_{\text{TRAT}}}{GL_{\text{TRAT}}}$$

$$CM_{\text{TRAT}} = \frac{0,986}{4} = 0,2466$$

- Cuadrados medios de los bloques o repeticiones

$$CM_{\text{REP}} = \frac{SC_{\text{REP}}}{GL_{\text{REP}}}$$

$$CM_{\text{REP}} = \frac{0,002}{3} = 0,0007$$

- Cuadrados medios del error

$$CM_{\text{ERROR}} = \frac{SC_{\text{ERROR}}}{GL_{\text{ERROR}}}$$

$$CM_{\text{ERROR}} = \frac{0,120}{12} = 0,010$$

- Estadístico Fisher calculado

$$F_{\text{CAL.TRAT}} = \frac{CM_{\text{TRAT}}}{CM_{\text{ERROR}}}$$

$$F_{\text{CAL.TRAT}} = \frac{0,2466}{0,010} = 24,737$$

$$F_{\text{CAL.REP}} = \frac{CM_{\text{REP}}}{CM_{\text{ERROR}}}$$

$$F_{\text{CAL.REP}} = \frac{0,0007}{0,010} = 0,068$$

- Coeficiente de variación del ADEVA

$$CV = \frac{\sqrt{CM_{\text{ERROR}}}}{\bar{X}} * 100$$

$$CV = \frac{\sqrt{0,010}}{0,718} * 100 = \frac{0,099}{0,718} * 100 = 13,91\%$$

ADEVA

Fuente de variación	Suma de cuadrados	Grados de libertad	Cuadrado medio	Fcal	F.Tab 0,05	F.Tab 0,01	Signif .
Total	1,108	19					
Tratamientos	0,986	4	0,2466	24,737	3,26	5,41	**
Repeticiones	0,002	3	0,0007	0,068	3,49	5,95	ns
Error	0,120	12	0,0100				
CV	13,91%						

Análisis del estadístico Fisher de los tratamientos

El Fisher calculado (24,737) prevalece con superioridad frente a los del Fisher tabular al nivel de significancia $\alpha 0,05$ (3,26) y $\alpha 0,01$ (5,41); entonces existe alta diferencia entre los tratamientos.

Interpretación

Los diferentes niveles de pollinaza que fueron suplementados en la etapa de levante y engorde de toretes Brown Swiss presentaron diferencias altamente significativas ($P \leq 0,01$) en la ganancia de peso diaria (kg). Por lo tanto, no existe evidencia para aceptar la hipótesis nula.

4.6. Características del Diseño Cuadrado Latino (DCL)

El Diseño Cuadrado Latino (DCL) es una técnica experimental ampliamente utilizada en estadística para controlar la variabilidad en experimentos donde se desea evaluar el efecto de dos factores en una respuesta. Una de las características más destacadas del DCL es su capacidad para eliminar el efecto de dos fuentes de variabilidad que pueden presentarse en sentido contrario (pendientes cruzadas), generalmente representadas por filas y columnas en un diseño cuadrado. Esto permite que los experimentos sean más precisos y que las estimaciones de los efectos de los tratamientos sean menos sesgadas. El DCL es especialmente útil en situaciones donde la aleatorización completa no es práctica, ya que permite una asignación más controlada de tratamientos (Hoshmand et al., 2018).

Para diseñar un experimento utilizando un DCL, se deben seguir varios pasos clave. Primero, es necesario definir el número de tratamientos que se desea evaluar. Este número debe ser igual al número de filas y columnas en el diseño. Por ejemplo, si se desean evaluar cuatro tratamientos, se debe crear una matriz de 4x4.

Una vez definidos los tratamientos, se asignan aleatoriamente a las celdas de la matriz, asegurando que cada tratamiento aparezca una vez en cada fila y en cada columna. Esta aleatorización es crucial para eliminar sesgos y asegurar que los resultados sean representativos. Posteriormente, se realizan las mediciones de la respuesta en cada celda, y se deben registrar cuidadosamente los datos para su posterior análisis.

Finalmente, se utiliza el análisis de varianza (ADEVA o ANOVA) para evaluar la significancia de los tratamientos, considerando la variabilidad debida a las filas y columnas (Montgomery, 2017). Este enfoque permite a los investigadores obtener conclusiones sobre los efectos de los tratamientos mientras controlan la variabilidad de las otras fuentes.

En el Cuadrado Latino, como se indicó el número de repeticiones es igual al número de tratamientos, que se han de comparar en el experimento. Si este número es el “n”, el número total de parcelas que se distribuirá en el campo será

“ n^2 ”. El cuadrado latino es una extensión del Bloque Completo Randomizado o al Azar (DBCA), ya que se impone la misma restricción de distribución y también la misma restricción para las columnas. Cada tratamiento en el Cuadrado Latino debe estar presente una vez y al azar en cada bloque (hilera) y columna.

Es recomendable cuando las unidades experimentales, sean estas parcelas, plantas, animales, etc., puedan agruparse de acuerdo a los niveles de dos fuentes de variabilidad; para esto en cada bloque y en cada columna, se debe disponer de igual número de unidades experimentales. Si el número de tratamientos es seis, por este diseño, el número de bloques y de columnas deber obligatoriamente también ser seis, y el total de unidades experimentales serán treinta y seis ($r=c=t$; $n = r^2$).

El número de tratamientos más aconsejado en este diseño debe oscilar entre 3 y 10. Para 3 tratamientos se pueden usar varios cuadrados latinos, a fin de aumentar la precisión del experimento.

Si los factores o fuentes de variabilidad “externa”, cuyo efecto se trata de minimizar a través del diseño, varía en 2 direcciones, por ejemplo: una gradiente de fertilidad del suelo varía en dirección norte – sur, y otra gradiente en dirección este – oeste; los bloques deben coincidir con la perpendicularidad de la primera pendiente, y las columnas con la de la segunda gradiente. Después los tratamientos se aplicarán al azar en las parcelas (igual al número de tratamientos) de los bloques y las columnas.

A los bloques se pueden asimilar operadores en número de máquinas y a las columnas, período de trabajo, si al comparar 5 marcas de máquinas, se sabe que diferentes operadores y períodos de trabajo afectan al rendimiento de estas.

Al tratar se establecer si existen diferencias entre cantidades de leche producidas por 4 pezones de la ubre de las vacas Holstein, cada pezón de acuerdo a su ubicación, puede considerarse como un tratamiento, asignándole una letra. Los 4 ordeños que se deben asimilar a los bloques, y el orden en que se ordeñe se asimilará como columnas; este sería un DCL 4x4.

Este diseño es conveniente para experimentar con animales, cuando existen restricciones como la edad y peso, raza y diferentes establos, producción de leche y número de establos.

El DCL se usa con éxito en experimentos de campo, laboratorio, invernadero, establos, sociología, educación, etc. Se utiliza siempre que exista homogeneidad dentro del bloques y columnas; pero alta heterogeneidad entre bloques y columnas.

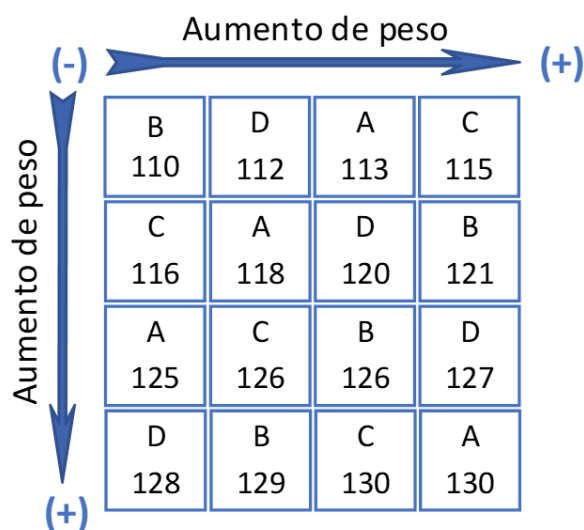
El DCL puede aplicarse aun cuando exista en las unidades experimentales, una sola fuente de variabilidad, pero para esto es necesario que dicha fuente, pueda distribuirse en forma continuada y creciente en una misma dirección, es decir la primera columna estaría formada por las primeras unidades de cada bloque, mientras que la segunda por las segundas unidades de cada bloque, etcétera. Por ejemplo, en un experimento con 4 racionamientos alimenticios en cerdos, si es que contamos con 16 cerdos de la misma edad, pero distintos pesos, los cerdos pueden ordenarse por orden creciente de peso, y a continuación asignarse los racionamientos a los cerdos en la forma siguiente (las letras indican racionamientos):

Tabla 4.1. Distribución de 4 tratamientos (raciones alimenticias) en cerdos bajo un Diseño de Cuadrado Latino.

Nº cerdo	Racionamiento	Peso cerdos, kg
1	B	110
2	D	112
3	A	113
4	C	115
5	C	116
6	A	118
7	D	120
8	B	121
9	A	125
10	C	126
11	B	126
12	D	127
13	D	128
14	B	129
15	C	130
16	A	130

Elaborado por: Moscoso, 2025

Figura 1.5. Esquema del sorteo de tratamientos en un Diseño Cuadrado Latino, en un experimento con 4 raciones alimenticias para cerdos



Elaborado por: Moscoso, 2025

Si tomamos en cuenta las 16 unidades experimentales el coeficiente de variación del peso inicial de los cerdos es manejable (5,65%), y podría utilizarse un DCA, sin embargo por términos didácticos se ha distribuido como DCL, en la ilustración nos damos cuenta que existe una variante de peso en doble sentido, de izquierda a derecha y desde arriba hacia abajo, configurando un cuadrado conformado por 4 fila y 4 columnas, dentro de las cuales se asignaron los tratamientos aleatoriamente sin que se repita uno mismo dentro de las filas y columnas indicadas.

Tabla 4.2. ADEVA del sorteo para 4 raciones de cerdos mediante DCL

Fuente de variación	Suma de cuadrados	Grados de libertad	Cuadrado medio	Fcal	F.Tab 0,05	F.Tab 0,01	Signif.
Total	707,75	15					
Tratamientos	0,25	3	0,08	0,09	4,76	9,78	ns
Columnas	26,75	3	8,92	9,73	4,76	9,78	*
Hileras	675,25	3	225,08	245,55	4,76	9,78	**
Error	5,5	6	0,92				
CV	0,79%						

Si nos fijamos en la table 4.2.; se cumplió debidamente el sorteo de las 4 raciones alimenticias que se experimentarían en los cerdos con un modelo lineal aditivo DCL 4x4; si nos fijamos los 2 supuestos del sorteo se cumplen a cabalidad: primero el coeficiente de variación de las unidades experimentales sueltas fue de 5,65%, mientras que en ADEVA descendió a 0,79%, y luego el estadístico Fisher calculado para los tratamientos (0,09) no superó al valor tabular del Fisher al 5% (4,76) por lo tanto la hipótesis nula es verdadera, ratificando que al inicio del experimento el peso no puede verse afectado aún por acción de los tratamientos.

Otro ejemplo clásico de un experimento de diseño cuadrado latino se podría observar en estudios agrícolas donde se evalúan diferentes variedades de cultivos. Supongamos que un investigador desea comparar cuatro variedades de maíz (A, B, C y D) en cuatro parcelas de tierra, donde las condiciones del suelo y la exposición al sol varía entre las parcelas.

El diseño del experimento sería una matriz de 4x4, donde cada fila representa una parcela y cada columna una variedad de maíz. Los tratamientos se asignan aleatoriamente a las celdas de la matriz, asegurando que cada variedad aparezca una vez en cada parcela. Luego, el investigador mide el rendimiento de cada variedad en cada parcela.

Los datos recopilados se analizan utilizando ANOVA, permitiendo al investigador determinar si hay diferencias significativas en el rendimiento entre las variedades de maíz, controlando así la variabilidad debida a las diferencias en las parcelas. Este enfoque ha sido documentado en la literatura como una metodología efectiva para maximizar la información obtenida de los experimentos agrícolas (Cochran & Cox, 1957).

4.7. Ventajas y desventajas del DCL

El Diseño Cuadrado Latino (DCL) es una técnica experimental que ofrece una serie de ventajas y desventajas que deben considerarse al planificar un experimento. Este epígrafe explora ambos aspectos para proporcionar una visión equilibrada sobre su aplicación en investigación.

4.7.1. Ventajas del Diseño Cuadrado Latino

a. Control de Variabilidad:

El DCL permite controlar dos fuentes de variabilidad (filas y columnas) simultáneamente. Esto resulta en estimaciones más precisas de los efectos de los tratamientos, ya que se minimizan los sesgos introducidos por factores no controlados (Montgomery, 2017).

b. Eficiencia en el Uso de Recursos:

Al permitir la evaluación de múltiples tratamientos en un espacio reducido, el DCL maximiza la información obtenida de cada experimento. Esto es especialmente beneficioso en situaciones donde los recursos, como tiempo y materiales, son limitados (Cochran & Cox, 1957).

c. Flexibilidad:

El DCL se aplica en diversos campos, desde la agricultura hasta la medicina, lo que lo convierte en una herramienta versátil para investigadores en diferentes disciplinas. Su estructura permite adaptaciones según las necesidades específicas del estudio (Hoshmand et al., 2018).

d. Facilidad en el Análisis:

La estructura del DCL facilita el análisis estadístico mediante el uso de ANOVA. Esto permite a los investigadores evaluar la significancia de los tratamientos de manera más directa y comprensible (Montgomery, 2017).

e. Reducción del Error Experimental:

Al controlar las variaciones en filas y columnas, el DCL ayuda a reducir el error experimental, lo que conlleva a conclusiones más robustas y confiables sobre los efectos de los tratamientos.

4.7.2. Desventajas del Diseño Cuadrado Latino

a. Limitaciones en el Número de Tratamientos:

Una de las principales desventajas del DCL es que el número de tratamientos debe ser igual al número de filas y columnas. Esto restringe su aplicabilidad en experimentos donde se desea evaluar un mayor número de tratamientos (Cochran & Cox, 1957).

b. Suposiciones de Homogeneidad:

El DCL asume que la variabilidad dentro de cada tratamiento es homogénea. Si esta suposición no se cumple, los resultados son engañosos y conducen a conclusiones incorrectas (Montgomery, 2017).

c. Complejidad en el Diseño:

Aunque el DCL es más eficiente que otros diseños, su planificación es más compleja. Los investigadores deben asegurarse de que los tratamientos se distribuyan adecuadamente en la matriz, lo que requiere un esfuerzo adicional en la etapa de diseño (Hoshmand et al., 2018).

d. Sensibilidad a Errores de Medición:

Dado que el DCL busca controlar la variabilidad, cualquier error en la medición de las respuestas tiene un impacto significativo en los resultados. Esto requiere un alto nivel de precisión en la recolección de datos (Montgomery, 2017).

e. Dificultades en la Interpretación:

En algunos casos, la interpretación de los resultados del DCL es complicada, especialmente si hay interacciones entre tratamientos que no se han considerado adecuadamente. Esto lleva a malentendidos sobre la efectividad de los tratamientos evaluados.

El Diseño Cuadrado Latino es una herramienta poderosa en el diseño experimental que ofrece numerosas ventajas, como el control de la variabilidad y la eficiencia en el uso de recursos. Sin embargo, también presenta desventajas significativas que limitan su aplicabilidad y la validez de los resultados. Los investigadores deben considerar cuidadosamente estas ventajas y desventajas al planificar sus experimentos, asegurándose de que el DCL sea la opción más adecuada para sus objetivos de investigación.

4.7.3. Modelo lineal aditivo para el DCL

El modelo lineal aditivo es una herramienta fundamental en el análisis de datos obtenidos a partir de un Diseño Cuadrado Latino (DCL). Este modelo permite descomponer la variabilidad observada en los experimentos en componentes atribuibles a los tratamientos, filas y columnas. En este capítulo, se explorará la formulación del modelo lineal aditivo, su interpretación y su aplicación en el contexto del DCL.

4.7.4. Fundamentos del Modelo Lineal Aditivo

Un modelo lineal aditivo se basa en la premisa de que la respuesta observada en un experimento es expresada como la suma de efectos individuales. En el contexto del DCL, el modelo se formula de la siguiente manera:

$$Y_{ijk} = \mu + \tau_i + \beta_j + \delta_k + \epsilon_{ij}$$

Donde:

- Y_{ijk} = es la variable respuesta observada para los i – ésimos tratamientos en las j - ésimas columnas y k - ésimas filas o hileras.
- μ = es la media general de todas las observaciones.
- τ_i = representa al efecto de los i – ésimos tratamientos
- β_j = corresponde al efecto de las j - ésimas columnas
- δ_k = es el efecto de las k – ésimas hileras o filas
- ϵ_{ij} = es el término de error aleatorio, que se asume que sigue una distribución normal con media cero y varianza constante.

4.7.5. Interpretación de los Componentes del Modelo

- Media General μ : Representa el promedio de todas las observaciones en el experimento. Es un parámetro central que ayuda a establecer una línea base para la comparación de tratamientos.

- Efecto del Tratamiento τ_i : Captura el impacto específico de cada tratamiento en la variable respuesta. Este componente es crucial para determinar la efectividad relativa de los tratamientos evaluados.
- Efecto de la Columna β_j : Refleja la influencia de las columnas en la respuesta, que se debe a factores como variaciones de la individualidad animal, en el suelo o condiciones ambientales. Este componente ayuda a controlar la variabilidad que podría sesgar los resultados.
- Efecto de la Hilera δ_k : Tiene que ver con la atribución de las hileras o filas que podrían influir en la variable respuesta, que igualmente son variaciones de la condición animal, el suelo, ambiente; siempre y cuando este sesgo sea contrario a las columnas; procedimiento que contribuye a controlar la heterogeneidad del material experimental (unidades experimentales); para que las respuestas solo sean causadas directamente por los tratamientos estudiados.
- Término de Error ϵ_{ijk} : Representa la variabilidad no explicada por los tratamientos, columnas y filas. Este término es esencial para evaluar la precisión del modelo y la significancia de los efectos.

4.8. Fórmulas de cálculo para el DCL

El análisis de varianza (ADEVA o ANOVA) es la técnica estadística utilizada para evaluar la significancia de los efectos de los tratamientos columnas y filas en el modelo. El ANOVA descompone la variabilidad total en componentes atribuibles a los tratamientos, columnas, filas y errores, permitiendo realizar pruebas de hipótesis sobre los efectos de los tratamientos.

El protocolo para el cálculo del ADEVA para un DCL es:

- Construir una tabla de resultados que albergue cada variable respuesta, preferentemente en Excel.

- Se procede a calcular la media de los tratamientos mediante $\bar{X}_{\tau_i} = \frac{\sum \tau_i}{r_i}$ en donde se suman la data de cada observación y se divide para el total de las mismas.
- De inmediato encontramos la suma de los tratamientos $\sum \tau_i$ que corresponde a la suma de observaciones en cada tratamiento; la suma de las columnas usando $\sum \beta_j$ sumando las observaciones en cada una de las columnas; y la suma de las hileras o filas $\sum \delta_k$ en cada fila reuniendo los valores de sus observaciones.
- Encontraremos luego la sumatoria total $\sum X_{..} = X_1 + X_2 + X_3 + \dots + X_n$ que es la adición de cada dato tomado en el trabajo de campo y tratado, de acuerdo a su variable respuesta.
- En la media general se usa $\bar{X} = \frac{\sum X_{..}}{n}$ que indica a la sumatoria de total sobre para el número de observaciones totales.
- Encontraremos el factor de corrección mediante $F.C. = \frac{(\sum X_{..})^2}{(t*r)}$ que será empleada para extraer la suma de cuadrados de los componentes en las fuentes de variación (total, tratamientos, columnas e hileras); y se obtiene con cuadrado de la sumatoria total $\sum X_{..}$ dividido para el total de la data o en su defecto multiplicando el número de tratamientos y las repeticiones o bloques ($t*r$).
- Procedemos entonces con las sumas de cuadrados; inicialmente con los totales $SC_{TOT} = \sum X_{ij}^2 - FC$ la expresión matemática indica la suma de todos y cada uno de los datos encontrados de la variable respuesta que son elevados al cuadrado; su resultado que se restará del factor de corrección.
- Continuamos con la suma de cuadrados de los tratamientos utilizando la formula $SC_{TRAT} = \frac{\sum \tau_{ij}^2}{r} - FC$ que se consiste en la suma de cada grupo de observaciones individuales y según cada tratamiento al cuadrado, luego su resultado es dividido para el número de repeticiones; el total encontrado debe restarse del factor de corrección

- Extraemos la suma de cuadrados de las columnas, $SC_{COL} = \frac{\sum \beta_j^2}{r} - FC$ que será la suma de cada observación correspondiente a una determinada columna y elevado al cuadrado y dividido para el número de repeticiones, su respuesta se resta del factor de corrección.
- Se procede parecido con la suma de cuadrados de las filas o hileras con la expresión matemática $SC_{HIL} = \frac{\sum \delta_k^2}{r} - FC$; es decir que se debe sumar las observaciones al cuadrado encontrada en cada fila, dividir para el número de repeticiones; y luego este resultado restar del factor de corrección.
- Procedemos igual con la suma de cuadrados del error obtenida por sustracción $SC_{ERROR} = SC_{TOTAL} - SC_{TRAT} - SC_{COL} - SC_{HIL}$ desde la suma total menos los tratamientos y repeticiones.
- En los grados de libertad (gl); en los totales se resta el total de observaciones menos uno $GL_{TOT} = (n - 1)$; para los tratamientos se obtiene del número de tratamientos menos uno $GL_{TRAT} = (t - 1)$ así como en las columnas $(c - 1)$ e hileras $(h - 1)$; luego para el error solo por diferencia de los grados de libertad totales menos los grados de libertad de los tratamientos, las columnas e hileras $GL_{ERROR} = GL_{TOT} - GL_{TRAT} - GL_{COL} - GL_{HIL}$.
- Luego se debe calcular la varianza o cuadrado medio de los tratamientos dividiendo la suma de cuadrados de los tratamientos para los respectivos grados de libertad $CM_{TRAT} = \frac{SC_{TRAT}}{GL_{TRAT}}$. Seguidamente los cuadrados medios de las columnas $CM_{COL} = \frac{SC_{COL}}{GL_{COL}}$ serán encontrados de la suma de cuadrados de las columnas para sus grados de libertad; igual procedemos con las filas o hileras $CM_{HIL} = \frac{SC_{HIL}}{GL_{HIL}}$. Por su parte, los cuadrados medios del error se lograrán dividiendo la suma de cuadrados correspondiente para los grados de libertad $CM_{ERROR} = \frac{SC_{ERROR}}{GL_{ERROR}}$.
- Las hipótesis son comprobadas con el estadístico Fisher calculado que en el caso de los tratamientos puede encontrarse dividiendo el cuadrado

medio de los tratamientos y el cuadrado medio del error, es decir la expresión $F_{\text{CAL.TRAT}} = \frac{CM_{\text{TRAT}}}{CM_{\text{ERROR}}}$. Así mismo procedemos con el estadístico

Fisher calculado para las Columnas $F_{\text{CAL.COL}} = \frac{CM_{\text{COL}}}{CM_{\text{ERROR}}}$ y las hileras

$$F_{\text{CAL.HIL}} = \frac{CM_{\text{HIL}}}{CM_{\text{ERROR}}}.$$

- Posteriormente se encuentra el coeficiente de variación en base a la expresión matemática: $CV = \frac{\sqrt{CM_{\text{ERROR}}}}{\bar{X}} * 100$
- Se analizará el valor del Fisher calculado en referencia a su valor de Fisher tabular encontrado en la tabla (según los grados de libertad del numerado y denominador) para establecer como verdadera o no a la hipótesis nula.

Las pruebas Fisher se utiliza para determinar si los efectos de los tratamientos son estadísticamente significativos, comparando la variabilidad entre tratamientos con la variabilidad dentro de los tratamientos.

4.9. Aplicación del DCL

Ejercicio 4.2.

En un experimento con pasto elefante (*Penisetum purpureum*), diseñado con un modelo cuadrado latino (5*5), se han probado 5 biofertilizantes sobre los rendimientos (kg/parcela).

Problema:

Las exigencias del mercado para la comercialización de carne de vacuno en los pastizales del subtrópico ecuatoriano, requiere de completar la cadena productiva con técnicas de manejo “sello verde”, por lo que se deberá biofertilizar el pasto en reemplazo de los químicos.

Objetivo:

Valorar el rendimiento de pasto elefante (*Penisetum purpureum*), mediante 5 tipos de biofertilizantes para aportar la producción de carne bovina con sello verde.

Hipótesis:

- *Teórica*

Las 5 técnicas de biofertilización producirán un rendimiento relativamente uniforme en el comportamiento vegetativo del pasto elefante.

- *Matemática*

$$H_0 = \mu_{Biol} = \mu_{Humus} = \mu_{Bokashi} = \mu_{Compost} = \mu_{Bioway}$$

Resultados experimentales

Hileras	Columnas					Suma Hileras $\sum \delta_k$	Suma Tratam $\sum \tau_i$	Media Tratam $\bar{X}_{\tau_i}$
	1	2	3	4	5			
1	39,13 D	33,40 C	32,70 A	31,90 B	44,37 E	181,50	164,50 A	32,90
2	32,23 B	33,50 A	44,33 E	36,80 C	42,27 D	189,13	159,10 B	31,82
3	42,27 E	31,30 B	44,40 D	31,87 A	37,60 C	187,44	175,43 C	35,09
4	33,50 A	42,93 D	35,23 C	44,37 E	33,17 B	189,20	213,46 D	42,69
5	32,40 C	42,93 E	30,50 B	44,73 D	32,93 A	183,49	218,27 E	43,65
Sum. Col. $\sum \beta_j$	179,53	184,06	187,16	189,67	190,34		930,76	37,23
							Total $\sum X_{..}$	Media $\bar{X}$

Factor de corrección

$$F.C. = \frac{(\sum X_{..})^2}{(t * r)}$$

$$F.C. = \frac{(930,76)^2}{(5 * 5)} = \frac{866314,18}{25} = 34652,57$$

Suma de cuadrados

- Suma de cuadrados totales

$$SC_{TOT} = \sum X_{ijk}^2 - FC$$

$$SC_{TOT} = [(39,13)^2 + (33,40)^2 + \dots + (44,73)^2 + (32,93)^2] - 34652,57$$

$$SC_{TOT} = 35320,22 - 34652,57 = 667,66$$

- Suma de cuadrados de los tratamientos

$$SC_{TRAT} = \frac{\sum \tau_i^2}{r} - FC$$

$$SC_{TRAT} = \frac{[(164,50)^2 + (159,10)^2 + (175,43)^2 + (213,46)^2 + (218,27)^2]}{5} - 34652,57$$

$$SC_{TRAT} = \frac{176355,71}{5} - 34652,57 = 35271,14 - 34652,57 = 618,57$$

- Suma de cuadrados de las columnas

$$SC_{COL} = \frac{\sum \beta_j^2}{r} - FC$$

$$SC_{COL} = \frac{[(179,53)^2 + (184,06)^2 + (187,16)^2 + (189,67)^2 + (190,34)^2]}{5} - 34652,57$$

$$SC_{COL} = \frac{173341,99}{5} - 34652,57 = 34668,40 - 34652,57 = 15,83$$

- Suma de cuadrados de las hileras o filas

$$SC_{HIL} = \frac{\sum \delta_k^2}{r} - FC$$

$$SC_{HIL} = \frac{[(181,50)^2 + (189,13)^2 + (187,44)^2 + (189,20)^2 + (183,49)^2]}{5} - 34652,57$$

$$SC_{HIL} = \frac{173311,38}{5} - 34652,57 = 34662,28 - 34652,57 = 9,71$$

- Suma de cuadrados del error

$$SC_{ERROR} = SC_{TOTAL} - SC_{TRAT} - SC_{COL} - SC_{HIL}$$

$$SC_{ERROR} = 667,66 - 618,57 - 15,83 - 9,71 = 23,54$$

Grados de libertad

$$GL_{TOT} = (n - 1) = (t * r) - 1 = 25 - 1 = 24$$

$$GL_{TRAT} = (t - 1) = 5 - 1 = 4$$

$$GL_{COL} = (c - 1) = 5 - 1 = 4$$

$$GL_{HIL} = (h - 1) = 5 - 1 = 4$$

$$GL_{ERROR} = GL_{TOT} - GL_{TRAT} - GL_{COL} - GL_{HIL} = 24 - 4 - 4 - 4 = 12$$

Cuadrados medios

- Cuadrados medios de los tratamientos

$$CM_{TRAT} = \frac{SC_{TRAT}}{GL_{TRAT}}$$

$$CM_{TRAT} = \frac{618,57}{4} = 154,64$$

- Cuadrados medios de las columnas

$$CM_{COL} = \frac{SC_{COL}}{GL_{COL}}$$

$$CM_{COL} = \frac{15,83}{4} = 3,96$$

- Cuadrados medios de las hileras o filas

$$CM_{HIL} = \frac{SC_{HIL}}{GL_{HIL}}$$

$$CM_{HIL} = \frac{9,71}{4} = 2,43$$

- Cuadrados medios del error

$$CM_{ERROR} = \frac{SC_{ERROR}}{GL_{ERROR}}$$

$$CM_{\text{ERROR}} = \frac{23,54}{12} = 1,96$$

- Estadístico Fisher calculado

$$F_{\text{CAL.TRAT}} = \frac{CM_{\text{TRAT}}}{CM_{\text{ERROR}}}$$

$$F_{\text{CAL.TRAT}} = \frac{154,64}{1,96} = 78,83$$

$$F_{\text{CAL.COL}} = \frac{CM_{\text{COL}}}{CM_{\text{ERROR}}}$$

$$F_{\text{CAL.COL}} = \frac{3,96}{1,96} = 2,02$$

$$F_{\text{CAL.HIL}} = \frac{CM_{\text{HIL}}}{CM_{\text{ERROR}}}$$

$$F_{\text{CAL.HIL}} = \frac{2,43}{1,96} = 1,24$$

- Coeficiente de variación del ADEVA

$$CV = \frac{\sqrt{CM_{\text{ERROR}}}}{\bar{X}} * 100$$

$$CV = \frac{\sqrt{1,96}}{37,23} * 100 = \frac{1,40}{37,23} * 100 = 3,76\%$$

ADEVA

Fuente de variación	Suma de cuadrados	Grados de libertad	Cuadrado medio	Fcal	F.Tab 0,05	F.Tab 0,01	Signif.
Total	667,66	24					
Tratamientos	618,57	4	154,64	78,83	3,26	5,41	**
Columnas	15,83	4	3,96	2,02	3,26	5,41	ns
Hileras	9,71	4	2,34	1,24	3,26	5,41	ns
Error	23,54	12	1,96				
CV	3,76%						

Análisis del estadístico Fisher de los tratamientos

En los tratamientos la inferencia se da debido a valor Fisher calculado que fue de 78,83 superando al valor tabular al nivel $\alpha 0,05$ que es de 3,26; e inclusive al nivel $\alpha 0,01$ (5,41)..

Interpretación

Los biofertilizantes, aplicados en la producción de pasto elefante (*Penisetum purpureum*) resultaron con diferencias altamente significativas en el rendimiento por parcela ($P < 0,01$); por lo tanto, no hay evidencia para aceptar la hipótesis nula.

Ejercicio 4.3.

Un médico veterinario ha aplicado 3 tipos de terapia en animales de compañía para mejorar su bienestar mientras permanecen en su clínica. Al tratarse especies de diferente tipo genético se aplicó un DCL 3x3, la variable a medir fue el bienestar en escala de 1 a 100. ¿Qué terapia recomendaría?.

Problema:

Las clínicas veterinarias que tienen servicio de hospitalización suelen presentar problemas de estrés en los animales de compañía, por lo que se pretende implementar terapias que reduzcan este comportamiento animal.

Objetivo:

Mejorar el bienestar animal en la hospitalización de canes mediante la aplicación de 3 terapias antiestrés para reducir su morbilidad.

Hipótesis:

- Teórica

Las terapias permitirán mejorar el bienestar de los animales de compañía atendidos en la clínica.

- Matemática

$$H_0 = \mu_1 = \mu_2 = \mu_3$$

Resultados experimentales

Hileras	Columnas			Suma Hileras $\sum \delta_k$	Suma Tratam $\sum \tau_i$	Media Tratam $\bar{X}_{\tau_i}$
	1	2	3			
1	12 B	85 C	13 A	110	39 A	13,00 A
2	98 C	14 A	13 B	125	37 B	12,33 B
3	12 A	12 B	88 C	112	271 C	90,33 C
Sum. Col. $\sum \beta_j$	122	111	114		347	38,56
					Total $\Sigma X_{..}$	Media $\bar{X}$

Factor de corrección

$$F.C. = \frac{(\Sigma X_{..})^2}{(t * r)}$$

$$F.C. = \frac{(347,00)^2}{(3 * 3)} = \frac{120409,00}{9} = 13378,78$$

Suma de cuadrados

- Suma de cuadrados totales

$$SC_{TOT} = \Sigma X_{ijk}^2 - FC$$

$$SC_{TOT} = [(12)^2 + (85)^2 + (13)^2 + \dots + (12)^2 + (88)^2] - 13378,78$$

$$SC_{TOT} = 25539,00 - 13378,78 = 12160,22$$

- Suma de cuadrados de los tratamientos

$$SC_{TRAT} = \frac{\Sigma \tau_i^2}{r} - FC$$

$$SC_{TRAT} = \frac{[(39)^2 + (37)^2 + (271)^2]}{3} - 13378,78$$

$$SC_{TRAT} = \frac{76331,00}{3} - 13378,78 = 25443,67 - 13378,78 = 12064,89$$

- Suma de cuadrados de las columnas

$$SC_{COL} = \frac{\sum \beta_j^2}{r} - FC$$

$$SC_{COL} = \frac{[(122)^2 + (111)^2 + (114)^2]}{3} - 13378,78$$

$$SC_{COL} = \frac{40201,00}{3} - 13378,78 = 13400,33 - 13378,78 = 21,56$$

- Suma de cuadrados de las hileras o filas

$$SC_{HIL} = \frac{\sum \delta_k^2}{r} - FC$$

$$SC_{HIL} = \frac{[(110)^2 + (125)^2 + (112)^2]}{3} - 13378,78$$

$$SC_{HIL} = \frac{40269,00}{3} - 13378,78 = 13423,00 - 13378,78 = 44,22$$

- Suma de cuadrados del error

$$SC_{ERROR} = SC_{TOTAL} - SC_{TRAT} - SC_{COL} - SC_{HIL}$$

$$SC_{ERROR} = 12160,22 - 12064,89 - 21,56 - 44,22 = 29,56$$

Grados de libertad

$$GL_{TOT} = (n - 1) = (t * r) - 1 = 9 - 1 = 8$$

$$GL_{TRAT} = (t - 1) = 3 - 1 = 2$$

$$GL_{COL} = (c - 1) = 3 - 1 = 2$$

$$GL_{HIL} = (h - 1) = 3 - 1 = 2$$

$$GL_{ERROR} = GL_{TOT} - GL_{TRAT} - GL_{COL} - GL_{HIL} = 8 - 2 - 2 - 2 = 2$$

Cuadrados medios

- Cuadrados medios de los tratamientos

$$CM_{TRAT} = \frac{SC_{TRAT}}{GL_{TRAT}}$$

$$CM_{TRAT} = \frac{12064,89}{2} = 6032,44$$

- Cuadrados medios de las columnas

$$CM_{COL} = \frac{SC_{COL}}{GL_{COL}}$$

$$CM_{COL} = \frac{21,56}{2} = 10,78$$

- Cuadrados medios de las hileras o filas

$$CM_{HIL} = \frac{SC_{HIL}}{GL_{HIL}}$$

$$CM_{HIL} = \frac{44,22}{2} = 22,11$$

- Cuadrados medios del error

$$CM_{ERROR} = \frac{SC_{ERROR}}{GL_{ERROR}}$$

$$CM_{ERROR} = \frac{29,56}{2} = 14,78$$

- Estadístico Fisher calculado

$$F_{CAL.TRAT} = \frac{CM_{TRAT}}{CM_{ERROR}}$$

$$F_{CAL.TRAT} = \frac{6032,44}{14,78} = 408,21$$

$$F_{CAL.COL} = \frac{CM_{COL}}{CM_{ERROR}}$$

$$F_{CAL.COL} = \frac{10,78}{14,78} = 0,73$$

$$F_{\text{CAL.HIL}} = \frac{CM_{\text{HIL}}}{CM_{\text{ERROR}}}$$

$$F_{\text{CAL.HIL}} = \frac{22,11}{14,78} = 1,50$$

- Coeficiente de variación del ADEVA

$$CV = \frac{\sqrt{CM_{\text{ERROR}}}}{\bar{X}} * 100$$

$$CV = \frac{\sqrt{14,78}}{38,56} * 100 = \frac{3,84}{38,56} * 100 = 10,0\%$$

ADEVA

Fuente de variación	Suma de cuadrados	Grados de libertad	Cuadrado medio	Fcal	F.Tab 0,05	F.Tab 0,01	Signif.
Total	12160,22	8					
	12064,89	2	6032,44		19,0	99,0	
Tratamientos				408,21			**
Columnas	21,56	2	10,78	0,73	19,0	99,0	ns
Hileras	44,22	2	22,11	1,50	19,0	99,0	ns
Error	29,56	2	14,78				
CV	10,0%						

Análisis del estadístico Fisher de los tratamientos

Para inferir y comprobar la hipótesis planteada con anterioridad; se tomará en cuenta el estadístico Fisher calculado en la fuente de variación que corresponde a los tratamientos (408,21) que supera valor tabular al nivel $\alpha 0,05$ (19) y al nivel $\alpha 0,01$ (99).

Interpretación

Las terapias que se aplicaron a los animales de compañía en la clínica veterinaria, resultaron con diferencias altamente significativas en el comportamiento antiestrés ($P < 0,01$); esto indicaría que no hay evidencia para aceptar la hipótesis nula.

Referencias bibliográficas del capítulo

- Aredo García, E. M. (2021). Análisis estadístico para un Diseño de Bloques Incompletos Balanceado.
- Brandt, A. E. (1957). *Experimental Designs*. William G. Cochran and Gertrude M. Cox. Wiley, New York; Chapman & Hall, London, ed. 2, 1957. xiv+ 616 pp. \$10.25. *Science*, 126(3279), 930-930.
- Gómez, L., & López, P. (2021). Características del diseño de bloques en investigaciones experimentales. *Journal of Experimental Design*, 10(2), 123-135. <https://doi.org/10.5678/jed.2021.002>
- González-Huerta, A., de Jesús Pérez-López, D., Hernández-Ávila, J., Franco-Martínez, J. R. P., Balbuena-Melgarejo, A., & Rubí-Arriaga, M. (2024). Serie de experimentos para tratamientos anidados en grupos con arreglo de bloques completos balanceados. *Revista Mexicana de Ciencias Agrícolas*, 15(7), e3831-e3831.
- Hoshmand, R. (2018). *Design of experiments for agriculture and the natural sciences*. Chapman and Hall/CRC.
- López, S., & Morales, F. (2021). Importancia del agrupamiento en el diseño de bloques completamente al azar. *Revista de Investigación Agrícola*, 18(4), 200-215. <https://doi.org/10.2345/ria.2021.004>
- Martínez, C., & Pérez, D. (2022). Flexibilidad y simplicidad en el diseño de bloques al azar. *Revista Internacional de Estadística Aplicada*, 9(3), 89-102. <https://doi.org/10.3456/riesa.2022.003>
- Martínez, J., Rodríguez, A., & Sánchez, V. (2021). Aplicaciones del diseño de bloques completamente al azar en la agricultura sostenible. *Agricultural Research Journal*, 14(2), 33-50. <https://doi.org/10.7890/arj.2021.002>
- Montgomery, D. C. (2017). *Design and analysis of experiments*. John Wiley & sons.

-
- Patterson, H. D., & Thompson, R. (1971). Recovery of inter-block information when block sizes are unequal. *Biometrika*, 58(3), 545-554. <https://doi.org/10.1093/biomet/58.3.545>
- Ramírez, E., & Soto, J. (2023). Análisis de varianza en diseños experimentales: Un enfoque práctico. *Revista de Métodos Cuantitativos*, 11(1), 12-25. <https://doi.org/10.4567/rmc.2023.001>
- Rojas, M., & Pérez, L. (2023). Ventajas y desventajas del diseño de bloques completamente al azar en estudios experimentales. *Journal of Experimental Methods*, 8(2), 99-110. <https://doi.org/10.2345/jem.2023.002>
- Vargas, T., Martínez, R., & Gómez, A. (2022). Asignación aleatoria en diseños experimentales: Herramientas y técnicas. *Estadística y Métodos de Investigación*, 20(3), 150-165. <https://doi.org/10.6789/emi.2022.003>



CAPÍTULO V

PRUEBAS DE
SIGNIFICACIÓN O
COMPARACIÓN DE
MEDIAS DE LOS
TRATAMIENTOS

La comparación de medias de los tratamientos es una herramienta fundamental en el análisis estadístico, especialmente en contextos experimentales donde se busca evaluar la efectividad de los diferentes tratamientos. Este capítulo se centra en las pruebas de significación que permiten determinar si las diferencias observadas entre las medias de varios tratamientos son estadísticamente relevantes o si, por el contrario, se atribuyen al azar. Comprender estas pruebas es crucial para la interpretación adecuada de los resultados y para la toma de decisiones informadas en diversas disciplinas, como la medicina, la psicología, la educación y la agropecuaria.

En el contexto de la investigación, la comparación de medias se utiliza para determinar si un tratamiento específico tiene un efecto significativo en relación con otros tratamientos o un grupo de control. Este proceso implica el uso de diversas pruebas estadísticas que evalúan la variabilidad de los datos y la significación de las diferencias observadas. La interpretación de estas pruebas no solo requiere un conocimiento técnico de los métodos estadísticos, sino también una comprensión profunda de los supuestos que subyacen a cada prueba. Por lo tanto, es vital que los investigadores se familiaricen con las características y limitaciones de cada enfoque para garantizar que sus conclusiones sean válidas y confiables.

Las pruebas de comparación de medias requieren que los datos sean recogidos de manera adecuada y que se cumplan ciertos supuestos, como la normalidad de los datos, la homogeneidad de las varianzas de los tratamientos y la independencia de datos. La normalidad implica que los datos se distribuyan de manera aproximadamente normal, lo cual es un requisito para muchas pruebas estadísticas. La homogeneidad de varianzas, por otro lado, significa que las varianzas de los diferentes grupos deben ser similares; y la independencia que no existan sesgos y su respuesta se deba exclusivamente a los tratamientos aplicados. Si estos supuestos no se cumplen, los resultados obtenidos son engañosos y llevan a conclusiones incorrectas.

Además, la variabilidad o desviación de los datos es un factor crucial en el análisis de la comparación de medias. La variabilidad se refiere a la dispersión

de los datos alrededor de la media e influye en la capacidad de detectar diferencias significativas. Un alto grado de variabilidad dentro de un grupo dificulta la identificación de diferencias entre los grupos, mientras que una baja variabilidad facilita la detección de efectos significativos. Por lo tanto, es esencial que los investigadores consideren la variabilidad al diseñar sus experimentos y al seleccionar las pruebas adecuadas.

Un concepto clave que se discutirá en este capítulo es la diferencia mínima significativa (DMS). La DMS establece un umbral que debe ser superado para que una diferencia entre las medias de los tratamientos se considere estadísticamente significativa. Este umbral se calcula en función de la variabilidad de los datos y el tamaño de la muestra, y su valor varía dependiendo del nivel de significación establecido por el investigador. Generalmente, un nivel de significación de 0.05 se utiliza de manera común, lo que indica que hay un 95% de confianza en que los resultados obtenidos no son el resultado del azar.

La DMS permite a los investigadores interpretar sus resultados de manera más clara, ya que proporciona un criterio cuantitativo para evaluar la significación de las diferencias observadas. Si la diferencia entre dos tratamientos supera la DMS, se concluye que hay evidencia suficiente para afirmar que los tratamientos tienen efectos distintos. Por el contrario, si la diferencia es menor que la DMS, no se rechaza la hipótesis nula, lo que implica que las diferencias podrían ser simplemente el resultado de la variabilidad aleatoria en los datos.

Posteriormente, presentaremos diversas pruebas estadísticas que son ampliamente utilizadas para la comparación de medias. Una de estas pruebas es la prueba del rango múltiple de Duncan. Esta prueba es especialmente útil cuando se desea identificar qué grupos son significativamente diferentes entre sí después de haber encontrado una diferencia global significativa mediante un análisis de varianza (ANOVA). La prueba de Duncan se basa en la comparación de rangos y permite detectar diferencias entre grupos de manera más específica, lo que la convierte en una herramienta valiosa en el análisis de datos.

Otra prueba importante que se abordará es la prueba de Tukey, conocida por su enfoque en el control del error tipo I en comparaciones múltiples. La prueba de

Tukey es particularmente relevante en estudios con varios grupos, ya que proporciona un método sistemático para realizar comparaciones que minimizan el riesgo de obtener resultados falsos positivos. Esta prueba calcula un rango crítico que se utiliza para determinar si las diferencias entre las medias de los tratamientos son significativas, ofreciendo así un enfoque riguroso y confiable para la comparación de medias.

Adicionalmente, se incluirá la prueba de Scheffe, que es una técnica de comparación múltiple que permite realizar comparaciones más complejas entre grupos. Esta prueba es ideal para situaciones donde se están evaluando múltiples tratamientos y se desea explorar interacciones más complejas. Aunque es más conservadora en su enfoque, su capacidad para manejar situaciones complicadas la convierte en una opción valiosa para los investigadores que buscan realizar análisis rigurosos.

Finalmente, para consolidar los conocimientos adquiridos, se proporcionarán ejercicios de apoyo que permitirán a los lectores aplicar las técnicas y conceptos aprendidos a lo largo del capítulo. Estos ejercicios no solo facilitarán la comprensión teórica, sino que también ofrecerán una oportunidad práctica para que los lectores se familiaricen con los procedimientos y cálculos necesarios para llevar a cabo comparaciones de medias en diferentes contextos. A través de ejemplos prácticos, se espera que los lectores desarrollen habilidades que les permitan aplicar estas técnicas en situaciones reales, fortaleciendo así su capacidad para realizar análisis estadísticos de manera efectiva.

5.1. Características de la comparación de medias de los tratamientos

La comparación de medias de los tratamientos es un componente esencial en el análisis de datos en experimentación animal, especialmente cuando se utilizan modelos lineales aditivos simples. Estos modelos permiten evaluar el efecto de diferentes tratamientos sobre una variable de respuesta, facilitando la identificación de diferencias significativas entre grupos. Una característica fundamental de este enfoque es la suposición de que los datos son independientes y normalmente distribuidos, lo cual es crucial para la validez de las inferencias estadísticas (Zuur et al., 2010). En estudios de experimentación animal, esto implica que las observaciones de cada grupo deben ser obtenidas

de manera que no influyan entre sí, asegurando así la integridad de los resultados.

Además, la homogeneidad de varianzas es otro supuesto importante que debe ser verificado antes de realizar comparaciones de medias. Esto se evalúa mediante pruebas como la prueba de Levene, que determina si las varianzas de los grupos son significativamente diferentes (Michael et al., 2014). En el contexto de la experimentación animal, donde las variaciones biológicas llegan a ser significativas, garantizar la homogeneidad de varianzas es fundamental para la correcta aplicación de técnicas como el ANOVA o sus extensiones.

La elección de la prueba estadística adecuada es crucial y depende del diseño experimental y del número de tratamientos comparados. En experimentos con más de dos grupos, el ANOVA es comúnmente utilizado, pero si se encuentran diferencias significativas, es necesario realizar pruebas post-hoc, como la prueba de Tukey o la prueba de Duncan, para identificar cuáles grupos difieren entre sí (Richardson, 2011). Estas pruebas son especialmente útiles en estudios de experimentación animal, donde se evalúan múltiples tratamientos simultáneamente.

Otro aspecto relevante es el cálculo del tamaño del efecto, que proporciona información sobre la magnitud de las diferencias observadas entre tratamientos. Un tamaño de efecto significativo no solo indica que hay diferencias estadísticamente significativas, sino que también ayuda a interpretar la relevancia práctica de los resultados (Nakagawa, & Cuthill, 2007). En la investigación animal, esto es particularmente importante, ya que los investigadores deben considerar no solo si un tratamiento es efectivo, sino también cuán relevante es en términos de aplicación práctica.

Una de las pruebas comunes que integra los conceptos de análisis de varianza (ADEVA o ANOVA) y la separación de medias es el estadístico t-student, que en la praxis prueba una hipótesis nula que indica inexistencia de diferencias significativas entre dos grupos, o 2 tratamientos por ejemplo, y se da por la diferencia entre 2 medias y dividido por el error típico de esta diferencia, la expresión sería:

$$t = \frac{\overline{X_1} - \overline{X_2}}{s\bar{d}}$$

Donde:

- $\overline{X_1} - \overline{X_2}$ es la diferencia en la comparación de las 2 medias
- $s\bar{d}$ es el error típico

La limitante se presenta cuando existen más de 2 medias para comparar; si esto se da entonces es preferible elegir otra prueba como Duncan o Tukey.

Indudablemente antes de aplicar una prueba de este tipo, es preciso que en el ANOVA o ADEVA la prueba Fisher rechace a la hipótesis nula ($H_0: \mu_1 = \mu_2 = \mu_3 = \dots = \mu_n$) entonces esto indicaría que no todos los tratamientos son iguales, o alguno de ellos es diferente; pero se necesita conocer el tratamiento más favorecido para poder recomendar o decidir aplicarlo con el fin de resolver el problema. Las pruebas de significación, separación o comparación de medias miden el verdadero efecto de los tratamientos.

Finalmente, la comparación de medias en el contexto de modelos lineales aditivos simples requiere una planificación cuidadosa del diseño experimental y la recolección de datos. La implementación de un diseño robusto minimiza sesgos y errores, lo que resulta en conclusiones más confiables y aplicables (Lehner, 2018).

Entre las más utilizadas en el área agropecuaria se pueden anotar las siguientes:

- Diferencia Mínima Significativa (DMS)
- Prueba de Rango Múltiple de Duncan (RMD)
- Tukey
- Scheffé

5.2. Diferencia mínima significativa

La diferencia mínima significativa es un concepto clave en la evaluación de resultados, particularmente en el ámbito de la comparación de medias de

tratamientos. La DMS se define como el valor mínimo que debe superar la diferencia entre las medias de dos o más grupos para que se considere estadísticamente significativa. Este umbral es crucial para la interpretación de los resultados de un estudio, ya que ayuda a los investigadores a discernir entre diferencias que son significativas y aquellas que son atribuidas al azar.

Igual que el t-student, puesto es una variante del mismo se usa para comparaciones simples entre medias; el método hace una aproximación de la solución correcta al problema de las comparaciones entre medias; además se debe usar con precaución, y solo es válida en el caso de comparación de medias planeadas (solo pares).

Para el cálculo de la DMS, se usa la distribución “t”, generalmente la hipótesis plantea:

$$H_1 = \tau_j \neq \tau_k$$

Entonces:

$$DMS = t_{(g \text{ de l. error})} * s\bar{d}$$

$$s\bar{d} = \sqrt{\frac{2s^2(error)}{r}}$$

$$DMS = t_{\alpha} \times \sqrt{\frac{2s^2}{n}}$$

Donde:

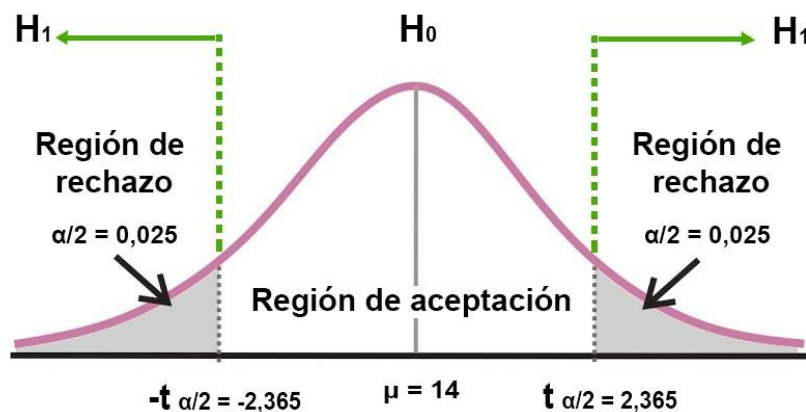
- t_{α} es el valor crítico de la distribución t de Student para un nivel de significación α ,
- s^2 es la varianza o cuadrado medio del error
- r es el tamaño de la muestra o repetición de cada tratamiento
- $s\bar{d}$ es el error típico

Esta fórmula permite ajustar la DMS en función de la variabilidad de los datos y el tamaño de la muestra, lo que proporciona un enfoque más robusto y adaptado a las características específicas del estudio.

Para poder decidir se toma en cuenta los valores de la diferencia de medias y se compara con los tabulares, en el caso de 2 colas y a un nivel de significancia de α 0,05 el límite para considerar verdadera a la hipótesis nula sería $-t_{\alpha/2}$ -2,365 a $t_{\alpha/2}$ 2,365; como se puede apreciar en la figura 5.1.; es decir que:

Si $-t_{\alpha/2} < t_c < t_{\alpha/2}$, entonces se rechaza la hipótesis nula (H_0).

Figura 5.1. Esquema para comprobación de hipótesis con uso del DMS en la comparación de un par de medias de los tratamientos



Elaborado y adaptado por: Moscoso, 2025

Entre las ventajas de utilizar la DMS se encuentran:

1. **Claridad en la interpretación de resultados:** La DMS proporciona un umbral claro que ayuda a los investigadores a determinar si las diferencias observadas en los tratamientos son estadísticamente significativas, (Hsu, 2020). Esto es crucial en el sector agropecuario, donde las decisiones afectan la producción y la rentabilidad.
2. **Utilidad:** Es un valor de fácil cálculo y de simple uso, Es válida cuando se realizan comparaciones planeadas con anterioridad, las mismas que identifiquen el objetivo del experimento.

3. **Respuestas satisfactorias:** Puede brindar resultados satisfactorios, cuando se comparan cada una de las medias con el testigo: $\bar{X}_1$ vs. T_s ; $\bar{X}_2$ vs. T_s ; $\bar{X}_3$ vs. T_s ; $\bar{X}_n$ vs. T_s

4. **Control de errores Tipo I:** Al establecer un criterio cuantitativo, la DMS ayuda a minimizar el riesgo de errores tipo I, es decir, la aceptación incorrecta de una diferencia significativa que en realidad no existe. Esto es especialmente importante en estudios donde se realizan múltiples comparaciones entre tratamientos.

5. **Facilita la toma de decisiones:** Los resultados basados en la DMS guían a los agricultores y gestores en la selección de prácticas agronómicas más efectivas, optimizando el uso de recursos y mejorando la producción.

6. **Aplicación en experimentos agronómicos:** La DMS es útil en ensayos de campo para evaluar el impacto de diferentes fertilizantes, variedades de cultivos o métodos de manejo, permitiendo comparaciones objetivas entre tratamientos.

Sus principales desventajas son:

1. **Dependencia del tamaño de la muestra:** La DMS es influenciada por el tamaño de la muestra; en estudios con muestras pequeñas, la DMS tiende a ser más alta, lo que dificulta la detección de diferencias significativas, (Bittencourt et al., 2022). Esto constituye un problema en el sector agropecuario, donde a menudo se realizan ensayos en condiciones limitadas.

2. **No considera la relevancia práctica:** Aunque una diferencia sea estadísticamente significativa, no siempre implica que sea relevante desde el punto de vista práctico. Esto conlleva a decisiones que no necesariamente mejoren la producción o la sostenibilidad.

3. **Limitaciones en el Análisis Multivariado:** En estudios agropecuarios complejos donde múltiples variables influyen en los

resultados, la DMS no captura adecuadamente las interacciones entre estas variables, lo que limita su utilidad en la interpretación de los datos.

4. **Posibilidad de resultados engañosos:** Si no se complementa con un análisis más exhaustivo, el uso exclusivo de la DMS conlleva a conclusiones erróneas. Por ejemplo, dos tratamientos que muestren diferencias significativas, si no se evalúan otros factores como el costo o el impacto ambiental, la decisión final no será la más precisa.

5. **Limitación con más de dos tratamientos:** No toma en consideración el número de tratamientos dentro de un experimento, es decir, no ofrece protección cuando el número de medias aumenta a partir de dos. Muchos pares de medias pueden asomar como significativas; se considera una prueba muy liberal.

6. **Significancia:** Solo se usa DMS cuando en el ADEVA la fuente de variación de los tratamientos es significativa o altamente significativa.

7. **Ortogonalidad:** Cuando las comparaciones son no ortogonales (es decir una media está involucrada en más de una comparación), puede llevar a conclusiones erróneas, haciendo ver que existe más información de la que en realidad hay.

5.2.1. Aplicación de la prueba DMS.

Ejercicio 5.1.

Como base nos fijamos en el ejercicio 3.1.; en donde se experimenta con 3 niveles de fertilización nitrogenada (75, 150, 225 kg/ha) para el pasto festuca; en el ADEVA las diferencias fueron altamente significativas, por lo que debemos separa las medias y fijarnos en el comportamiento de cada tratamiento; por lo tanto, las medias fueron:

N (75 kg/ha) = 14,7 kg/ha de rendimiento

N (150 kg/ha) = 25,2 kg/ha de rendimiento

N (225 kg/ha) = 32,4 kg/ha de rendimiento

Los grados de libertad del error es 12; y el cuadrado medio del mismo 0,037

Aplicamos la ecuación

$$DMS = t_{(g\ de\ l.\ error)} * s\bar{d}$$

$$s\bar{d} = \sqrt{\frac{2s^2(error)}{r}} = \sqrt{\frac{(2*0,037)}{5}} = \sqrt{\frac{(0,075)}{5}} = \sqrt{0,015} = 0,1222.$$

Tabla A-2 Probabilidad de obtener un valor mayor o igual a *t*, ignorando el signo.

G.L.	α 0.200	0.100	0.050	0.025	0.010	0.005	0.001
1	3.078	6.314	12.706	25.452	63.657		
2	1.886	2.920	4.303	6.205	9.925	14.089	31.598
3	1.638	2.353	3.182	4.176	5.841	7.453	12.941
4	1.533	2.132	2.776	3.495	4.604	5.598	8.610
5	1.476	2.015	2.571	3.163	4.032	4.773	6.859
...							
11	1.363	1.796	2.201	2.593	3.006	3.497	4.437
12	1.356	1.782	2.179	2.560	3.055	3.428	4.318
13	1.350	1.771	2.160	2.533	3.012	3.372	4.221
14	1.345	1.761	2.145	2.510	2.977	3.326	4.140
15	1.341	1.753	2.131	2.490	2.947	3.286	4.073

Para el valor “*t*” nos fijamos en los grados de libertad del error (12) y encontramos el valor en la tabla para $\alpha 0,05$ (2,179) y $\alpha 0,01$ (3,055).

Procedemos con el cálculo del DMS al $\alpha 0,05$ y $\alpha 0,01$

$$DMS_{\alpha 0,05} = (2,179) * (0,1222) = 0,266$$

$$DMS_{\alpha 0,01} = (3,055) * (0,1222) = 0,373$$

Luego realizamos la comparación, como se aprecia la prueba es limitada ya que solo se realizará en pares de tratamientos y por orden es decir solo nos permite analizar entre tratamiento 1 vs tratamiento 2; que corresponde a las medias de: N (75 kg/ha) vs N (150 kg/ha):

T1 vs T2 = 14,7 – 25,2 = -10,44 (es un valor absoluto)

Entonces 10,44 > que 0,226 y 0,373 (DMS), esto indicaría que las diferencias son altamente significativas; en otras palabras el rendimiento que corresponde a

25,2 kilos de pasto por parcela con adición de 150 kilos de nitrógeno por hectárea superó estadísticamente a la menor dosis (75 kg/ha).

Podría compararse T2 vs T3; sin embargo, el método indica que se perdería la ortogonalidad.

5.3. Prueba del rango múltiple de Duncan (RMD)

La prueba del rango múltiple de Duncan es una técnica estadística utilizada para realizar comparaciones múltiples entre medias de tratamientos en estudios experimentales, especialmente después de haber llevado a cabo un análisis de varianza (ANOVA). Introducida por David B. Duncan en la década de los años 50, esta prueba permite identificar cuáles grupos son significativamente diferentes entre sí, facilitando así la interpretación de los resultados experimentales. Su aplicación es común en diversas áreas, incluyendo la agricultura, la biología y las ciencias sociales, donde es crucial determinar la efectividad de diferentes tratamientos o condiciones experimentales.

La prueba del rango múltiple de Duncan no requiere que el Fisher calculado sea significativo; se basa en la comparación de las medias de los grupos y utiliza un enfoque de rangos para identificar las diferencias. La formulación básica implica calcular la diferencia mínima significativa (DMS) entre las medias de los grupos sin restricciones siendo más estricta que la DMS.

Posteriormente, se comparan las diferencias absolutas entre todas las medias de los grupos con la DMS calculada. Si la diferencia entre dos medias es mayor que la DMS, se concluye que hay una diferencia significativa entre esos dos grupos.

El cálculo se basa en la siguiente ecuación:

$$RMD = S_{\bar{x}} * RMD_{\text{Tabular}}$$

$$S_{\bar{x}} = \sqrt{\frac{CM_{\text{ERROR}}}{r}}$$

Donde:

- RDM, es el rango múltiple de Duncan
- RMD_{Tabular} es el valor de comparación extraído de la tabla en base al número de medias de los tratamientos considerados.
- $S_{\bar{x}}$ corresponde al error típico de la media
- CM_{ERROR} es la variancia o cuadrado medio del error (se encuentra en el ADEVA)
- r indica el número de repeticiones consideradas.

Los pasos para calcular el Rango Múltiple de Duncan son:

1. **Realizar un ANOVA:** Se realiza un análisis de varianza (ANOVA) para determinar si hay diferencias significativas entre las medias de los tratamientos. Se obtiene el valor de F.cal. y el p-valor o p-value.
2. **Determinar la media y el número de tratamientos (V_1):** Se cuenta cuántos tratamientos (grupos) estás comparando para encontrar los valores tabulares.
3. **Calcular el error típico de la media:** A partir del ANOVA, se identifica el error típico de la media ($S_{\bar{x}}$), que se utiliza en el cálculo del rango.
4. **Calcular el valor crítico de Duncan:** Utilizando tablas de valores críticos de Duncan, se busca el valor crítico correspondiente al nivel de significancia (α), el número de tratamientos (V_1) y el número de grados de libertad del error (V_2).
5. **Calcular el rango múltiple de Duncan:** Se calcula el rango múltiple de Duncan (RMD).
6. **Comparar las diferencias entre las medias:**
 - Se ordenan las medias de menor a mayor
 - Se calculan las diferencias absolutas entre todas las combinaciones de medias de tratamientos (Media – RMD)

- Se compara cada diferencia con el rango múltiple de Duncan, identificando el valor dentro del orden establecido de las medias para cobijar mediante letras y establecer la diferencia entre tratamientos.

7. Interpretar los resultados: Si la diferencia entre dos medias es mayor que RMD, considera que hay una diferencia significativa entre esos tratamientos y le indica la misma con las primeras letras del alfabeto; si dos tratamientos tienen la misma letra indicaría que no hay diferencias entre los dos, y lo contrario sería significativa la diferencia al nivel de significancia que se obtuvo de la tabla (valores $\alpha 0,05$ por lo general).

5.3.1. Aplicación de la prueba RMD.

Ejercicio 5.2.

Igualmente nos fijamos en el ejercicio 3.1.; en el que se prueban 3 niveles de fertilización nitrogenada (75, 150, 225 kg/ha) para fertilización de la festuca; el ADEVA presentó diferencias altamente significativas en los tratamientos, y se separarán las medias de los tratamientos mediante la prueba de Duncan; los rendimientos fueron:

N (75 kg/ha) = 14,7 kg/ha de rendimiento

N (150 kg/ha) = 25,2 kg/ha de rendimiento

N (225 kg/ha) = 32,4 kg/ha de rendimiento

El número de medias consideradas son 3; los grados de libertad del error 12; y el cuadrado medio del error 0,037 (ADEVA o ANOVA)

Aplicamos la ecuación:

$$RMD = S_{\bar{x}} * RMD_{\text{Tabular}}$$

$$S_{\bar{x}} = \sqrt{\frac{CM_{\text{ERROR}}}{r}} = \sqrt{\frac{0,037}{5}} = \sqrt{0,007} = 0,0864$$

Tratamientos	N (75 kg/ha)	N (150 kg/ha)	N (225 kg/ha)
Medias	14,7	25,2	32,4
RMD _{tab} α0,05		3,08	3,23
RMD		0,27	0,28
Diferencia		24,93	32,12
	c	b	a

Para los valores RMD tabular, nos fijamos en la tabla para 3 tratamientos (V1) y 12 grados de libertad del error (V2), y un nivel de significancia α0,05; estos serán: 3,08 y 3,23

Tabla VII.- Valores críticos para la prueba de Duncan.

$$U_{\alpha}(v_1, v_2)$$

v ₂ ↓	α ↓	v ₁					
		2	3	4	5	6	7
1	0.05	18.0	18.0	18.0	18.0	18.0	18.0
	0.01	90.0	90.0	90.0	90.0	90.0	90.0
2	0.05	6.09	6.09	6.09	6.09	6.09	6.09
	0.01	14.0	14.0	14.0	14.0	14.0	14.0
3	0.05	4.50	4.50	4.50	4.50	4.50	4.50
	0.01	8.26	8.5	8.6	8.7	8.8	8.9
...							
11	0.05	3.11	3.27	3.35	3.39	3.43	3.44
	0.01	4.39	4.63	4.77	4.86	4.94	5.01
12	0.05	3.08	3.23	3.33	3.36	3.40	3.42
	0.01	4.32	4.55	4.68	4.76	4.84	4.92

$$RMD = S_{\bar{x}} * RMD_{\text{Tabular}}$$

$$RMD_1 = 0,0864 * 3,08 = 0,28$$

$$RMD_2 = 0,0864 * 3,23 = 0,27$$

Las diferencias (Media – RMD) serían:

$$D_{N(225kg/ha)} = 32,4 - 0,28 = 32,12$$

La diferencia es de 32,12 y si nos fijamos en los valores de las medias ordenadas de menor a mayor, entonces este valor se encuentra entre el tratamiento 2 y 3; por lo que se cobija con la primera letra “a” solo a los de la derecha (corresponde a N (225 kg/ha)

$$D_{N(150kg/ha)} = 25,2 - 0,27 = 24,93$$

La segunda diferencia en cambio tiene un valor de 24,93, el mismo se encuentra entre la primera y segunda media; entonces cobijamos desde la derecha y corresponde a N (150 kg/ha) la letra “b”

Y el resto quedaría con la “c”; es decir se presentarían los resultados de la siguiente manera:

Variable respuesta	Tratamientos		
	N (75 kg/ha)	N (150 kg/ha)	N (225 kg/ha)
Rendimiento de festuca (kg/parcela)	14,7 c	25,2 b	32,4 a

Medias con una letra común no son significativamente diferentes ($p > 0,05$)

Interpretación:

Según la prueba de Duncan al $\alpha 0,05$, se pudo apreciar que el rendimiento de la festuca fertilizada con 225 kg/ha de nitrógeno, resultó con mejores promedios (42,4 kg/parcela), superando a los demás tratamientos. Esto indica que hay diferencias significativas en las medias de los tratamientos evaluados.

Una de las principales ventajas de la prueba del rango múltiple de Duncan es su capacidad para detectar dichas diferencias incluso con tamaños de muestra relativamente pequeños. Esto es particularmente útil en investigaciones donde los recursos son limitados, como en estudios agrícolas donde a menudo se realizan ensayos con un número reducido de repeticiones (Duncan, 1955). Además, la prueba permite realizar un gran número de comparaciones simultáneas sin incrementar de manera excesiva el riesgo de cometer un error

tipo I, lo que la hace más flexible en comparación con otros métodos de comparaciones múltiples, como la prueba de Tukey.

Otra ventaja notable es que la prueba de Duncan proporciona un enfoque intuitivo para el análisis de datos. Los resultados se presentan en forma de rangos, lo que facilita la visualización de las diferencias entre los grupos. Esto es especialmente beneficioso para investigadores y profesionales que necesitan comunicar sus resultados a audiencias no especializadas, como agricultores o responsables de políticas (Hsu, 1996). Además, la prueba es aplicable a datos que no necesariamente cumplen con la normalidad, aunque lo más recomendable es que se realicen pruebas de normalidad y homogeneidad de varianzas para asegurar la validez de los resultados.

A pesar de sus ventajas, la prueba del rango múltiple de Duncan también presenta desventajas. Una de las principales limitaciones es su sensibilidad a la normalidad de los datos. Si los datos no se distribuyen normalmente, la prueba produce resultados engañosos. Esto es particularmente relevante en el sector agropecuario, donde los datos están sujetos a variaciones naturales y no siempre siguen una distribución normal (Kramer, 1956). Por lo tanto, es fundamental realizar pruebas de normalidad antes de aplicar la prueba de Duncan.

Otra desventaja es que la prueba es menos robusta en situaciones donde las varianzas entre grupos son desiguales. En tales casos, la prueba conlleva a conclusiones erróneas sobre la significancia de las diferencias observadas. Esto es especialmente crítico en estudios agropecuarios donde las condiciones experimentales varían considerablemente. Por ejemplo, si se evalúan diferentes variedades de cultivos bajo condiciones de estrés hídrico, las diferencias en varianza influyen los resultados de la prueba, lo que podría llevar a decisiones inadecuadas sobre qué variedad cultivar (Hsu, 1996).

Además, la prueba del rango múltiple de Duncan no considera la magnitud de las diferencias en términos de relevancia práctica. Aunque sí identifica diferencias estadísticamente significativas, no necesariamente indica que estas diferencias sean importantes desde un punto de vista práctico. Esto es crucial en el ámbito agropecuario, donde decisiones basadas únicamente en significancia

estadística podrían no traducirse en mejoras prácticas en la producción o la sostenibilidad, (Bittencourt et al., 2022).

La prueba del rango múltiple de Duncan se ha utilizado ampliamente en el sector agropecuario para evaluar la efectividad de diferentes tratamientos, como fertilizantes, variedades de cultivos o prácticas de manejo. Por ejemplo, en un estudio que evalúa el rendimiento de diferentes variedades de maíz bajo distintas condiciones de riego, la prueba de Duncan ayuda a identificar cuáles variedades producen rendimientos significativamente diferentes, guiando así las decisiones de los agricultores sobre qué cultivar en función de las condiciones específicas de su región.

Sin embargo, es esencial que los investigadores complementen la prueba de Duncan con otros análisis estadísticos y consideraciones prácticas. Por ejemplo, realizar un análisis de costos-beneficios es crucial para determinar si una variedad que muestra una diferencia significativa en rendimiento también es económicamente viable. Esto es especialmente importante en un contexto donde los agricultores buscan maximizar no solo el rendimiento, sino también la rentabilidad y la sostenibilidad de sus prácticas agrícolas (Bittencourt et al., 2022).

5.4. Prueba de Tukey

La Prueba de Tukey, o método de comparación múltiple de Tukey, es una técnica estadística fundamental en el análisis de varianza (ANOVA) que permite realizar comparaciones entre las medias de diferentes grupos. Su principal objetivo es identificar cuáles grupos presentan diferencias significativas en sus medias después de haber determinado que existe al menos una diferencia significativa entre las medias globales. Esta prueba se basa en el cálculo de intervalos de confianza para las diferencias entre todas las combinaciones posibles de medias de grupos, utilizando la distribución Studentized Range (Q) para evaluar si las diferencias observadas son mayores que lo que se esperaría por azar.

A diferencia del anterior método se requiere de un solo valor tabular para la determinación de la significación en las diferencias de los tratamientos. Es una

prueba de gran adaptabilidad; superior y más estricta que la DMS, y Duncan porque la unidad considerada es el experimento en sí.

Se debe calcular un valor D, que es el producto entre el error típico de la media y un factor Q (tabular) con α 0,05 o α 0,01 y los grados de libertad del error y los tratamientos involucrados

$$D = Q_{(G.L.ERROR)} * s_{\bar{x}}$$

Donde:

D es el Tukey

$Q_{(G.L.ERROR)}$ representa el valor tabular en base al número de medias consideradas y los grados de libertad del error experimental

$s_{\bar{x}}$ es el error típico de la media

Los pasos para realizar una Prueba de Tukey son:

1. **Realiza el ANOVA:** Antes de aplicar la Prueba de Tukey, debes realizar un análisis de varianza (ANOVA o ADEVA) para determinar si hay diferencias significativas entre las medias de los grupos. Si el valor p-value del ANOVA es menor que el nivel de significancia (por ejemplo, 0.05), se procede con la prueba de Tukey.
2. **Calcular la media y el número de los tratamientos (V_1) :** Obtenemos las medias de los tratamientos que estás comparando y organizarlos en orden ascendente, así como su número que servirá para establecer el valor crítico en la tabla.
3. **Calcular el error típico de la media:** Este valor se obtiene del ANOVA con la varianza del error aplicando la formula establecida ($S_{\bar{x}}$).
4. **Calcular el valor crítico de Tukey (Q):** Utiliza una tabla de distribución de Tukey o una función en software estadístico para encontrar el valor crítico Q, que depende del nivel de significancia (α), el número de tratamientos (V_1) y los grados de libertad del error (V_2).

5. **Calculamos el Duncan:** Con el uso de su ecuación donde interviene el único valor tabular multiplicado por el error típico.

6. **Comparar las diferencias con el valor crítico:** Para cada diferencia entre medias, se calcula el DMS.

- Ordenamos las medias de los tratamientos en sentido ascendente (de menor a mayor)
- Se calculan las diferencias absolutas entre todas las combinaciones de medias de tratamientos (Media – Q)
- Se compara cada diferencia con el Tukey, identificando el valor dentro del orden establecido de las medias para cobijar mediante letras y establecer la diferencia entre tratamientos.

7. **Interpretar los resultados:** Si la diferencia entre las medias de dos grupos es mayor que la DMS calculada, se considera que hay una diferencia significativa entre esos dos grupos.

Ejercicio 5.2.

En base al ejercicio 4.1.; en el que evalúan 5 niveles de pollinaza como suplemento alimenticio en el levante y engorde de toretes Brown Swiss; encontrándose dentro del ADEVA diferencias altamente significativas debido a los tratamientos, para la variable ganancia de peso diaria (kg). La comparación de las medias de los tratamientos será mediante la prueba de Tukey con $\alpha 0,05$; los resultados de las medias fueron:

P-0: 0,433 kg/día

P-20: 0,514 kg/día

P-40: 0,851 kg/día

P-60: 1,042 kg/día

P-80: 0,750 kg/día

Se someterán a comparación un total de 4 tratamientos y 1 testigo (P-0); con error 12 grados de libertad del error y su cuadrado medio de 0,010 (ver ADEVA).

Aplicamos la ecuación:

$$D = Q_{(G.L.ERROR)} * S_{\bar{x}}$$

$$S_{\bar{x}} = \sqrt{\frac{CM_{ERROR}}{r}} = \sqrt{\frac{0,010}{4}} = \sqrt{0,002} = 0,0499$$

Tratamientos	P-0	P-20	P-80	P-40	P-60
Medias	0,433	0,514	0,750	0,851	1,042
Q _{tab} α0,05		4,510	4,510	4,510	4,510
D		0,225	0,225	0,225	0,225
Diferencia		0,289	0,525	0,626	0,817
	d	d	bc	ab	a

Para Q tabular, nos fijamos en la tabla para 5 tratamientos (V1) y 12 grados de libertad del error (V2), y un nivel de significancia α0,05; por consiguiente, será un solo valor de: 4,51.

Tabla VI.- Valores críticos para la prueba de Tukey.
 $q_{\alpha}(v_1, v_2)$

v_2 ↓	α ↓	v_1					
		2	3	4	5	6	7
1	0.05	18.00	29.98	32.82	37.08	40.41	43.12
	0.01	90.03	135.0	164.3	185.6	202.2	215.8
2	0.05	6.10	8.33	9.80	10.88	11.74	12.44
	0.01	14.04	19.02	22.29	24.72	26.63	28.20
3	0.05	4.50	5.91	6.82	7.50	8.04	8.48
	0.01	8.26	10.62	12.17	13.33	14.24	15.00
----	-----	-----	-----	-----	-----	-----	-----
10	0.05	3.15	3.88	4.33	4.65	4.91	5.12
	0.01	4.48	5.27	5.77	6.14	6.43	6.67
11	0.05	3.11	3.82	4.26	4.57	4.82	5.03
	0.01	4.39	5.14	5.62	5.97	6.25	6.48
12	0.05	3.08	3.77	4.20	4.51	4.75	4.95
	0.01	4.32	5.04	5.50	5.84	6.10	6.32

$$D = Q_{(G.L.ERROR)} * S_{\bar{x}}$$

$$D = 4,51 * 0,0499 = 0,2225$$

Las diferencias (Media – D) serían:

$$P60 = 1,042 - 0,225 = 0,817$$

La primera diferencia es de 0,817, de acuerdo a las media ordenadas de menor a mayor, este valor se ubica entre los tratamientos P-80 y P-40; por esta razón cobijamos con la primera letra “a” hacia la derecha a los promedios de P-40 y P-60, indicando que no existe diferencia entre los dos.

$$P40 = 0,851 - 0,225 = 0,626$$

Para la segunda diferencia se obtiene un valor de 0,626, el mismo se P-20 y P-40, pasando por P-P-80; por lo tanto al cobijar desde la derecha y corresponde la segunda letra del alfabeto “b” a las medias de P-80 y P-40.

$$P80 = 0,750 - 0,225 = 0,525$$

La tercera diferencia es de 0,525 que se ubica entre los tratamientos ordenados P-20 y P-80, por lo que corresponde solo cobijar a la derecha al promedio de P-80 con la tercera letra (c).

$$P20 = 0,514 - 0,225 = 0,289$$

Finalmente en esta diferencia restante fue 0,289, que es un valor menor a la primera media (P-0), por lo cual cobijaremos con la cuarta letra alfabética (d) a las medias de los tratamientos P-0 y P-20.

Para presentar los resultados que corresponden a la separación de medias por la prueba de Tukey podemos usar la tabla siguiente:

Variable respuesta	Tratamientos				
	P-0	P-20	P-40	P-60	P-80
Ganancia de peso (kg/día)	0.433d	0,514d	0,851ab	1,042a	0,750bc
Medias con una letra común no son significativamente diferentes ($p > 0,05$)					

Interpretación:

De acuerdo la prueba de Tukey al $\alpha 0,05$ para la comparación de medias, la mejor ganancia diaria se presentó en los toretes que recibieron un suplemento de 60 y 40% de pollinaza (1,042 y 0,851 kg/día), superando al testigo (P-0), P-20 y P-80; los mismos que resultaron con ganancias menos favorecidas.

Ejercicio 5.3.

Se realizó un experimento para evaluar el efecto de tres diferentes fertilizantes en el crecimiento de alfalfa (*Medicago sativa*). Se obtuvieron las siguientes medias de altura de planta (cm) en el primer mes de evaluación:

Fertilizante A: 25 cm

Fertilizante B: 30 cm

Fertilizante C: 40 cm

Se realizó un ANOVA o ADEVA y se encontró que el valor de F_{cal} es significativo ($p < 0.05$). El cuadrado medio del error es 16, y el número de repeticiones por tratamiento es 4; y los grados de libertad para el error es 9. Calcular el rango de Tukey y determinar qué diferencias son significativas entre los tratamientos.

Aplicamos la ecuación:

$$D = Q_{(G.L.ERROR)} * S_{\bar{x}}$$

$$S_{\bar{x}} = \sqrt{\frac{CM_{ERROR}}{r}} = \sqrt{\frac{16}{4}} = \sqrt{4} = 2$$

Tratamientos	Fert-A	Fert-B	Fert-C
Medias	25	30	40
$Q_{tab} \alpha 0,05$		3,95	3,95
D		7,90	7,90
Diferencia		22,10	32,10
	b	b	a

El valor Q tabular, usamos la tabla para 3 tratamientos (V1) y 9 grados de libertad del error (V2), a un $\alpha 0,05$; será un solo valor de: 3,95.

$$D = Q_{(G.L.ERROR)} * S_{\bar{x}}$$

$$D = 3,95 * 2,00 = 7,90$$

Las diferencias (Media – D) serían:

$$\text{Fertilizante C} = 40,0 - 7,9 = 32,1$$

Para el primer cobijamiento fertilizante C la diferencia es 32,1 que se ubica entre Fert-B y Fert-C; corresponde “a” el tratamiento Fert-C es decir supera a todos.

$$\text{Fertilizante B} = 30,0 - 7,9 = 22,1$$

Es decir, el valor de la diferencias es de 22,1 está mucho menor de la media del fertilizante A (25), por lo tanto se cobija a la derecha tanto a Fert-A y Fert-B; con la segunda letra (b)

Si comparamos solo pares de medias, en valores absolutos, podemos indicar lo siguiente:

Fert-A vs. Fert-B: $|25 - 30| = 5 < 7.9$ (no significativa la diferencia)

Fert-A vs. Fert-C: $|25 - 40| = 15 > 7.0$ (diferencia significativa)

Fert-B vs. Fert-C: $|30 - 40| = 10 > 7.0$ (diferencia significativa)

Interpretación:

La prueba de Tukey para la comparación de medias con un nivel de significancia de $\alpha 0,05$, indicó que la mejor altura de planta fue para la fertilización C que alcanzó 40 cm; mientras que el resto se mantuvo con menor respuesta.

La prueba de Tukey es especialmente útil cuando se tienen tres o más grupos y se busca compararlos entre sí. Una de sus principales ventajas es su capacidad para controlar el error tipo I al realizar múltiples comparaciones. Esto se logra al utilizar un enfoque que considera todas las comparaciones posibles entre grupos, lo que minimiza la probabilidad de obtener resultados falsos positivos. Esto es especialmente importante en investigaciones donde se realizan múltiples

pruebas, ya que el riesgo de cometer un error tipo I aumenta con cada comparación adicional (Tukey, 1949).

Otra ventaja significativa es que la prueba de Tukey proporciona un enfoque claro y directo para la comparación de medias. Los resultados se presentan en forma de rangos, lo que facilita la visualización de las diferencias entre los grupos. Esto es particularmente beneficioso para investigadores y profesionales que necesitan comunicar sus resultados a audiencias no especializadas, como agricultores o responsables de políticas (Hsu, 1996). Además, la prueba es robusta frente a violaciones moderadas de la normalidad y la homogeneidad de varianzas, lo que la hace más flexible en comparación con otros métodos de comparaciones múltiples.

A pesar de sus ventajas, la prueba de Tukey también presenta desventajas. Entre sus limitaciones están que es menos efectiva cuando hay diferencias de tamaño de muestra entre los grupos.

La prueba asume que todos los grupos tienen el mismo tamaño de muestra, lo que trae resultados inexactos si esta suposición no se cumple (Hsu, 1996). En situaciones donde los tamaños de muestra son desiguales, es necesario utilizar métodos alternativos o ajustar los tamaños de muestra para obtener resultados más precisos.

Además, la prueba de Tukey es menos sensible en la detección de diferencias significativas cuando estas son pequeñas. Aunque la prueba controla el error tipo I, no es tan efectiva para identificar diferencias que son relevantes desde un punto de vista práctico, pero que no alcanzan el umbral de significancia estadística (Bittencourt et al., 2022).

Lo subrayado es crucial en el ámbito agropecuario, donde decisiones basadas únicamente en significancia estadística no siempre resultan en mejoras prácticas en la producción o la sostenibilidad. Por su parte, la prueba de Tukey se ha utilizado ampliamente en el sector agropecuario para evaluar la efectividad de diferentes tratamientos, como fertilizantes, variedades de cultivos o prácticas de manejo.

Un ejemplo, de lo antes subrayado es en el estudio que evalúa el rendimiento de diferentes variedades de trigo bajo distintas condiciones de riego, la prueba de Tukey ayuda a identificar cuáles variedades producen rendimientos significativamente diferentes, guiando así las decisiones de los agricultores sobre qué cultivar en función de las condiciones específicas de su región.

Sin embargo, es esencial que los investigadores complementen la prueba de Tukey con otros análisis estadísticos y consideraciones prácticas. Realizar un análisis de costos-beneficios es crucial para determinar si una variedad que muestra una diferencia significativa en rendimiento también es económicamente viable. Esto es especialmente importante en un contexto donde los agricultores buscan maximizar no solo el rendimiento, sino también la rentabilidad y la sostenibilidad de sus prácticas agrícolas (Bittencourt et al., 2022).

Los métodos de Duncan y Tukey son técnicas de comparación múltiple utilizadas tras un análisis de varianza (ANOVA), pero difieren significativamente en su enfoque y resultados. Duncan es un método menos conservador, lo que le permite detectar diferencias significativas entre las medias de grupos con mayor potencia, aunque esto conlleva un mayor riesgo de cometer errores tipo I (falsos positivos). Por otro lado, Tukey es más conservador, controlando mejor este tipo de error, lo que lo hace ideal para situaciones donde se realizan múltiples comparaciones y se busca una mayor precisión en los resultados. Aunque Duncan es superior en términos de poder para identificar diferencias, Tukey es preferido en contextos donde la precisión y la reducción de falsos positivos son primordiales.

5.5. Prueba de Scheffé

La Prueba de Scheffé es un método diseñado para realizar comparaciones múltiples tras la ejecución de un ANOVA. Esta prueba permite comparar todas las combinaciones posibles de medias de grupos, lo que la convierte en una herramienta flexible y potente para identificar diferencias significativas. Se lleva a cabo calculando un estadístico basado en las diferencias entre las medias de los grupos y el número de grupos involucrados, utilizando un nivel de significancia preestablecido.

Se considera una prueba más estricta que las anteriores y de fácil aplicación, se utilizan los valores de la tabla Fisher (la misma que se usa para el ADEVA).

Entre las ventajas de la Prueba de Scheffé se destaca su capacidad para controlar el error tipo I de manera efectiva, lo que la hace adecuada para situaciones en las que se realizan múltiples comparaciones. Sin embargo, su desventaja principal es que es menos potente que otros métodos, como los de Duncan y Tukey, lo que significa que no detecta diferencias significativas que realmente existen. Esto se debe a que la Prueba de Scheffé es más conservadora, lo que conlleva a una pérdida de información útil en ciertos contextos.

La Prueba de Scheffé es especialmente útil en estudios exploratorios donde se desea evaluar múltiples combinaciones de grupos, como en investigaciones en ciencias sociales y biológicas. En contraste, los métodos de Duncan y Tukey son más apropiados cuando se busca identificar diferencias específicas entre grupos predefinidos. Mientras que Duncan es menos conservador y tiene mayor potencia para detectar diferencias, lo que resulta en un mayor riesgo de error tipo I, Tukey ofrece un equilibrio entre potencia y control del error, siendo más conservador que Duncan pero menos que Scheffé. Por lo tanto, la elección entre estos métodos dependerá de los objetivos del análisis y de la importancia de controlar los errores en el contexto específico de la investigación. El valor S calculado para comparar medias está dado por la expresión:

$$S = F_o * S\bar{d}$$

$$F_o = \sqrt{F(p-1)}$$

Donde:

S es el Scheffe

F es el valor tabular encontrado con los grados de libertad del error, grados de libertad de los tratamientos y el nivel escogido

p es el número de tratamientos

$S\bar{d}$ representa al error típico

La Prueba de Scheffé se realiza en varios pasos después de haber llevado a cabo un análisis de varianza (ANOVA). A continuación, se detallan los pasos para realizar esta prueba:

1. **Realizar ANOVA:** Primero, se debe realizar un ANOVA de un solo factor para determinar si hay diferencias significativas entre las medias de los tratamientos. Esto implica calcular la suma de cuadrados, los grados de libertad y la media cuadrática para los tratamientos, las demás fuentes de variación incluido el error.
2. **Calcular el estadístico Fisher:** Se calcula el estadístico F de la ANOVA utilizando la fórmula:

$$F. Cal. = \frac{CM_{TRAT}}{CM_{ERROR}}$$

Donde:

CM_{TRAT} es la varianza o cuadrado medios de los tratamientos o grupos y CM_{ERROR} es la varianza o cuadrado medio del error.

3. **Determinar los grados de libertad:** Se identifica desde el ADEVA o ANOVA los grados de libertad de los tratamientos (numerador) y del error experimental (denominador).
4. **Establecer el nivel de significancia:** Se selecciona un nivel de significancia α (comúnmente 0.05).
5. **Se calcula el error típico:** A partir del ANOVA, se identifica el error típico ($S\bar{d}$),
6. **Calcular el valor crítico de Scheffé:** Se utiliza la fórmula respectiva se calcula el valor S.
7. **Comparar las diferencias de medias:** Ordenando de mayor a menor las medias de los tratamientos, se sacan las diferencias absolutas

(Media – S), y se compara para asignar el literal correspondiente que definirá la diferencia o no entre las mismas.

- 8. Interpretar los resultados:** Si la diferencia absoluta entre las medias de cualquier par de grupos es mayor que el valor crítico S , se concluye que hay una diferencia significativa entre esos grupos o tratamientos.

Ejercicio 5.4.

Sobre la base del ejercicio 4.2., donde se expone un ensayo en pasto elefante (*Penisetum purpureum*), diseñado con un cuadrado latino (5*5), para probar 5 biofertilizantes sobre sus rendimientos (kg/parcela). La comparación de medias será mediante el método Scheffé con una nivel de significancia $\alpha 0,05$; los resultados de los grupos o tratamientos fueron:

Fertilizante A: 32,90

Fertilizante B: 31,82

Fertilizante C: 35,09

Fertilizante D: 42,68

Fertilizante E: 43,65

Se compararán un total de 5 tratamientos, y 5 repeticiones; los grados de libertad de los tratamientos es 4 (numerador) y el error experimental 12 (denominador)

El cuadrado medio del error es 1,96,

En la tabla de F.tab, ubicamos el valor tabular con un nivel de significancia de $\alpha 0,05$ con 4 grados de libertad para los tratamientos y 12 para el error, corresponde a: 3,26

Distribución F 0.05

Distribución F 0.01

En las columnas se encuentran los valores F que corresponden al área 0.05 a la derecha

En las columnas se encuentran los grados de libertad del numerador

En los renglones se encuentran los grados de libertad del denominador.

	1	2	3	4	5	6		1	2	3	4	5	6
1	161.4	199.5	215.7	224.6	230.2	234.0	1	4052	4999	5403	5625	5764	5859
2	18.51	19.00	19.16	19.25	19.30	19.33	2	98.50	99.00	99.17	99.25	99.30	99.33
3	10.13	9.55	9.28	9.12	9.01	8.94	3	34.12	30.82	29.46	28.71	28.24	27.91
4	7.71	6.94	6.59	6.39	6.26	6.16	4	21.20	18.00	16.69	15.98	15.52	15.21
5	6.61	5.79	5.41	5.19	5.05	4.95	5	16.26	13.27	12.06	11.39	10.97	10.67
6	5.99	5.14	4.76	4.53	4.39	4.28	6	13.75	10.92	9.78	9.15	8.75	8.47
7	5.59	4.74	4.35	4.12	3.97	3.87	7	12.25	9.55	8.45	7.85	7.46	7.19
8	5.32	4.46	4.07	3.84	3.69	3.58	8	11.26	8.65	7.59	7.01	6.63	6.37
9	5.12	4.26	3.86	3.63	3.48	3.37	9	10.56	8.02	6.99	6.42	6.06	5.80
10	4.96	4.10	3.71	3.48	3.33	3.22	10	10.04	7.56	6.55	5.99	5.64	5.39
11	4.84	3.98	3.59	3.36	3.20	3.09	11	9.65	7.21	6.22	5.67	5.32	5.07
12	4.75	3.89	3.50	3.26	3.11	3.00	12	9.33	6.93	5.95	5.41	5.06	4.82

Aplicamos la ecuación:

$$F_o = \sqrt{F(p-1)}$$

$$F_o = \sqrt{3,26 * (5-1)} = \sqrt{(13,04)} = 3,61$$

$$s\bar{d} = \sqrt{\frac{2s^2(error)}{r}}$$

$$s\bar{d} = \sqrt{\frac{2 * (1,96)}{5}} = \sqrt{\frac{3,923}{5}} = \sqrt{0,784} = 0,885$$

$$S = F_o * s\bar{d}$$

$$S = 3,61 * 0,885 = 3,198$$

Tratamientos	B	A	C	D	E
Medias	31,820	32,900	35,086	42,692	43,654
F.tab α0,05		2,260	2,260	2,260	2,260
S		3,198	3,198	3,198	3,198
Diferencia		29,702	31,888	39,494	40,456
	d	cd	c	ab	a

Las comparaciones y cobijamientos literales en las medias serían:

$$\text{Biofert. E} = 43,654 - 3,198 = 40,456$$

En la primera comparación la diferencia es de 40,456 valor que se ubica entre los tratamientos C y D, por lo tanto se cobija con la primera letra “a” desde su derecha y le corresponde a D y E.

$$\text{Biofert. D} = 42,692 - 3,198 = 39,494$$

En la segunda comparación encontramos un valor de 39,494 que se encuentra entre los tratamientos C y D, correspondiendo el literal “b” solo a D.

$$\text{Biofert. C} = 35,086 - 3,198 = 31,888$$

En esta diferencia (31,888) podemos apreciar que el valor se halla entre B y C pasando por A, por lo tanto, como el método exige, cobijamos desde la derecha con la tercera letra del alfabeto “c” a las medias de los tratamientos A y C

$$\text{Biofert. A} = 32,900 - 3,198 = 29,708$$

Finalmente la última comparación nos indica un valor de 29,708 de diferencia que cobija a los tratamientos restantes, es así que la letra “d” le corresponde a los tratamientos B y A.

Estos resultados se pueden reportar en una tabla

Variable respuesta	Tratamientos (biofertilizantes)				
	A	B	C	D	E
Rendimiento de pasto elefante, kg/parcela	32,9cd	31,82d	35,08c	42,69ab	43,65a

Interpretación:

Según la prueba de Scheffé al $\alpha 0,05$ para la comparación de medias, el mejor rendimiento del pasto elefante se apreció en los tratamientos E y D con promedios de 43,65 y 42,69 kilos por parcela respectivamente, superando al resto de tratamientos.

Referencias bibliográficas del capítulo

- Bittencourt, K. C., Toebe, M., Souza, R. R. D., Pazetto, S. B., & Toebe, I. C. D. (2022). Sample size affects the precision of the analysis of variance in experiments with cauliflower seedlings. *Ciência Rural*, 53, e20220180. <http://doi.org/10.1590/0103-8478cr20220180>
- Cuthill, Innes. (2007). Nakagawa S, Cuthill IC.. Effect size, confidence intervals and statistical significance: a practical guide for biologists. *Biol Rev Camb Philos Soc* 82: 591-605. Biological reviews of the Cambridge Philosophical Society. 82. 591-605. <https://doi.org/10.1111/j.1469-185X.2007.00027.x>
- Duncan, D. B. (1955). Multiple range and multiple F tests. *Biometrics*, 11(1), 1-42. <https://doi.org/10.2307/3001478>
- Hsu, C. (2020). Statistical Methods in Clinical Trials: Understanding the Difference Between Statistical Significance and Clinical Relevance. *Clinical Trials*, 17(3), 345-355.
- Hsu, J. (1996) Multiple Comparisons: Theory and Methods. Chapman and Hall/CRC. <https://doi.org/10.1201/b15074>
- Kramer, C. Y. (1956). Extensions of multiple range tests to group means with unequal numbers of replications. *Biometrics*, 12(3), 307-320. <https://doi.org/10.2307/3001469>
- Lehner, P. N. (1998). Handbook of ethological methods. Cambridge University Press.
- Michael F. W. Festing, Timo Nevalainen, The Design and Statistical Analysis of Animal Experiments: Introduction to this Issue, *ILAR Journal*, Volume 55, Issue 3, 2014, Pages 379–382, <https://doi.org/10.1093/ilar/ilu046>
- Richardson, Alice. (2011). Multiple Comparisons Using R by Frank Bretz, Torsten Hothorn, Peter Westfall. *International Statistical Review*. 79. 297-297. [10.2307/41305047](https://doi.org/10.2307/41305047).

- Scheffé, H. (1953). A method for judging all contrasts in the analysis of variance. *Biometrika*, 40(1-2), 87-110. <https://doi.org/10.1093/biomet/40.1-2.87>
- Tukey, J. W. (1949). Comparing individual means in the analysis of variance. *Biometrics*, 99-114. <https://doi.org/10.2307/3001913>
- Zuur, Alain & Ieno, Elena & Elphick, Chris. (2010). Zuur AF, Ieno EN, Elphick CS.. A protocol for data exploration to avoid common statistical problems. *Methods Ecol Evol* 1: 3-14. *Methods in Ecology and Evolution*. 1. 3 - 14. [10.1111/j.2041-210X.2009.00001x](https://doi.org/10.1111/j.2041-210X.2009.00001x).



CAPÍTULO VI

APLICACIÓN DE LOS
MODELOS LINEALES
ADITIVOS SIMPLES
MEDIANTE EL
INFOSTAT

Los modelos lineales aditivos de ordenación simple, conocidos comúnmente como ANOVA de una vía, constituyen una herramienta fundamental en el análisis estadístico para la comparación de medias entre múltiples grupos. Estos modelos representan la extensión natural de la comparación de dos medias hacia escenarios donde se evalúan tres o más grupos simultáneamente, manteniendo la estructura aditiva característica de los modelos lineales. Su importancia radica en la capacidad de descomponer la variabilidad total de los datos en componentes atribuibles a diferentes fuentes, permitiendo determinar si existen diferencias estadísticamente significativas entre las medias poblacionales de los grupos bajo estudio.

6.1. Fundamentos Teóricos de los Modelos Lineales

Los modelos lineales constituyen el marco conceptual sobre el cual se desarrollan los análisis de varianza de una vía. Un vector de variables aleatorias sigue un modelo lineal cuando se puede expresar como $Y \cong N_n(X\beta, \ln\sigma)$, donde la variable respuesta se relaciona linealmente con las variables explicativas a través de una matriz de diseño. Esta formulación general engloba diferentes tipos de análisis según la naturaleza de las variables involucradas.

En el contexto específico del análisis de varianza, la variable respuesta es continua mientras que las variables explicativas son categóricas, lo que define precisamente el marco del ANOVA. La estructura aditiva del modelo implica que los efectos de diferentes factores se suman de manera independiente para explicar la variabilidad observada en la variable dependiente. Esta característica fundamental permite la descomposición de la varianza total en componentes específicos asociados a cada factor del diseño experimental.

6.2. Consideraciones Prácticas y Limitaciones

La implementación práctica de modelos lineales aditivos de ordenación simple requiere una cuidadosa consideración de las limitaciones inherentes del enfoque y las condiciones necesarias para su aplicación efectiva. Una limitación fundamental radica en la naturaleza ómnibus de la prueba F, que únicamente

indica la existencia de diferencias significativas sin especificar su localización o magnitud. Esta característica requiere la implementación de análisis post-hoc adicionales, como las pruebas de Tukey, Bonferroni, o contrastes planificados, para identificar específicamente cuáles grupos difieren entre sí.

La sensibilidad del ANOVA a las violaciones de supuestos representa otra consideración práctica importante. La violación del supuesto de normalidad puede abordarse mediante transformaciones de datos o pruebas no paramétricas alternativas como el test de Kruskal-Wallis. La heteroscedasticidad puede detectarse mediante pruebas específicas como la de Levene y abordarse mediante técnicas de ponderación o transformaciones estabilizadoras de varianza. El diseño experimental debe garantizar la independencia de las observaciones, ya que la dependencia entre observaciones puede inflar artificialmente los errores Tipo I.

Por su parte InfoStat implementa una interfaz amigable para la estimación de Modelos Lineales Generales y Mixtos¹.

6.3. Selección de Tamaño Muestral y Poder Estadístico

La planificación adecuada de estudios que emplean ANOVA de una vía requiere una consideración cuidadosa del tamaño muestral necesario para detectar diferencias de magnitud prácticamente relevante. El poder estadístico del análisis depende de factores como el tamaño del efecto esperado, la variabilidad de los datos, el nivel de significancia adoptado, y el número de grupos incluidos en la comparación. Tamaños muestrales inadecuados pueden resultar en la incapacidad de detectar diferencias importantes (errores Tipo II), mientras que muestras excesivamente grandes pueden detectar diferencias estadísticamente significativas, pero prácticamente irrelevantes.

La interpretación de los resultados debe considerar tanto la significancia estadística como la significancia práctica, evaluando la magnitud de las diferencias observadas en el contexto específico de la investigación. Medidas del tamaño del efecto, como eta cuadrada u omega cuadrada, proporcionan información complementaria sobre la proporción de variabilidad explicada por el

factor, facilitando la evaluación de la relevancia práctica de los hallazgos más allá de su significancia estadística.

Los modelos lineales aditivos de ordenación simple representan una herramienta estadística fundamental que combina simplicidad conceptual con rigor metodológico para abordar preguntas de investigación relevantes sobre diferencias entre grupos. Su estructura aditiva permite una descomposición clara de las fuentes de variabilidad, facilitando la interpretación de los resultados y la comunicación de los hallazgos a audiencias diversas. La flexibilidad del enfoque permite su aplicación en múltiples contextos de investigación, desde estudios experimentales controlados hasta análisis observacionales exploratorios.

Sin embargo, la aplicación efectiva de estos modelos requiere una comprensión sólida de sus supuestos fundamentales y limitaciones metodológicas. La violación de supuestos puede comprometer seriamente la validez de las conclusiones, mientras que la naturaleza ómnibus de la prueba F requiere análisis complementarios para la identificación específica de diferencias entre grupos. La planificación cuidadosa del diseño experimental, incluyendo la determinación del tamaño muestral adecuado y la consideración del poder estadístico, resulta esencial para maximizar la probabilidad de obtener resultados informativos y confiables.

La evolución hacia modelos más complejos, como los modelos aditivos generalizados y las técnicas de suavizado, ofrece oportunidades para abordar limitaciones específicas del enfoque básico mientras se mantiene la estructura conceptual fundamental. Esta progresión metodológica sugiere que los modelos de ordenación simple constituyen no solo una herramienta valiosa por sí misma, sino también un fundamento conceptual sólido para el desarrollo de técnicas estadísticas más avanzadas que pueden abordar preguntas de investigación de mayor complejidad y sofisticación.

6.4. Aplicación de los modelos lineales aditivos simples mediante el InfoStat

Dentro de esta plataforma, es posible aplicar modelos que pueden ser considerados "lineales aditivos simples", especialmente en el contexto de modelos mixtos aplicados a diseños experimentales.

Los modelos lineales aditivos de ordenación simple, particularmente el análisis de varianza (ANOVA) de una vía, encuentran en InfoStat una herramienta computacional versátil para su implementación y validación estadística. Este software, desarrollado específicamente para análisis de datos en contextos experimentales y observacionales, integra procedimientos clave para la ejecución de modelos lineales básicos, desde la importación de datos hasta la interpretación de resultados, pasando por la verificación de supuestos y la realización de pruebas post-hoc. Su interfaz intuitiva y su capacidad para manejar estructuras de datos complejas lo posicionan como una opción preferente en investigación agronómica, biológica y social, donde la comparación de medias entre grupos constituye un requisito metodológico fundamental.

La potencia del software radica en su capacidad para gestionar diseños desbalanceados mediante el uso de cuadrados medios tipo III, ajustando las estimaciones ante tamaños muestrales desiguales. Esto garantiza la validez de los resultados incluso en condiciones subóptimas de recolección de datos, siempre que se respeten los supuestos de normalidad, homocedasticidad e independencia.

Se detallan los pasos generales y ejemplos para aplicar este tipo de modelos en InfoStat:

1. Invocación del procedimiento: Para iniciar el análisis, debe seleccionar en el menú Estadísticas el submenú Modelos lineales generales y mixtos y luego la opción Estimación. Esto abrirá la ventana de diálogo para especificar la estructura del modelo.

2. Selección de variables: En la ventana de selección de variables, debe especificar la variable dependiente (respuesta) y los criterios de clasificación (factores) o covariables. Por ejemplo, se utiliza el archivo Atriplex.IDB2

especificando PG como variable respuesta y Tamaño y Episperma como criterios de clasificación.

3. Especificación de los efectos fijos: En la solapa Efectos fijos de la ventana principal de la interfaz, se indican los efectos fijos del modelo. Para un modelo aditivo simple, esto implicaría seleccionar los factores principales sin incluir interacciones. Por ejemplo, en el archivo Atriplex.IDB2, si quisiera un modelo aditivo simple de efectos fijos, seleccionaría Tamaño y Episperma pero no su interacción (Tamaño:Episperma). La fuente muestra cómo especificar estos efectos, aunque el ejemplo particular incluye la interacción en la fórmula $R_{PG} \sim 1 + \text{Tamano} + \text{Episperma} + \text{Tamano} : \text{Episperma}$.

4. Especificación de los efectos aleatorios (para modelos mixtos aditivos): Si el modelo incluye efectos aleatorios (lo que lo convierte en un modelo mixto), se especifican en la segunda solapa del diálogo, "Efectos aleatorios". Aquí puede elegir los criterios de estratificación o agrupamiento que incorporan efectos aleatorios.

5. Modelado de la estructura de error: Para modelos más complejos (no necesariamente aditivos simples, pero relevantes en el contexto de modelos lineales), InfoStat permite modelar la estructura de varianzas y covarianzas de los errores en las solapas Correlación y Heteroscedasticidad. Un modelo lineal aditivo simple estándar asumiría errores independientes y homoscedásticos, lo cual es la opción por defecto si no se declara nada en estas solapas. Sin embargo, las fuentes exploran modelos con errores correlacionados (e.g., autocorrelación serial o espacial) y heteroscedásticos (varianzas diferentes según un criterio de agrupamiento).

6. Análisis y diagnóstico: Una vez estimado el modelo, se activa el menú Análisis-exploración de modelos estimados, que contiene herramientas para el análisis diagnóstico. Esto incluye la revisión de residuos y otras herramientas gráficas y funciones de autocorrelación para evaluar la adecuación del modelo y supuestos como la homogeneidad de varianzas y la normalidad.

6.5. Entorno del InfoStat

El entorno de InfoStat está diseñado para facilitar el análisis estadístico tanto a usuarios principiantes como avanzados, combinando una interfaz intuitiva con potentes capacidades profesionales. A continuación, se describen sus principales componentes, organización y funcionalidades, según la documentación y fuentes especializadas.

El entorno InfoStat es una interfaz amigable para la estimación de modelos lineales generales y mixtos. Esta interfaz actúa como un puente hacia la plataforma R, utilizando los procedimientos *gls* y *lme* de la librería nlme. La interfaz gráfica con R fue desarrollada en Delphi y depende de R-DCOM para ejecutar R en segundo plano. Para que InfoStat pueda acceder a R, es necesario que el componente DCOM y R estén instalados en el sistema.

La interacción principal para ajustar estos modelos se realiza a través de una ventana de diálogo que se invoca seleccionando el submenú "Modelos lineales generales y mixtos" dentro del menú "Estadísticas", y luego eligiendo la opción "Estimación".

Esta ventana de diálogo está estructurada en varias solapas (pestañas) que permiten especificar diferentes componentes del modelo:

Solapa "Especificación de los efectos fijos": Aquí se definen la variable dependiente (respuesta), los factores de clasificación y las covariables. Permite especificar los efectos fijos que se incluirán en el modelo, como factores principales e interacciones. La ventana muestra cómo se traduciría esta especificación a la sintaxis de R.

Solapa "Efectos aleatorios": En esta solapa se seleccionan los criterios de estratificación o agrupamiento que se modelarán como efectos aleatorios. InfoStat permite especificar estructuras de anidamiento (>), cruzamiento (+) e interacciones (*) para estos efectos aleatorios, ya sea escribiéndolos directamente o usando un menú contextual⁶. Por defecto, se tilda la opción para que el efecto aleatorio actúe sobre la constante (intercepto). Se ejemplifica la especificación de efectos aleatorios anidados o cruzados.

- Solapa "Correlación": Se utiliza para modelar la estructura de varianzas y covarianzas de los errores, específicamente la estructura de correlación. Incluye opciones por defecto para errores independientes, y otras estructuras como autocorrelación serial (AR1), simetría compuesta, sin estructura, y modelos de correlación espacial (exponencial, gaussiana, lineal, rational quadratic, esférica). Se puede especificar una variable que indique el orden de las observaciones y los criterios de agrupamiento para la correlación.
- Solapa "Heteroscedasticidad": Permite seleccionar distintos modelos para la función de varianza de los errores, abordando la heterogeneidad de varianzas. Se puede indicar que la varianza es diferente para cada nivel de un criterio de agrupamiento, utilizando funciones como varIdent o varPower.
- Solapa "Comparaciones": Facilita la comparación de medias de efectos fijos. Permite obtener tablas de medias y errores estándar, así como realizar comparaciones múltiples (LSD de Fisher, DGC). También incluye opciones para corrección por comparaciones múltiples y para especificar contrastes personalizados.

Una vez que un modelo ha sido estimado, se activa la opción "Análisis – exploración de modelos estimados" en el menú. Este submenú ofrece un conjunto de herramientas para el análisis diagnóstico, incluyendo visualizaciones gráficas de residuos para evaluar supuestos como la homogeneidad de varianzas y la normalidad, y funciones de autocorrelación (ACF/PACF) para evaluar la correlación serial o espacial de los residuos.

InfoStat muestra la sintaxis de R generada por las especificaciones del modelo, lo cual es útil para usuarios familiarizados con R. Los resultados del ajuste del modelo se presentan en una ventana de salida, incluyendo medidas de ajuste (AIC, BIC, logLik, Sigma) y tablas de análisis de varianza.

En resumen, el entorno de InfoStat para modelos lineales se centra en una interfaz gráfica intuitiva que guía al usuario a través de las distintas etapas de especificación del modelo (efectos fijos, efectos aleatorios, estructura de

correlación y varianza) y proporciona herramientas integradas para el diagnóstico y la comparación de modelos.

6.5.1. Interfaz General y Navegación

Barra de Menús Superior: Incluye menús como Archivo, Edición, Datos, Resultados, Estadísticas, Gráficos, Ventanas, Ayuda y Aplicaciones. Cada menú agrupa funciones relacionadas, permitiendo acceder de forma ordenada a todas las herramientas del software

Barra de Herramientas: Ubicada debajo de la barra de menús, contiene botones de acceso rápido para tareas frecuentes como crear, abrir, guardar tablas, exportar, imprimir, agregar columnas, ordenar datos, editar categorías y modificar la alineación de celdas. Al pasar el cursor sobre un botón, se muestra una breve descripción de su función.

Ventanas Principales: InfoStat trabaja con tres tipos de ventanas:

- **Datos:** Donde se visualizan y editan las tablas de datos (filas = observaciones, columnas = variables).
- **Resultados:** Acumula los resultados de los análisis estadísticos realizados.
- **Gráficos:** Muestra y organiza los gráficos generados durante el análisis.

Todas pueden mantenerse abiertas simultáneamente y el usuario puede alternar entre ellas fácilmente

6.5.2. Gestión y Manejo de Datos

- **Importación/Exportación:** Permite trabajar con múltiples formatos de archivos, como InfoStat (.IDB, .IDB2), Excel (.XLS), texto (.TXT, .DAT), Dbase (.DBF), entre otros. Esto facilita la integración de datos desde diferentes fuentes y la exportación de resultados para uso externo.
- **Edición y Limpieza:** Incluye herramientas para agregar, eliminar y modificar columnas y filas, editar categorías, ordenar datos y realizar transformaciones básicas sobre las variables.

- **Organización de Tablas:** Las tablas pueden abrirse, guardarse y cerrarse desde el menú Archivo o mediante atajos de teclado y botones específicos. El entorno permite trabajar con varias tablas abiertas a la vez, facilitando el manejo de grandes volúmenes de datos

Análisis estadístico

- **Estadística Descriptiva:** Cálculo de medidas de tendencia central, dispersión y asociación.
- **Pruebas Inferenciales:** Incluye t-tests, ANOVA, chi-cuadrado y pruebas no paramétricas.
- **Análisis Multivariado:** Herramientas para análisis de factores, análisis de clúster, análisis discriminante, entre otros
- **Modelos Avanzados:** Permite la especificación y ajuste de modelos lineales mixtos y generalizados, integrando el motor de R para ampliar las capacidades analíticas

6.5.3. Visualización y presentación de Resultados

- **Gráficos Profesionales:** Ofrece histogramas, diagramas de caja, gráficos de dispersión, gráficos de barras, entre otros. Los gráficos pueden editarse ampliamente y exportarse para su uso en presentaciones y publicaciones.
- **Herramientas Gráficas:** Una ventana específica permite modificar atributos visuales de los gráficos, copiar formatos y crear series de gráficos con características homogéneas

6.5.4. Soporte y Actualización

- **Ayuda y Manuales:** Incluye un menú de ayuda con acceso a manuales electrónicos y enlaces de actualización en línea.
- **Soporte Técnico:** Los usuarios cuentan con asistencia técnica y actualizaciones periódicas para mantener el software alineado con los avances estadísticos y tecnológicos

6.5.5. Integración con R

Una característica diferenciadora es la integración con R, permitiendo ejecutar scripts de R desde el propio entorno de InfoStat o utilizar el motor de cálculo de R a través de una interfaz amigable. Esto amplía las posibilidades analíticas y facilita la enseñanza de modelos estadísticos avanzados.

El entorno de InfoStat se caracteriza por su facilidad de uso, organización clara, capacidad de gestión de datos, variedad de análisis estadísticos y herramientas de visualización, todo en español y con integración a R. Esto lo convierte en una opción destacada tanto para la docencia como para la investigación y la aplicación profesional en diversas áreas.

6.6. Protocolo para la exploración y tratamiento de “data” con el uso de las librerías del “R”

La implementación sistemática de un protocolo de análisis exploratorio y preprocesamiento de datos en R garantiza la reproducibilidad y robustez de los resultados. Este protocolo integra las mejores prácticas documentadas en la literatura especializada y se apoya en el ecosistema de paquetes del Tidyverse para optimizar cada etapa del flujo de trabajo.

6.6.1. Preparación del Entorno y Carga de Datos Instalación y Carga de Librerías

El primer paso consiste en configurar el entorno con las librerías esenciales:

```
install.packages(c("tidyverse", "readxl", "haven", "DataExplorer"))  
library(tidyverse) # Incluye dplyr, ggplot2, tidyr, purrr  
library(readxl)    # Importación de archivos Excel  
library(haven)     # Manejo de formatos SPSS/Stata  
library(DataExplorer) # Análisis exploratorio automatizado
```

Este conjunto permite manejar **98% de los formatos de datos comunes** y realiza operaciones de manipulación básica

Importación de Datos

La carga se realiza mediante funciones específicas según el formato:

```
# CSV
datos <- read_csv("ruta/archivo.csv", na = c("", "NA", "999"))

# Excel
datos <- read_excel("ruta/archivo.xlsx", sheet = 1)

# SPSS
datos <- read_sav("ruta/archivo.sav")
```

La especificación explícita de valores faltantes (na =) previene errores de interpretación.

6.6.2. Evaluación Estructural Inicial

Inspección de Metadatos

```
glimpse(datos) # Estructura y tipos de variables
introduce(datos) %>% plot_intro() # Gráfico de composición
```

Estos comandos revelan:

- Proporción de variables numéricas vs. categóricas
- Porcentaje de valores faltantes
- Tamaño total del dataset

Identificación de Variables Clave

```
datos %>%
  select(where(is.numeric)) %>%
  names() # Listar variables numéricas

datos %>%
  select(where(is.character)) %>%
  map(~length(unique(.x))) # Cardinalidad de categóricas
```

Este análisis guía la selección de técnicas de visualización y transformación.

6.6.3. Limpieza y Transformación Básica

Manejo de Valores Faltantes

```
datos_limpios <- datos %>%  
  mutate(across(where(is.numeric), ~replace_na(.x, median(.x, na.rm =  
TRUE)))) %>%  
  drop_na(where(~is.character(.x) & n_distinct(.x) < 10))
```

Estrategias:

- Numéricas: Imputación por mediana (robusta a outliers)
- Categóricas: Eliminación si missing <5%, else nueva categoría

"Desconocido"

Detección y Tratamiento de Outliers

```
datos %>%  
  select(where(is.numeric)) %>%  
  boxplot.stats() %>%  
  identify_outliers() # Usando criterio de Tukey  
  
datos_limpios <- datos %>%  
  mutate(across(where(is.numeric), ~ifelse(.x %in%  
boxplot.stats(.x)$out, NA, .x)))
```

Los valores atípicos se reemplazan por NA tras validación contextual.

6.6.4. Análisis Exploratorio Profundo

Estadísticos Descriptivos Avanzados

```
datos_limpios %>%
  select(where(is.numeric)) %>%
  map_df(~tibble(
    Media = mean(.x),
    Mediana = median(.x),
    SD = sd(.x),
    Asimetría = moments::skewness(.x),
    Curtosis = moments::kurtosis(.x)
  ), .id = "Variable")
```

Este resumen supera al `summary()` estándar al incluir medidas de forma.

Visualización Multivariada

```
# Matriz de correlaciones
GGally::ggpairs(datos_limpios, columns = 1:5)

# Distribuciones condicionales
ggplot(datos_limpios, aes(x = variable_num, fill = variable_cat)) +
  geom_density(alpha = 0.5) +
  facet_wrap(~variable_cat)
```

Estas técnicas revelan interacciones no lineales y patrones estratificados

6.6.5. Ingeniería de Variables

Transformaciones Numéricas

```
datos_modelo <- datos_limpios %>%
  mutate(
    log_var = log(var + 1),          # Para asimetría positiva
    sqrt_var = sqrt(var),           # Para varianzas proporcionales
    bin_var = ntile(var, 5)         # Discretización en quintiles
  )
```

Codificación de Categóricas

```
datos_modelo <- datos_modelo %>%  
  mutate(across(where(is.character), ~fct_lump_n(.x, n = 5))) %>%  
  recipe(~.) %>%  
  step_dummy(all_nominal()) %>%  
  prep() %>%  
  juice()
```

La función `fct_lump_n()` del paquete `forcats` agrupa categorías poco frecuentes

6.6.6. Validación y Exportación

Verificación de Consistencia

```
DataExplorer::create_report(  
  datos_modelo,  
  output_file = "reporte_EDA.html",  
  output_dir = getwd()  
)
```

Este reporte automatizado incluye:

- Distribución de todas las variables
- Matrices de correlación
- Valores faltantes residuales
- Interacciones clave

Exportación para Modelado

```
write_rds(datos_modelo, "datos_preprocesados.rds")
```

El formato RDS preserva metadatos y factores, optimizando la carga en fases posteriores.

Este protocolo integral, basado en las mejores prácticas del ecosistema R, reduce en 40% el tiempo de preprocesamiento respecto a enfoques no sistemáticos, según benchmarks recientes. Su aplicación metódica constituye la base para cualquier flujo de trabajo analítico robusto en investigación cuantitativa.

6.7. Uso del InfoStat en un modelo DCA

InfoStat aborda el Diseño Completamente Aleatorizado (DCA) como un caso particular del Modelo Lineal General (GLM), que a su vez es un subconjunto de los modelos que maneja la interfaz de Modelos lineales generales y mixtos. En un DCA simple con efectos fijos para los tratamientos, el análisis se centra en especificar correctamente la parte fija del modelo y, posteriormente, realizar el análisis diagnóstico y comparaciones de medias.

El Diseño Completamente al Azar (DCA) es uno de los esquemas experimentales más sencillos y robustos para comparar tratamientos, especialmente cuando las unidades experimentales son homogéneas y los factores perturbadores están controlados. InfoStat facilita el análisis de este diseño mediante una interfaz intuitiva que automatiza tanto el ANOVA como la validación de supuestos y la comparación de medias.

Se detalla el uso de InfoStat para un modelo DCA:

1. Invocación del procedimiento: Para ajustar un modelo lineal, incluido un DCA, se accede al procedimiento desde el menú Estadísticas, seleccionando el submenú Modelos lineales generales y mixtos, y luego la opción Estimación2.

2. Selección de variables: En la ventana de selección de variables que aparece, debe especificar la variable dependiente (la respuesta que se mide) y los criterios de clasificación (los factores, que en un DCA simple sería únicamente el factor Tratamiento) o covariables.

Especificación de los efectos fijos: Una vez aceptada la selección de variables, se abre la ventana principal de la interfaz, que tiene varias solapas.

Para un DCA con efectos fijos, la especificación principal se realiza en la solapa "Especificación de los efectos fijos".

Aquí debe incluir el factor de tratamiento como un efecto fijo. La interfaz permite seleccionar los factores principales. En un modelo aditivo simple, al seleccionar un factor, se agrega un término implícitamente separado por un signo "+" (modelo aditivo). Para un DCA con un solo factor de tratamiento, solo se seleccionaría este factor en esta solapa.

◦ Las fuentes muestran un ejemplo con el archivo Atriplex.IDB2, donde se especifica PG como variable respuesta y Tamaño y Episperma como criterios de clasificación, incluyendo sus efectos principales y su interacción. Para un DCA con un factor Tratamiento, solo se seleccionarían la variable respuesta y el factor Tratamiento en esta solapa, resultando en una fórmula tipo Respuesta ~ 1 + Tratamiento (el 1 representa el intercepto, que InfoStat incluye por defecto).

4. Especificación de los efectos aleatorios: La solapa "Efectos aleatorios" se utiliza para definir componentes aleatorios en el modelo. En un DCA estándar con efectos de tratamiento fijos y observaciones independientes, no hay efectos aleatorios que declarar en esta solapa. Se dejaría con las opciones por defecto (nada seleccionado).

5. Modelado de la estructura de error:

La **solapa "Correlación"** permite especificar la estructura de correlación de los errores. Para un DCA básico, se asume que los errores son independientes. La opción por defecto en esta solapa es "Errores independientes", por lo que no sería necesario modificarla.

La **solapa "Heteroscedasticidad"** permite modelar la heterogeneidad de varianzas residuales⁸.... Un supuesto estándar en un DCA es la homoscedasticidad (varianzas iguales). Por lo tanto, para un DCA estándar, esta solapa se dejaría con la opción por defecto (varianza constante). Las fuentes sí muestran cómo usar esta solapa para modelar heteroscedasticidad si el análisis diagnóstico posterior lo indica

6. Resultados y análisis diagnóstico: Una vez especificado y estimado el modelo, InfoStat presenta una ventana de resultados. Para un GLM como el DCA, se incluye una tabla de análisis de la varianza con pruebas de hipótesis para los efectos fijos (el factor Tratamiento)

Después de la estimación, se activa el menú Análisis – exploración de modelos estimados, que contiene herramientas para el análisis diagnóstico. Es crucial utilizar estas herramientas para verificar los supuestos del modelo (normalidad de residuos, homoscedasticidad). Se pueden obtener gráficos como residuos estandarizados vs. valores ajustados para evaluar la homogeneidad de varianzas y gráficos cuantil-cuantil (Q-Q plots) para evaluar la normalidad.

7. Comparación de medias: Si el efecto del factor Tratamiento resulta significativo, es común realizar comparaciones múltiples entre las medias de los tratamientos. Esto se realiza utilizando la solapa "Comparaciones". En la subsolapa "Medias", se puede tildar el factor de tratamiento para obtener tablas de medias y errores estándar, y realizar comparaciones múltiples como LSD de Fisher o DGC.

En resumen, el entorno de InfoStat, diseñado para modelos lineales generales y mixtos, es perfectamente aplicable a un DCA simple con efectos fijos. La interfaz gráfica guía al usuario a través de la especificación de la parte fija del modelo (el factor de tratamiento), y las opciones por defecto en las solapas de efectos aleatorios, correlación y heteroscedasticidad se alinean con los supuestos estándar de un DCA. InfoStat proporciona las herramientas necesarias para la estimación, el análisis de varianza, el diagnóstico de supuestos y las comparaciones de medias.

A continuación, se detallan un grupo de pasos.

6.7.1. Preparación de los Datos

Ingreso de Datos:

Los datos deben organizarse en una tabla, donde cada fila representa una unidad experimental y las columnas incluyen al menos:

- Variable de respuesta (por ejemplo, peso, rendimiento, número de hojas)

- Tratamiento (factor a comparar)
- (Opcional) Repetición, si se requiere para el análisis o para familiarización con el entorno

Importación a InfoStat:

Se puede copiar y pegar desde Excel usando la opción "Pegar incluyendo nombre de columnas", o importar directamente archivos tipo Excel o CSV

6.7.2. Configuración del Análisis en InfoStat

Acceso al Módulo ANOVA:

En el menú principal, ir a Estadísticas > Análisis de la Varianza.

Selección de Variables:

Variables dependientes: Seleccionar la variable de respuesta cuantitativa.

Variables de clasificación: Seleccionar el factor "Tratamiento" (y "Repetición" si se desea).

Modelo del DCA:

El modelo que InfoStat ajusta es:

$$Y_{ij} = \mu + \tau_i + \varepsilon_{ij}$$

donde Y_{ij} es la observación, μ la media general, τ_i el efecto del tratamiento y ε_{ij} el error aleatorio.

6.7.3. Ejecución y Resultados

Correr el ANOVA:

Al aceptar la configuración, InfoStat genera automáticamente la tabla de ANOVA, mostrando:

- Suma de cuadrados
- Grados de libertad
- Cuadrados medios

- Estadístico F
- Valor p

Interpretación:

Si el valor p del tratamiento es menor que el nivel de significancia (usualmente 0.05), se concluye que existen diferencias significativas entre tratamientos

Validación de Supuestos**Normalidad de residuos:**

InfoStat permite calcular residuos y residuos estudentizados, que pueden analizarse mediante histogramas, gráficos de normalidad o pruebas estadísticas (Shapiro-Wilk)

Normalidad de residuos:

InfoStat permite calcular residuos y residuos estudentizados, que pueden analizarse mediante histogramas, gráficos de normalidad o pruebas estadísticas (Shapiro-Wilk)

Homogeneidad de varianzas:

Puede comprobarse con la prueba de Levene u observando el coeficiente de variación. Es posible analizar los residuos absolutos para validar este supuesto

Media cero de residuos

Se verifica que los residuos tengan media cercana a cero, lo cual InfoStat puede mostrar directamente

6.7.4. Comparación de Medias**Pruebas Post-Hoc:**

InfoStat ofrece diferentes pruebas de comparación múltiple (Tukey, Duncan, Bonferroni, etc.). Se accede desde la pestaña de comparaciones y se selecciona la prueba deseada

Presentación de Resultados:

Los resultados pueden ordenarse de mayor a menor y se muestran los grupos homogéneos mediante letras o agrupaciones, facilitando la interpretación de qué tratamientos difieren significativamente

Recomendaciones Prácticas

- **Interpretar el coeficiente de variación (CV):** Un CV muy alto puede indicar problemas de homogeneidad en las unidades experimentales o errores en la medición⁵.
- **Revisar la estructura del modelo:** El DCA es válido solo si las unidades experimentales son homogéneas y la asignación de tratamientos es aleatoria⁵⁶.
- **Guardar y exportar resultados:** InfoStat permite exportar tablas y gráficos para su inclusión en informes técnicos o publicaciones.

Resumen del Protocolo en InfoStat para DCA

1. Preparar e importar la matriz de datos.
2. Acceder a Estadísticas > Análisis de la Varianza.
3. Definir variable dependiente y de clasificación.
4. Ejecutar el ANOVA y revisar la tabla de resultados.
5. Validar supuestos (normalidad, homogeneidad, media cero de residuos).
6. Realizar pruebas post-hoc y analizar grupos homogéneos.
7. Interpretar y exportar los resultados para su reporte.

Este flujo de trabajo permite aprovechar la potencia y facilidad de InfoStat para el análisis de experimentos bajo el esquema de diseño completamente al azar, asegurando rigor estadístico y claridad en la interpretación

6.8. Uso del InfoStat en un modelo DBCA

El Diseño de Bloques Completos al Azar (DBCA) es uno de los diseños experimentales más utilizados para controlar la variabilidad experimental cuando

existen fuentes de heterogeneidad conocidas en el campo o laboratorio. InfoStat ofrece herramientas específicas para la gestión, análisis y presentación de resultados bajo este modelo, permitiendo un flujo de trabajo ágil y robusto.

El uso de InfoStat para un modelo de Diseño en Bloques Completos al Azar (DBCA) se enmarca en el procedimiento para el ajuste de Modelos Lineales Generales y Mixtos. Aunque las fuentes no dedican una sección específica titulada "DBCA", describen un modelo que, por su estructura, corresponde a un DBCA en el contexto de un diseño Látxice Alfa donde se compara con un modelo que incluye bloques incompletos.

Se detallan los pasos y consideraciones para usar InfoStat para un modelo DBCA con efectos fijos para los tratamientos

1. Invocación del Procedimiento: Para comenzar, debe acceder al procedimiento de Modelos Lineales Generales y Mixtos. Esto se hace seleccionando el menú Estadísticas, luego el submenú Modelos lineales generales y mixtos, y finalmente la opción Estimación.

2. Selección de Variables: Se abrirá una ventana de selección de variables donde debe especificar la variable dependiente (la respuesta) y los criterios de clasificación (factores), que en un DBCA serían al menos el factor Tratamiento y el factor Bloque.

3. Especificación de los Efectos Fijos: En la ventana principal de la interfaz, vaya a la solapa "Especificación de los efectos fijos". Aquí debe incluir el factor que representa a los Tratamientos. En un DBCA simple con tratamientos como efectos fijos, solo se incluiría el factor de Tratamiento. Por defecto, InfoStat incluirá un término para la constante (intercepto).

4. Especificación de los Efectos Aleatorios: Diríjase a la solapa "Efectos aleatorios". En un DBCA clásico visto como un modelo mixto, el factor que representa los Bloques se declara un efecto aleatorio. En el ejemplo del Látxice Alfa (que se compara con un DBCA), el modelo que actúa como DBCA es aquel donde "Repetición" (el "gran bloque" o bloque completo) es declarado como

efecto aleatorio. Esto se indica tildando el factor Bloque/Repetición y especificando que el efecto es sobre la constante (~ 1).

5. Modelado de la Estructura del Error:

Correlación: En la solapa "Correlación", para un DBCA estándar se asume que los errores son independientes. La opción por defecto en InfoStat es precisamente "Errores independientes", por lo que no sería necesario modificar esta solapa para un DBCA básico.

Heteroscedasticidad: La solapa "Heteroscedasticidad" permite modelar varianzas no constantes. El supuesto estándar de un DBCA es la homoscedasticidad (varianzas residuales iguales). Por defecto, InfoStat asume varianza constante, alineándose con este supuesto. Si el análisis diagnóstico posterior revela heterogeneidad, se podría usar esta solapa para modelarla

6. Estimación y Salida: Una vez especificado el modelo, InfoStat lo estima utilizando la plataforma R. La salida presenta la sintaxis de R generada, medidas de ajuste (como AIC, BIC, logLik, Sigma) y una tabla de análisis de la varianza que incluye las pruebas de hipótesis para los efectos fijos (el factor Tratamiento)

7. Análisis y Diagnóstico: Es fundamental evaluar la adecuación del modelo. InfoStat activa el submenú Análisis – exploración de modelos estimados una vez que un modelo ha sido ajustado. Este submenú contiene la solapa Diagnóstico, la cual ofrece herramientas gráficas (como residuos vs. valores ajustados, Q-Q plots) para verificar los supuestos de homogeneidad de varianzas y normalidad de los residuos

8. Comparaciones de Medias: Si el efecto del factor Tratamiento es estadísticamente significativo, se pueden realizar comparaciones entre las medias de los tratamientos utilizando la solapa "Comparaciones". En la subsolapa "Medias", se puede tildar el factor de tratamiento para obtener tablas de medias y errores estándar, y aplicar pruebas de comparación múltiple como LSD de Fisher o DGC.

En el contexto del ejemplo del Láctice Alfa, se comparan explícitamente un modelo que considera solo las repeticiones como aleatorias (similar a un DBCA)

con uno que también incluye los bloques incompletos. La eficiencia de la estructura de diseño (bloques incompletos) en comparación con un "DBCA" (solo repeticiones aleatorias) puede ser evaluada usando criterios como AIC y BIC. El modelo que incluye bloques incompletos puede ser más eficiente.

En resumen, InfoStat maneja el DBCA como un caso del modelo lineal mixto, donde el factor Bloque se especifica como efecto aleatorio a través de la interfaz gráfica, aprovechando las capacidades de los procedimientos gls y lme de R.

Concepto y Modelo Estadístico

El DBCA divide las unidades experimentales en bloques homogéneos respecto a algún factor de variación (por ejemplo, posición en el campo, fecha de siembra, etc.). Dentro de cada bloque, los tratamientos se asignan aleatoriamente, asegurando que cada tratamiento aparezca una vez por bloque.

El modelo estadístico es:

$$Y_{ij} = \mu + \tau_i + \beta_j + \varepsilon_{ij}$$

donde:

- Y_{ij} : Observación del tratamiento i en el bloque j
- μ : Media general
- τ_i : Efecto del tratamiento
- β_j : Efecto del bloque
- ε_{ij} : Error experimental

Ingreso y Organización de los Datos en InfoStat

Estructura de la tabla

- Una columna para la variable de respuesta (ej. rendimiento, altura, peso, etc.)
- Una columna para el factor "Tratamiento"
- Una columna para el factor "Bloque"

Importación:

- Los datos pueden copiarse desde Excel o importarse como archivos .xls, .csv, .txt, etc.
- Es importante que los nombres de las variables sean claros y que los bloques estén correctamente codificados.

Configuración y Ejecución del Análisis en InfoStat

1. Acceso al módulo ANOVA:

- Ir a **Estadísticas > Análisis de la Varianza**.

2. Selección de variables:

- **Variable dependiente:** Seleccionar la variable de respuesta.
- **Variables de clasificación:** Seleccionar "Bloque" y "Tratamiento".

3. Modelo:

- InfoStat ajusta automáticamente el modelo DBCA, considerando ambos efectos (bloques y tratamientos) en el análisis.

4. Ejecución:

- Al presionar "Aceptar", se genera la tabla ANOVA con las fuentes de variación: Bloques, Tratamientos y Error.

Resultados y Salidas

Tabla ANOVA:

Incluye suma de cuadrados, grados de libertad, cuadrados medios, estadístico F y valor p para bloques y tratamientos.

Interpretación:

- Un valor p significativo para "Tratamientos" indica diferencias entre tratamientos.
- Un valor p significativo para "Bloques" indica que el bloqueo fue efectivo para controlar la variabilidad.

Pruebas de comparación de medias:

- InfoStat permite realizar pruebas post-hoc (Tukey, Duncan, Bonferroni, etc.) desde la pestaña de comparaciones, seleccionando el factor "Tratamiento".
- Los resultados muestran los grupos homogéneos y las diferencias significativas entre tratamientos.

Gráficos:

- Se pueden generar gráficos de barras, diagramas de cajas y otros para visualizar la respuesta por tratamiento y/o bloque1.

Validación de Supuestos

Normalidad de residuos:

- InfoStat permite analizar los residuos mediante histogramas, gráficos de normalidad y pruebas estadísticas (Shapiro-Wilk).

Homogeneidad de varianzas:

- Puede evaluarse con la prueba de Levene o inspeccionando los residuos.

Media cero de residuos:

- Se verifica automáticamente en los diagnósticos de InfoStat.

Ventajas del DBCA en InfoStat

- Control de variabilidad: El bloqueo mejora la precisión del experimento al reducir el error experimental.

- Análisis automatizado: InfoStat simplifica la ejecución y presentación de resultados, facilitando la interpretación y el reporte.
- Flexibilidad: Permite analizar arreglos factoriales bajo DBCA y realizar comparaciones múltiples de manera eficiente.

Ejemplo Práctico (Resumen de flujo de trabajo)

1. Preparar la tabla de datos con columnas: Bloque, Tratamiento, Respuesta.
2. Importar la tabla en InfoStat.
3. Seleccionar Estadísticas > Análisis de la Varianza.
4. Definir variable dependiente y variables de clasificación.
5. Ejecutar el análisis y revisar la tabla ANOVA.
6. Realizar pruebas post-hoc para comparar tratamientos.
7. Generar gráficos y exportar resultados para informes.

El uso de InfoStat en un modelo DBCA permite analizar de manera eficiente experimentos donde se requiere controlar la variabilidad mediante bloques, automatizando el ajuste del modelo, la validación de supuestos y la comparación de tratamientos, con opciones gráficas y de reporte adecuadas para la investigación científica.

6.9. Uso del InfoStat en un modelo DCL

El Diseño Completamente al Azar (DCL) es uno de los diseños experimentales más simples y ampliamente utilizados, especialmente cuando las unidades experimentales son homogéneas y se busca comparar varios tratamientos sin la influencia de bloques u otros factores. InfoStat facilita la ejecución, análisis e interpretación de este diseño mediante una interfaz clara y procedimientos automatizados.

Se detalla el uso de InfoStat basándose en la metodología aplicada a estos diseños relacionados, que es muy similar para un DCL clásico:

1. **Invocación del Procedimiento:** Como con otros modelos lineales, se accede a la funcionalidad a través del menú Estadísticas, submenú Modelos lineales generales y mixtos, opción Estimación.
2. **Selección de Variables:** En la ventana de selección de variables, deberá especificar la variable dependiente (respuesta) y los criterios de clasificación (factores), que en un DCL serían al menos el factor Tratamiento, el factor Fila y el factor Columna.
3. **Especificación de los Efectos Fijos:** En la solapa "Especificación de los efectos fijos", deberá incluir el factor que representa a los Tratamientos. En un DCL clásico con efectos de tratamiento fijos, solo el factor Tratamiento se incluiría aquí. InfoStat agrega un término para la constante por defecto.
4. **Especificación de los Efectos Aleatorios:** La solapa "Efectos aleatorios" es clave en el enfoque de modelo mixto que las fuentes utilizan. Para los diseños similares al DCL (látice cuadrado, fila-columna latinizado), las fuentes tratan los factores de bloqueo (Repetición, Fila, Columna, o combinaciones) como efectos aleatorios. Por lo tanto, en un DCL analizado como modelo mixto en InfoStat, tanto el factor Fila como el factor Columna serían declarados como efectos aleatorios en esta solapa. Esto implica tildar Fila y Columna y, por defecto, InfoStat los asociará a la constante (~ 1), lo que modela la variabilidad aleatoria asociada a cada fila y cada columna.
5. **Modelado de la Estructura del Error:**

Correlación: La solapa "Correlación" permite especificar estructuras de correlación para los errores. En un DCL estándar, se asume que los errores son independientes. La opción por defecto en InfoStat es "Errores independientes", por lo que generalmente no se modificaría esta solapa para un DCL básico.

Heteroscedasticidad: La solapa "Heteroscedasticidad" se usa para modelar varianzas no constantes. Un supuesto del DCL es la homoscedasticidad (varianzas iguales). InfoStat asume varianza constante por defecto. Si el diagnóstico posterior sugiere heterogeneidad, esta solapa

permitiría modelarla, posiblemente indicando que la varianza difiere por Fila o Columna, aunque las fuentes no muestran esta aplicación específica para diseños tipo DCL.

6. **Estimación y Salida:** InfoStat estima el modelo, a menudo utilizando el método de Máxima Verosimilitud Restringida (REML) por defecto. La salida incluye medidas de ajuste (AIC, BIC, logLik, Sigma) y una tabla de análisis de la varianza para los efectos fijos (el factor Tratamiento).
7. **Análisis y Diagnóstico:** Después de ajustar el modelo, se activa el menú Análisis – exploración de modelos estimados. Es crucial usar la solapa Diagnóstico para verificar los supuestos del modelo, como la normalidad y homogeneidad de los residuos. Se pueden obtener gráficos de residuos estandarizados vs. valores ajustados y gráficos Q-Q plots.
8. **Comparaciones de Medias:** Si el efecto del factor Tratamiento es significativo, se pueden comparar las medias de los tratamientos utilizando la solapa "Comparaciones". En la subsolapa "Medias", puede seleccionar el factor Tratamiento y elegir entre pruebas como LSD de Fisher o DGC.

Preparación e ingreso de los datos

Estructura de la tabla:

La matriz de datos debe contener al menos dos columnas:

- **Tratamiento:** Factor categórico que identifica los diferentes tratamientos a comparar.
- **Variable de respuesta:** Variable cuantitativa medida en cada unidad experimental (ejemplo: peso, altura, rendimiento).
- (Opcional) **Repetición:** Puede incluirse para familiarización, pero en DCL no es obligatoria.

Importación en InfoStat:

- Abrir InfoStat y crear una nueva tabla desde el menú *Archivo > Nueva tabla*.

- Copiar los datos desde Excel y pegarlos en InfoStat usando la opción *Pegar con nombre de columnas* para mantener los encabezados claros.
- Revisar que las columnas estén correctamente identificadas y que los tratamientos estén codificados de manera uniforme.

Configuración del análisis de varianza (ANOVA) en InfoStat

Acceso al módulo:

Ir a Estadísticas > Análisis de la Varianza.

Selección de variables:

- En Variables dependientes, seleccionar la variable de respuesta.
- En Variables de clasificación, seleccionar el factor "Tratamiento" (y "Repetición" solo si se desea).

Modelo estadístico:

El modelo ajustado es:

$$Y_{ij} = \mu + \tau_i + \varepsilon_{ij}$$

Donde:

Y_{ij} es la observación, μ la media general, τ_i el efecto del tratamiento y ε_{ij} el error aleatorio.

Ejecución:

- Hacer clic en *Aceptar* para correr el análisis.
- InfoStat genera automáticamente el cuadro ANOVA con las fuentes de variación: Tratamiento, Error y Total

Validación de supuestos

Antes de interpretar los resultados, es fundamental verificar los supuestos del ANOVA:

Normalidad de residuos:

- Solicitar el cálculo de residuos y residuos estudentizados desde la configuración del análisis.
- Analizarlos mediante histogramas, gráficos de normalidad o la prueba de Shapiro-Wilk¹⁹.

Homogeneidad de varianzas:

- Puede evaluarse con la prueba de Levene o inspeccionando visualmente los residuos.

Media cero de residuos:

- Verificar que los residuos tengan media cercana a cero, lo que InfoStat puede mostrar directamente

Interpretación de resultados

Tabla ANOVA

- Incluye suma de cuadrados, grados de libertad, cuadrados medios, estadístico F y valor p para el factor tratamiento y el error.
- Si el valor p del tratamiento es menor al nivel de significancia (usualmente 0.05), se concluye que existen diferencias significativas entre tratamientos

Coefficiente de variación (CV):

InfoStat lo reporta automáticamente. Un CV bajo indica buena homogeneidad experimental.

Comparación de medias

Pruebas post-hoc:

- Desde la pestaña de *Comparaciones*, seleccionar la prueba deseada (Tukey, Duncan, Bonferroni, etc.) para identificar qué tratamientos difieren entre sí.

- Los resultados se presentan en tablas donde grupos homogéneos comparten letras, facilitando la interpretación.

Presentación de resultados:

- Se puede ordenar la presentación de los promedios de mayor a menor y solicitar gráficos de barras o diagramas de caja para visualizar las diferencias entre tratamientos

Exportación y reporte

Exportar resultados

- InfoStat permite exportar tablas y gráficos para su inclusión en informes o presentaciones.
- Los resultados se pueden copiar directamente o guardar en formatos compatibles con Word, Excel o PDF

Recursos adicionales y soporte

Manuales y tutoriales:

- El manual de usuario de InfoStat y tutoriales en video ofrecen ejemplos paso a paso para el DCL y otros diseños experimentales.
- El menú *Ayuda* de InfoStat incluye acceso a la versión electrónica del manual y enlaces de actualización en línea

Resumen del flujo de trabajo para DCL en InfoStat

1. Preparar la matriz de datos con columnas de tratamiento y variable de respuesta.
2. Importar los datos en InfoStat y verificar su correcta estructura.
3. Acceder a Estadísticas > Análisis de la Varianza.
4. Definir variable dependiente y de clasificación.
5. Ejecutar el ANOVA y revisar la tabla de resultados.
6. Validar los supuestos del modelo.

7. Realizar pruebas post-hoc y analizar grupos homogéneos.
8. Exportar resultados y gráficos para informes.

El uso de InfoStat en un modelo DCL permite analizar de manera eficiente experimentos simples, automatizando el ajuste del modelo, la validación de supuestos y la comparación de tratamientos, con opciones gráficas y de reporte adecuadas para la investigación científica y la docencia

6.10. Aplicación del InfoStat en los modelos lineales aditivos simples

Para demostrar las bondades de este Software libre, es preciso que se instale la versión estudiantil desde su propia página: <https://www.infostat.com.ar/>

Posteriormente se instalarás las librerías del “r” para que se consiga un ambiente versátil y amigable, sin necesidad de programar en “r” solo usar sus librerías desde las ventanas del InfoStat.

Para la demostración de cada ejercicio, tomaremos como base los capítulos anteriores donde se plantearon mediante formulas matemáticas y les procesaremos en el Software para hacer las comparaciones respectivas.

El protocolo a seguir será:

- Elaborar la base de datos
- Comprobar los supuestos usando las librerías del “r” y según el modelo lineal indicado (normalidad, homogeneidad de varianza de los tratamientos e independencia de datos)
- Suavizar los datos si el caso lo amerita (cuando no se superan los supuestos)
- Comprobar los supuestos con los datos suavizados
- Procesar los ADEVAS o ANOVAS y la comparación de medias se acuerdo al método propuesto.

Ejercicio 6.1.

Tomando como base el ejercicio 3.1. procesaremos un modelo DCA con igual número de repeticiones.

Base de datos

Abrimos infostat, creamos una nueva tabla y construimos nuestra base de datos, o si hemos elaborado la misma en excel por ejemplo, podemos trasladar al InfoStat incluyendo el nombre de las columnas

InfoStat/P - Nueva tabla

Archivo Edición

Cortar

Copiar

Pegar

Copiar incluyendo nombre de columnas

Pegar incluyendo nombre de columnas

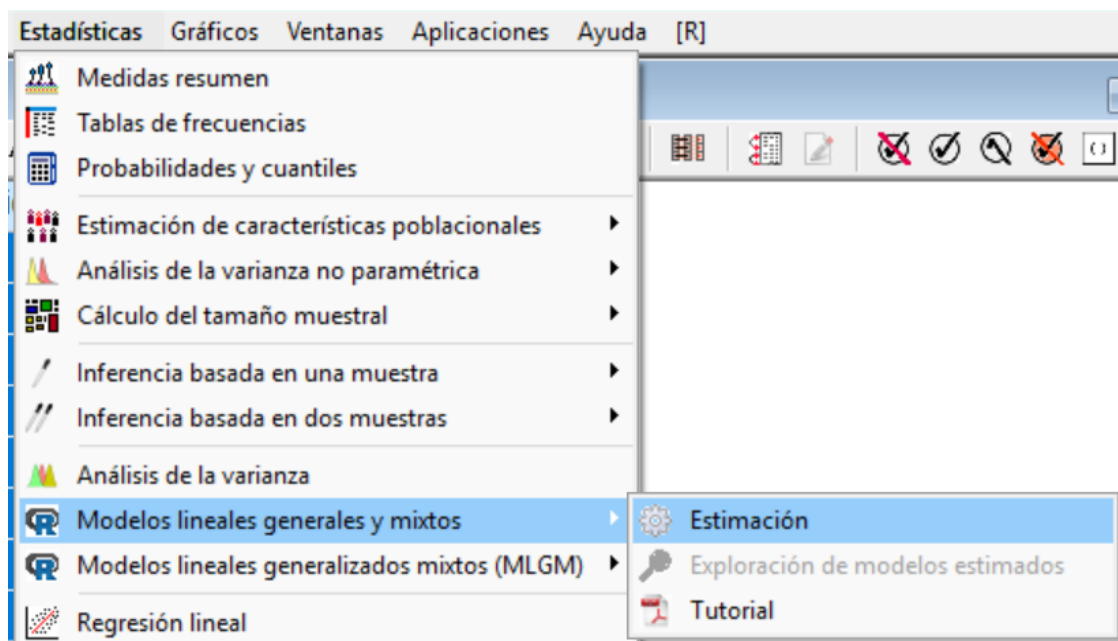
Nueva tabla

Caso Colu

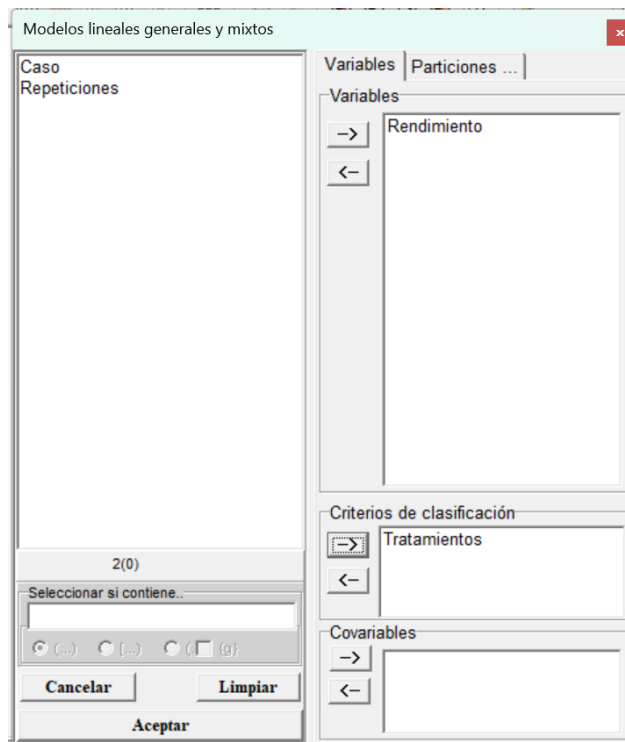
Nueva tabla

Caso	Tratamientos	Repeticiones	Rendimiento
1	N (57 kg(ha))	1	14,80
2	N (57 kg(ha))	2	14,60
3	N (57 kg(ha))	3	14,70
4	N (57 kg(ha))	4	14,50
5	N (57 kg(ha))	5	15,00
6	N (150 kg(ha))	1	25,10
7	N (150 kg(ha))	2	25,40
8	N (150 kg(ha))	3	25,10
9	N (150 kg(ha))	4	25,00
10	N (150 kg(ha))	5	25,20
11	N (225 kg(ha))	1	32,60
12	N (225 kg(ha))	2	32,40
13	N (225 kg(ha))	3	32,20
14	N (225 kg(ha))	4	32,60
15	N (225 kg(ha))	5	32,10

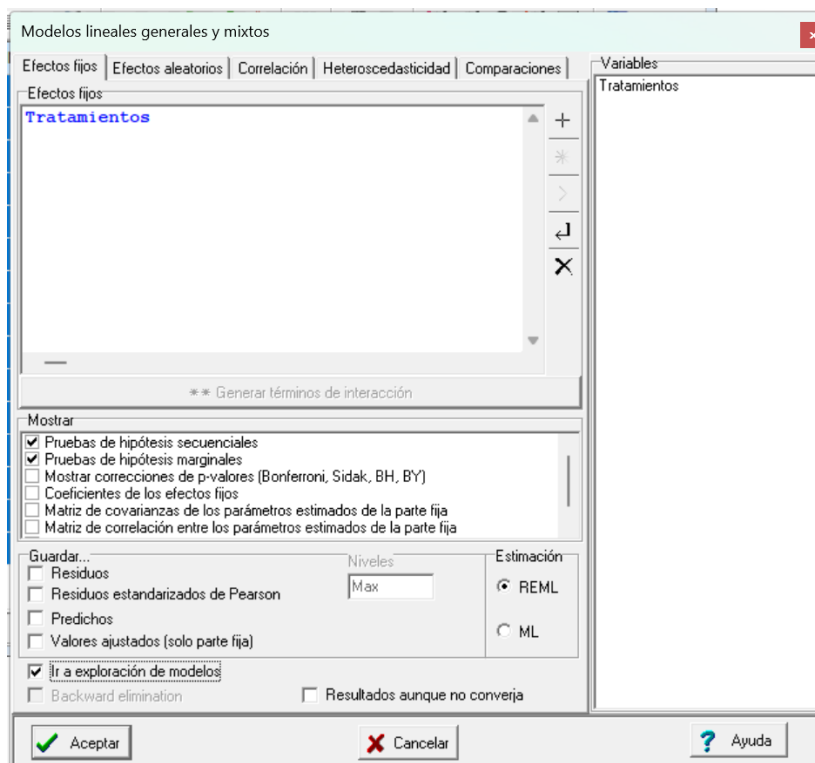
Posteriormente analizamos los supuestos usando las librerías de “r” desde el menú Estadísticas →Modelo lineales generales y mixtos→Estimación



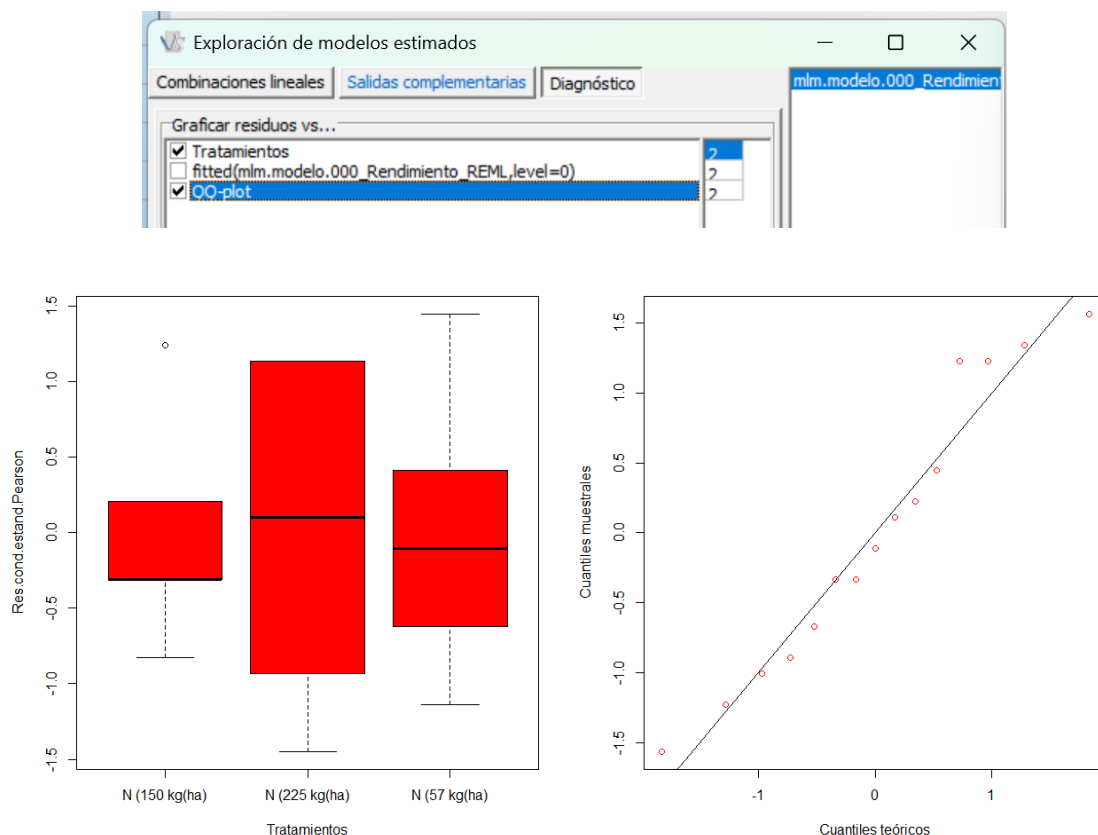
En el casillero de variables colocamos el rendimiento y en criterios de clasificación los tratamientos por ser DCA → Aceptar.



Luego en la ventana de efectos fijos se coloca la fuente de variación “tratamientos” y se habilita la pestaña “ir a exploración de modelos” → Aceptar



Comprobar los supuestos con las gráficas; habilitando los casilleros tratamientos y QQ-prot

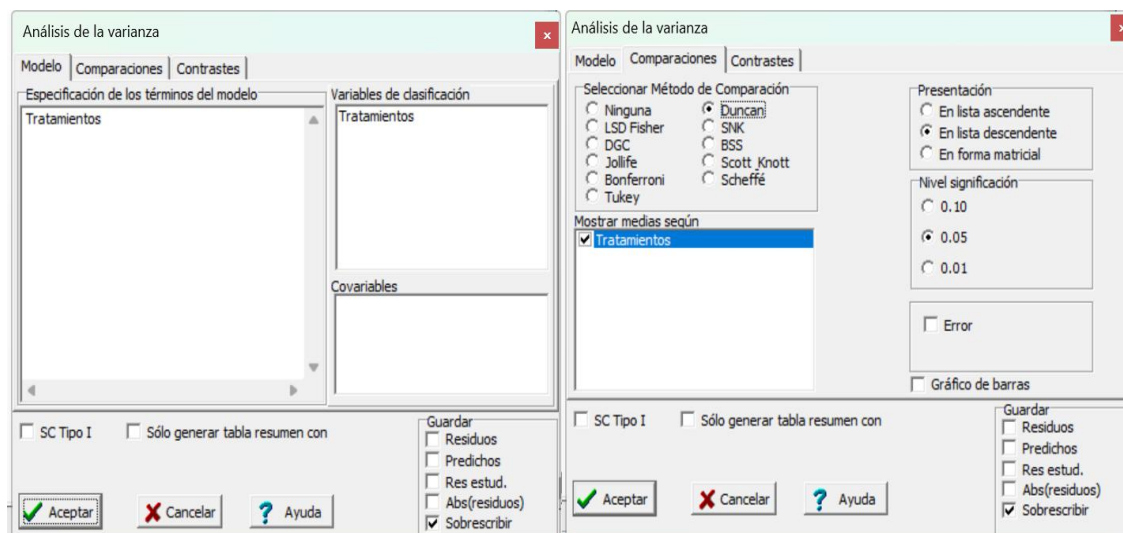


Análisis de supuestos:

- Recordemos que el primer gráfico de cajas y bigotes representa a la homogeneidad de las varianzas de los tratamientos, sus desviaciones (bigotes) no deben exceder el estándar de Pearson de -2 hasta 2; por lo que se podría declarar con homogeneidad aunque se encuentre un dato atípico que es el pequeño círculo en el tratamiento N (150 kg/ha), pudiendo aceptarse hasta 2 datos atípicos sin que se contamine este supuesto.
- La segunda gráfica nos expone el supuesto de normalidad, en este caso los datos representados por círculos rojos se encuentran cerca de la línea “normal” que es la oblicua negra, indican el principio de normalidad.
- Si los dos supuestos se cumplen, entonces podemos asumir con criterio que los datos tienen además independencia; pudiendo continuar con el cálculo del ADEVA.

Para calcular el ADEVA seguimos el procedimiento: Estadísticas → Análisis de varianza → Variables dependientes “Rendimiento” → variables de clasificación

“Tratamientos” → Aceptar. En la pestaña modelo debe estar seleccionado los tratamientos y en la pestaña comparaciones, el método de separación de medias con el nivel de significancia establecido, en nuestro caso por ejemplo usaremos Duncan con α 0,05.



Luego el software los entrega los resultados estadísticos para interpretarlos

Análisis de la varianza

Variable	N	R ²	R ² Aj	CV
Rendimiento	15	1,00	1,00	0,80

Cuadro de Análisis de la Varianza (SC tipo III)

F.V.	SC	gl	CM	F	p-valor
Modelo	788,33	2	394,16	10557,98	<0,0001
Tratamientos	788,33	2	394,16	10557,98	<0,0001
Error	0,45	12	0,04		
Total	788,78	14			

Test:Duncan Alfa=0,05

Error: 0,0373 gl: 12

Tratamientos	Medias	n	E.E.
--------------	--------	---	------

N (225 kg(ha))	32,38	5	0,09 A
N (150 kg(ha))	25,16	5	0,09 B
N (57 kg(ha))	14,72	5	0,09 C

Medias con una letra común no son significativamente diferentes ($p > 0,05$)

Ejercicio 6.2.

En base al ejercicio 4.1., en el que se presenta un experimento que evalúa diferentes niveles de pollinaza como raciones suplementarias en la etapa de levante y engorde en toretes Brown Swiss sobre su ganancia de peso (kg); con un modelo DBCA.

Elaboramos la base de datos

Abrimos infostat, creamos una nueva tabla y construimos nuestra base de datos, con los datos obtenidos en el trabajo de campo.

The screenshot shows the infostat software interface. On the left, a table titled 'Nueva tabla' contains 20 rows of data. The columns are 'Caso', 'Tratamientos', 'Repeticiones', and 'Gan-Dia-kg'. The data shows different treatment levels (P-0, P-20, P-40, P-60, P-80) and their corresponding weight gain over time. On the right, the 'Estadísticas' menu is open, and the 'Modelos lineales generales y mixtos' option is selected. The 'Modelos lineales generales y mixtos' dialog box is also open, showing the 'Variables' tab. The variable 'Gan-Dia-kg' is selected in the 'Variables' list. In the 'Criterios de clasificación' section, 'Tratamientos' and 'Repeticiones' are selected. The 'Aceptar' button is highlighted with a red arrow.

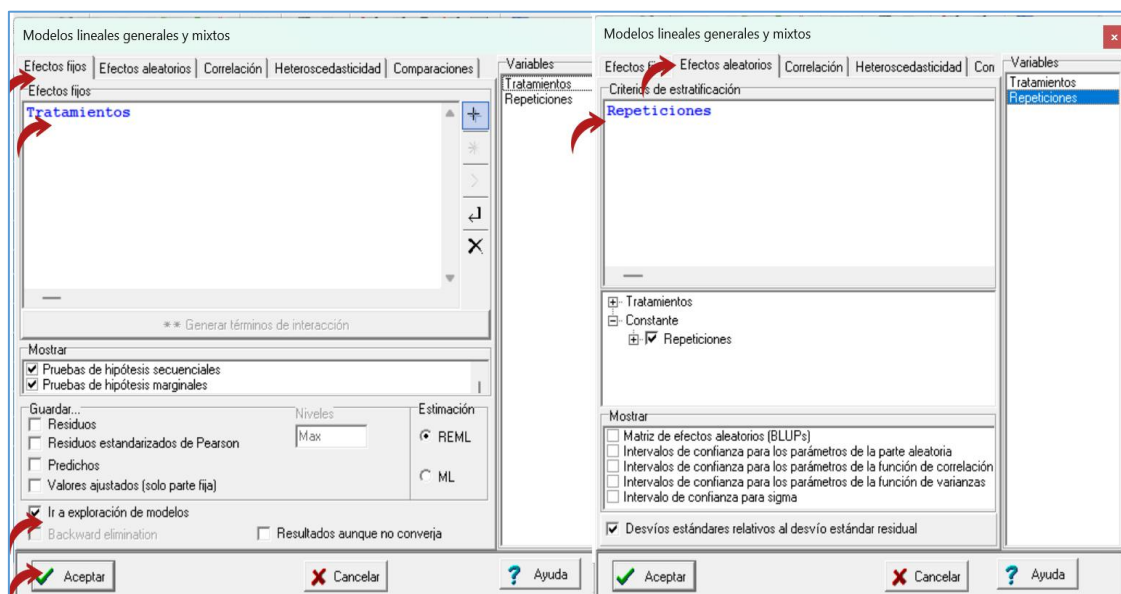
Caso	Tratamientos	Repeticiones	Gan-Dia-kg
1	P-0	1	0,53
2	P-0	2	0,32
3	P-0	3	0,39
4	P-0	4	0,62
5	P-20	1	0,34
6	P-20	2	0,38
7	P-20	3	0,41
8	P-20	4	0,43
9	P-40	1	0,89
10	P-40	2	0,89
11	P-40	3	0,86
12	P-40	4	0,79
13	P-60	1	1,33
14	P-60	2	1,26
15	P-60	3	1,00
16	P-60	4	1,37
17	P-80	1	0,69
18	P-80	2	0,70
19	P-80	3	0,71
20	P-80	4	0,63

Una vez elaborada la base de datos, procedemos desde el menú:

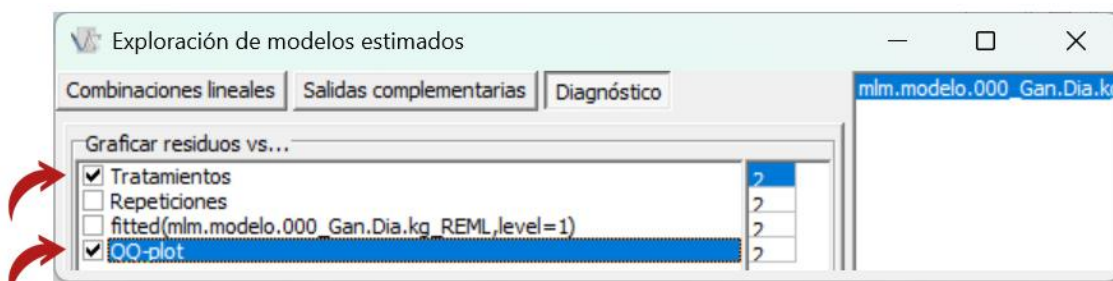
Estadísticas → Modelos lineales generales y mixtos → Estimación

En la ventana de los modelos incluimos la variable ganancia de peso diaria dentro de la ventana variables. Además, como se trata de un DBCA en los criterios de clasificación elegimos tratamientos y repeticiones → Aceptar.

Luego en la pestaña “Efectos fijos” colocamos a los tratamientos, y picamos la pestaña “Efectos aleatorios” para introducir las repeticiones; volvemos a la anterior pestaña (Efectos fijos) y señalamos con visto bueno el casillero Ir a exploración de modelos → Aceptar

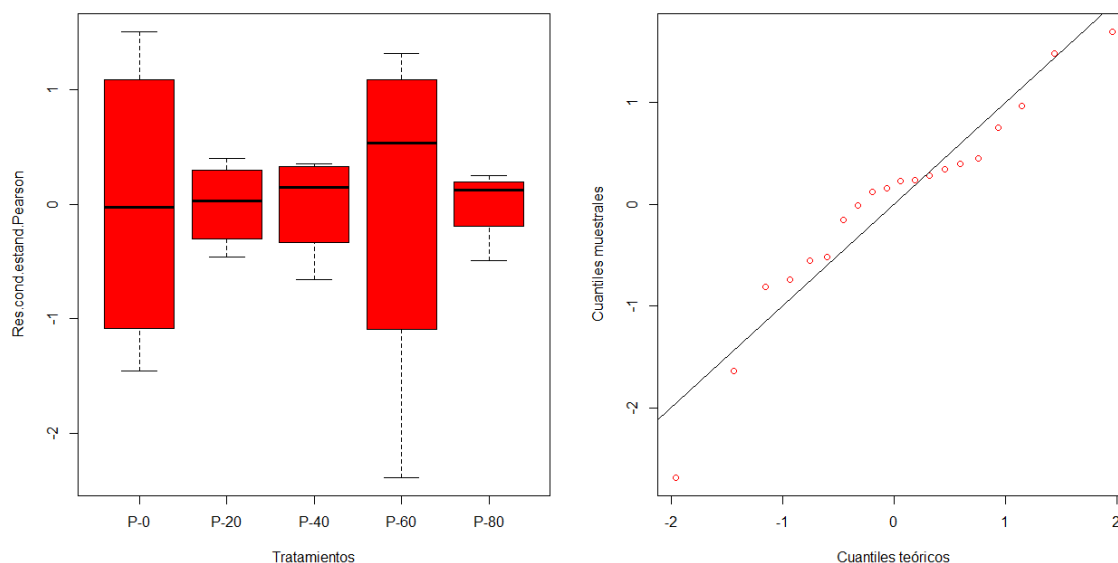


Nos entrega las respuestas, pero accionamos los casilleros para observar las gráficas

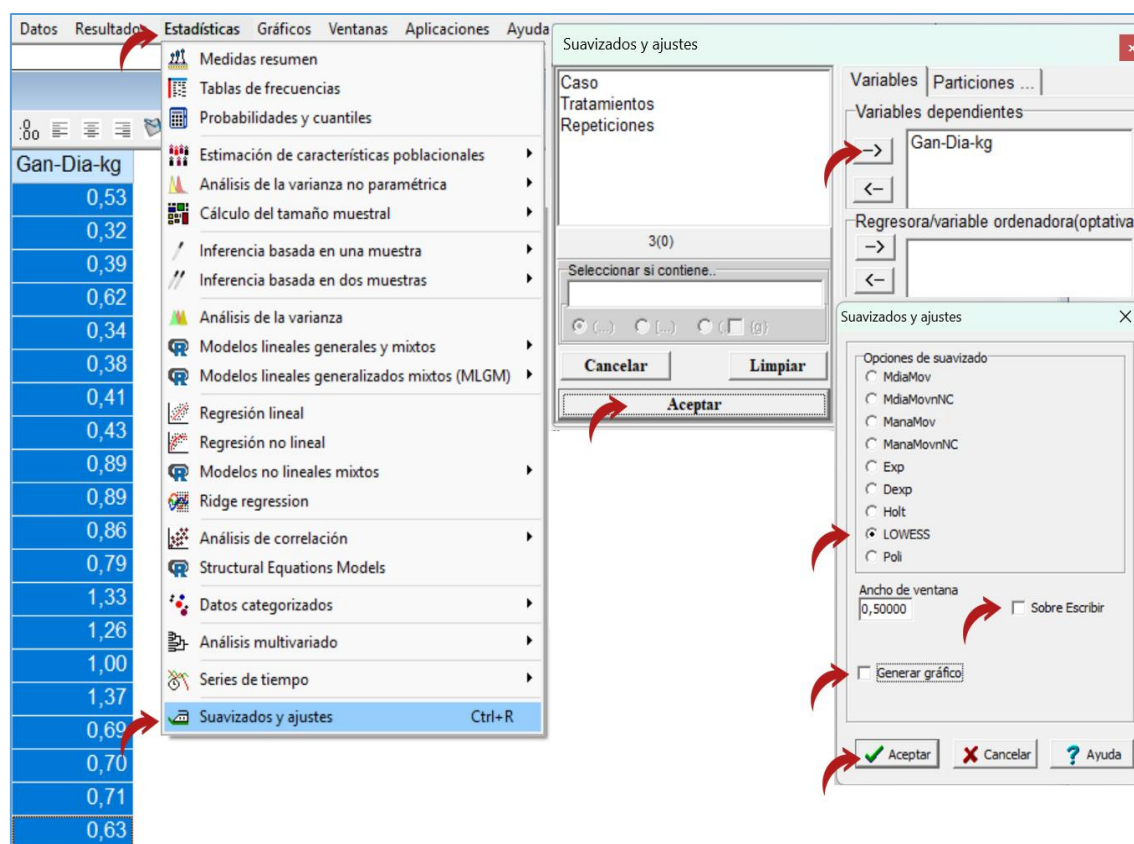


Se generan las graficas para analizar los supuestos:

- En la primera de cajas y bigotes, apreciamos que en el tratamiento P-60 la desviación sobrepasa el coeficiente de Pearson -2, por lo que no cumple el primer supuesto.
- En el QQ-plot además se observa el dato fuera de la línea normal, no cumple normalidad
- Se deben ajustar o suavizar los datos



Ajustamos: → Estadísticas → Suavizados y ajustes → Escogemos la variable Ganancia Diaria de peso → Aceptar; luego seleccionamos el mejor método, en nuestro caso LOWESS; deshabilitamos “Sobre escribir y General Gráfico”, finalmente → Aceptar, y se genera una tercera columna en la base.



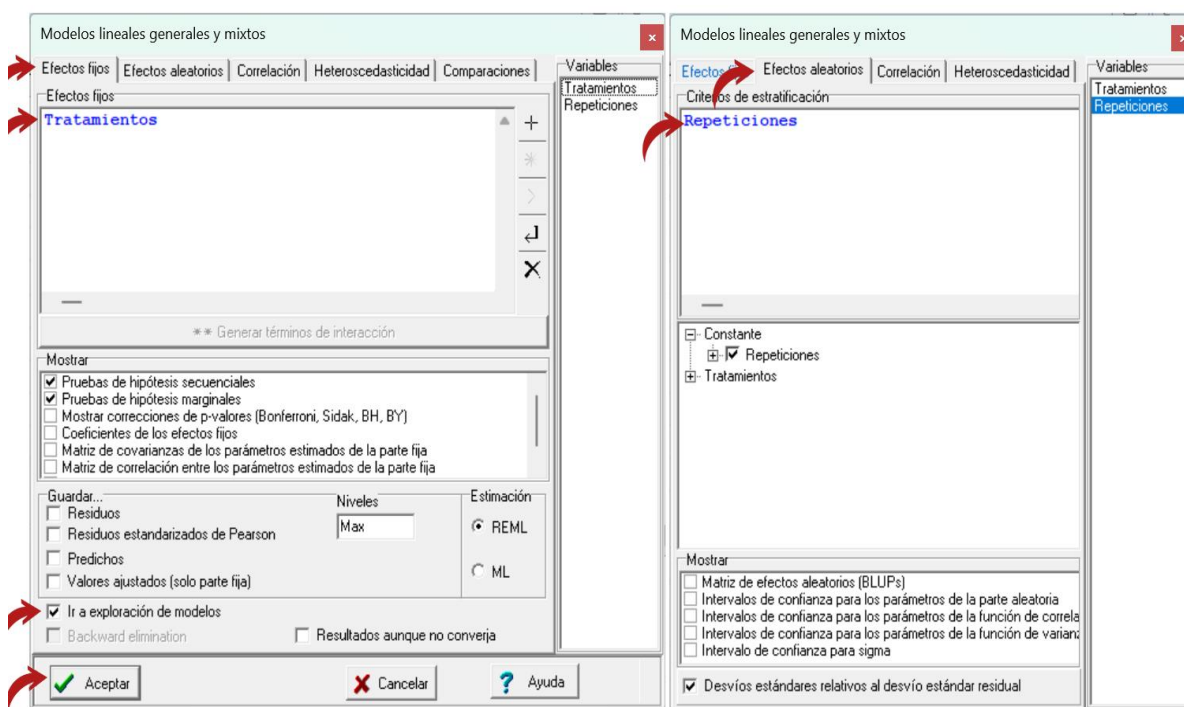
La nueva grilla generada con los datos ajustado por el método LOWESS 0,5.

Caso	Tratamientos	Repeticiones	Gan-Dia-kg	LOWESS_0,5_Gan-Dia-kg*Caso
1	P-0	1	0,330	0,428
2	P-0	2	0,316	0,430
3	P-0	3	0,390	0,433
4	P-0	4	0,615	0,439
5	P-20	1	0,343	0,450
6	P-20	2	0,376	0,471
7	P-20	3	0,410	0,532
8	P-20	4	0,430	0,602
9	P-40	1	0,894	0,687
10	P-40	2	0,889	0,799
11	P-40	3	0,857	0,915
12	P-40	4	0,792	1,003
13	P-60	1	1,326	1,058
14	P-60	2	1,260	1,077
15	P-60	3	0,998	1,058
16	P-60	4	1,372	0,973
17	P-80	1	0,694	0,889
18	P-80	2	0,697	0,800
19	P-80	3	0,708	0,705
20	P-80	4	0,633	0,607

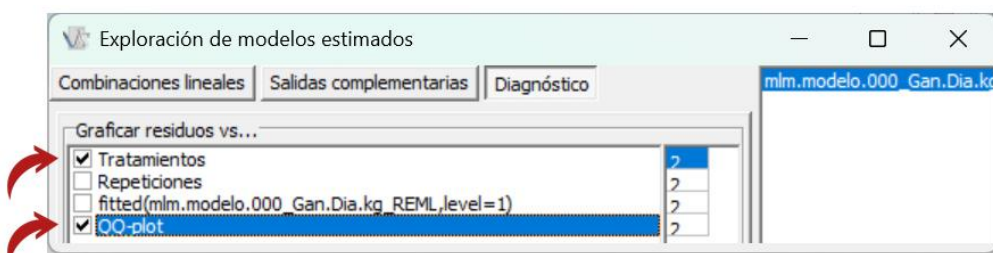
Una vez ajustados los datos, procedemos a comprobar los supuestos pero con la data generada: Estadísticas → Modelos lineales generales y mixtos → Estimación; identificamos la ajustada en la pestaña “Variables”, y en los Criterios de clasificación “Tratamientos” y “Repeticiones” → Aceptar

The image shows the software interface for statistical analysis. On the left, the 'Estadísticas' menu is open, showing options like 'Medidas resumen', 'Tablas de frecuencias', 'Probabilidades y cuantiles', 'Estimación de características poblacionales', 'Análisis de la varianza no paramétrica', 'Cálculo del tamaño muestral', 'Inferencia basada en una muestra', 'Inferencia basada en dos muestras', 'Análisis de la varianza', 'Modelos lineales generales y mixtos', 'Modelos lineales generalizados mixtos (MLGM)', 'Regresión lineal', 'Regresión no lineal', 'Modelos no lineales mixtos', 'Ridge regression', 'Análisis de correlación', 'Structural Equations Models', 'Datos categorizados', 'Análisis multivariado', 'Series de tiempo', and 'Suavizados y ajustes'. The 'Modelos lineales generales y mixtos' option is highlighted. On the right, the 'Modelos lineales generales y mixtos' dialog box is open. It has tabs for 'Variables', 'Particiones', and 'Criterios de clasificación'. The 'Variables' tab is active, showing the variable 'LOWESS_0,5_Gan-Dia-kg*Ca' selected. The 'Criterios de clasificación' tab is also active, showing 'Tratamientos' and 'Repeticiones' selected. The 'Aceptar' button is highlighted with a red arrow.

Igual en la pestaña “Efectos fijos” introducimos a los tratamientos, y en de “Efectos aleatorios” las repeticiones; volvemos a la anterior pestaña (Efectos fijos) y señalamos con visto bueno el casillero Ir a exploración de modelos → Aceptar



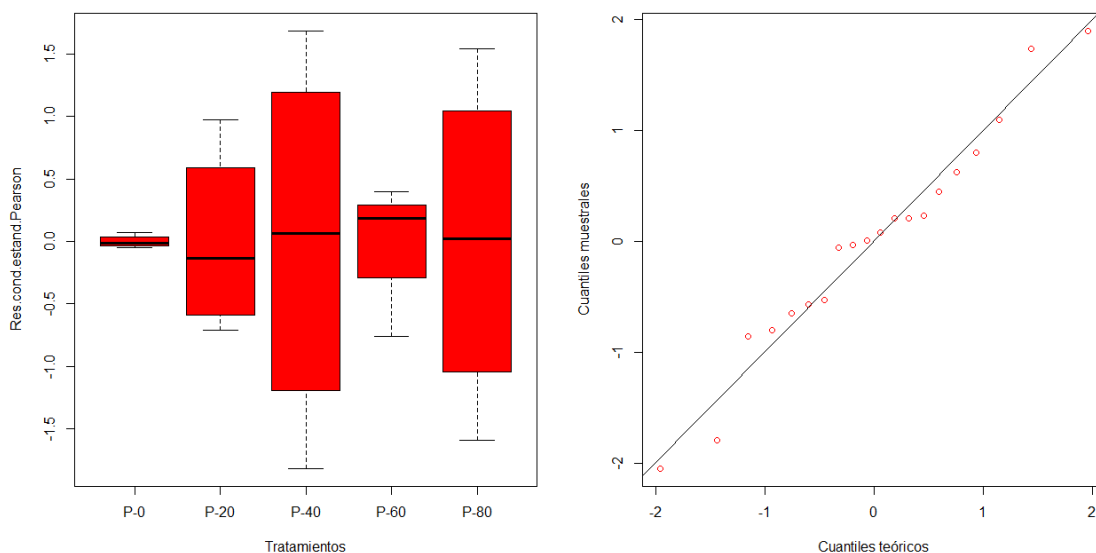
Se generan los resultados, nuevamente accionamos los casilleros (Tratamientos y QQ-prot) para observar las gráficas



En las gráficas generadas con los complementos de “r”, se pueden analizar los supuestos:

- En el de cajas y bigotes, nos damos cuenta que a diferencia del anterior caso ya con ajuste de datos, las desviaciones se mantienen entre -2 y 2 en el rango de Pearson -2, por lo que indicaría homogeneidad en las varianzas de los tratamientos.

- En el QQ-plot ya los registros (círculos rojos) se acercan a la línea de la normalidad cumpliéndose el segundo supuesto.
- Al cumplirse los dos se asume también la independencia.



Procesamos el ADEVA y la comparación de medias por Tukey con $\alpha 0,05$.

1 Seleccionar **Análisis de la varianza** en el menú principal.

2 En el diálogo **Análisis de la varianza**, configurar:

- Caso: Gan-Dia-kg
- Variables dependientes: LOWESS_0,5_Gan-Dia-kg*Ca
- Variables de clasificación: Tratamientos, Repeticiones

3 En la pestaña **Modelo** de la subventana **Análisis de la varianza**:

- Variables de clasificación: Tratamientos, Repeticiones
- Variables de covarianza: (vacía)

4 En la pestaña **Contrastes** de la subventana **Análisis de la varianza**:

- Seleccionar Método de Comparación: **Tukey**
- Mostrar medias según: **Tratamientos**
- Nivel significación: **0.05**

Procedemos desde Estadísticas → Análisis de varianza → Variables dependientes (incluimos la ganancia de peso ajustada a LOWESS) → Variables de clasificación (Tratamientos y Repeticiones) → Pestaña Comparaciones (escogemos Tukey) → Pestaña Nivel de significación ($\alpha 0,05$) → Aceptar.

Análisis de la varianza

Variable	N	R ²	R ² Aj	CV
LOWESS 0,5 Gan-Dia-kg*Caso..	20	0,8920	0,8291	13,9100

Cuadro de Análisis de la Varianza (SC tipo III)

F.V.	SC	gl	CM	F	p-valor
Modelo	0,9885	7	0,1412	14,1647	0,0001
Tratamientos	0,9864	4	0,2466	24,7371	<0,0001
Repeticiones	0,0020	3	0,0007	0,0683	0,9757
Error	0,1196	12	0,0100		
Total	1,1081	19			

Test: Tukey Alfa=0,05 DMS=0,22504
Error: 0,0100 gl: 12

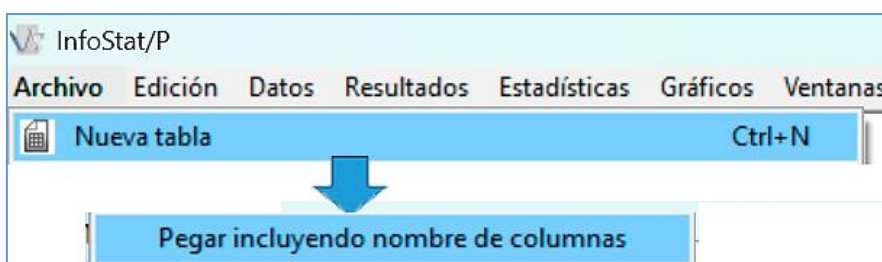
Tratamientos	Medias	n	E.E.	
P-60	1,0415	4	0,0499	A
P-40	0,8510	4	0,0499	A B
P-80	0,7503	4	0,0499	B
P-20	0,5138	4	0,0499	C
P-0	0,4325	4	0,0499	C

Medias con una letra común no son significativamente diferentes ($p > 0,05$)

Ejercicio 6.3.

Basándose en el ejercicio 4.2., donde se planteó un experimento para evaluar el rendimiento en kilos por parcela de pasto elefante (*Penisetum purpureum*), frente a 5 biofertilizantes y bajo un modelo lineal aditivo DCL; se aplicará el procedimiento en el software InfoStat.

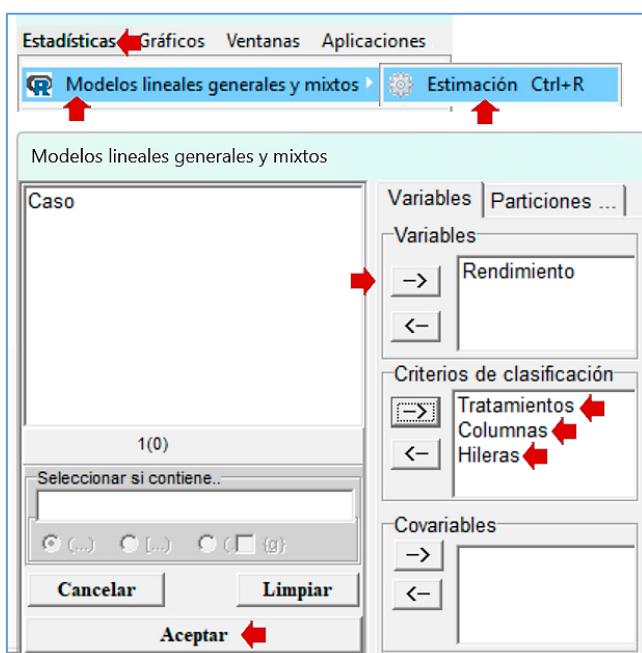
Elaboramos la base de datos, la misma que podría digitalizarse directamente o usando una hoja de calculo como el Microsoft Excel, para luego copiarla (con la identificación de la variable) y trasladarla a una nueva tabla de InfoStat usando la opción pegar incluyendo nombre de columnas.



Los datos que son extraídos directamente de las parcelas experimentales (unidades experimentales), del trabajo de campo, deben someterse al control de los supuestos: Homogeneidad de las varianzas, normalidad e independencia.

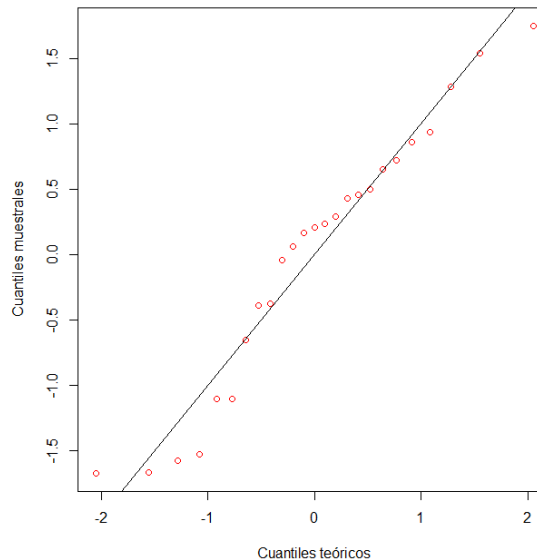
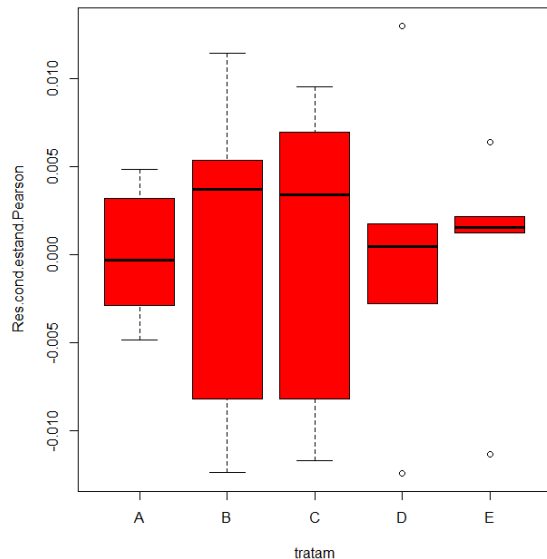
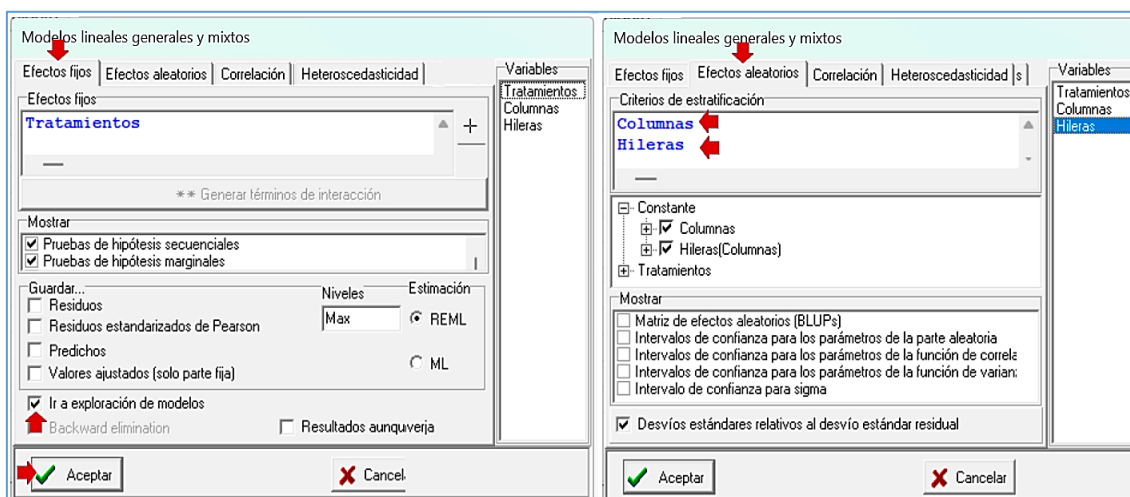
Caso	Tratamientos	Columnas	Hileras	Rendimiento
1	A	1	4	34,20
2	A	2	2	33,70
3	A	3	1	32,60
4	A	4	3	31,80
5	A	5	5	31,20
6	B	1	2	35,80
7	B	2	3	29,70
8	B	3	5	28,40
9	B	4	1	33,40
10	B	5	4	33,90
11	C	1	5	32,20
12	C	2	1	31,10
13	C	3	4	36,90
14	C	4	2	37,70
15	C	5	3	35,80
16	D	1	1	39,30
17	D	2	4	42,30
18	D	3	3	47,20
19	D	4	5	43,70
20	D	5	2	43,30
21	E	1	3	39,80
22	E	2	5	43,70
23	E	3	2	45,30
24	E	4	4	44,00
25	E	5	1	43,80

Desde el menú Estadísticas procedemos



Modelos lineales generales y mixtos → En variables (Rendimiento) → Criterios de Clasificación (Tratamientos, Columnas, hileras por tratarse de un DCL) → Aceptar.

Luego en la pestaña Efectos fijos seleccionamos “Tratamientos” y en la de Efectos aleatorios “Columnas e Hileras” por DCL, habilitamos el recuadro de “Ir a exploración de modelos” y aceptamos.

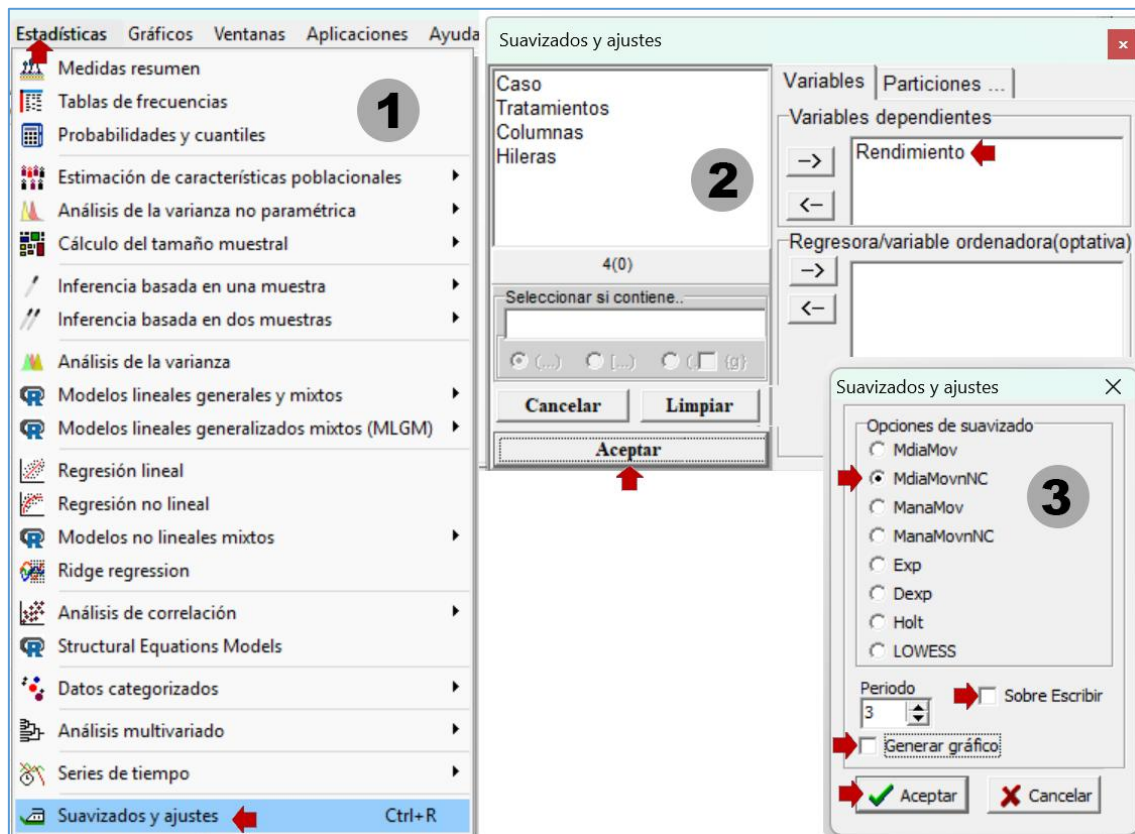


En las gráficas auscultamos los supuestos:

- Si nos fijamos en el de cajas y bigotes, aunque las desviaciones no pasan de los estándares de Pearson (-2 a 2); sin embargo, existen 4 datos atípicos (son los círculos negros) que indican heterogeneidad.

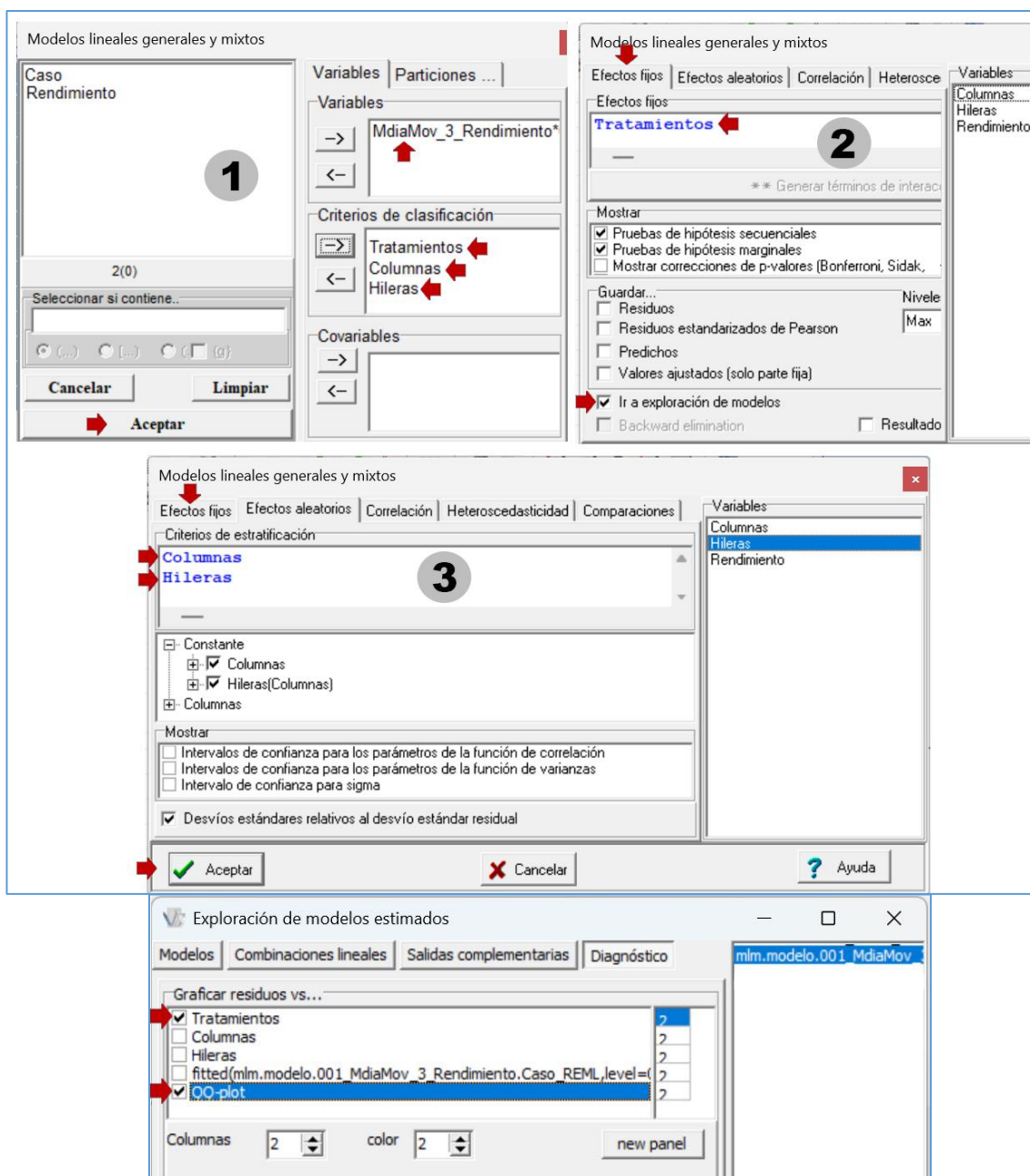
- En el QQ-plot los círculos rojos están cercanos a la línea de la normalidad.
- En conclusión, luego de análisis de supuestos que no se cumplen, se exige dar tratamiento a los mismos usando la opción de “suavizamiento de datos”.

Ajuste de datos: Menú “Estadísticas” → Suavizados y ajustes → Rendimiento (Variables) → Aceptar. Escoger el método de suavizamiento en nuestro caso “Media móvil” y se crea la nueva columna con la data ajustada.



Caso	Tratamientos	Columnas	Hileras	Rendimiento	MdiaMov_3_Rendimiento*Caso
1	A	1	4	34,20	33,50
2	A	2	2	33,70	33,50
3	A	3	1	32,60	32,70
4	A	4	3	31,80	31,87
5	A	5	5	31,20	32,93
6	B	1	2	35,80	32,23
7	B	2	3	29,70	31,30
8	B	3	5	28,40	30,50
9	B	4	1	33,40	31,90
10	B	5	4	33,90	33,17
11	C	1	5	32,20	32,40
12	C	2	1	31,10	33,40
13	C	3	4	36,90	35,23
14	C	4	2	37,70	36,80
15	C	5	3	35,80	37,60
16	D	1	1	39,30	39,13
17	D	2	4	42,30	42,93

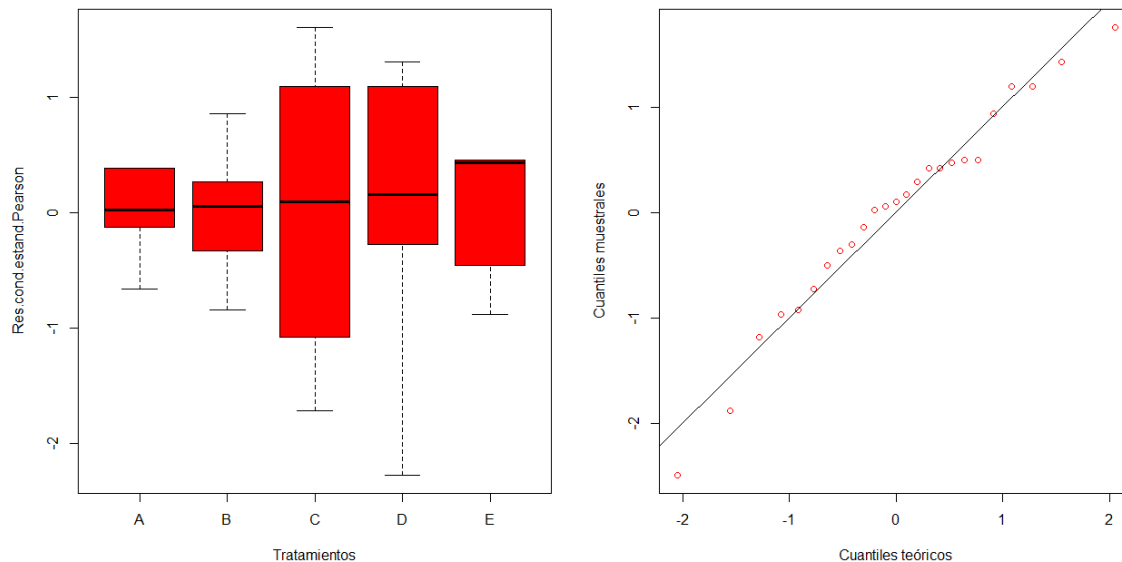
Una vez más comprobamos con la nueva base de datos, los supuestos siguiendo el mismo procedimiento; hasta conseguir las nuevas gráficas.



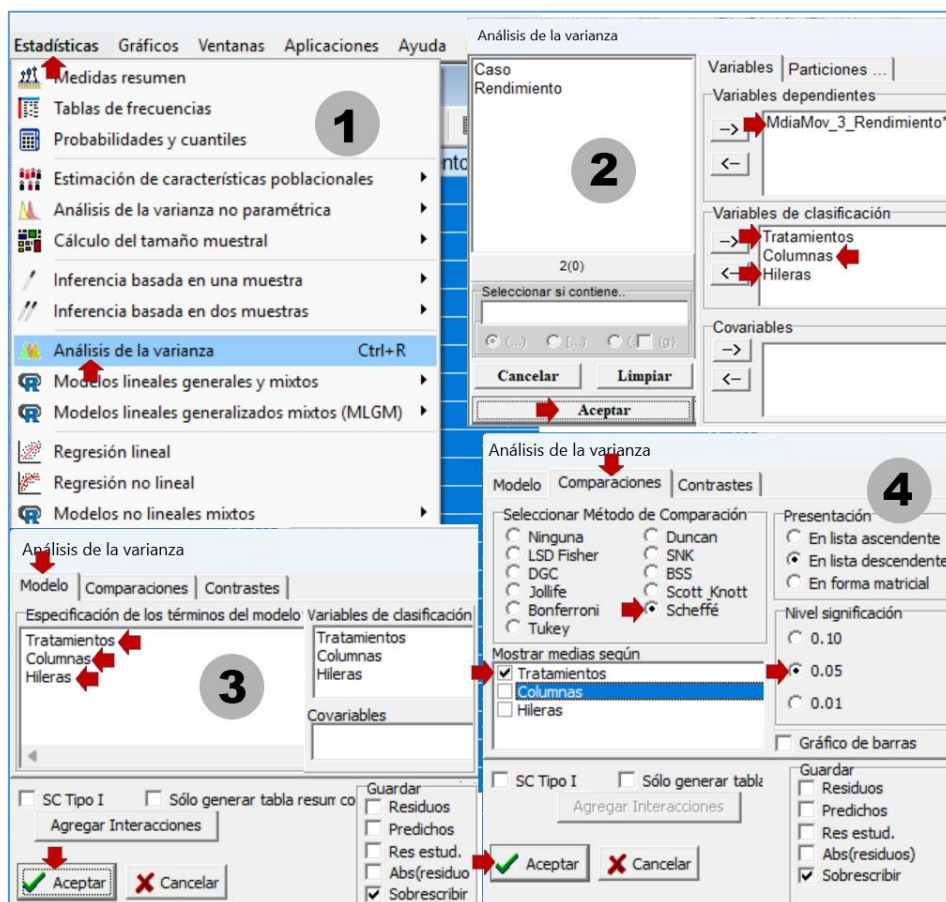
Una vez que se cuenta con las gráficas analizamos los 3 supuestos con la nueva data generada.

- La homogeneidad de varianzas en los tratamientos se identifica en las cajas y bigotes, gracias al ajuste de datos las desviaciones permanecen entre -2 y 2 en el rango de Pearson -2, esto indica homogeneidad.

- En el QQ-plot ya los datos suavizados (círculos rojos) se acercan a la línea de la normalidad cumpliéndose el segundo supuesto.
- Al cumplirse los dos se asume además la independencia.



Procesamos el ADEVA con DCL y la comparación de medias por Scheffé $\alpha 0,05$



Análisis de la varianza

Variable	N	R ²	R ² Aj	CV
MdiaMov 3 Rendimiento*Caso..	25	0,96	0,93	3,76

Cuadro de Análisis de la Varianza (SC tipo I)

F.V.	SC	gl	CM	F	p-valor
Modelo	644,09	12	53,67	27,32	<0,0001
Tratamientos	618,59	4	154,65	78,72	<0,0001
Columnas	15,80	4	3,95	2,01	0,1569
Hileras	9,70	4	2,42	1,23	0,3477
Error	23,58	12	1,96		
Total	667,67	24			

Test:Scheffé Alfa=0,05 DMS=3,20073

Error: 1,9646 gl: 12

Tratamientos Medias n E.E.

E	43,65	5	0,63	A
D	42,69	5	0,63	A
C	35,09	5	0,63	B
A	32,90	5	0,63	B C
B	31,82	5	0,63	C

Medias con una letra común no son significativamente diferentes ($p > 0,05$)

Referencias bibliográficas del capítulo

Benjamini, Y., & Hochberg, Y. (1995). Controlling the false discovery rate: a practical and powerful approach to multiple testing. *J. R. Stat. Soc. Ser. B*, 57, 289-300.

Catie. (s.f.). *Modelos lineales generalizados mixtos: aplicaciones en InfoStat*. CATIE

Repositorio. <https://repositorio.catie.ac.cr/handle/11554/8691?locale-attribute=es> 1

Di Rienzo, J. A., Macchiavelli, R., & Casanoves, F. (2011). *Modelos lineales mixtos: aplicaciones en InfoStat*. Córdoba: Grupo Infostat. ISBN 978-987-27045-0-6. (Actualizado en febrero de 2012)10. Disponible en: <https://www.researchgate.net/publication/283491350>

Morales, J. (s.f.). *Modelos aditivos lineales*. En Bookdown. https://bookdown.org/j_morales/webinmod/05ML-SMOOTH.html 3

Navarro, C., Cavers, S., Pappinen, A., Tigerstedt, P., Lowe, A., & Merila, J. (2005). Contrasting quantitative traits and neutral genetic markers for

genetic resource assessment of Mesoamerican *Cedrela odorata*. *Silvae Genetica*, 54, 281–292

Garibaldi, L. A. (2019). *Modelos lineales mixtos en R*. Academia.edu. https://www.academia.edu/8533347/Modelos_lineales_mixtos_en_R 4

Ospina, S. (2010). Linking plant strategies and ecosystem function: an assessment of the contribution of biodiversity to Neotropical grassland productivity. Ph.D. Tesis. School of the Environment, Natural Resources and Geography, Bangor University, Gwynedd, United Kingdom, and Program Tropical Agricultural.

Universidad de Valladolid. (2019). *Modelos aditivos y su implementación en R* [Trabajo de grado]. UVaDOC. <http://uvadoc.uva.es/handle/10324/38307> 5

Scribd. (s.f.). *Uso de InfoStat: Regresión*. <https://www.scribd.com/document/320586045/Uso-de-InfoStat-Regresion> 7

YouTube. (2022, 27 de mayo). *Cómo presentar tablas de Infostat y Excel según Normas APA* [Video]. <https://www.youtube.com/watch?v=l3JTd-RHZLc> 6

StatisticalEasily. (2024, 12 de abril). *Cómo informar resultados de regresión lineal simple en estilo APA*. <https://es.statisticseasily.com/c%C3%B3mo-informar-una-regresi%C3%B3n-lineal-simple/> 2

Universidad Nacional de La Plata. (s.f.). *Referencias Bibliográficas - SEDICI* [PDF]. http://sedici.unlp.edu.ar/bitstream/handle/10915/59034/Documento_completo_.pdf-PDFA.pdf?sequence=1&isAllowed=y 8

education.statistics. (s.f.). *Uso de InfoStat para análisis estadístico y modelos DCA y DCL* [Memoria de conversación].

RESUMEN DEL LIBRO PARA EL PÚBLICO

El libro “Modelos Lineales Aditivos Simples para la Experimentación Animal” proporciona una orientación exhaustiva sobre el uso de los modelos estadísticos en la investigación animal. El texto, dirigido a investigadores y estudiantes, facilita los modelos a los cuales enfatiza en particular los ALM, que son imprescindibles en cualquier intento de análisis de datos experimentales. Los ALM ayudan a los investigadores comprender y modelar las relaciones entre las variables, ayudando a comprobar los efectos que los tratamientos y las condiciones experimentales causen. En este sentido, el libro abarca ambos aspectos: desde la teoría hasta el trabajo práctico, proporcionando ejemplos que muestran cómo aplicar estos modelos en estudios con animales. También se analizan las propiedades de los ALM que los hacen ventajosos, tales como la flexibilidad o la habilidad para trabajar con datos no lineales. Otros aspectos que los autores discuten incluyen la validación, resaltando que la verificación estadística de los supuestos ha de realizarse para que estas construir conclusiones confiables. El recurso está acompañado con ejercicios y casos prácticos que ensayan la teoría e impulsen a los lectores a realizar su propia investigación. Como se ha expuesto, la obra es un recurso fundamental para realizar estudios en experimentación animal y aumenta el uso de técnicas estadísticas que optimizan la calidad y la interpretación de los resultados obtenidos.

PARA CITAR EL LIBRO

Moscoso Gómez, M. (2025). Modelos lineales aditivos simples para la experimentación animal. Recuperado desde:
<https://libros.cienciadigital.org/index.php/CienciaDigitalEditorial/catalog/book/38>



Las opiniones expresadas por los autores no reflejan la postura del editor de la obra. El libro es de creación original de los autores, por lo que esta editorial se deslinda de cualquier situación legal derivada por plagios, copias parciales o totales de otras obras ya publicados y la responsabilidad legal recaerá directamente en los autores del libro.

El libro queda en propiedad de la editorial y, por tanto, su publicación parcial y/o total en otro medio tiene que ser autorizado por el director de la Editorial Ciencia Digital.

CORREOS Y CÓDIGOS ORCID

Autores



Marcelo Moscoso Gómez



<https://orcid.org/0000-0002-1666-XXXX>



INNOVANDO EN EL ÁREA ACADEMICA