

UNIVERSIDAD NACIONAL AUTONOMA DE MEXICO

FACULTAD DE CIENCIAS

"ANALISIS DE RESIDUALES EN DATOS CON DOBLE CLASIFICACION"

1605

T E S I S

QUE PARA OBTENER

EL TITULO DE :

A C T U A R I O

P R E S E N T A

VICTOR ALFREDO BUSTOS Y DE LA HERRERA

DICIEMBRE

1974



Universidad Nacional
Autónoma de México



UNAM – Dirección General de Bibliotecas
Tesis Digitales
Restricciones de uso

DERECHOS RESERVADOS ©
PROHIBIDA SU REPRODUCCIÓN TOTAL O PARCIAL

Todo el material contenido en esta tesis esta protegido por la Ley Federal del Derecho de Autor (LFDA) de los Estados Unidos Mexicanos (México).

El uso de imágenes, fragmentos de videos, y demás material que sea objeto de protección de los derechos de autor, será exclusivamente para fines educativos e informativos y deberá citar la fuente donde la obtuvo mencionando el autor o autores. Cualquier uso distinto como el lucro, reproducción, edición o modificación, será perseguido y sancionado por el respectivo titular de los Derechos de Autor.

A MIS PADRES

NATALIA

ANTONIO

CON INFINITO AGRADECIMIENTO

A MI NOVIA EVANGELINA
CON TODO MI AMOR

A MIS HERMANOS

JOSE ANTONIO

MARIA EUGENIA

LETICIA

LOURDES

SERGIO

Agradezco al Dr. Ignacio Méndez Ramírez la dirección de este trabajo así como la proposición del tema. Asimismo agradezco al CIMAS la ayuda prestada en el tiempo en que este trabajo se llevó a cabo y las facilidades brindadas con respecto al uso de tiempo de computadora. También a los Sres. Act. Miguel Cervera Flores, Dr. Federico O'Reilly Torno, Act. Santiago Rincón Gallardo, Act. Lucio Pérez Rodríguez la revisión de este trabajo.

I N D I C E

	Pág.
1.- <u>CAPITULO 1</u>	1
Objetivo	
2.- <u>CAPITULO 2</u>	3
Introducción	
3.- <u>CAPITULO 3</u>	23
El Modelo general propuesto por el Dr. Mandel.	
4.- <u>CAPITULO 4</u>	32
Breve descripción de las Pruebas de Hipótesis usadas.	
5.- <u>CAPITULO 5</u>	46
El programa de computadora.	
6.- <u>CAPITULO 6</u>	56
Resultados y Discusión	
7.- <u>CAPITULO 7</u>	121
Resumen y Conclusiones	
8.- <u>APENDICES</u>	
Apéndice I. Listado de los programas usados	A 1
Apéndice II. Tablas	B 1
9.- <u>BIBLIOGRAFIA</u>	

CAPITULO 1

1. OBJETIVO

El objetivo del presente trabajo es investigar las propiedades distribucionales del error en un modelo de diseño de experimentos con dos criterios de clasificación y una sola observación por celda, cuando los efectos de interacción existen y son estimados mediante el método propuesto por el Dr. John Mandel, así el modelo propuesto es:

$$Y_{ij} = \mu + \alpha_i + \beta_j + \sum_{k=1}^K \theta_k u_{ki} v_{kj} + \epsilon_{ij}$$

Donde i se refiere al nivel del primer criterio en el que fué obtenida la información; j al nivel del segundo criterio,

$$\sum_{k=1}^K \theta_k u_{ki} v_{kj}$$

es la representación de la interacción; K es el número de términos de la forma $\theta_k u_{ki} v_{kj}$, este número no se fija de antemano sino que ha de determinarse de alguna forma; ϵ_{ij} es un error aleatorio, el cual se supone debe cumplir con las siguientes propiedades:

- a) $E(\epsilon_{ij}) = 0$
- b) $E(\epsilon_{ij}^2) = \sigma^2$
- c) $E(\epsilon_{ij} \epsilon_{i'j'}) = 0$

El valor de K se determina mediante el examen de las propiedades de los ϵ_{ij} . Así se pretende encontrar el valor de K y en consecuencia la forma real de la relación que cumplen las observaciones con los criterios de acuerdo a los cuales se han clasificado las mismas. Es decir, una vez que las propiedades de los errores se cumplan satisfactoriamente diremos entonces que el modelo propuesto explica satisfactoriamente los datos. El criterio usado por el Dr. Mandel para fijar el valor de K consiste en ordenar los elementos $\hat{\theta}_k$ en forma decreciente, e incluir en el modelo solamente aquellos términos de la forma $\theta u_{ki} v_{kj}$ cuyo valor del cuadrado medio correspondiente es significativamente mayor que aquellos que le suceden en orden de magnitud.

Sin embargo no hace el Dr. Mandel referencia alguna a las hipótesis acerca de la distribución de los errores las cuales permitirán extraer información extra acerca de las observaciones además de las inferencias que acerca de los parámetros sea posible hacer.

Se hace a continuación una breve exposición del modelo propuesto por el Dr. Mandel y se trabaja un criterio para determinar la forma real del modelo a ser ajustado a los datos. Se discuten los métodos usados y se presenta la discusión de los resultados obtenidos al aplicar todos los métodos a varios grupos de datos.

CAPITULO 2

2.- INTRODUCCION

MODELOS

A lo largo de la historia, el hombre siempre ha intentado relacionar los fenómenos que observa con las causas de estos; así nués, ha propuesto relaciones ó modelos entre los fenómenos y sus posibles explicaciones.

Acerca de lo anterior se transcribe lo que al respecto señalan los autores Hillier y Lieberman (4):

"Los modelos son parte integral de la vida común y corriente. Ejemplos simples de modelos, pueden ser representaciones tales como aviones a escala, retratos, globos terráqueos, etc. Por otra parte, los modelos juegan un papel muy importante en la ciencia o en los negocios; así nués, el hombre ha creado también modelos del átomo, modelos de la estructura genética en los seres vivos, modelos matemáticos para describir leyes físicas de movimiento o reacciones químicas, gráficas y sistemas contables industriales. Tales modelos son de inmenso valor para abstraer la esencia del objeto por cuyas características nos preguntamos, mostrando interrelaciones y facilitando el análisis.

MODELOS MATEMATICOS

"Los modelos matemáticos son también representaciones idealizadas, pero son expresados en términos de símbolos matemáticos y expresiones. Leyes físicas tales como:

$$F = m a$$

o

$$E = m c^2$$

Son ejemplos con los que estamos familiarizados".

Más adelante esos mismos autores escriben:

Los modelos matemáticos tienen muchas ventajas sobre una descripción verbal del problema. Una ventaja obvia es que el modelo matemático describe el problema en una forma mucho más concisa. Esto tiende a hacer la estructura total del problema comprensible y ayuda a encontrar importantes relaciones causa-efecto. De esta manera indica más claramente cuáles datos adicionales resultan relevantes para el análisis".

Una observación crucial para el desarrollo de la presente tesis es hecha por los autores antes mencionados, en el siguiente párrafo:

"Un modelo matemático es necesariamente una idealización abstracta del problema, y generalmente se requiere de suposiciones y aproximaciones que simplifican el modelo haciéndolo manejable. De ahí que se debe tener mucho cuidado de asegurarse que el modelo sigue siendo una representación válida del fenómeno. El criterio adecuado para juzgar la validez de un modelo es determinar si predice o no los efectos relativos de los cursos alternativos de acción con precisión suficiente para dar lugar a una decisión adecuada. De ahí que no es necesario incluir

detalles irrelevantes o factores que tienen aproximadamente el mismo efecto para todos los cursos alternativos de acción".

En su libro "Introducción a la metodología estadística" el Dr. Ignacio Méndez Ramírez (10) hace las siguientes observaciones acerca de los modelos matemáticos:

"Los modelos matemáticos, como los puramente simbólicos, deberán ser probados en la práctica para que puedan considerarse útiles; este proceso de pruebas es también, como parte del método científico, un proceso helicoidal que consta de los siguientes pasos:

"1.- Problemas y situaciones del mundo real se trasladan al modelo como problemas y situaciones matemáticas. Esto se hace mediante una abstracción; así de todas las características del mundo real, se representan en el modelo aquellas que se consideren relevantes para el fenómeno en estudio. Así en las leyes que ligan la masa de los cuerpos con otras variables, usualmente no interesa la forma, color, etc., de los cuerpos."

"2.- El modelo matemático debe ser funcional matemáticamente, en el sentido de que un problema del mundo real pueda ser trasladado a un problema matemático y que este último pueda ser resuelto manejando la lógica y postulados matemáticos inherentes al modelo."

"3.- Las soluciones matemáticas obtenidas por el modelo se interpretan términos de soluciones plausibles de los problemas rea-

les."

En la misma referencia se presenta el siguiente diagrama de funcionamiento del modelo:

MUNDO REAL

MUNDO SIMBOLICO

DATOS - - - - -	>	MODELO
PROBLEMAS - - - - -	>	PROBLEMAS AL MODELO
SOLUCION A PROBLEMAS < - -		SOLUCION MATEMATICA

TAXONOMIA DE MODELOS MATEMATICOS

Dentro de los modelos matemáticos podemos establecer la siguiente taxonomía:

- 1.- Modelos determinísticos
- 2.- Modelos aleatorios

Los primeros pretenden, y a veces lo consiguen, establecer relaciones causa-efecto entre ciertas características del fenómeno en estudio y los resultados observados, así por ejemplo, en el modelo:

$$E. = m c^2$$

basta con conocer la masa de un cuerpo (m) para determinar la energía (E) que es posible liberar de ese cuerpo. Obsérvese que modelos como el anterior dependen en la mayor parte de los casos, de constantes o parámetros que las más de las veces es necesario determinar: así el ejemplo arriba citado depende de el parámetro c^2 (la velocidad de la luz, en centímetros/segundo al cuadrado). La determinación de estos parámetros puede ser hecha

en base a experimentación, a repetición del experimento variando las condiciones en que este es realizado y obteniendo distintas mediciones de los efectos; y así a partir del modelo determinar los parámetros del mismo.

Los modelos aleatorios (o estocásticos o probabilísticos) por su parte no pretenden establecer relaciones causa-efecto entre los resultados de un experimento y las condiciones en que este fué realizado. Este tipo de modelos incluye un efecto o error aleatorio debido a las condiciones no controladas en las que fué realizado el experimento, proporcionando a su vez una idea de la distribución espacial de este error. Se puede decir que un modelo aleatorio es un modelo determinístico aumentado con una explicación de las fluctuaciones o variaciones observadas al realizar el experimento en condiciones más o menos homogéneas.

Es decir, supóngase que las condiciones $A, B, C, D, E, \dots, Z$ tienen una mayor o menor influencia con respecto a los resultados de un cierto fenómeno. Entonces el modelo aleatorio puede ser planteado como sigue:

$$Y_{ij\dots s} = f(A_i, B_j, \dots, Z_r) + \psi_s$$

Donde A_i significa: el criterio A en su nivel i.

B_j significa: el criterio B en su nivel j, etc.

Y ψ_{ij} es el error o la explicación aleatoria de las fluctuaciones observadas al realizar repetidas veces el experimento mante-

niendo A, B, ..., Z en los mismos niveles; ψ_{ij} también puede ser un error de medición.

Ahora bien puede suponerse que los criterios G, H, I, ..., Z tienen pequeña influencia en los resultados del experimento, lo que nos permitiría reescribir el modelo como:

$$Y_{ijk...s} = f_1(A_i, B_j, \dots, F_k) + \epsilon_s$$

Donde ϵ_s incluye tanto al error aleatorio ψ_s como la "pequeña" influencia de los criterios G, H, ..., Z.

En este caso es también de observarse que la forma analítica de la función f_1 depende también de un cierto número de parámetros, mismos que a su vez han de ser determinados de alguna manera (estimados). Sin embargo en este tipo de modelos, con el objeto de explicar las desviaciones aleatorias, se debe contar también con un modelo acerca de la distribución de las variables aleatorias

ϵ_s , la forma de la distribución de las variables aleatorias nos indicará a su vez cuantos parámetros es preciso determinar.

La consideración de un elemento aleatorio en el modelo da a su vez un carácter aleatorio a las observaciones, y en consecuencia es preciso determinar la forma de su distribución y los parámetros de los que esta depende.

Con el objeto de aclarar lo antes mencionado se incluye el siguiente ejemplo:

El rendimiento de trigo en una parcela no puede ser predicho con precisión. Muchos de los factores que afectan este rendimiento

son conocidos, pero la ecuación que relaciona estas cantidades es obscura. Se sabe que la temperatura X_1 , lluvia X_2 , cantidad de luz solar X_3 y muchos otros factores afectan el rendimiento Y . Aunque todos los factores que afectan este rendimiento no se conoce y aunque la relación funcional tampoco se conoce, es útil sin embargo suponer que existe un número finito de factores $X_1, X_2, \dots, X_n$ y una función g tal que el rendimiento Y puede ser exactamente determinado mediante:

$$Y = g(X_1, \dots, X_n)$$

Sin embargo si consideramos que ciertos factores tales como la distancia entre el sol y la tierra en el momento de la siembra, la cantidad de rayos gamma absorbidos por las plantas durante su crecimiento, etc., tienen poca influencia sobre el resultado final, las variables $X_6, \dots, X_n$ podrían ser ignoradas y su influencia podría ser considerada como aleatoria, así el modelo podría ser expresado como:

$$Y = g_1(X_1, \dots, X_5) + \varepsilon$$

Donde ε sigue una distribución normal con media cero y dispersión σ común. Así el modelo completo quedaría expresado como:

$$Y = g(X_1, \dots, X_5) + \varepsilon$$

$$\varepsilon \sim N(0, \sigma^2)$$

MODELO LINEAL

Una parte importante de los modelos estocásticos la componen los modelos llamados lineales; una definición de estos modelos esta dada en el libro de Graybill (12).

"Definición 5.2- Por un modelo lineal entenderemos una ecuación que incluye variables aleatorias, variables matemáticas y parámetros, y que es lineal en los parámetros y en las variables aleatorias".

En el ejemplo anterior podemos suponer:

$$Y = \varepsilon_1(X_1, \dots, X_5) + \varepsilon = \alpha_0 + \alpha_1 X_1 + \alpha_2 X_1^2 + \alpha_3 X_2 + \alpha_4 X_3 X_4 + \alpha_5 \log X_5 + \varepsilon$$

En este caso podemos decir que nuestro modelo es lineal en el sentido de la definición arriba mencionada.

Una clasificación de estos modelos está dada en el libro de Graybill; aparece reproducida en seguida:

"Modelo 1.- Supongamos que $X_1, X_2, \dots, X_n$ son variables matemáticas conocidas, Y una variable aleatoria, $\alpha_0, \alpha_1, \dots, \alpha_n$ son parámetros desconocidos y ε una variable aleatoria no observable con media 0. Bajo estas condiciones el modelo:

$$Y = \alpha_0 + \alpha_1 X_1 + \dots + \alpha_n X_n + \varepsilon$$

se define como el modelo 1.

"Modelo 2.- Modelos funcionalmente relacionados con variables sujetas a errores de medida. Supóngase que las variables matemáticas $Y, X_1, X_2, \dots, X_n$ son no observables (medidas con error), pero que $Y, X_1, X_2, \dots, X_n$ pueden ser observadas, don-

$$\text{de } y=Y + \varepsilon, \quad x_i = X_i + \varepsilon_i$$

Suponga además que existe una relación funcional entre las variables matemáticas, dada por:

$$Y = \alpha_0 + \alpha_1 X_1 + \dots + \alpha_n X_n$$

Substituyendo tenemos:

$$y = \alpha_0 + \alpha_1 x_1 + \dots + \alpha_n x_n + \varepsilon - \alpha_1 \varepsilon_1 - \dots - \alpha_n \varepsilon_n$$

Cuando las especificaciones anteriores se cumplen, llamemos al modelo Modelo 2.

"Modelo 3.- Modelos de Regresión. Suponga $y, x_1, \dots, x_n$ son conjunto de variables aleatorias conjuntamente distribuidas tales que el valor generado en la distribución condicional de y dados

$$x_i = X_i \quad (i=1, 2, \dots, n) \text{ está dado por } \alpha_0 + \alpha_1 X_1 + \dots + \alpha_n X_n.$$

Podemos entonces escribir:

$$y_x = \alpha_0 + \alpha_1 X_1 + \dots + \alpha_n X_n + \varepsilon$$

Cuando las condiciones arriba apuntadas se cumplen, definiremos el modelo como el Modelo 3.

"Modelo 4.- Modelos de Diseños Experimentales. Sea y una variable aleatoria, y sean $X_1, \dots, X_n$ variables que toman valores cero ó uno. Si $\alpha_1, \alpha_2, \dots, \alpha_n$ son parámetros desconocidos, entonces:

$$y = \alpha_1 X_1 + \alpha_2 X_2 + \dots + \alpha_n X_n + \varepsilon$$

Lo llamaremos Modelo 4. Este modelo es un caso especial del Modelo 1, pero dada su importancia, será discutido por separado.

Este modelo será escrito como sigue:

$$y_{ij \dots k} = \mu + \alpha_i + \beta_j + \dots + \epsilon_{ij \dots k}$$

Donde $\alpha_i, \beta_j, \dots$, son parámetros desconocidos y $\epsilon_{ij \dots k}$ es una variable aleatoria.

"Modelo 5.- Modelos de componentes de varianza. Sean

$a_1, \dots, a_p$ variables aleatorias no observables de una distribución con medio 0 y varianza σ_a^2 ; sean $b_{11}, b_{12}, \dots, b_{np}$ variables aleatorias no observables de una distribución con media 0 y varianza σ_b^2 . Sea Y_{ij} una variable aleatoria tal que:

$$y_{ij} = \mu + a_i + b_{ij}$$

Donde μ es un parámetro desconocido. Modelos como este serán llamados Modelo 5."

Modelo de Diseños Experimentales

En el presente trabajo estaremos interesados en estudiar casos especiales del Modelo 4 de Diseños Experimentales.

Un ejemplo de este modelo es el siguiente (Graybill, (12,p.223)).

"Suponga que un fabricante de focos, buscando incrementar la duración de sus productos, ha desarrollado dos recubrimientos químicos para los filamentos de los focos. Si el filamento no es tratado con cubrimiento alguno, los focos durarán un promedio de μ (desconocido) hrs. El fabricante supone que, si trate un filamento con el cubrimiento 1, este aumentará la vida del filamento en τ_1 (desconocido) hrs., y que si usa el cubrimiento 2 este

aumentará la duración de τ_2 (desconocido) hrs.

Ahora bien, a él le gustaría encontrar τ_1 y τ_2 el número de horas que los cubrimientos 1 y 2 respectivamente prolongan la vida de los focos (τ_1 y τ_2 pueden ser cero o negativos), y $\tau_1 - \tau_2$, la diferencia entre ambos tratamientos. Para evaluar estas constantes el fabricante podría producir un foco con cubrimiento 1 y probarlo para observar su duración, podría entonces hacer un foco con cubrimiento 2 y observar su duración. El modelo puede ser escrito:

$$Y_1 = \mu + \tau_1 + \epsilon_1$$

$$Y_2 = \mu + \tau_2 + \epsilon_2$$

Donde Y_1 es el número de horas que dura el foco con cubrimiento 1, y ϵ_1 es un error aleatorio debido a todos los errores no controlados, tales como voltaje variable, imperfecciones en la manufactura de los distintos componentes, temperatura, humedad, etc.

"Ahora bien, el fabricante probablemente no estará contento en sacar conclusiones basadas en la observación de solamente un foco por recubrimiento. Supongamos que decide fabricar 6 focos y cubrir 3 con cubrimiento 1 y 3 con cubrimiento 2. El modelo puede ser escrito como:

$$Y_{11} = \mu + \tau_1 + \epsilon_{11}$$

$$Y_{12} = \mu + \tau_1 + \epsilon_{12}$$

$$Y_{13} = \mu + \tau_1 + \varepsilon_{13}$$

$$Y_{21} = \mu + \tau_2 + \varepsilon_{21}$$

$$Y_{22} = \mu + \tau_2 + \varepsilon_{22}$$

$$Y_{23} = \mu + \tau_3 + \varepsilon_{23}$$

Que expresado en forma matricial quedaría:

$$\begin{bmatrix} Y_{11} \\ Y_{12} \\ Y_{13} \\ Y_{21} \\ Y_{22} \\ Y_{23} \end{bmatrix} = \begin{bmatrix} 1 & 1 & 0 \\ 1 & 1 & 0 \\ 1 & 1 & 0 \\ 1 & 0 & 1 \\ 1 & 0 & 1 \\ 1 & 0 & 1 \end{bmatrix} \begin{bmatrix} \mu \\ \tau_1 \\ \tau_2 \end{bmatrix} + \begin{bmatrix} \varepsilon_{11} \\ \varepsilon_{12} \\ \varepsilon_{13} \\ \varepsilon_{21} \\ \varepsilon_{22} \\ \varepsilon_{23} \end{bmatrix}$$

$$\text{o } \underline{Y} = \underline{X}\underline{\beta} + \underline{\varepsilon}$$

Es de hacerse notar que para este tipo de modelos el rango de la matriz X es siempre menor que el número de columnas.

Entonces la matriz $X'X$ que aparece frecuentemente no tiene inversa.

A el modelo anterior se le conoce como modelo de un diseño experimental con un criterio de clasificación (cubrimiento).

Ahora bien, en el presente trabajo estaremos interesados en los modelos de diseño de experimentos con dos clasificaciones y una sola observación por celda. Para este tipo de modelos nuestras observaciones pueden ser conjuntadas en una tabla como sigue:

CRITERIO 1

	1	2	3	...	n
1	Y_{11}	Y_{12}	Y_{13}	...	Y_{1n}
2	Y_{21}	Y_{22}	Y_{23}	...	Y_{2n}
<u>CRITERIO 2</u>	3	Y_{31}	Y_{32}	...	Y_{3n}
	.	.	.	...	.
	.	.	.	...	.
	.	.	.	...	.
	m	Y_{m1}	Y_{m2}	Y_{m3}	Y_{mn}

Donde Y_{ij} es el resultado obtenido al realizar el experimento con el criterio 1 en su nivel J , y el criterio 2 en su nivel i .

n es el total de niveles considerados para el criterio 1.

m es el total de niveles considerados para el criterio 2.

Si lo que nos interesa medir es, por ejemplo, la absorción de rayos ultravioleta, los que hemos llamado criterios pueden ser, intensidad de emisión y distancia entre emisor y receptor.

Si en cambio nuestras observaciones se refieren a la densidad de soluciones acuosas del alcohol etílico, los criterios pueden ser temperatura (distintos niveles) y concentración.

Para este tipo de datos el modelo lineal general que se ajusta es el siguiente:

$$Y_{ij} = \mu + \alpha_i + \beta_j + \eta_{ij} + \varepsilon_{ij}$$

$$i = 1, \dots, m \quad j = 1, \dots, n \quad (1)$$

$$E(\varepsilon_{ij}) = 0 \quad ; \quad E(\varepsilon_{ij}^2) =$$

$$E(\varepsilon_{ij} \varepsilon_{i'j'}) = 0$$

Es decir, se supone que las observaciones son explicadas por:

- μ un efecto común a todas las observaciones cuyo valor es desconocido.
- α_i el efecto de que el criterio 2 se encuentre en su nivel i . Es desconocido también, pero es función de i solamente.
- β_j el efecto de tener el criterio 1 en su nivel j . Este efecto es también conocido; es función de j solamente.
- η_{ij} el efecto de interacción de los niveles i y j de los criterios 2 y 1 respectivamente. Este efecto es desconocido y es función de i y de j simultáneamente. Su forma no es aditiva ya que todo lo que podía considerarse aditivo ha sido extraído por α_i y β_j .
- ε_i error aleatorio debido a las características no controladas en el experimento y aquellas variaciones naturales en el fenómeno que se estudia.

La interacción queda definida, según Graybill como:

"Definición.- $f(X, Z)$ será definida como una función sin interacción si y sólo si existen funciones $g(X)$ y $h(Z)$ tales que:

$$f(X, Z) = g(X) + h(Z)"$$

Debe hacerse notar que la estructura para Y_{ij} dada por el modelo (1), se simplifica enormemente si el efecto de interacción

η_{ij} es igual a cero; en este caso el modelo (1) se simplifica

a:

$$Y_{ij} = \mu + \alpha_i + \beta_j + \varepsilon_{ij}$$

$$i = 1, \dots, n$$

$$j = 1, \dots, m$$

(2)

$$E(\varepsilon_{ij}) = 0, \quad E(\varepsilon_{ij}^2) = \sigma^2, \quad E(\varepsilon_{ij} \varepsilon_{i'j'}) = 0$$

Este modelo sin interacción se conoce también como el modelo estrictamente aditivo y tiene la ventaja de descomponer la función de dos variables $E(Y_{ij})$ en funciones de una sola variable (α_i, β_j) .

La confusión de ambos modelos acarrea graves problemas a las inferencias y pruebas de hipótesis que se hacen con respecto a los parámetros del modelo. Bancroft (1) menciona algunos de los problemas a que lleva el uso del modelo equivocado: "Si efectos de interacción están presentes en la población, pero el modelo escogido fué un modelo aditivo estrictamente (como el (2)) entonces cualquier suma de cuadrados de los efectos principales es solamente un estimador aproximado de la magnitud poblacional de este efecto. Si por otro lado, no hay interacción entre los criterios en la población, pero se sumuso que un modelo con interacción era el adecuado entonces ambas sumas de cuadrados de los efectos principales serán insesgadas pero ineficientemente estimadas".

En vista de lo anterior fué necesario desarrollar criterios para determinar la existencia o no de efectos de interacción en la población y en base a ellos elegir el modelo adecuado a ser ajusta-

do a los datos. Así Tukey (19) desarrolló un método para probar la hipótesis $H_0: \eta_{ij} = 0, i=1, \dots, m, j=1, \dots, n$, bajo la suposición de una forma especial de interacción.

Este método es explicado en Graybill (12), Rao (13). El método consiste brevemente en:

1.- Suma de cuadrados debida a no aditividad:

$$N_{ss} = \frac{\sum_{ij} Y_{ij} (\bar{Y}_{i.} - \bar{Y}_{..}) (\bar{Y}_{.j} - \bar{Y}_{..})^2}{\sum_i (\bar{Y}_{i.} - \bar{Y}_{..})^2 \sum_j (\bar{Y}_{.j} - \bar{Y}_{..})^2}$$

2.- Suma de cuadrados debida a balance (residuo):

$$R_{ss} = \sum_{ij} (Y_{ij} - \bar{Y}_{i.} - \bar{Y}_{.j} + \bar{Y}_{..})^2 - N_{ss}$$

Si $\eta_{ij} = 0$ para todos los valores de i y de j entonces N_{ss}/σ^2 se distribuye como $\chi^2(1)$. R_{ss}/σ^2 se distribuye como

$\chi^2((m-1)(n-1)-1)$ y es independiente de N_{ss}/σ^2 . Entonces bajo la hipótesis $H_0: \eta_{ij} = 0$ para todas i y j , la cantidad:

$$F_c = (N_{ss}/R_{ss}) ((m-1)(n-1)-1)$$

Se distribuye como $F(1, (m-1)(n-1)-1)$. Si $F_c > F_\alpha$ la hipótesis

$H_0: \eta_{ij} = 0$ puede ser rechazada al nivel α .

La derivación de la prueba anterior puede ser resumida como sigue:

Si se supone que $\eta_{ij} = \lambda \alpha_i \beta_j$

La prueba $H_0: \eta_{ij} = 0$, para todo valor de i y de j puede ser escrita como:

Además, sea λ igual a cero o no se tiene:

$$E(\bar{Y}_{i.} - \bar{Y}_{..}) = \alpha_i$$

$$E(\bar{Y}_{.j} - \bar{Y}_{..}) = \beta_j$$

$$E(\bar{Y}_{..}) = \mu$$

De manera que estimadores insesgados de α_i , β_j , μ están dados por:

$$\hat{\alpha}_i = \bar{Y}_{i.} - \bar{Y}_{..}$$

$$\hat{\beta}_j = \bar{Y}_{.j} - \bar{Y}_{..}$$

$$\hat{\mu} = \bar{Y}_{..}$$

Además:

$$E(Y_{ij} - \bar{Y}_{i.} - \bar{Y}_{.j} + \bar{Y}_{..}) = \lambda \alpha_i \beta_j$$

y si se supone que α_i y β_j son conocidas, se puede dar un estimador de mínimos cuadrados para λ , se obtiene:

$$\hat{\lambda} = \frac{\sum_i \sum_j \alpha_i \beta_j (Y_{ij} - \alpha_i - \beta_j - \bar{Y}_{..})}{\sum_i \sum_j \alpha_i^2 \beta_j^2}$$

Se sustituyen los estimadores $\hat{\alpha}_i$ y $\hat{\beta}_j$ por α_i y β_j respectivamente, en la expresión para $\hat{\lambda}$, se obtiene:

$$\hat{\lambda} = \frac{\sum_i \sum_j (Y_{ij}) (\bar{Y}_{i.} - \bar{Y}_{..}) (\bar{Y}_{.j} - \bar{Y}_{..})}{\sum_i (\bar{Y}_{i.} - \bar{Y}_{..})^2 \sum_j (\bar{Y}_{.j} - \bar{Y}_{..})^2}$$

con una varianza, para $\hat{\alpha}_i$ y $\hat{\beta}_j$ dados, igual a:

$$\frac{m^2 n^2 \sigma^2 \sum_i \sum_j (\bar{Y}_{i.} - \bar{Y}_{..})^2 (\bar{Y}_{.j} - \bar{Y}_{..})^2}{\left[\frac{1}{m} \sum_i \bar{Y}_{i.}^2 - \frac{1}{mn} \bar{Y}_{..}^2 \right]^2 \left[\frac{1}{n} \sum_j \bar{Y}_{.j}^2 - \frac{1}{mn} \bar{Y}_{..}^2 \right]^2}$$

o equivalentemente:

$$\frac{m n \sigma^2}{\left[\frac{1}{m} \sum_i \bar{Y}_{i.}^2 - \frac{1}{mn} \bar{Y}_{..}^2 \right] \left[\frac{1}{n} \sum_j \bar{Y}_{.j}^2 - \frac{1}{mn} \bar{Y}_{..}^2 \right]}$$

con la estadística:

$$N_{ss} = \frac{m n \left[\sum_i \sum_j Y_{ij} (\bar{Y}_{i.} - \bar{Y}_{..}) (\bar{Y}_{.j} - \bar{Y}_{..}) \right]^2}{\left[\frac{1}{m} \sum_i \bar{Y}_{i.}^2 - \frac{1}{mn} \bar{Y}_{..}^2 \right] \left[\frac{1}{n} \sum_j \bar{Y}_{.j}^2 - \frac{1}{mn} \bar{Y}_{..}^2 \right]^2} \sim \chi^2_{(1)}$$

pero:

$$\frac{\sum_i \sum_j (Y_{ij} - \bar{Y}_{i.} - \bar{Y}_{.j} + \bar{Y}_{..})^2}{\sigma^2} \sim \chi^2_{((m-1)(n-1)-1)}$$

Siguiéndose de aquí la construcción de R_{ss} y F_c además de la prueba.

Lo anterior es un resumen de los teoremas dados en Graybill

(12, pp.324-331).

Debe hacerse énfasis en el hecho de que la prueba de Tukey supone una forma particular para el efecto de interacción η_{ij} . Como se verá en la siguiente sección esta forma para el efecto de interacción queda incluida como un caso muy particular de la forma propuesta por Mandel (8,9,), para expresar el efecto de interacción (un sólo término multiplicativo con $\theta = \lambda$, $u_i = \alpha_i$ para toda i , $v_j = \beta_j$ para toda j). ^{1/}

La necesidad de suponer una forma especial para el efecto de interacción se debe a que los estimadores, obtenidos mediante el método de mínimos cuadrados, sin hacer ninguna suposición extra acerca de los parámetros del modelo, ajustan perfectamente a nuestras observaciones, lo cual da como resultado que la suma de cuadrados debida al error sea cero.

Esta suma de cuadrados dividida entre sus grados de libertad es usada como el denominador de la estadística que se usa para las pruebas de hipótesis y como estimador de la varianza del error σ^2 . Es obvio que si se toma el valor cero ninguna prueba de hipótesis puede ser llevada a cabo.

^{1/} Si en cada celda se tiene más de una observación, no es necesario hacer suposición alguna acerca de la forma que debe tener el efecto de interacción, es decir, la prueba de la hipótesis $H_0: \eta_{ij} = 0$ para todos los valores de i y de j , puede ser llevada a cabo sin suponer ninguna forma particular para η_{ij} .

En el caso en que no se tiene evidencia alguna, por experiencia anterior o por una muestra independiente de la actual, será preciso aplicar una prueba de este tipo para determinar el modelo que se ajusta mejor a la evidencia disponible (la muestra). Sin embargo la aplicación previa al análisis de los datos de una prueba tal, alterará los niveles de significancia de las pruebas que se realicen con esta misma muestra y con respecto a los parámetros del modelo usado. En efecto, Bancroft (1, p11) propone que la prueba de noaditividad sea realizada a un nivel de significancia del 25% si las pruebas siguientes han de realizarse al 5% para mantenerse en un nivel próximo al 5%.

Sin embargo, asegura, esto disminuirá la potencia de las últimas pruebas.

Es de esperarse que la anterior presentación dé un punto de vista suficientemente amplio para la mejor comprensión de lo que sigue.

CAPITULO 3

3.- EL MODELO GENERAL PROPUESTO POR MANDEL (8,9)

Antes de plantear la derivación y estimación del modelo, se incluyen algunas consideraciones que resultan ser interesantes.

A diferencia de otros métodos de estimación de la interacción entre renglones y columnas, el presente método se propone atacar el problema de la interacción de una manera completamente sistemática, así se parte este término del modelo en tantos términos individuales como sean requeridos por los datos. No se hace ninguna suposición previa sobre linealidad. En buena medida son los mismos datos lo que generan el número de términos para representar la interacción en el modelo.

Este método se deriva del método de componentes principales usado para el análisis de factores en análisis estadístico multivariado. El estudio de esta técnica se debe principalmente a Harold Hotelling (14).

Por otro lado, mientras que el método presentado aquí se enoja en resultados matemáticos conocidos, el punto de vista es esencialmente novedoso.

En particular, mientras que practicamente todas las discusiones acerca de este método asignan funciones asimétricas a los renglones y a las columnas de la tabla, el presente enfoque trata a los renglones y a las columnas de la misma manera, quedando el papel de renglón o columna sujeto a una elección subjetiva de

los criterios que se han mencionado con anterioridad.

Además el presente método puede ser usado para probar aditividad, de efectos de los criterios; para ver esto es preciso hacer ciertas suposiciones acerca de la forma particular en que se expresa el efecto de interacción.

Supóngase que el efecto de interacción (η_{ij} en el modelo) puede ser partido y expresado como sigue:

$$\eta_{ij} = \theta_1 u_{1i} v_{1j} + \theta_2 u_{2i} v_{2j} + \dots + \theta_k u_{ki} v_{kj} + \varepsilon_{ij} \quad (3.1)$$

$$\text{con } E(\varepsilon_{ij}) = 0, \quad E(\varepsilon_{ij}^2) = \sigma^2, \quad E(\varepsilon_{ij} \varepsilon_{i'j'}) = 0$$

De tal modo que la función de 2 variables η_{ij} se expresa como una combinación de funciones de una sola variable (u_{ri}, v_{rj}) nuevamente, lo cual es realmente una simplificación.

Sin pérdida de generalidad se imponen las siguientes restricciones:

$$\sum_i u_{i1} = \sum_i u_{i2} = \dots = \sum_j v_{j1} = \sum_j v_{j2} = \dots = 0 \quad (3.2)$$

$$\sum_i u_{i1}^2 = \sum_i u_{i2}^2 = \dots = \sum_j v_{j1}^2 = \sum_j v_{j2}^2 = \dots = 1$$

El problema esencial consiste en encontrar estimadores de los nuevos parámetros $\theta_1, \theta_2, \dots, u_{i1}, u_{i2}, \dots, v_{j1}, v_{j2}, \dots$ y para σ , la desviación estándar del error aleatorio.

La estimación de los parámetros se lleva a cabo haciendo uso del método de mínimos cuadrados. Se sabe que un estimador de mínimos cuadrados de el efecto de interacción, lo constituye el resi-

dual d_{ij} que esta dado por la siguiente relación:

$$d_{ij} = Y_{ij} - \hat{\mu} - \hat{\alpha}_i - \hat{\beta}_j$$

donde

$$\hat{\mu} = \bar{Y}_{..} = \frac{1}{mn} \sum_i \sum_j Y_{ij}$$

$$\hat{\alpha}_i = \bar{Y}_{i.} - \bar{Y}_{..} = \frac{1}{n} \sum_j Y_{ij} - \bar{Y}_{..}$$

$$\hat{\beta}_j = \bar{Y}_{.j} - \bar{Y}_{..} = \frac{1}{m} \sum_i Y_{ij} - \bar{Y}_{..}$$

Con el objeto de (estimar) ajustar la interacción supóngase primero que:

$$Y_{ij} = \theta u_i v_j + \epsilon_{ij}$$

Donde u_i y v_j están sujetos a las restricciones (3.2) los estimadores de mínimos cuadrados se obtienen minimizando la siguiente expresión:

$$G = \sum_i \sum_j (d_{ij} - \theta u_i v_j)^2 - \lambda_1 (\sum_i u_i^2 - 1) - \lambda_2 (\sum_j v_j^2 - 1) - 2\mu_1 (\sum_i u_i) - 2\mu_2 (\sum_j v_j)$$

Donde $\lambda_1, \lambda_2, \mu_1, \mu_2$, son multiplicadores de Lagrange. Después de llevar a cabo los cálculos respectivos se obtienen las siguientes relaciones:

$$u_i = \theta^{-1} \sum_j d_{ij} v_j \quad (3.4)$$

$$v_j = \theta^{-1} \sum_i d_{ij} u_i \quad (3.5)$$

sustituyendo (3.5) en (3.4) y definiendo

$$S_{it} = \sum_j d_{ij} d_{tj}$$

obtenemos

$$\theta^2 u_t = \sum_i u_i s_{it}$$

En notación matricial (3.6) y (3.7) se pueden expresar como:

$$S = D D' \quad , \quad S \in M_{m \times m} \quad , \quad D \in M_{m \times n} \quad (3.6)'$$

donde D es la matriz de residuales y apostrofe representa la transpuesta. Entonces (3.7) se convierte en:

$$\theta^2 \underline{u} = S \underline{u} \quad , \quad \underline{u} \in R^m \quad (3.7)'$$

De la anterior relación se desprende que θ^2 es un eigenvalor de la matriz S y $\underline{u}$, el eigenvector asociado. De manera similar puede demostrarse que:

$$\theta^2 \underline{v} = (D'D) \underline{v} \quad , \quad \underline{v} \in R^n$$

y así la minimización de G se logra tomando como θ^2 el eigenvalor más grande de la matriz S . Más aún, puede probarse el hecho de que, habiendo escogido como θ^2 el máximo del conjunto de eigenvalores, si se desea minimizar con respecto a θ_2, u_{2i}, v_{2j} la siguiente expresión.

$$\sum_i \sum_j ((d_{ij} - \theta_1 u_{1i} v_{1j}) - \theta_2 u_{2i} v_{2j})^2$$

con u_{2i}, v_{2j} sujetos a restricciones similares a (3.2) se tendrá que el valor mínimo de esa expresión se logra tomando a θ_2 como el segundo más grande eigenvalor de la matriz S original con $\underline{u}_2$, el eigenvector asociado, y con $\underline{v}_2$ en situación similar. Así pues, basta con obtener el conjunto total de eigenvalores de la matriz $S = DD'$, el conjunto $\underline{u}, \underline{u}_2, \dots$ de eigenvectores asocia-

dos, para tener los estimadores de mínimos cuadrados para todos los parámetros de (3.1).

Por otra parte, ya que los residuales se obtienen a partir de los estimadores obtenidos por el método de mínimos cuadrados, se tiene que la matriz $DD' = S$ tiene a lo más, rango igual a $\min(m, n) - 1$; es decir, habrá a lo más $\min(m, n) - 1$ eigenvalores distintos de cero y en consecuencia se tendrá a lo más $\min(m, n) - 1$ términos del tipo $\theta u_i v_j$.

El método para determinar el número de términos de la forma $\theta u_i v_j$ que se han de incluir en el modelo, y que fué desarrollado por el Dr. Mandel (8,9) es interesante e ilustrativo. Es en cierta forma una generalización del criterio usado en el método de componentes principales para determinar cuáles son las combinaciones lineales de las componentes del vector de observaciones que explican la mayor varianza de las mismas observaciones. Este método es explicado por Mandel de la siguiente manera: "Supóngase que calculamos todos los términos del tipo $\theta u_i v_j$, de modo que ϵ_{ij} es hecho igual a cero. "Entonces obtenemos lo siguiente:

$$\sum_i \sum_j d_{ij}^2 = \theta_1^2 + \theta_2^2 + \theta_3^2 + \dots \quad (3.8)$$

"Esto se obtiene a partir de la condición de ortogonalidad de los eigenvectores, de donde todos los productos cruzados son cero, y de la relación (3.1).

"La ecuación (3.8) corresponde a una partición de la suma de cuadrados de la interacción. Entonces se ve uno tentado a intentar un tratamiento del tipo de análisis de varianza. Sin embargo, las θ^2 no son formas cuadráticas de las observaciones originales Y_{ij} ^{1/}. A pesar de lo anterior, es posible formular el problema en un lenguaje de análisis de varianza en virtud de las siguientes consideraciones.

Supóngase primero que las Y_{ij} son una muestra aleatoria de una población $N(0,1)$. Entonces las cantidades θ_1^2 , θ_2^2 , ... formarán cada una poblaciones estadísticas particulares. Sean:

$$M_1 = E_N(\theta_1^2), \quad M_2 = E_N(\theta_2^2), \quad \dots$$

"Donde E es el operador esperanza y N está por una distribución normal estandar ($N(0,1)$).

"Si la población normal tuviera una varianza igual a σ^2 en vez de igual a 1, entonces las cantidades $M_1, M_2, \dots$, sería simplemente multiplicadas por σ^2 .

"Entonces, los cocientes de los valores θ_i^2 obtenidos a partir de una población $N(0, \sigma^2)$ entre los valores M correspondientes, son todos estimadores de σ^2 .

"Si lo que tenemos dado ahora es una matriz de observaciones Y_{ij} donde los términos de interacción no son sino errores normales

^{1/} Si el modelo es estrictamente aditivo, la suma de cuadrados de los residuos estima insensiblemente a un múltiplo de la varianza de ϵ_{ij} .

(con distribución de probabilidad normal, entonces los cocientes de los valores θ^2 obtenidos a partir de estos datos, entre los correspondientes valores de M obtenidos de una matriz de observaciones $N(0,1)$ de las mismas dimensiones son simplemente estimadores de σ^2 .

"Yendo un poco más adelante con este argumento, y haciendo un razonamiento de tipo heurístico, podemos esperar que si el modelo real contiene K términos del tipo $\theta u_i v_j$, entonces los correspondientes K valores de θ^2 serán inflados por los efectos sistemáticos de estos términos, mientras que los términos restantes $\theta^2_{k+1}, \theta^2_{k+2} \dots$ sólo serán estimadores de:

$$M_{k+1} \sigma^2, M_{k+2} \sigma^2, \dots$$

"Luego

$$\frac{\theta^2_{k+1}}{M_{k+1}}, \quad \frac{\theta^2_{k+2}}{M_{k+2}}, \quad \frac{\theta^2_{k+3}}{M_{k+3}}, \quad \dots$$

serán estimadores de σ^2 .

"Es de estagarera en que uno se puede hacer un juicio sobre el número k de términos del tipo $\theta u_i v_j$ que han de ser considerados en el modelo".

Haciendo uso de una simulación Monte-Carlo fueron determinados los valores de $M_1, M_2, M_3, \dots$ para distintos valores de m y n . Valiéndose de estos valores se planteó una analogía de las tablas de análisis de varianza.

Esta se representa a continuación:

Fuente	Grados de Libertad	Suma de Cuadrados	Cuadrados Medios
Total	mn	$\sum_i \sum_j y_{ij}^2$	
α_i	$m-1$	$\sum_i \frac{y_{i.}^2}{n} - \frac{y_{..}^2}{mn}$	
β_j	$n-1$	$\sum_j \frac{y_{.j}^2}{m} - \frac{y_{..}^2}{mn}$	
η_{ij}	$(m-1)(n-1)$	$\sum_i \sum_j d_{ij}^2$	$d_{ij}^2 / (m-1)(n-1)$
$\theta_{1i}^u v_{1j}$	M_1	$\hat{\theta}_1^2$	$\hat{\theta}_1^2 / M_1$
$\theta_{2i}^u v_{2j}$	M_2	$\hat{\theta}_2^2$	$\hat{\theta}_2^2 / M_2$
$\theta_{3i}^u v_{3j}$	M_3	$\hat{\theta}_3^2$	$\hat{\theta}_3^2 / M_3$
$\vdots$	$\vdots$	$\vdots$	$\vdots$
$\vdots$	$\vdots$	$\vdots$	$\vdots$
$\vdots$	$\vdots$	$\vdots$	$\vdots$

Ahora se examinan los cuadrados medios correspondientes a la partición de η_{ij} en la suma de términos $\theta_{ui}^u v_{uj}$. En este momento no se tiene una distribución teórica para estos cuadrados medios, y las pruebas exactas de significancia no se pueden realizar. Sin embargo, aún la interpretación intuitiva de los cuadrados medios nos llevará a determinar fácilmente el número de términos multiplicativos que han de ser considerados en el modelo.

Así, si ninguno de los cuadrados medios relacionado con términos θ_{ij} resulta ser drásticamente mayor que los demás, por lo expuesto antes, ningún término multiplicativo será considerado en el modelo que se ha de ajustar a los datos que se estudien. Es decir el modelo es un modelo estrictamente aditivo:

$$Y_{ij} = \mu + \alpha_i + \beta_j + \varepsilon_{ij}$$

Esta puede ser una alternativa a la prueba de aditividad de Tukey; sin embargo como se apunta arriba se descoroce la distribución de $\hat{\theta}^2$ y en consecuencia el nivel real de significancia de la prueba.

Si en cambio, resulta ser que solamente el primer cuadrado medio, es decir, el correspondiente a θ_{ij} es significativamente mayor que los correspondientes a los términos semejantes posteriores. Así el modelo sería

$$Y_{ij} = \mu + \alpha_i + \beta_j + \theta_{ij} + \varepsilon_{ij}$$

que incluye como caso particular la suposición de Tukey para su prueba de no aditividad

$$Y_{ij} = \mu + \alpha_i + \beta_j + \lambda \alpha_i \beta_j + \varepsilon_{ij}$$

Así podría continuarse usando los valores de la tabla de análisis de varianzas.

CAPITULO 4

4.- BREVE DESCRIPCION DE LAS PRUEBAS DE HIPOTESIS USADAS

En el presente trabajo se hace uso de un criterio esencialmente distinto, para determinar la forma del modelo que ha de ser ajustado a los datos, a aquel presentado en la sección anterior y debido a esto se expone a continuación:

Se planteó en la sección 2 la forma general de un modelo estadístico, que era:

$$Y_{AB...Z} = f(A, B, \dots, Z) + \epsilon_{AB...Z}$$

Donde E es una variable aleatoria cuya distribución cumple con las siguientes propiedades:

- 1) $E(\epsilon_{AB...Z}) = 0$
- 2) $E(\epsilon_{AB...Z}^2) = \sigma^2$ (constante e independiente de $A, B, \dots, Z$)
- 3) $\epsilon_{AB...Z}$ se distribuye normalmente
- 4) $E(\epsilon_{AB...Z} \epsilon_{A'B'...Z'}) = 0$

Es decir, tenemos las suposiciones de:

- 2') Homogeneidad de Varianzas
- 3') Normalidad
- 4') Independencia de los errores

Lo que para este modelo en particular puede ser expresado como:

$$Y_{ij} = \mu + \alpha_i + \beta_j + \sum_{r=1}^k \alpha_r \alpha_{ri} \alpha_{rj} + \epsilon_{ij}$$

con

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

$$E(\varepsilon_{ij} \varepsilon_{i'j'}) = 0$$

Esta última definición de el modelo, en particular la parte que se refiere a las propiedades que deben satisfacer los errores aleatorios asociados con cada observación, permiten atacar el problema de la determinación del modelo más adecuado (número de términos de la forma $\theta_{ij} v_j$ que expresan la interacción), para representar el fenómeno que generó los datos. Es decir, el criterio que se usará en adelante será:

Aumentar a un modelo del tipo $Y_{ij} = \mu + \alpha_i + \beta_j$ tantos términos de la forma $\theta_{ij} v_j$ como sea necesario hasta lograr que las propiedades que deben cumplir los errores sean satisfechas.

Aquí cabe hacer un paréntesis para preguntarse por el objeto de llevar a cabo tarea semejante. La razón por la que se hace lo anterior es lo siguiente:

a) La hipótesis de distribución normal de los errores lleva implícitas varias hipótesis más, como son:

- 1) La distribución es simétrica con respecto a su media.
- 2) El coeficiente de Kurtosis ($E(X - \mu)^4 / \sigma^4 - 3$) es cero, es decir, la gráfica de la función de densidad no es ni rebatida ni picuda.
- 3) La media es independiente de la varianza, es decir, una

translación lineal del origen no afecta la forma de la gráfica de la distribución.

4) Los estimadores de media y varianza ($\bar{X}$, S^2) resumen toda la información obtenible a partir de la muestra (son suficientes minimales). Además son estocásticamente independientes. El hecho de que esta hipótesis no se cumple, en general, afectará fuertemente las inferencias que se hagan con respecto a la varianza y en menor grado las que se hagan con respecto a la media. Los efectos más graves de dicho incumplimiento son el que tampoco se cumpla alguna de las hipótesis que se enumeran arriba (Sesgo, Kurtosis, etc.), se obtendrán demasiados resultados significativos en algunas pruebas de hipótesis hechas con respecto a los parámetros del modelo y se afectará la potencia de alguna de estas pruebas.

Para probar la hipótesis de normalidad se han desarrollado distintos criterios, ampliamente difundidos en la literatura. En el presente trabajo se hace uso de cuatro criterios para probar algunas de las hipótesis, tanto de normalidad como de simetría y kurtosis. Dos de estas pruebas fueron usadas para probar normalidad y estas son:

I) Prueba de Kolmogorov-Smirnov.

II) Prueba de Ji-cuadrada asintótica.

Además se hizo uso de una prueba para simetría y otra para kurtosis.

La primera de las pruebas para normalidad fué usada cuando el número de celdas en las tablas en que se resumió la información era menor o igual a sesenta. La razón de hacer esto fué que en el programa de computadora que se hizo para llevar a cabo las pruebas de hipótesis, la rutina o procedimiento para obtener integrales numéricas, y en consecuencia la que realiza pruebas de Kolmogorov-Smirnov, lleva mucho tiempo de proceso.

Además para muestras grandes, la prueba de Ji-cuadrada asintótica es suficientemente buena.

Se intenta a continuación explicar el funcionamiento de estas pruebas.

La prueba de Ji-cuadrada asintótica.— En esta prueba los datos se agrupan en clases para formar una distribución de frecuencias, además se calculan la $\bar{X}$, la media, y S, la desviación estándar; a partir de estas cantidades se ajusta una distribución normal considerando $\mu = \bar{X}$ y $\sigma = S$.

Para cada clase se encuentra la siguiente cantidad:

$$\frac{(f_i - F_i)^2}{F_i} = \frac{(\text{OBSERVADOS} - \text{ESPERADOS})^2}{\text{ESPERADOS}}$$

Siendo el criterio de prueba en que a continuación se presenta

$$D^2 = \sum_i (f_i - F_i)^2 / F_i$$

Donde la suma corre sobre todas las clases. Si todos los da-

tos realmente pertenecen a una distribución normal, entonces D^2 se distribuye aproximadamente como una Ji-cuadrada con $(R-1)$ grados de libertad con R , el número de clases usadas para el cálculo de D^2 . Si en cambio los datos pertenecen a alguna otra distribución, la diferencia entre f_i , la frecuencia observada, y F_i , la frecuencia esperada, será grande; consecuentemente el valor de D^2 será grande. Lo que causa que si se obtienen valores grandes de D^2 , la hipótesis de normalidad sea rechazada.

El teorema en que se propone el uso de esta prueba (Cramér (2)), requiere que el número esperado F_i no sea demasiado pequeño. Estos pueden ocurrir en las clases de las colas. Como regla práctica para trabajar esos casos extremos se aconseja construir las clases de los extremos de manera tal que los valores esperados en ellas sean al menos uno, siempre que casi todas las restantes tengan al menos un valor esperado de 5.

La prueba de Ji-cuadrada no propone ninguna hipótesis alternativa para la distribución de los errores. Las tablas usadas en esta prueba se encuentran reproducidas en la tabla I del apéndice II.

Kolmogorov-Smirnov.— La distribución muestral

$$P_n(x) = \frac{j}{n} \quad \text{para} \quad X_{(j)} < x \leq X_{(j+1)}, \quad j=0,1,\dots,n$$

con $X_{(0)} = -\infty$ y $X_{(n+1)} = \infty$, en general resultará ser diferente de la función de distribución poblacional. Pero si esa diferencia resulta ser grande podemos usarla como una base para rechazar la hipótesis de que $F(x)$ es realmente la distribución de la población de la cual fué extraída la muestra. Es decir, la diferencia existente entre la distribución empírica y la establecida en la hipótesis debe servir como una estadística muy adecuada para determinar si la distribución que se supuso es efectivamente la poblacional.

De hecho es la diferencia $|F_n(x) - F(x)|$ la que se usa para esta prueba. Sin embargo determinar esta diferencia para todos los valores de X , y decidir cual de ellas usar para la prueba puede resultar difícil; por lo que la estadística de Kolmogorov-Smirnov que llamaremos D_n , porque depende del tamaño de muestra, se toma como:

$$D_n = \sup_{-\infty < x < \infty} |F_n(x) - F(x)|$$

Esta estadística, bajo H_0 , tiene una distribución independiente de la distribución $F(x)$, que es la que se propone como hipótesis nula.

Para decidir si la hipótesis H_0 se rechaza o no, el valor de D_n se compara con valores que se encuentran tabulados (P.ej.

Lindgren (16), P.436). Esta tabla, por que la distribución de D_n es independiente de $F(x)$, es válida para cualquiera que sea

nuestra elección de $F(x)$.

Por otra parte, la hipótesis alternativa, como en el caso de la prueba de Ji-cuadrada asintótica, no se especifica. Esto impide conocer con exactitud la potencia de la prueba. Debe notarse que hasta el momento se ha hablado de "La distribución $F(x)$ " o "La hipótesis H_0 ", es decir, se ha hablado de establecer una hipótesis simple; sin embargo, para el presente trabajo, la hipótesis nula es: "La distribución es normal con media cero" por lo que las tablas usuales no son útiles en este caso ya que la distribución no se ha especificado completamente. Por lo anterior se hará uso de las tablas obtenidas por Illiefors (7).

Estas tablas resultan ser más restrictivas que las tablas usuales, en el sentido de proporcionar valores críticos más pequeños para un mismo tamaño de muestra y un mismo nivel de significancia. Además estas tablas fueron construidas para hipótesis del tipo: "La distribución es normal", con media y varianza no especificadas. Las tablas se encuentran reproducidas en la tabla 2 del Apéndice número II, al final del volumen.

Prueba de Simetría.— Medida de la simetría de una población dada por el valor promedio de $(X-\mu)^3$ ($E(X-\mu)^3/\sigma^3$), tomado sobre toda la población. Si los valores de X se reparten que μ se encuentran agrupados cerca de la media, pero los valores de X se

vores que μ se encuentran muy alejados de este valor, entonces la medida de simetría tomará un valor positivo ya que la contribución de los valores positivos de $(X-\mu)^3$ predominará sobre la de los valores de $(X-\mu)^3$ negativos. También se puede dar el caso de poblaciones con medida de simetría negativa.

Este coeficiente involucra las unidades (al cubo), en que fueron medidas las observaciones; para evitar esta propiedad no deseable, esto es, hacer el coeficiente independiente de las unidades, se divide este coeficiente entre σ^3 . De lo anterior resulta el llamado coeficiente de asimetría el cual es en ocasiones denotado por γ_1 .

El estimador muestral de este coeficiente, llamado g_1 , se calcula como sigue: Sean:

$$m_3 = \sum (x_i - \bar{x})^3 / n$$

$$m_2 = \sum (x_i - \bar{x})^2 / n$$

entonces

$$g_1 = \frac{m_3}{m_2^{3/2}}$$

Obsérvese que en el cálculo de m_2 , la varianza muestral, el denominador es n y no $(n-1)$, lo que facilita los cálculos posteriores.

Si la muestra se toma de una población normal entonces g_1 se distribuye aproximadamente normal con media cero y desviación

estandar $(6/n)^{1/2}$.

La suposición de que σ_1 es normal es suficientemente buena para un tamaño de muestra mayor que 50 (para tamaños de muestra entre 25 y 200, los límites de significancia para pruebas de una cola, están dados para .05 y .01 en la tabla 3 del apéndice II).

PRUEBA DE KURTOSIS

Otro tipo de medida que es útil para caracterizar una población normal es la de Kurtosis (picudez). En general la medida de Kurtosis se define como el promedio de $(X - \mu)^4$ dividido entre σ^4 . Para una población normal este cociente toma el valor 3. Si el cociente excede el valor 3 eso indica que hay una gran concentración de los valores de X alrededor de la media; consecuencia de lo anterior es que los valores alejados de la media tienen poca o ninguna importancia. Es decir la gráfica de la curva tiene una forma "picuda" o "puntiguda". Por otra parte, en caso de que el coeficiente resulte tomar un valor menos que 3, la gráfica de la curva tiene una forma "achata" o "más achata" que la curva normal. Este coeficiente se denota usualmente por γ_2 .

El estimador muestral del coeficiente de Kurtosis está dado por

$$\mu_2 = \frac{m_4}{m_2^2}$$

donde

$$m_4 = \frac{\sum_i (x_i - \bar{x})^4}{n}$$

es el cuarto valor muestral con respecto a la media.

Desgraciadamente la distribución de g_2 está lejos de parecerse a la normal para tamaños de muestra menores que 1000. Para tamaños de muestra entre 50 y 1000, la tabla 4 del apéndice II nos da valores de los límites de significancia para niveles de significancia del 5% y del 1%.

b) La hipótesis de independencia de los errores, en caso de no satisfacerse puede traer problemas en, por ejemplo, inferencias hechas acerca de las medias poblacionales. Se sabe con una correlación positiva de los errores vicia los pruebas de t y de F produciendo muchos resultados significativos.

Esta hipótesis plantea la no existencia de la relación entre diferentes observaciones, la cual no se toma en cuenta por los términos del modelo que representan los efectos. Si los individuos u observaciones se toman al azar y se miden por separado, la suposición es muy razonable; sin embargo pueden aparecer casos en los que esta suposición no se cumple.

Se han desarrollado varios métodos para probar tal hipótesis como son la de Theil y Nacar, la de Durbin-Watson, la de Von Neumann; en el presente trabajo, va que en general las hipóte-

sis que se mide que cumplan las observaciones no se satisfacen (como en Theil y Nagar) o no son concluyentes en algunos casos (Durbin Watson), se hará uso de la prueba Von Neuman para la cual se han tabulado los extremos de significancia. La tabla aparece reproducida en el apéndice II como la tabla 5.

Brevemente esta prueba se lleva a cabo como sigue:

El criterio que se usa para probar la hipótesis de independencia, es en este caso el cociente entre el cuadrado medio de las diferencias sucesivas y varianza, es decir:

$$V = \frac{\sum_{t=2}^n (e_t - e_{t-1})^2 / (n-1)}{\sum_{t=1}^n (e_t - \bar{e})^2 / n} = \frac{\delta^2}{S^2}$$

Para tamaños de muestra (n) grandes, δ^2/S^2 se distribuye aproximadamente como normal con

$$E(V) = \frac{2n}{n-1} ; \quad \text{Var}(V) = \frac{4n^2 (n-2)}{(n+1) (n-1)^3}$$

entonces, para n grande, una prueba para autocorrelación puede llevarse a cabo mediante la comparación de los valores de la razón de Von-Neumann con los límites de la región crítica elegida para una normal con media y varianza dadas arriba, sin embargo es importante enfatizar que las fórmulas para media y varianza son ciertas solamente si los errores se distribuyen independientemente, y esto no se cumple para los residuales obtenidos

por mínimos cuadrados, aún cuando los errores reales si se distribuyen independientemente.

Como se dijo anteriormente, para tamaños de muestra pequeños ($4 \leq n \leq 60$), se usan los límites de significancia para niveles de .01 y .05 dados en la tabla 5 del apéndice II.

c) La hipótesis de igualdad de varianzas de los errores o como usualmente se le conoce, la homoscedasticidad, es realmente una suposición fuerte. Es más difícil de ser satisfecha que, por ejemplo, la independencia de los errores; sin embargo el análisis de varianza formal está basado en esta suposición. De hecho, una característica esencial de varias técnicas es la estimación de esta, supuestamente común, varianza.

Muchas circunstancias pueden viciar esta suposición. Es frecuente que distintos métodos de medición produzcan variaciones de diferentes magnitudes; y valores esperados mayores estén comúnmente asociados con desviaciones estándar mayores. Así pues es necesario considerar posibles variaciones en las desviaciones estándar antes aún de aplicar el análisis de varianza a cualquier conjunto de datos.

El incumplimiento de esta suposición tiende a dar, mediante la prueba de F, demasiados resultados significativos. Cuando en un modelo de diseño de experimentos se supone la igualdad de varianzas se combina la evidencia muestral con el fin de obtener un estimador ponderado de la varianza, obtenido a partir de

las observaciones muestrales; este estimador tiene la propiedad de dar intervalos de confianza más estrechos para la media y diferencia entre distintas medias, y en consecuencia dar pruebas de significancia más potentes. Si la igualdad de varianzas no se satisface, el estimador ponderado de la varianzas no es adecuado como estimador de la varianzas de las observaciones en cada grupo.

A la desigualdad de las varianzas se le ha dado el nombre de heteroscedasticidad, acerca de la cual se ha concluido que si se supone un modelo con heteroscedasticidad, cuando esta realmente existe, es posible obtener un mejor estimador del vector de coeficientes en el modelo. Dicho estimador es un mejor estimador lineal e insesgado para el vector de coeficientes, y es obtenido mediante los llamados mínimos cuadrados generalizados.

El criterio para probar la violación de esta hipótesis, utilizado en este trabajo, fué propuesto por Cochran y hace uso de la siguiente estadística

$$C = \frac{\max_t \sum_{i=1}^{n_t} (x_{ti} - \bar{x}_t)^2}{\sum_{t=1}^k \sum_{i=1}^{n_t} (x_{ti} - \bar{x}_t)^2} = \frac{\max_t S_t}{\sum_t S_t}$$

Para esta prueba, el libro de Johnson y Leone (5) proporciona una tabla de los límites de significancia para el 1% de significancia. Esta tabla es la número 6 del anéndice II.

Esta prueba se puede aplicar solamente cuando el cálculo de todo S_t esté basado en el mismo número de observaciones.

El criterio para rechazar la hipótesis de homoscedasticidad es el siguiente:

Rechace si $C >$ valor tabulado.

De otra manera no rechace.

CAPITULO 5

5.- EL PROGRAMA DE COMPUTADORA

Para el cálculo de los residuales, de los estimadores y la utilización de las pruebas de hipótesis a las que se ha hecho mención a lo largo de este trabajo, se desarrolló un programa de computadora en lenguaje ALGOL (un listado de dicho programa aparece en el anéndice I).

Para el desarrollo de este programa se procuró que las rutinas (procedures) que realicen tal o cual prueba fuesen lo más independientes posible de las restantes, con el objeto de facilitar su utilización en otros programas; aunque esto haya provocado una reducción fuerte en la eficiencia del programa total.

En la presente sección se hace un breve resumen de dicho programa con el fin de que la gente interesada pueda hacer uso del mismo.

Debido a la estructura que debe guardar un programa en ALGOL, se hará referencia únicamente de las declaraciones llamadas (procedures) procedimientos.

Procedimientos LEEMATRIZ E IMPRIMEMATRIZ.- Son rutinas de entrada y salida de datos. LEEMATRIZ, como su nombre lo indica, lee un arreglo de datos dados en forma matricial. IMPRIMEMATRIZ a su vez imprime matrices; este último procedimiento fué usado para la impresión de las matrices de observaciones, de residuos, de eigenvectores y de diferencias en el procedimiento Kolmogorov-

Smirnov.

Procedimiento ESTIMAR.- Este procedimiento calcula los estimadores de mínimos cuadrados de los parámetros que representan los efectos principales (μ , α_i , β_j) si se considera que el modelo es del tipo.

$$Y_{ij} = \mu + \alpha_i + \beta_j + \varepsilon_{ij}$$

los estimadores se encuentran en las salidas del programa bajo "residuos 1".

Procedimiento RESIDUO1.- Este procedimiento hace el cálculo de los residuos $d_{ij} = Y_{ij} - \hat{\mu} - \hat{\alpha}_i - \hat{\beta}_j$, donde, como es usual ($\hat{}$) representa estimador. Estos residuos son expresados en forma de matriz y se encuentran impresos en cada salida del programa.

Procedimiento ORDENA.- Este procedimiento arregla, según su magnitud, los elementos de una matriz. Fue usado en el procedimiento KOTCOGSHIR para obtener la estadística de orden de los residuos.

Procedimiento INTEGRA.- La función de este procedimiento es calcular integrales numéricas. Tiene como una de sus declaraciones al procedimiento INTEGRAL el cual es el que realmente hace el cálculo de las integrales. La función que ha de ser integrada la es proporcionada en el "Define" $F(x,s)$, donde x es el argumento y s es la varianza de la función; la función que se usó es la

de densidad normal con media cero y varianza s (la media es una de las hipótesis del modelo). Es además un procedimiento (integral) iterativo con el fin de lograr que el valor de la integral obtenida sea "bueno" (se permitió un error hasta 10^{-5}).

Procedimiento KOIMOOSMIR.- Este procedimiento lleva a cabo la prueba de la hipótesis de normalidad (con valor esperado igual a cero). En virtud de que la función consiste en rechazar o no rechazar una hipótesis, le fué asignado un tipo (BOOLEAN); es decir, el procedimiento "toma" el valor (false) cero si la hipótesis es rechazada, y el valor uno (true) si no.

En este caso la prueba fué hecha a cinco distintos niveles de significancia (20%, 15%, 10%, 5%, 1%), siendo el tamaño de muestra el total de elementos de la matriz de residuos, y entendiendo por rechazar la hipótesis, el que esta sea rechazada a uno de los cinco niveles de significancia.

Un estimador de la varianza lo da el procedimiento SUPUCUAD el cual se explica a continuación.

El procedimiento KOIMOOSMIR trabaja como se indica a continuación:

Dada una matriz de residuos, estos son ordenados de menor a mayor mediante el procedimiento ORDENA. Después se calcula la función de distribución mediante el procedimiento INTEGRA; esta función es comparada con la función de distribución empírica pero sólo en los valores de los residuos, ya que las mayores diferen-

Este estimador es

$$S_1^2 = \frac{\sum_{i=1}^m \sum_{j=1}^n d_{ij}^2}{(m-1)(n-1)}$$

Para los modelos que incluyen k términos de la forma $\theta u_i v_j$, el estimador de la varianza se obtuvo a partir de la siguiente fórmula:

$$S_k^2 = \frac{\sum_{i=1}^m \sum_{j=1}^n d_{ij}^2}{(m-1)(n-1) - M_1 \dots - M_k}$$

Los valores de M_j fueron obtenidos de las tablas reportadas por Mandel en los artículos en los que expone su método de estimación de los efectos de interacción; estos valores como se acun- to anteriormente fueron obtenidos mediante una simulación Monte-Carlo.

Procedimiento HOMOSCED.- Este es otro procedimiento con tipo (Booleano) ya que fué hecho para probar la hipótesis de igualdad de varianzas (homoscedasticidad). La prueba usada en este procedimiento es la propuesta por Cochran. Como se explicó antes dicha prueba consiste en calcular el cociente del segundo momento muestral con respecto a su media para una de las subpoblaciones entre la suma de los segundos momentos muestrales con respecto a la media de cada una de las subpoblaciones. Tanto los renglones como las columnas de la matriz fueron tomados co-

mo las subpoblaciones que se usan en la prueba y en cada caso fué realizada la prueba de homocedasticidad para el nivel de significancia del 1% por ser este el único nivel para el cual se tabularon los límites de significancia. Se consideró no rechazada la hipótesis sólo en el caso en que simultáneamente no se rechazaba la hipótesis por renglones y por columnas; en este caso el procedimiento tomaba el valor true, de otro modo era false el valor.

Procedimiento JICUAD.- Este procedimiento realiza la prueba de la hipótesis de normalidad con media cero, por lo que también fué asignado el carácter Booleano. Como se dijo antes, este procedimiento fué usado cuando el número de elementos de la matriz de residuo resultó ser mayor o igual que 60. Para llevar a cabo la prueba de la hipótesis se hizo uso del estimador de la desviación estandar de las desviaciones, que se calculó con el procedimiento SUMCUAD. Las observaciones (elementos de la matriz) fueron divididos en 6 clases equiprobables (es decir con probabilidad $1/6$ cada una de ellas) y por lo tanto con el mismo valor esperado, lo que motivó que las longitudes de los intervalos de clase fueran distintas, aunque resulten ser simétricos. Ya que la hipótesis que se desea probar es la de normalidad con valor esperado igual a cero, se usaron los límites de significancia obtenidos en tablas de la distribución Ji-cuadrada con $mn-2$ grados de libertad.

Procedimiento TRANSPON.- Es bien simple la función de este procedimiento, la cual consiste en encontrar la matriz transpuesta de otra matriz que le es dada como entrada.

Procedimiento MULTIPLICAMATRIZ.- Como su nombre claramente lo indica, este procedimiento encuentra la matriz producto de dos matrices que son dadas como entrada para el procedimiento.

El uso dado a los procedimientos anteriores (TRANSPON, MULTIPLICAMATRIZ), fué el encontrar la matriz $DD' = S$, y fueron usados dentro del procedimiento ESTIMAU.

Procedimiento SINETRIA.- Este es otro procedimiento del tipo Booleano ya que fué usado para probar la hipótesis de que la distribución de los residuos era simétrica. Ya que los límites de significancia tabulados están dados para pruebas de una cola para los niveles de 5% y 1%, la prueba de hipótesis de simetría se llevo a cabo para niveles de significancia del 10% y del 2% (la distribución del estimador es simétrica). Como en los casos anteriores se consideró no rechazada la hipótesis sólo en el caso en que la hipótesis fuera no rechazada a ambos niveles, lo cual, por el acortamiento de la longitud del intervalo de no rechazo, implica que la hipótesis no fué rechazada con un nivel de significancia del 10%.

Procedimiento KURTOSIS.- Como se indicó al explicar las pruebas de hipótesis usadas, la prueba KURTOSIS fué usada como una prueba paralela para probar la hipótesis de normalidad. Por lo ante-

rior, a este procedimiento también le fué asignado el tipo Booleano.

Los niveles de significancia usados fueron del 10% y del 2%.

La hipótesis se considera no rechazada si el valor de la estadística calculada se encontraba entre los límites de significancia dados en las tablas para el nivel de significancia más grande; en ese caso el procedimiento toma el valor (true) uno y consecuentemente (false) cero si la hipótesis se rechazó a ese nivel.

Procedimiento NONSYMEIGENVALUESANDEIGENVECTORS.- Este es un procedimiento que fué tomado de la cinta Bibilomet de la biblioteca de la máquina. Fué usado para encontrar los eigenvalores y los eigenvectores de la matriz $DD'=S$, los cuales se recordará, son básicos para el método de estimación de los sumandos en que es dividida el efecto de interacción. Este fué usado también dentro del procedimiento ESTIMAU.

Procedimiento ESTIMAU.- En este procedimiento, haciendo uso de los procedimientos TRANSPON, MULTIPLICAMATRIZ, NONSYMEIGENVALUESANDEIGENVECTORS y ORDENA, se calcula a cada paso la matriz D' , después la matriz S y después se encuentran los eigenvalores y los eigenvectores de la matriz S ; de entre los eigenvectores de la matriz S se escoge el mayor y el eigenvector asociado a dicho eigenvalor. A este procedimiento se llega solamente si alguna de

las hipótesis fué rechazada.

Procedimiento ESTIMAV.- Haciendo uso del eigenvector encontrado en el procedimiento ESTIMAU y el valor propio asociado a dicho vector, se calcula el vector $\underline{v}$ cuyas componentes son los valores v_j que se usan en el cálculo de los términos de la forma $\theta_{ij} v_j$; este cálculo se lleva a cabo haciendo uso de la relación

$$v_j = \theta^{-1} \sum_i d_{ij} u_i$$

que fué introducida al comentar el método de Mandel.

Procedimiento RESIDUOS2.- En este procedimiento se calculan los nuevos residuos resultantes de suponer un modelo de la forma

$$Y_{ij} = \mu + \alpha_i + \beta_j + \theta_1^u u_{1i} v_{1j} + \dots + \theta_r^u u_{ri} v_{rj} + \varepsilon_{ij}$$

donde el número de términos (r) que explican el efecto de interacción se incrementa en uno cada vez que se llega a este procedimiento, es decir, cada vez que estadísticamente alguna de las suposiciones hechas acerca de la distribución del error aleatorio no se cumple.

El programa, brevemente dicho, trabaja de la siguiente manera:

Una matriz de observaciones es leída mediante LEEBTRIZ. Información adicional como son las dimensiones de la matriz, también se proporcionan al programa, además de límites de significancia para las distintas pruebas, etc.

Para esta matriz se calcula la suma de cuadrados de sus componentes (SUMQUAD), se estima el medio general y los efectos principales

los por renglón y por columna (ESTIMAB) y se calculan los residuos (RESIDU01) que se obtendrían si se supone el siguiente modelo:

$$Y_{ij} = \mu + \alpha_i + \beta_j + \varepsilon_{ij}$$

Para la matriz de residuos resultantes se calcula la suma de cuadrados de sus componentes y, si el número de elementos de la matriz es menor o igual que 60, se prueban sobre estos residuos las hipótesis de:

- 1) Homoscedasticidad (HOMOSCED)
- 2) Normalidad (KOLMOGSMIR)
- 3) Independencia de los residuos (AUTOCORR)
- 4) Simetría o Sesgo ((SIMETRIA)
- 5) Kurtosis (KURTOSIS)

Si el número de elementos de la matriz es mayor que 60 entonces se realizan las mismas pruebas, con la única diferencia de que la prueba, de normalidad es realizada mediante JIQUAD.

Si ninguna de las hipótesis es rechazada entonces el programa termina y estos nos indica que el modelo

$$Y_{ij} = \mu + \alpha_i + \beta_j + \varepsilon_{ij}$$

es suficientemente bueno para explicar el comportamiento de las observaciones.

Si en cambio, alguna de las hipótesis es rechazada, se procede a encontrar un término de la forma $\theta u_i v_j$ con el objeto de ajustar

a los datos un modelo como el siguiente:

$$Y_{ij} = \mu + \alpha_i + \beta_j + \theta_{1i} u_{1i} v_{1j} + \epsilon_{ij} \quad (2)$$

(ESTIMAU, ESTIMAV), entonces se calcula la nueva matriz de residuos (RESIDUOS2). Sobre estos nuevos residuos se prueban una vez más todas las hipótesis.

Si ninguna hipótesis es rechazada se concluye que el modelo (2) es suficientemente bueno. Si alguna hipótesis es rechazada entonces se procede a encontrar un nuevo término de la forma $\theta_2 u_{2i} v_{2j}$ para lo cual se hace uso de la nueva matriz de residuos. Este procedimiento se espera termine al incluir 2 o 3 términos multiplicativos para explicar la interacción.

Para probar la efectividad del modelo y de la explicación del efecto de interacción como una suma de términos multiplicativos, se corrió el programa sobre varias matrices de datos. Los resultados se discuten en la siguiente sección.

CAPITULO 6

6.- RESULTADOS OBTENIDOS Y DISCUSION

El objeto de esta sección es probar las suposiciones hechas acerca de la distribución del error en dos tipos de grupos de datos. Primero sobre varios grupos de datos reales tomados de varias publicaciones. Después se realizan todas las pruebas sobre un grupo de 25 matrices de datos las cuales se generaron mediante modelos con interacción y para las cuales se simulaban variables aleatorias normales con media cero y distintas desviaciones estandar. Los grupos de datos fueron generados con distintos efectos de interacción. Para ello se hizo uso del siguiente modelo:

$$Y_{ij} = 2(i-1) + 3j + ((i-1)(j-3))^k + e_{ij}$$

Donde a k le fueron asignados valores 1, 3, 5, 7 y 9; además para cada valor de k se generaron 5 poblaciones normales con media cero y desviación estandar igual a $1/20$, $6/20$, $11/20$, $16/20$, $21/20$, lo que dió como resultado 25 matrices de datos, todos ellos independientes entre sí. (El programa usado para generar estas poblaciones está en el anéndice I).

1) El primer grupo de datos a los que se aplicó todo el juego de pruebas fué tomado del artículo en el que el Dr. Mandel expone su método de estimación. La matriz de observaciones está constituida por mediciones de la densidad de soluciones acuosas del alcohol etílico a 7 niveles de temperatura y a 6 nive-

les de concentración. La matriz aparece representada en la tabla E 1.

TABLA E 1

DENSIDAD DE SOLUCIONES AGUOSAS DE ALCOHOL ETILICO

CONCENTRACION	10	15	20	25	30	35	40
30.086	.95965	.95672	.95369	.95052	.94725	.94387	.94039
39.988	.94241	.93885	.93521	.93150	.92772	.92387	.91994
49.961	.92170	.91784	.91392	.90993	.90588	.90178	.89758
59.976	.89932	.89528	.89120	.88704	.88284	.87857	.87423
70.012	.87598	.87184	.86764	.86337	.85906	.85468	.85024
80.036	.85188	.84764	.84336	.83903	.83464	.83020	.82569

Para estos datos los estimadores de la media general y de los efectos de renglones y de columna están dados por:

$$\begin{aligned} \hat{\mu} &= 0.8968 \\ \hat{\alpha}_1 &= 0.0535 & \hat{\beta}_1 &= 0.0117 \\ \hat{\alpha}_2 &= 0.0346 & \hat{\beta}_2 &= 0.0079 \\ \hat{\alpha}_3 &= 0.0130 & \hat{\beta}_3 &= 0.0041 \\ \hat{\alpha}_4 &= -0.0098 & \hat{\beta}_4 &= 0.0001 \\ \hat{\alpha}_5 &= -0.0325 & \hat{\beta}_5 &= -0.0039 \\ \hat{\alpha}_6 &= -0.0578 & \hat{\beta}_6 &= -0.0079 \\ & & \hat{\beta}_7 &= -0.0121 \end{aligned}$$

Estos estimadores están sujetos a la restricción no estimable:

$$\sum \hat{\alpha}_i = \sum \hat{\beta}_j = 0$$

Para la primera matriz de residuos los resultados de las pruebas fueron los siguientes:

Homoscedasticidad.- Esta hipótesis fué rechazada cuando las subpoblaciones fueron formadas por los renglones de la matriz (se dirá en adelante, "rechazada (o no rechazada) por renglones (o por columnas)"). Pero la misma hipótesis fué aceptada por columnas (el primer rechazo era suficiente para llevar a cabo todas las pruebas una vez más).

Normalidad.- La prueba de Kolmogorov-Smirnov para normalidad no rechazó esta hipótesis a ninguno de los cinco niveles (0.2, 0.15, 0.10, 0.05, 0.01) a los que fué probada.

Autocorrelación.- La prueba de Von Neuman rechazó esta hipótesis a los niveles de 0.10 y 0.02.

Simetría.- La hipótesis de simetría de los residuos no fué rechazada ni al nivel de significancia de 0.10 ni a nivel de 0.02.

Kurtosis.- Esta hipótesis fué rechazada al nivel 0.10 pero en cambio no lo fué al nivel 0.02.

Como resultado del rechazo de varias de las hipótesis, en particular de la autocorrelación (para la cual la razón de Von Neuman dió un valor de 0.861023 siendo la cota inferior más pequeña (para 10%) igual a 1.5408), se hace necesaria la obtención de al menos un término para explicar la interacción.

Los estimadores del primer término del efecto de interacción están dados por:

$$\hat{\sigma}_1^2 = 0.000026108$$

$$\hat{u}_1 = 0.78475$$

$$\hat{v}_1 = -0.59028$$

$$\hat{u}_2 = 0.23149$$

$$\hat{v}_2 = -0.37743$$

$$\hat{u}_3 = -0.05214$$

$$\hat{v}_3 = -0.17315$$

$$\hat{u}_4 = -0.22048$$

$$\hat{v}_4 = 0.01832$$

$$\hat{u}_5 = -0.33470$$

$$\hat{v}_5 = 0.20187$$

$$\hat{u}_6 = -0.40893$$

$$\hat{v}_6 = 0.37602$$

$$\hat{v}_7 = 0.54466$$

Los vectores $\underline{u}$ y $\underline{v}$ han sido normalizados, es decir:

$$\sum_i \hat{u}_i^2 = \sum_j \hat{v}_j^2 = 1$$

Para los residuos resultantes al ajustar un modelo con un término de interacción, las hipótesis no rechazadas, con los respectivos niveles de significancia, fueron los siguientes:

Homoscedasticidad	(Por columnas al 0.1)
Normalidad	(0.20, 0.15, 0.10, 0.05, 0.01)
Autocorrelación	(0.02)
Simetría	(0.10, 0.02).
Kurtosis	(0.10, 0.02).

Mientras que las hipótesis que se rechazaron y los niveles de significancia para los que se rechazaron, fueron:

Homoscedasticidad	(Por renglones 0.01)
Autocorrelación	(0.10)

El rechazo de estas dos hipótesis hizo necesario la inclusión de

un término de la forma $\theta u_i v_j$ y los estimadores para estos valores están dados por:

$$\hat{\theta}_2^2 = 0.000000003$$

$$\hat{u}_1 = -0.45933$$

$$\hat{v}_1 = 0.51188$$

$$\hat{u}_2 = 0.72448$$

$$\hat{v}_2 = 0.00852$$

$$\hat{u}_3 = 0.32933$$

$$\hat{v}_3 = -0.33764$$

$$\hat{u}_4 = -0.04954$$

$$\hat{v}_4 = -0.46830$$

$$\hat{u}_5 = -0.32126$$

$$\hat{v}_5 = -0.35555$$

$$\hat{u}_6 = -0.22368$$

$$\hat{v}_6 = 0.12990$$

$$\hat{v}_7 = 0.51118$$

$$\text{con } \sum_i \hat{u}_i^2 = \sum_j \hat{v}_j^2 = 1.$$

Para los nuevos residuos las hipótesis no rechazadas son las siguientes:

Homoscedasticidad (Por renglones y por columnas, 0.01)

Normalidad (0.10, 0.05, 0.01)

Simetría (0.10, 0.02)

Kurtosis (0.10, 0.02)

Siendo las hipótesis rechazadas, las siguientes:

Normalidad (0.20, 0.15)

Autocorrelación (0.10, 0.02)

De lo anterior se desprende el que un modelo que incluya sólo dos términos multiplicativos no cumple con dos hipótesis centrales

(aunque la normalidad se rechaza con niveles altos de α , probabilidad de error tipo I, por lo que no es tan grave este aspecto), hechas acerca de la distribución conjunta de los errores, sin embargo según el criterio del Dr. John Mandel, este modelo explica suficientemente bien los datos del ejemplo.

Si se aumenta un término multiplicativo a la explicación del efecto de interacción, siendo los estimadores de los parámetros los siguientes:

$$\hat{\theta}_3^2 = (1.47535854) 10^{-10}$$

$$\hat{u}_1 = 0.00425$$

$$\hat{v}_1 = -0.23143$$

$$\hat{u}_2 = -0.39407$$

$$\hat{v}_2 = 0.36943$$

$$\hat{u}_3 = 0.79460$$

$$\hat{v}_3 = -0.13611$$

$$\hat{u}_4 = -0.26636$$

$$\hat{v}_4 = 0.27288$$

$$\hat{u}_5 = 0.13845$$

$$\hat{v}_5 = -0.62633$$

$$\hat{u}_6 = -0.32686$$

$$\hat{v}_6 = 0.53833$$

$$\hat{v}_7 = -0.13677$$

Entonces las hipótesis no rechazadas son:

Homoscedasticidad (Por renglones, por columnas, 0.01)

Normalidad (0.20, 0.15, 0.10, 0.05, 0.01)

Autocorrelación (0.02)

Simetría (0.10, 0.02)

Kurtosis (0.02)

Y las hipótesis rechazadas

Autocorrelación (0.10)

Kurtosis (0.10)

Es decir, dos de las hipótesis centrales se rechazaron pero sólo a uno de los niveles a los que fueron probados recuérdese que ocurrió lo mismo con un sólo término para la interacción, con lo cual un modelo que incluya uno o tres términos multiplicativos mejora el ajuste de los datos a las hipótesis planteadas acerca de la distribución del error. Entonces los modelos que mejor se ajustan a los datos según los criterios establecidos en el presente trabajo son los siguientes:

$$Y_{ij} = \mu + \alpha_i + \beta_j + \theta_1 u_{1i} v_{1j} + \epsilon_{ij}$$

$$Y_{ij} = \mu + \alpha_i + \beta_j + \sum_{k=1}^3 \theta_k u_{ki} v_{kj} + \epsilon_{ij}$$

Para discriminar entre ambos puede usarse el criterio de Mandel, a partir del cual, el modelo que incluye tres términos es mejor. Aquí terminaría el ajuste del modelo, sin embargo si por fines de estudio teórico o por curiosidad se ocurriese el aumentar un término multiplicativo a la explicación del efecto de interacción, entonces se tendría que las hipótesis rechazadas serían las siguientes:

Homocedasticidad (Por renglones 0.01)

Normalidad (0.20, 0.15, 0.10).

Kurtosis (0.10, 0.02)

Y en consecuencia, las no rechazadas serían:

Homoscedasticidad (Por columnas 0.01)

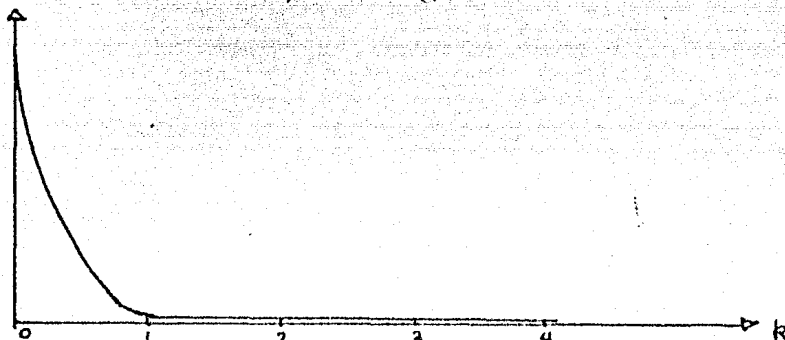
Normalidad (0.05, 0.01)

Autocorrelación (0.10, 0.02)

Simetría (0.10, 0.02)

Además es interesante considerar lo ocurrido con la suma de cuadrados de los residuos; para ello simplemente se plotea la suma de cuadrados contra el número de términos multiplicativos incluidos en el modelo, en la gráfica G 1.

SCD



Gráfica G 1.- Reducción en suma de cuadrados contra número de sumandos del efecto de interacción.

2) El siguiente grupo de datos fué tomado de el artículo publicado por D.H. Levcock (Tron. Agriculture, Trin, Vol 32 No. 2, Abril 1955), en el cual se describe el efecto de las formas de las parcelas en la reducción del error en experimentos de té. Para este artículo se formaron 54 parcelas en un territorio rectangular con curvas de fertilidad definidas, las cuales se presentan en la tabla E 2 con los rendimientos de cada parcela incluidos.

TABLA N. 2

RENDIMIENTOS DE TR. CON CURVAS DE FERTILIDAD
DEL CAMPO EXPERIMENTAL

28.5	27.1	30.4	27.6	25.1	23.2
32.7	20.5	22.3	24.8	22.1	22.8
31.7	23.0	23.9	29.4	22.5	29.0
26.7	21.9	23.0	25.3	26.9	29.4
25.8	21.4	20.0	23.7	31.1	37.3
25.1	20.2	21.2	22.2	27.4	34.7
28.1	25.3	21.4	25.1	28.5	35.7
30.6	37.0	52.7	41.9	43.9	45.1
41.4	49.7	49.2	47.4	47.8	43.7

La matriz de observaciones dió como resultado los siguientes estimadores de los parámetros:

$$\hat{\mu} = 30.2667$$

$$\hat{\alpha}_1 = -3.2833$$

$$\hat{\beta}_1 = -0.2000$$

$$\hat{\alpha}_2 = -6.0667$$

$$\hat{\beta}_2 = -2.9222$$

$$\hat{\alpha}_3 = -3.6833$$

$$\hat{\beta}_3 = -0.9222$$

$$\hat{\alpha}_4 = -4.7333$$

$$\hat{\beta}_4 = 0.5556$$

$$\hat{\alpha}_5 = -3.7167$$

$$\hat{\beta}_5 = 0.3222$$

$$\hat{\alpha}_6 = -3.4667$$

$$\hat{\beta}_6 = 3.1667$$

$$\hat{\alpha}_7 = -2.0167$$

$$\hat{\alpha}_8 = 11.6000$$

$$\hat{\alpha}_9 = 16.2667$$

A la matriz de residuos que da como resultado el uso de estos estimadores se aplicaron las pruebas de hipótesis, rechazando las

siguientes:

Normalidad (0.20, 0.15, aunque es a niveles altos de α).

Y no rechazadas las restantes hipótesis. De lo anterior es posible suponer que un modelo que incluya únicamente los términos aditivos es, para todos los fines prácticos, suficientemente bueno. Para este modelo, la suma de cuadrados del error es igual a 901.03555.

Si se aumenta un término multiplicativo para explicar los efectos de la interacción, si alguno de estos efectos existe, los estimadores de los parámetros involucrados en este término son:

$$\hat{\sigma}_1^2 = 453.1528$$

$$\hat{u}_1 = 0.26127$$

$$\hat{v}_1 = -0.52572$$

$$\hat{u}_2 = -0.16968$$

$$\hat{v}_2 = 0.17493$$

$$\hat{u}_3 = -0.18439$$

$$\hat{v}_3 = 0.72779$$

$$\hat{u}_4 = -0.10770$$

$$\hat{v}_4 = 0.01233$$

$$\hat{u}_5 = -0.34007$$

$$\hat{v}_5 = 0.01437$$

$$\hat{u}_6 = -0.23988$$

$$\hat{v}_6 = -0.40371$$

$$\hat{u}_7 = -0.28756$$

$$\hat{u}_8 = 0.65282$$

$$\hat{u}_9 = 0.40918$$

Un modelo con un término multiplicativo para explicación de los efectos de interacción ajustado a los datos originales dio como resultado que las hipótesis rechazadas fueran las siguientes:

Normalidad (0.20)

Es decir ahora los errores son "más normales" que los obtenidos con un modelo estrictamente aditivo. Por otro lado la suma de cuadrados del error alcanzó un valor de sólo 447.88275, que es menor que el 50% de la suma de cuadrados obtenida con un modelo estrictamente aditivo, luego la inclusión de ese término de interacción mejora el ajuste del modelo.

Para estos datos la inclusión de un nuevo término de interacción sólo empeora las propiedades que tenían los errores con el modelo anterior, como lo demuestra el rechazo de las hipótesis de:

Normalidad	(0.20, 0.15, 0.10, 0.05)
Autocorrelación	(0.10, 0.02)
Simetría	(0.10, 0.02)
Kurtosis	(0.10, 0.02)

Los estimadores de este nuevo término de la explicación de la interacción son:

$$\hat{\theta}_2^2 = 318.9956$$

$$\hat{u}_1 = 0.40551$$

$$\hat{v}_1 = 0.63069$$

$$\hat{u}_2 = 0.50841$$

$$\hat{v}_2 = 0.21988$$

$$\hat{u}_3 = 0.29973$$

$$\hat{v}_3 = 0.04704$$

$$\hat{u}_4 = -0.11220$$

$$\hat{v}_4 = 0.09778$$

$$\hat{u}_5 = -0.43815$$

$$\hat{v}_5 = -0.34488$$

$$\hat{u}_6 = -0.2619$$

$$\hat{v}_6 = -0.65051$$

$$\hat{u}_7 = -0.18909$$

$$\hat{u}_8 = -0.42210$$

$$\hat{u}_9 = 0.11220$$

Aunque se ha ganado en suma de cuadrados del error logrando reducir el valor hasta 128.8871, se han perdido casi todas las propiedades deseables para los errores.

Si para el mismo grupo de datos se construye la tabla de análisis de la varianza propuesta por el Dr. Mandel, esta tendría el aspecto siguiente:

FUENTE	GRADOS DE LIBERTAD	SUMA DE CUADRADOS	CUADRADOS MEDIOS
$\theta_1^{u_{1i} v_{1j}}$	18.6	453.1528	48.4428
$\theta_2^{u_{2i} v_{2j}}$	10.89	318.9956	41.1269
$\theta_3^{u_{3i} v_{3j}}$	6.24	128.8871	20.6549

De aquí es fácil ver que, de acuerdo al criterio propuesto por Mandel, se tienen dos alternativas:

A) No incluir ningún término para explicación de la interacción ya que los cuadrados medios tienen el mismo orden de magnitud.

Con esto se tendría una suma de cuadrados igual a 901.03555.

B) Incluir al menos tres términos para la explicación de la interacción. Esta alternativa daría como resultado errores poco apropiados para realizar inferencias y pruebas de hipótesis, por contradecir casi todas las hipótesis hechas acerca de la distribución del error.

De acuerdo a las propiedades de los errores y la suma de cuadrados del error el "mejor" modelo es el que incluye únicamente

θ_1, u_{1i}, v_{1j} , lo que no coincide en ninguno de los casos a los que conduce el criterio de Mandel.

3) El programa también se aplicó a una serie de matrices que representan los rendimientos promedio de árboles frutales en libras por parcela para los años 1921, 1923, 1925 y 1927. Estos rendimientos fueron reportados por Parker-Batchelor en Hilgardia (Variation in yields of Fruit Trees, Hilgardia, October 1932 Vol. 7, No. 32).

Las matrices usadas representan sólo una parte del experimento total y son las siguientes:

1921

18	34	27	28	15
22	35	19	22	40
29	36	35	24	40
14	36	34	24	31
23	31	20	14	16
30	16	25	28	12
12	23	32	32	13
19	30	27	21	14
15	33	27	32	20
16	22	24	28	15
15	31	19	28	28

1923

50	76	60	72	63
66	58	72	63	62
74	65	52	56	79
52	60	42	66	63
65	73	65	73	55
71	59	72	79	55
46	75	85	86	51
57	51	83	74	75
47	61	64	75	60
62	62	72	81	68
66	107	80	91	86

1925

121	138	125	128	145
116	135	136	125	151
129	132	136	114	171
95	126	121	140	164
115	120	128	124	145
113	134	129	145	141
107	158	149	149	131
119	143	143	134	123
105	133	130	122	135
106	132	136	155	135
109	141	148	144	153

1927

108	104	107	103	128
128	152	169	129	144
136	153	156	114	183
107	153	146	129	157
145	162	174	135	168
134	168	179	134	132
159	186	186	145	146
154	167	170	144	140
159	180	145	152	164
152	177	187	164	167
135	180	185	186	163

Para la matriz que representa los rendimientos en el año de 1921 la aplicación del programa dió los siguientes resultados:

Los estimadores de los parámetros fueron:

$$\hat{\mu} = 24.6364$$

$$\hat{\alpha}_1 = -0.2634$$

$$\hat{\beta}_1 = -5.2727$$

$$\hat{\alpha}_2 = 2.9636$$

$$\hat{\beta}_2 = 5.0909$$

$$\hat{\alpha}_3 = 8.1636$$

$$\hat{\beta}_3 = 1.7273$$

$$\hat{\alpha}_4 = 3.3636$$

$$\hat{\beta}_4 = 0.9091$$

$$\hat{\alpha}_5 = -3.86634$$

$$\hat{\beta}_5 = -2.4545$$

$$\hat{\alpha}_6 = -2.4363$$

$$\hat{\alpha}_7 = -2.2364$$

$$\hat{\alpha}_8 = -2.4364$$

$$\hat{\alpha}_9 = 0.7636$$

$$\hat{\alpha}_{10} = -3.6364$$

$$\hat{\alpha}_{11} = -0.4364$$

A los residuos resultantes de substituir los valores anteriores de los estimadores, se les aplicaron todas las pruebas de las hipótesis, lo que dió como resultado el que ninguna de las hipótesis se rechazara; en otras palabras, se supone que un modelo estrictamente aditivo es adecuado para ser ajustado a los datos. El valor de la suma de cuadrados de los residuos fué igual a 185.70909, con un estimador de la varianza igual a 43.3939272.

Para la matriz que representa los rendimientos en el año de 1923 los resultados variaron únicamente en la estimación de los parámetros del modelo aditivo y en el valor de la suma de cuadrados de los residuos. Los estimadores obtenidos se presentan a continuación.

$$\hat{\mu} = 67.0545$$

$$\hat{\alpha}_1 = -2.8545$$

$$\hat{\alpha}_2 = -2.8545$$

$$\hat{\alpha}_3 = -1.8545$$

$$\hat{\beta}_1 = -7.4182$$

$$\hat{\beta}_2 = 0.8545$$

$$\hat{\beta}_3 = 0.8545$$

$$\hat{\alpha}_4 = -9.4545$$

$$\hat{\beta}_4 = 7.1273$$

$$\hat{\alpha}_5 = -0.8545$$

$$\hat{\beta}_5 = -1.4182$$

$$\hat{\alpha}_6 = 0.1455$$

$$\hat{\alpha}_7 = 1.5455$$

$$\hat{\alpha}_8 = 0.9455$$

$$\hat{\alpha}_9 = -5.6545$$

$$\hat{\alpha}_{10} = 1.9455$$

$$\hat{\alpha}_{11} = 18.9455$$

Y la suma de cuadrados de los residuos aumentó hasta el valor 4635.309090, dando un estimador de la varianza igual a 115.88272; sin embargo ninguna de las hipótesis distribucionales del error fué rechazada, ajustándose en consecuencia solamente un modelo aditivo.

Para la matriz que representa los rendimientos para el año 1925 tampoco fué rechazada, ninguna de las hipótesis hechas al modelo, por lo que el modelo estrictamente aditivo se ajustó a los datos. Lo anterior trajo como consecuencia que la suma de cuadrados del error aumentara aún más hasta alcanzar el valor 5351.30909, dando un estimador de la varianza de los errores igual a 133.782727.

Los estimadores obtenidos fueron los siguientes:

$$\hat{\mu} = 132.4000$$

$$\hat{\alpha}_1 = -1.0$$

$$\hat{\alpha}_2 = 0.2$$

$$\hat{\alpha}_3 = 4.0$$

$$\hat{\alpha}_4 = -3.2$$

$$\hat{\alpha}_5 = -6.0$$

$$\hat{\alpha}_6 = 0.0$$

$$\hat{\alpha}_7 = 6.4$$

$$\hat{\alpha}_8 = 9.3132(10)^{-10}$$

$$\hat{\alpha}_9 = -7.4$$

$$\hat{\alpha}_{10} = 0.4$$

$$\hat{\alpha}_{11} = 6.6$$

$$\hat{\beta}_1 = -20.1273$$

$$\hat{\beta}_2 = 3.2364$$

$$\hat{\beta}_3 = 2.2364$$

$$\hat{\beta}_4 = 2.1455$$

$$\hat{\beta}_5 = 12.5091$$

Sin embargo para la matriz que presenta los rendimientos frutales para el año de 1927, los resultados cambiaron notablemente. Los estimadores de los parámetros para el modelo lineal aditivo resultaron ser los siguientes:

$$\hat{\mu} = 151.4545$$

$$\hat{\alpha}_1 = -41.4545$$

$$\hat{\alpha}_2 = -7.0545$$

$$\hat{\alpha}_3 = -3.0545$$

$$\hat{\alpha}_4 = -13.0545$$

$$\hat{\alpha}_5 = 5.3455$$

$$\hat{\alpha}_6 = -2.0545$$

$$\hat{\alpha}_7 = 12.9455$$

$$\hat{\beta}_1 = -13.5455$$

$$\hat{\beta}_2 = 10.5455$$

$$\hat{\beta}_3 = 12.5455$$

$$\hat{\beta}_4 = -11.9091$$

$$\hat{\beta}_5 = 2.3636$$

$$\hat{\alpha}_8 = 3.5455$$

$$\hat{\alpha}_9 = 8.5455$$

$$\hat{\alpha}_{10} = 17.9455$$

$$\hat{\alpha}_{11} = 18.3455$$

Para los residuos obtenidos al sustituir los estimadores arriba anotados la suma de cuadrados fué igual a 8022.43636. Al aplicar las pruebas a estos residuos, la hipótesis de independencia de los mismos fué rechazada únicamente a el nivel de significancia del 0.10 (el valor de la razón de Von Neuman fué igual a 2.510922 que es mayor que el límite superior del intervalo de no rechazo el cual tiene el valor de 2.4819). Lo anterior trabajo como consecuencia que se aumentara un término multiplicativo para intentar explicar el efecto de interacción, los estimadores para los parámetros de dicho término estuvieron dados por:

$$\hat{\theta}_1^2 = 3899.55294$$

$$\hat{u}_1 = 0.37193$$

$$\hat{v}_1 = 0.09339$$

$$\hat{u}_2 = -0.09818$$

$$\hat{v}_2 = -0.2308$$

$$\hat{u}_3 = 0.57935$$

$$\hat{v}_3 = -0.42886$$

$$\hat{u}_4 = 0.19704$$

$$\hat{v}_4 = -0.26174$$

$$\hat{u}_5 = 0.14902$$

$$\hat{v}_5 = 0.823$$

$$\hat{u}_6 = -0.39709$$

$$\hat{u}_7 = -0.33477$$

$$\hat{u}_8 = -0.23751$$

$$\hat{u}_9 = 0.1783$$

$$\hat{u}_{10} = -0.12004$$

$$\hat{u}_{11} = -0.28806$$

La suma de cuadrados para los nuevos residuos fué reducida hasta 4122.88342 al ajustar el nuevo modelo, con lo que se obtuvo un estimador de la varianza común de los errores igual a 103.072085. Para estos residuos ninguna de las hipótesis fué rechazada con lo que según el criterio desarrollado aquí el modelo que satisfactoriamente explica los datos originales es el dado por:

$$Y_{ij} = \mu + \alpha_i + \beta_j + \theta u_i v_j + \epsilon_{ij}$$

Es decir, es notable la pérdida en eficiencia del modelo lineal aditivo para explicar los rendimientos de los árboles frutales al aumentar la edad de los mismos. Así al corregir la falta de independencia se logra reducir la suma de cuadrados del error aproximadamente a la mitad.

4) El siguiente grupo de datos al que se le aplicó el método de estimación de Mandel fué publicado por R.J. Garber, T.C. Mollvane, y E.M. Hoover (18 pp. 286-298). El experimento realizado consistió en sembrar maíz en 1927, avena en 1928 y trigo en 1929, en el mismo campo experimental. Los rendimientos, para su comparación, fueron expresados como desviaciones con respecto a la producción promedio de todo el campo, para el mismo tipo de cul-

tivo. Esto es lo que se presenta en la tabla E 4, aunque sólo se presenta una parte del total del campo utilizado en el experimento.

TABLA E 4.- DESVIACIONES PORCIENTO CON RESPECTO A LA MEDIA

- 0.7	6.7	2.0	- 8.1	5.4	- 8.1	0.1	2.4	-11.4
-11.0	16.6	4.2	-11.6	-13.9	-21.6	-10.9	3.4	9.6
- 8.5	10.4	5.3	5.4	7.9	- 3.7	6.6	-13.8	- 3.2
- 8.5	2.9	-13.8	-13.7	-12.2	- 0.4	14.6	-20.2	-2.1
- 3.3	- 8.3	- 6.7	-15.2	- 6.3	10.6	5.1	-11.7	- 0.1
- 9.3	- 2.1	-12.7	-17.7	-34.0	-17.7	-26.5	-12.5	-10.1
-27.6	-32.1	-17.7	-17.0	-29.8	-18.7	-11.4	2.9	-15.6
-26.4	-13.1	-22.2	-20.9	-28.8	-26.1	-19.1	-16.6	- 2.1
-10.9	0.7	- 0.3	- 6.0	-18.8	-31.4	-31.7	-33.9	-21.2
13.9	-10.5	- 9.5	3.8	-20.1	-12.6	-41.1	-18.5	-15.1
- 1.3	18.4	11.0	2.4	- 2.1	- 1.3	-13.8	- 8.1	7.9

MAIZ

-25.9	39.3	32.2	25.3	- 0.9	12.8	24.1	3.0	22.0
-38.7	23.2	24.7	21.0	17.4	- 7.6	- 3.0	18.6	32.3
-35.7	38.7	7.9	41.8	11.0	2.7	13.1	-0 .6	15.2
-42.7	-23.8	2.4	12.8	7.3	24.1	18.0	-19.2	-13.7
-31.7	- 7.0	- 2.1	-16.2	8.2	17.4	- 7.0	-26.2	-15.9
-31.4	-23.8	8.8	8.8	-29.0	-26.3	-49.4	-33.8	-39.0
-58.8	-30.5	- 9.1	5.5	-14.6	-20.1	- 7.9	- 7.9	-29.6
-60.4	-28.7	-11.3	- 0.9	-26.0	-33.5	-16.8	-15.5	- 4.6
-22.3	-22.9	- 1.8	- 1.8	-16.2	-40.9	-34.8	-48.5	-25.0
-18.3	-28.7	-15.5	3.4	-18.9	-21.0	-24.1	-18.0	-16.2
-15.9	- 7.3	-17.7	6.1	- 9.1	- 5.2	-26.2	- 4.6	- 0.9

AVENA

10.8	68.6	61.9	29.4	31.4	34.0	51.0	28.9	14.9
- 0.5	25.1	49.0	30.4	12.9	12.4	17.0	21.6	44.3
- 3.6	45.4	33.5	34.5	35.6	12.9	18.6	26.3	25.8
2.6	2.1	18.6	26.3	20.1	16.0	17.0	-12.4	8.8
- 9.8	3.1	9.8	32.0	40.2	2.6	-20.6	-37.1	-12.4
4.1	10.8	11.9	23.2	-15.5	-33.5	-31.4	-32.0	-3.1
-25.8	-22.2	- 3.1	9.8	7.7	-21.6	7.7	-11.9	-26.3
-42.3	-16.5	-1.0	16.0	- 1.5	-18.6	- 1.5	- 0.3	- 5.2
2.6	- 5.7	-12.9	-16.5	-17.0	-30.9	-18.6	-22.7	-13.9
-10.8	-24.7	-17.5	-17.0	-24.2	-23.7	-10.3	-14.4	- 5.7
4.6	2.6	14.9	-10.8	- 6.7	- 8.2	-14.9	3.6	22.7

TRIGO

La primera matriz esta formada por las desviaciones en porcentaje, de la media total del rendimiento de maíz obtenida.

Los restantes representan lo mismo pero para avena y trigo respectivamente.

Para el caso del maíz se obtuvo que los parámetros del modelo estrictamente aditivo eran los siguientes:

$$\hat{\mu} = -8.7212$$

$\hat{\alpha}_1 = 7.5768$	$\hat{\beta}_1 = 0.3394$
$\hat{\alpha}_2 = 4.8101$	$\hat{\beta}_2 = 7.7758$
$\hat{\alpha}_3 = 9.4323$	$\hat{\beta}_3 = 3.2303$
$\hat{\alpha}_4 = 2.7879$	$\hat{\beta}_4 = -0.2424$
$\hat{\alpha}_5 = 4.7323$	$\hat{\beta}_5 = -5.1606$
$\hat{\alpha}_6 = -7.1222$	$\hat{\beta}_6 = -3.1879$
$\hat{\alpha}_7 = -9.8343$	$\hat{\beta}_7 = -2.9242$
$\hat{\alpha}_8 = -10.7566$	$\hat{\beta}_8 = -2.7879$
$\hat{\alpha}_9 = -8.3343$	$\hat{\beta}_9 = 2.9576$

$$\hat{\alpha}_{10} = -3.4677$$

$$\hat{\alpha}_{11} = 10.1768$$

La suma de cuadrados de los residuos obtenidos, después de sustituir los estimadores arriba apuntados, resultó ser igual a 8498.5422; la aplicación de las pruebas de hipótesis a la matriz de residuos únicamente resultó en rechazo de la hipótesis de independencia de los residuos a un nivel de significancia del 10%. El que esta hipótesis haya sido rechazada motivó la inclusión de un término para explicar el efecto de interacción. Los estimadores de los parámetros que componen dicho término fueron:

$$\hat{\theta}_1^2 = 3497.86699$$

$$u_1 = -0.06258$$

$$u_2 = 0.0382$$

$$u_3 = -0.06717$$

$$u_4 = -0.37186$$

$$u_5 = -0.33709$$

$$u_6 = 0.12015$$

$$u_7 = -0.37387$$

$$u_8 = -0.15874$$

$$u_9 = 0.45068$$

$$u_{10} = 0.55102$$

$$u_{11} = 0.21757$$

$$\hat{v}_1 = 0.38530$$

$$\hat{v}_2 = 0.27183$$

$$\hat{v}_3 = 0.25858$$

$$\hat{v}_4 = 0.33753$$

$$\hat{v}_5 = -0.00916$$

$$\hat{v}_6 = -0.22594$$

$$\hat{v}_7 = -0.70058$$

$$\hat{v}_8 = -0.20624$$

$$\hat{v}_9 = -0.11131$$

La inclusión de este término trajo como resultado una reducción de la suma de cuadrados, la que alcanzó el valor de 5000.6763417;

además ninguna hipótesis fue rechazada con lo que se supone que el modelo que explica los datos de una manera satisfactoria, de acuerdo a los criterios establecidos, es un modelo que incluye un sólo término como explicación del efecto de interacción existente. Este hecho hizo suponer que existía cierta interacción entre los renglones y las columnas en que fueron arregladas las parcelas lo cual sería apoyado por los resultados obtenidos en los otros grupos de datos. Así pues se procedió a realizar un análisis semejante en los restantes grupos de datos. Para el rendimiento de avena los estimadores obtenidos para los parámetros del modelo estrictamente aditivo fueron:

$$\hat{\mu} = -8.5707$$

$$\begin{aligned}\hat{\alpha}_1 &= 23.2263 \\ \hat{\alpha}_2 &= 18.3374 \\ \hat{\alpha}_3 &= 19.0263 \\ \hat{\alpha}_4 &= 4.7040 \\ \hat{\alpha}_5 &= -0.3737 \\ \hat{\alpha}_6 &= -16.4404 \\ \hat{\alpha}_7 &= -10.6515 \\ \hat{\alpha}_8 &= -13.2848 \\ \hat{\alpha}_9 &= -15.2293 \\ \hat{\alpha}_{10} &= -8.9071 \\ \hat{\alpha}_{11} &= -0.4071\end{aligned}$$

$$\begin{aligned}\hat{\beta}_1 &= -26.1384 \\ \hat{\beta}_2 &= 2.0707 \\ \hat{\beta}_3 &= 10.2525 \\ \hat{\beta}_4 &= 18.1889 \\ \hat{\beta}_5 &= 2.2253 \\ \hat{\beta}_6 &= -1.2111 \\ \hat{\beta}_7 &= -1.7929 \\ \hat{\beta}_8 &= -5.3111 \\ \hat{\beta}_9 &= 1.7162\end{aligned}$$

La suma de cuadrados de los residuos obtenidos de la aplicación de estos valores fué igual a 16561.013939. Las pruebas de las hipótesis resultaron todas en no rechazo de ninguna de las hi-

nótesis. De lo anterior se concluye que el comportamiento de los errores obtenidos es bueno si se supone un modelo estrictamente aditivo, sin embargo el tamaño de la suma de cuadrados de los residuos es aún muy grande lo que, como se verá más adelante, puede ser indicio de que existe un efecto de interacción pequeño.

La matriz de datos que reporta el rendimiento de trigo para el siguiente año tiene el inconveniente de haber sido tomada en un año en que los fenómenos meteorológicos tuvieron grandes cambios, lo que afectó los rendimientos; esto puede explicar la diferencia tan grande en el comportamiento de los datos. Mediante la aplicación del programa el modelo que, según los criterios usados, mejor se ajusta los datos es un modelo que contiene cuatro términos multiplicativos para explicar el efecto de interacción. La historia del proceso de este grupo de datos es la siguiente. Los estimadores de los parámetros del modelo estrictamente aditivo son:

$$\hat{\mu} = -14.9394$$

$\hat{\alpha}_1 =$	51.7061	$\hat{\beta}_1 =$	7.3485
$\hat{\alpha}_2 =$	38.5172	$\hat{\beta}_2 =$	20.5121
$\hat{\alpha}_3 =$	40.2828	$\hat{\beta}_3 =$	27.0020
$\hat{\alpha}_4 =$	25.9505	$\hat{\beta}_4 =$	122.5061
$\hat{\alpha}_5 =$	15.3061	$\hat{\beta}_5 =$	23.0939
$\hat{\alpha}_6 =$	7.6616	$\hat{\beta}_6 =$	10.3576
$\hat{\alpha}_7 =$	5.4172	$\hat{\beta}_7 =$	17.5667

$$\begin{aligned}\hat{\alpha}_8 &= 6.0616 \\ \hat{\alpha}_9 &= -0.1273 \\ \hat{\alpha}_{10} &= -189.8384 \\ \hat{\alpha}_{11} &= -1.5384\end{aligned}$$

$$\begin{aligned}\hat{\beta}_8 &= 9.2121 \\ \hat{\beta}_9 &= 17.4121\end{aligned}$$

La suma de cuadrados de los residuos es igual a 2628175.64077.

Las hipótesis rechazadas fueron:

Homoscedasticidad	(Por renglones y por columnas al 0.01)
Normalidad	(Al 0.10, 0.05, 0.01, usando Ji cuadrada)
Simetría	(0.10, 0.02)
Kurtosis	(0.10, 0.02).

Con la inclusión de un término multiplicativo para la interacción se tuvo que los estimadores fueron los siguientes:

$$\hat{\sigma}_1^2 = 2615029.9695$$

$\hat{u}_{11} = -0.082246$	$\hat{v}_{11} = 0.12103$
$\hat{u}_{12} = -0.09146$	$\hat{v}_{12} = 0.10367$
$\hat{u}_{13} = -0.02299$	$\hat{v}_{13} = 0.10392$
$\hat{u}_{14} = -0.0968$	$\hat{v}_{14} = -0.94253$
$\hat{u}_{15} = -0.10751$	$\hat{v}_{15} = 0.11736$
$\hat{u}_{16} = -0.10728$	$\hat{v}_{16} = 0.1261$
$\hat{u}_{17} = -0.09939$	$\hat{v}_{17} = 0.12173$
$\hat{u}_{18} = -0.10308$	$\hat{v}_{18} = 0.12729$
$\hat{u}_{19} = -0.08593$	$\hat{v}_{19} = 0.12152$
$\hat{u}_{110} = 0.95212$	
$\hat{u}_{111} = -0.02622$	

La suma de cuadrados de los residuos fue igual a 13145.6713.

Esto representa una reducción muy grande relativa al modelo adi-

tivo.

Las hipótesis rechazadas fueron:

Normalidad (0.10, 0.05)

Independencia (0.10, 0.02)

Con el segundo término multiplicativo para la interacción los resultados obtenidos fueron los siguientes. Los estimadores de los parámetros incluídos resultaron ser:

$$\hat{\theta}_2^2 = 4019.9959$$

$$\hat{u}_{21} = -0.27793$$

$$\hat{u}_{22} = 0.0162$$

$$\hat{u}_{23} = -0.16643$$

$$\hat{u}_{24} = -0.07734$$

$$\hat{u}_{25} = 0.00575$$

$$\hat{u}_{26} = 0.64464$$

$$\hat{u}_{27} = -0.36355$$

$$\hat{u}_{28} = -0.37$$

$$\hat{u}_{29} = 0.41973$$

$$\hat{u}_{210} = 0.00167$$

$$\hat{u}_{211} = 0.16725$$

$$\hat{v}_{21} = 0.73429$$

$$\hat{v}_{22} = 0.16557$$

$$\hat{v}_{23} = -0.0045$$

$$\hat{v}_{24} = -0.00421$$

$$\hat{v}_{25} = -0.26657$$

$$\hat{v}_{26} = -0.22283$$

$$\hat{v}_{27} = -0.44413$$

$$\hat{v}_{28} = -0.21808$$

$$\hat{v}_{29} = 0.26046$$

La suma de cuadrados de los residuos resultó ser 9125.67539.

Las hipótesis rechazadas fueron:

Homoscedasticidad (Por renglones)

Normalidad (0.10, 0.05)

Simetría (0.10)

Independencia (0.10)

Kurtosis (0.10, 0.02)

Para el tercer término multiplicativo los estimadores de los parámetros fueron:

$$\hat{\theta}_3^2 = 3700.86987$$

$$\hat{u}_{31} = -0.14125$$

$$\hat{v}_{31} = 0.16714$$

$$\hat{u}_{32} = -0.3991$$

$$\hat{v}_{32} = -0.0407$$

$$\hat{u}_{33} = -0.09276$$

$$\hat{v}_{33} = -0.05772$$

$$\hat{u}_{34} = 0.21373$$

$$\hat{v}_{34} = -0.00293$$

$$\hat{u}_{35} = 0.81293$$

$$\hat{v}_{35} = 0.71261$$

$$\hat{u}_{36} = -0.00574$$

$$\hat{v}_{36} = 0.24791$$

$$\hat{u}_{37} = 0.07648$$

$$\hat{v}_{37} = -0.21491$$

$$\hat{u}_{38} = -0.16384$$

$$\hat{v}_{38} = 0.51162$$

$$\hat{u}_{39} = -0.05057$$

$$\hat{v}_{39} = -0.29978$$

$$\hat{u}_{310} = 0.01493$$

$$\hat{u}_{311} = -0.26481$$

La suma de cuadrados de los residuos fué igual a 5424.80551.

Las hipótesis rechazadas fueron:

Homoscedasticidad (Por renglones y por columnas)

Kurtosis (0.10, 0.02)

Con el cuarto y último término multiplicativo los estimadores de los parámetros fueron:

$$\hat{\theta}_4^2 = 2814.99131$$

$$\hat{u}_{41} = 0.6184$$

$$\hat{v}_{41} = -0.34011$$

$$\hat{u}_{42} = 0.08638$$

$$\hat{v}_{42} = 0.76162$$

$$\hat{u}_{43} = 0.29633$$

$$\hat{v}_{43} = 0.39283$$

$$\hat{u}_{44} = -0.25594$$

$$\hat{v}_{44} = -0.0199$$

$$\hat{u}_{45} = 0.08617$$

$$\hat{v}_{45} = -0.09921$$

$$\hat{u}_{46} = 0.27431$$

$$\hat{v}_{46} = -0.07093$$

$$\hat{u}_{47} = -0.31234$$

$$\hat{v}_{47} = -0.2449$$

$$\hat{u}_{48} = -0.15347$$

$$\hat{v}_{48} = -0.1528$$

$$\hat{u}_{49} = -0.16208$$

$$\hat{v}_{49} = -0.2266$$

$$\hat{u}_{410} = -0.00156$$

$$\hat{u}_{411} = -0.47619$$

La suma de cuadrados de los residuos fué igual a 2609.4142.

Esta vez ninguna de las hipótesis fué rechazada.

Parecería que si el terreno tuviese alguna tendencia de fertilidad todos los modelos ajustados a los distintos grupos de datos deberían parecerse, sin embargo es notable la diferencia existente entre el tercero y los dos restantes; como se explicó antes aquel grupo de datos fué tomado en un año en que los cambios climatológicos resultaron particularmente serios, por lo que el tercer modelo resulta de poca aplicabilidad a los rendimientos de otros años. Además la configuración muy particular del experimento (1o. Maíz, 2o. Avena, 3o. Trigo), puede tener

influencia en los resultados obtenidos del mismo. Es decir parece ser necesario que el experimento se repita en las mismas condiciones para poder afirmar algo positivo o tomar otra parte de la superficie cultivada en el mismo experimento con el objeto de poder hacer alguna consideración al respecto. Esto es una indicación clara de que los rendimientos siguen un patrón de fertilidad del suelo pero este patrón cambia al considerar diferentes cultivos y diferentes condiciones ambientales.

5) Otro grupo de datos a los que se aplicó el análisis fué tomado del libro de V.G. Panse y P.V. Sukhatme (20), y consiste de un ensayo de uniformidad ("cosecha de un terreno con un cultivo particular con un tratamiento uniforme, dividiendo el terreno en pequeñas unidades, y cosechando y anotando separadamente la producción de cada una de estas unidades"), con აღო-
dón efectuado en el Institute of Plant Industry, de Indore.

La matriz de datos que se presenta a continuación representa sólo una parte del experimento.

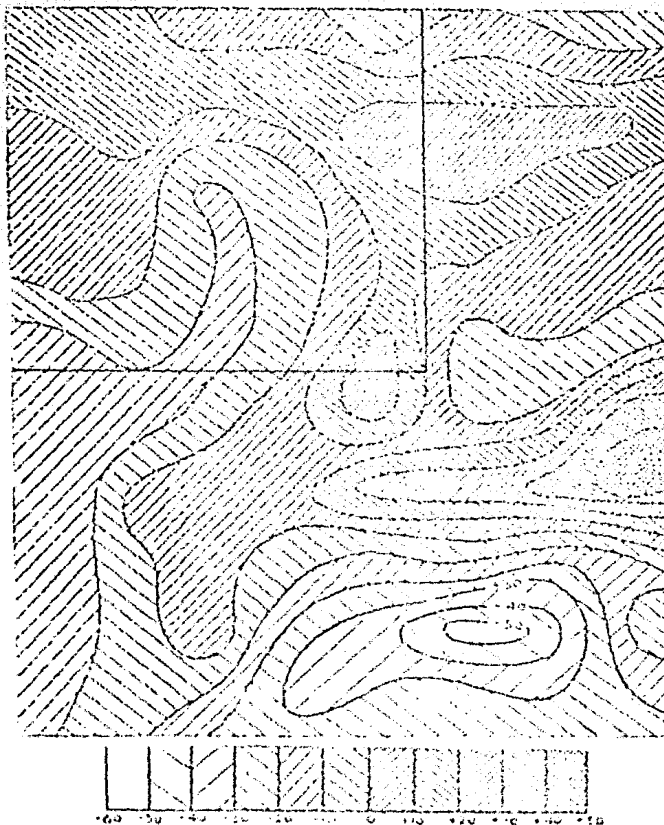
Esta matriz se reproduce en la tabla E 5.

TABLA E 5

ENSAYO DE UNIFORMIDAD CON ALGODON MALVI, 1933-34, INSTITUTE OF PLANT INDUSTRY, INDORE

93	95	111	116	97	82	80	131	80	102	97	95	87	136	118	89	57	106	118	97
49	63	50	115	83	78	69	93	85	84	64	69	51	44	115	61	54	54	36	88
89	90	95	129	97	133	72	115	97	76	91	106	40	147	119	66	122	63	72	34
85	91	115	107	99	77	76	63	70	71	59	78	58	85	61	104	84	104	59	56
63	95	68	69	101	120	96	100	70	110	88	75	101	61	106	113	93	62	92	109
86	82	113	74	117	28	83	81	109	91	53	108	83	74	129	89	79	89	75	114
62	85	112	30	79	92	73	103	89	77	67	88	97	50	101	111	100	96	66	113
94	76	119	75	125	96	95	74	90	63	68	66	96	66	85	74	160	148	84	110
45	96	115	130	99	77	104	116	64	77	32	88	67	76	72	87	59	123	67	59
51	79	111	79	84	84	67	63	66	40	56	80	58	68	52	76	87	72	91	81
81	80	72	102	104	81	42	76	85	80	56	74	93	88	60	111	128	124	117	118
74	51	66	86	71	73	98	87	61	111	81	61	73	61	63	85	115	80	82	76
35	89	112	88	105	76	76	78	72	92	59	81	90	52	82	68	88	94	68	106
94	64	101	90	69	101	87	70	63	75	63	88	74	81	68	87	82	80	64	105
35	55	88	87	73	119	60	76	60	73	77	44	83	82	74	52	97	65	59	75
82	83	102	80	88	57	81	55	75	69	70	84	59	73	77	81	98	74	95	110
46	54	83	60	57	75	73	50	66	31	56	66	43	70	59	35	90	62	89	96
71	68	39	80	73	86	64	89	101	88	63	83	74	78	74	50	126	98	104	109
58	52	85	99	72	96	67	86	37	71	62	48	63	64	106	90	91	86	83	104
65	45	77	78	93	60	94	82	72	81	84	64	75	97	114	109	111	74	107	86

En el mismo libro se presenta un "mapa de contorno de fertilidad", donde resulta obvio que tanto los renglones como las columnas interactúan, es decir los efectos principales por renglón no son constantes, ocurriendo lo mismo en las columnas. El mapa aparece reproducido a continuación



Mapa de contorno de fertilidad de los datos del ensayo de uniformidad de siembra.

En él se marcan los límites del terreno considerado en el grupo de datos al que fué aplicado el programa.

Los estimadores de los parámetros que representan los efectos principales son los siguientes:

$$\hat{\mu} = 83.7275$$

$$\hat{\alpha}_1 = 15.6225$$

$$\hat{\alpha}_2 = -13.4775$$

$$\hat{\alpha}_3 = 8.4225$$

$$\hat{\alpha}_4 = -3.6275$$

$$\hat{\alpha}_5 = 5.9225$$

$$\hat{\alpha}_6 = 4.1225$$

$$\hat{\alpha}_7 = 3.6725$$

$$\hat{\alpha}_8 = 9.4725$$

$$\hat{\alpha}_9 = -1.0775$$

$$\hat{\alpha}_{10} = -11.4775$$

$$\hat{\alpha}_{11} = 4.8725$$

$$\hat{\alpha}_{12} = -5.9775$$

$$\hat{\alpha}_{13} = -0.6775$$

$$\hat{\alpha}_{14} = 29.3725$$

$$\hat{\alpha}_{15} = -9.5275$$

$$\hat{\alpha}_{16} = -4.0775$$

$$\hat{\alpha}_{17} = -20.6775$$

$$\hat{\beta}_1 = -11.3275$$

$$\hat{\beta}_2 = -8.9775$$

$$\hat{\beta}_3 = 8.3225$$

$$\hat{\beta}_4 = 7.4725$$

$$\hat{\beta}_5 = 5.5725$$

$$\hat{\beta}_6 = 0.3225$$

$$\hat{\beta}_7 = -5.9275$$

$$\hat{\beta}_8 = 0.6725$$

$$\hat{\beta}_9 = -8.1275$$

$$\hat{\beta}_{10} = -5.6725$$

$$\hat{\beta}_{11} = -16.4275$$

$$\hat{\beta}_{12} = -6.4275$$

$$\hat{\beta}_{13} = 22.8225$$

$$\hat{\beta}_{14} = -6.0775$$

$$\hat{\beta}_{15} = 3.0225$$

$$\hat{\beta}_{16} = -1.8275$$

$$\hat{\beta}_{17} = 12.3225$$

$$\hat{\alpha}_{18} = -2.8275$$

$$\hat{\beta}_{18} = 3.9725$$

$$\hat{\alpha}_{19} = -7.7275$$

$$\hat{\beta}_{19} = -2.3275$$

$$\hat{\alpha}_{20} = -0.3275$$

$$\hat{\beta}_{20} = 8.5725$$

La suma de cuadrados de los residuos resultó ser igual a 536191.4025.

Las hipótesis rechazadas fueron:

Normalidad (Ji-Cuadrada, 0.10, 0.05, 0.01)

Simetría (0.10, 0.02)

Kurtosis (0.10, 0.02)

El rechazo de estas hipótesis motivó la inclusión de un término del tipo $\theta_{u_i v_j}$; los estimadores de los parámetros fueron los siguientes:

$$\hat{\theta}_1^2 = 411689.2813222$$

$$\hat{u}_{11} = -0.05857$$

$$\hat{v}_{11} = -0.02827$$

$$\hat{u}_{12} = -0.06914$$

$$\hat{v}_{12} = -0.06505$$

$$\hat{u}_{13} = -0.12441$$

$$\hat{v}_{13} = -0.033$$

$$\hat{u}_{14} = -0.07018$$

$$\hat{v}_{14} = -0.05683$$

$$\hat{u}_{15} = -0.01744$$

$$\hat{v}_{15} = -0.07832$$

$$\hat{u}_{16} = -0.04653$$

$$\hat{v}_{16} = -0.01896$$

$$\hat{u}_{17} = -0.01918$$

$$\hat{v}_{17} = -0.03262$$

$$\hat{u}_{18} = -0.03224$$

$$\hat{v}_{18} = -0.07201$$

$$\hat{u}_{19} = -0.06247$$

$$\hat{v}_{19} = -0.06870$$

$$\hat{u}_{110} = -0.05765$$

$$\hat{u}_{111} = -0.03053$$

$$\hat{u}_{112} = -0.04341$$

$$\hat{u}_{113} = -0.02386$$

$$\hat{u}_{114} = 0.96903$$

$$\hat{u}_{115} = -0.02016$$

$$\hat{u}_{116} = 0.06785$$

$$\hat{u}_{117} = -0.06674$$

$$\hat{u}_{118} = -0.04982$$

$$\hat{u}_{119} = -0.05600$$

$$\hat{u}_{120} = -0.05287$$

$$\hat{v}_{110} = -0.04824$$

$$\hat{v}_{111} = -0.05495$$

$$\hat{v}_{112} = -0.0348$$

$$\hat{v}_{113} = 0.97114$$

$$\hat{v}_{114} = -0.05267$$

$$\hat{v}_{115} = -0.07982$$

$$\hat{v}_{116} = -0.0347$$

$$\hat{v}_{117} = -0.0669$$

$$\hat{v}_{118} = -0.05473$$

$$\hat{v}_{119} = -0.07341$$

$$\hat{v}_{120} = -0.01617$$

Con una suma de cuadrados igual a 124502.120655.

La única hipótesis rechazada fué la siguiente:

Kurtosis (0.10, 0.02)

Lo que propició la inclusión de otro término multiplicativo para la explicación de la interacción; los estimadores de los parámetros considerados en este término fueron:

$$\hat{\theta}_2^2 = 27708.3075459$$

$$\hat{u}_{21} = 0.33403$$

$$\hat{u}_{22} = 0.3212$$

$$\hat{u}_{23} = 0.506$$

$$\hat{u}_{24} = 0.06267$$

$$\hat{v}_{21} = -0.06925$$

$$\hat{v}_{22} = 0.06354$$

$$\hat{v}_{23} = -0.02216$$

$$\hat{v}_{24} = 0.34258$$

$\hat{u}_{25} = 0.04373$	$\hat{v}_{25} = -0.08131$
$\hat{u}_{26} = -0.0429$	$\hat{v}_{26} = 0.12726$
$\hat{u}_{27} = -0.06595$	$\hat{v}_{27} = 0.01489$
$\hat{u}_{28} = -0.46116$	$\hat{v}_{28} = 0.32595$
$\hat{u}_{29} = 0.2585$	$\hat{v}_{29} = -0.05737$
$\hat{u}_{210} = -0.09003$	$\hat{v}_{210} = 0.10497$
$\hat{u}_{211} = -0.33146$	$\hat{v}_{211} = 0.08738$
$\hat{u}_{212} = -0.05363$	$\hat{v}_{212} = 0.10279$
$\hat{u}_{213} = -0.08419$	$\hat{v}_{213} = 0.01272$
$\hat{u}_{214} = 0.05844$	$\hat{v}_{214} = 0.29172$
$\hat{u}_{215} = 0.098$	$\hat{v}_{215} = 0.26634$
$\hat{u}_{216} = -0.17952$	$\hat{v}_{216} = -0.05868$
$\hat{u}_{217} = -0.16371$	$\hat{v}_{217} = -0.45641$
$\hat{u}_{218} = -0.19248$	$\hat{v}_{218} = -0.30781$
$\hat{u}_{219} = 0.00732$	$\hat{v}_{219} = -0.24742$
$\hat{u}_{220} = -0.02487$	$\hat{v}_{220} = -0.43273$

La suma de cuadrados de los residuos alcanzó el valor de 97793.8131.

Con hipótesis rechazadas, las siguientes:

Normalidad (0.10, 0.05, 0.01)

Kurtosis (0.10)

Los estimadores del tercer término multiplicativo fueron los siguientes:

$$\hat{\sigma}_3^2 = 21307.0868892$$

$\hat{u}_{31} = 0.05993$	$\hat{v}_{31} = -0.5498$
$\hat{u}_{32} = -0.01911$	$\hat{v}_{32} = 0.24252$
$\hat{u}_{33} = -0.29119$	$\hat{v}_{33} = 0.48614$
$\hat{u}_{34} = 0.29149$	$\hat{v}_{34} = 0.08163$
$\hat{u}_{35} = -0.20065$	$\hat{v}_{35} = 0.20885$
$\hat{u}_{36} = 0.37377$	$\hat{v}_{36} = -0.37978$
$\hat{u}_{37} = 0.17063$	$\hat{v}_{37} = 0.08479$
$\hat{u}_{38} = 0.05414$	$\hat{v}_{38} = -0.03151$
$\hat{u}_{39} = 0.48417$	$\hat{v}_{39} = 0.00633$
$\hat{u}_{310} = 0.1156$	$\hat{v}_{310} = -0.09113$
$\hat{u}_{311} = -0.06157$	$\hat{v}_{311} = -0.31913$
$\hat{u}_{312} = -0.21736$	$\hat{v}_{312} = 0.18168$
$\hat{u}_{313} = 0.21659$	$\hat{v}_{313} = -0.00604$
$\hat{u}_{314} = -0.00588$	$\hat{v}_{314} = -0.20317$
$\hat{u}_{315} = -0.28581$	$\hat{v}_{315} = -0.05452$
$\hat{u}_{316} = 0.07974$	$\hat{v}_{316} = 0.16103$
$\hat{u}_{317} = -0.12516$	$\hat{v}_{317} = -0.39549$
$\hat{u}_{318} = -0.24209$	$\hat{v}_{318} = 0.29578$
$\hat{u}_{319} = -0.10623$	$\hat{v}_{319} = -0.2043$
$\hat{u}_{320} = -0.19101$	$\hat{v}_{320} = -0.0087$

La suma de cuadrados de los residuos obtenidos al sustituir estos estimadores descendió hasta 76486.766222.

La hipótesis rechazada fué:

Normalidad (0.10, 0.05)

Entonces se incluyó un cuarto término, con lo que se logró que ninguna de las hipótesis fuese rechazada; los estimadores de los parámetros que componen dicho término fueron los siguientes:

$$\hat{\theta}_4^2 = 18400.91374743$$

$\hat{u}_{41} = -0.03637$	$\hat{v}_{41} = 0.12066$
$\hat{u}_{42} = -0.37919$	$\hat{v}_{42} = 0.06741$
$\hat{u}_{43} = 0.40973$	$\hat{v}_{43} = 0.27358$
$\hat{u}_{44} = 0.30708$	$\hat{v}_{44} = 0.18766$
$\hat{u}_{45} = -0.38116$	$\hat{v}_{45} = 0.0039$
$\hat{u}_{46} = -0.36171$	$\hat{v}_{46} = 0.24898$
$\hat{u}_{47} = -0.16988$	$\hat{v}_{47} = -0.09739$
$\hat{u}_{48} = 0.26277$	$\hat{v}_{48} = -0.11825$
$\hat{u}_{49} = 0.18194$	$\hat{v}_{49} = -0.06469$
$\hat{u}_{410} = 0.22148$	$\hat{v}_{410} = -0.34737$
$\hat{u}_{411} = 0.12148$	$\hat{v}_{411} = -0.07114$
$\hat{u}_{412} = -0.09167$	$\hat{v}_{412} = -0.00633$
$\hat{u}_{413} = -0.06233$	$\hat{v}_{413} = -0.00305$
$\hat{u}_{414} = 0.04045$	$\hat{v}_{414} = 0.36386$
$\hat{u}_{415} = 0.20218$	$\hat{v}_{415} = -0.42572$
$\hat{u}_{416} = -0.04222$	$\hat{v}_{416} = -0.2031$
$\hat{u}_{417} = 0.12217$	$\hat{v}_{417} = 0.26236$

$$\hat{u}_{418} = -0.04135$$

$$\hat{v}_{418} = 0.235$$

$$\hat{u}_{419} = -0.11613$$

$$\hat{v}_{419} = -0.00646$$

$$\hat{u}_{420} = -0.18726$$

$$\hat{v}_{420} = -0.42039$$

Con una suma de cuadrados, para los residuos, igual a 58085.812476.

Comparando este resultado con el que se hubiera obtenido por el método de Mandel que ve que el número de términos incluido en el modelo es cuatro veces mayor; el estimador de la varianza para el modelo que incluye cuatro términos sería de 868.89, mientras que el obtenido con el modelo que incluye un solo término sería de 1863.39. Los valores requeridos para la aplicación del criterio de Mandel aparecen ordenados en la tabla de análisis de la varianza incluida en seguida.

TABLA 5.1

ANÁLISIS DE LA VARIANZA PARA LOS TÉRMINOS DE LA INTERACCIÓN

FUENTE	SUMA DE CUADRADOS	GRADOS DE LIBERTAD	CUADRADOS MEDIOS
$\theta_{111}^{u_{11}v_{1j}}$	411689.28	74.25	5558.13
$\theta_{221}^{u_{21}v_{2j}}$	26708.31	62.24	429.11
$\theta_{331}^{u_{31}v_{3j}}$	21307.09	53.50	398.26
$\theta_{441}^{u_{41}v_{4j}}$	18400.91	48.13	382.31
errores	58085.81	66.85	868.89

6) Los siguientes grupos de datos son los obtenidos mediante una simulación en la computadora. El mecanismo generador de números aleatorios de la computadora produce números o variables aleatorias independientes que tienen una distribución uniforme en el intervalo $(0,1)$. Haciendo uso de un teorema que dice "sean U y V variables aleatorias independientes y uniformemente distribuidas en el intervalo $(0,1)$, entonces las variables aleatorias:

$$\begin{aligned} X_1 &= (-2 \log_e U)^{1/2} \sin (2\pi V) \\ X_2 &= (-2 \log_e U)^{1/2} \cos (2\pi V) \end{aligned}$$

Son independientes y tienen una distribución normal con media cero y varianza uno", y de las variables aleatorias generadas en la computadora, se generaron variables aleatorias con distribución normal con media cero y desviación estándar σ . Estos errores aleatorios fueron sumados a distintos modelos (distintos en el efecto de interacción), y a esos datos simulados se les aplicó el programa de computadora.

Como se indicó anteriormente, el modelo usado en la simulación fué el siguiente:

$$Y_{ij} = 2(i-1) + 3j + \left((i-3)(j-1) \right)^n + \varepsilon_{ij}$$

$$i=1, \dots, 5$$

$$j=1, \dots, 4$$

$$\text{con } \varepsilon_{ij} \sim N(0, \sigma^2)$$

A r se le asignaron los valores 1,3,5,7,9 y σ tomó los valores 1/20, 6/20, 11/20, 16/20, 21/20; es decir, se generaron 25 matrices de datos.

Los resultados pueden resumirse en una tabla donde se indica cuáles de las hipótesis hechas acerca de los errores se rechazaron habiendo incluido k términos para explicación del efecto de interacción.

Además se incluyen los valores de la suma de cuadrados de los residuos y el valor del estimador de σ^2 , para el mismo número de términos.

σ	r	NO. DE TERMINOS	SUMA DE CUADRADOS	$\hat{\sigma}^2$	$\hat{\sigma}$	HIPOTESIS RECHAZADAS
.05	1	0	50.62431	2.057	-	Ninguna
.30	1	0	48.1089	2.0022	-	Ninguna
.55	1	0	52.5376	2.092	-	Ninguna
.80	1	0	79.90938	2.942	-	Ninguna
1.05	1	0	23.95009	1.8291	-	Ninguna
.05	3	0	61147.2634	71.383	-	N5, K2
		1	0.0204	.074	61143.243	I1
		2	0.00486	.073	.0155411	II, I2, K1
		3	$(1.457) \times 10^{-17}$	-	.0048584	III, N5, S2

σ	r	NO. DE TERMINOS	SUMA DE CUADRADOS	$\hat{\sigma}^2$	$\hat{\sigma}$	HIPOTESIS RECHAZADAS
0.30	3	0	60920.57	71.25	-	N5, K2
		1	0.407855	.333	60920.1628	I2
		2	.04155284	.25	0.3663	Ninguna
.55	3	0	60885.3151	71.23	-	N5, K2
		1	1.55388	.65	60883.1612	Ninguna
.80	3	0	61371.52956	71.51	-	N5, K2
		1	7.67	1.442	61363.8595	I 2
		2	0.466178	0.84	7.203824	N4, HR
		3	1.1778×10^{-16}	-	0.466178	HC, N5, S2
1.05	3	0	61192.4799	71.40	-	N5, K2
		1	1.94122	.72	61190.5387	Ninguna
.05	5	0	33703930.187	1675.90	-	HR, HC, N5, S2, K2
		1	4.587266	1.1178	33703925.599	HR, N3, K2
		2	0.00068326	0.1211	4.5776	N5, I2, K1
		3	$(3.3615) \times 10^{-15}$	-	0.00068	N5
.30	5	0	84112214.767	2647.51	-	HC, N5, K2

σ	r	NO. DE TERMINOS	SUMA DE CUADRADOS	$\hat{\theta}^2$	$\hat{\sigma}$	HIPOTESTIS RECHAZADAS
		1	0.142344261	.1969	84112214.625	I2
		2	0.027256766	.2032	0.115087	Ninguna
.55	5	0	84132634.0647	2647.75	-	HC, N5, K2
		1	0.993415382	.52	84132633.072	N3
		2	0.2100978	.564	0.783317	N4, I2
		3	$(4.3831) \times 10^{-14}$	-	0.2100978	HC, N5, S2
.80	5	0	84120425.8584	2647.56	-	HC, N5, K2
		1	1.404441	.618	84120424.454	N3, S1, K2
		2	.2044	.556	1.20003735	I2
		3	$(7.0726) \times 10^{-14}$	-	0.204404	HC, N5, S2, K2
1.05	5	0	84132634.0647	2647.83	-	HC, N5, K2
		1	0.993415	.520	84132633.072	N3, S1
		2	.2100978	.564	0.7833176	N4, I2
		3	$(4.3831) \times 10^{-14}$	-	0.21009779	HC, N5, S2
.05	7	0	$(1.1328) \times 10^{11}$	$0(10^5)$	-	HC, N5, K2
		1	10.746	1.71	$0(10^{11})$	-

σ	r	NO. DE TERMINOS	SUMA DE CUADRADOS	$\hat{\theta}^2$	$\hat{\theta}$	HIPOTESIS RECHAZADAS
.30	7	0	$(1.1328) \times 10^{11}$	$0(10^5)$	-	HC, N5, K2
		1	13.19405	1.89	$0(10^{11})$	Ninguna
.55	7	0	$(1.13331) \times 10^{11}$	$0(10^5)$	-	HC, N5, K2
		1	2.145226	.764	$0(10^{11})$	N4
		2	.008469	.1132	2.13675787	N1, I2
		3	$(2.8896) \times 10^{-11}$	-	0.0084689	N5, S2
.80	7	0	$(1.1333) \times 10^{11}$	$0(10^5)$	-	HC, N5, K2
		1	2.544475	.832	$0(10^{11})$	Ninguna
1.05	7	0	$(1.13332) \times 10^{11}$	$0(10^5)$	-	HC, N5, K2
		1	2.6159744	.844	$0(10^{11})$	N5, I2,
		2	.57813344	.935	2.037841	HR, N5, S2, K2
		3	$(2.4488) \times 10^{-10}$	-	0.57813344	N5, S2
.05	9	0	$(1.497964) \times 10^{14}$	$0(10^6)$	-	HC, N5, K2
		1	0.016686518	0.0674	$0(10^{14})$	N5, S1, K1
		2	0.00232039	0.0592	.014366129	N2, I2, K1
		3	0.000000038	-	.00232035	N5, S2

σ	r	NO. DE TERMINOS	SUMA DE CUADRADOS	$\hat{\sigma}^2$	$\hat{\sigma}$	HIPOTESIS RECHAZADAS
.30	9	0	$(1.4979648) \times 10^{14}$	$0(10^6)$	-	HC, N5, K2
		1	0.507979325	.37203	$0(10^{13})$	N5
		2	0.125681986	.4363	0.3822997	N3, I2
		3	0.00000003	-	0.125682	HC, N5, S2
.55	9	0	$(1.49496487) \times 10^{14}$	$0(10^6)$	-	HC, N5, K2
		1	18.639114	2.2536	$0(10^{13})$	HR, N2, S2, K2
		2	0.450466909	.8261	18.188647	N3, I2
		3	0.000000042	-	0.4504668	HC, N5, S2
.80	9	0	$(1.497964456) \times 10^{14}$	$0(10^6)$	-	HC, N5, K2
		1	10.123865964	1.6609	$0(10^{13})$	Ninguna
1.05	9	0	$(1.49796476) \times 10^{14}$	$0(10^6)$	-	HC, N5, K2
		1	8.6271067	1.533	$0(10^{13})$	N4, K1
		2	0.310009369	.6853	8.3170972	I2
		3	0.000000003	-	0.3100093	HC, N5, S2, K2

Donde por ejemplo KR quiere decir Kurtosis en sus R niveles de significancia mayores. Por otra parte $0(10^r)$ representa un número cuyo orden de magnitud era de 10^r .

Cómo puede observarse de la tabla, cuando los efectos de interacción son pequeños, por ejemplo para $r=1$ no se detectó ningún efecto de interacción (puede haber ocurrido en los rendimientos de avena del ejemplo 4), en contraste con los efectos de interacción grandes, por ejemplo para $r=9$.

Con el objeto de aclarar un poco más los puntos de interés se desarrollan en seguida los resultados obtenidos por las simulaciones.

Como se indica en el cuadro que resume los resultados, en el caso en el que el modelo era el siguiente:

$$Y_{ij} = 2(i-1) + 3j + (i-3)(j-1) + \varepsilon_{ij}$$

$$i=1, \dots, 5$$

$$j=1, \dots, 4$$

Con $\varepsilon_{ij} \sim N(0, 1/400)$

El modelo ajustado fué un modelo estrictamente aditivo.

Dando por resultado una suma de cuadrados del error igual a 50.62431 con un estimador de la desviación estándar dado por 2.057 cuando la desviación estándar real era de $0.05 = 1/20$; es decir la varianza real fué sobrestimada y el modelo real no fué descubierto ya que no detectó el efecto de interacción, y es esto lo que produce sobre estimaciones de σ_1^2 .

Algo semejante ocurrió con los modelos para los cuales las desviaciones estándar de los errores eran 0.30, 0.55, 0.80, 1.05

σ	r	NO. DE TERMINOS	SUMA DE CUADRADOS	$\hat{\sigma}^2$	$\hat{\sigma}$	HIPOTESIS RECHAZADAS
.30	9	0	$(1.4979648) \times 10^{14}$	$0(10^6)$	-	HC, N5, K2
		1	0.507979325	.37203	$0(10^{13})$	N5
		2	0.125681986	.4363	0.3822997	N3, I2
		3	0.0000003	-	0.125682	HC, N5, S2
.55	9	0	$(1.49496487) \times 10^{14}$	$0(10^6)$	-	HC, N5, K2
		1	18.639114	2.2536	$0(10^{13})$	HR, N2, S2, K2
		2	0.450466909	.8261	18.188647	N3, I2
		3	0.000000042	-	0.4504668	HC, N5, S2
.80	9	0	$(1.497964456) \times 10^{14}$	$0(10^6)$	-	HC, N5, K2
		1	10.123865964	1.6609	$0(10^{13})$	Ninguna
1.05	9	0	$(1.49796476) \times 10^{14}$	$0(10^6)$	-	HC, N5, K2
		1	8.6271067	1.533	$0(10^{13})$	N4, K1
		2	0.310009369	.6853	8.3170972	I2
		3	0.00000003	-	0.3100093	HC, N5, S2, K2

Donde por ejemplo ER quiere decir Kurtosis en sus R niveles de significancia mayores. Por otra parte $0(10^r)$ representa un número cuyo orden de magnitud era de 10^r .

con los estimadores respectivos dados por 2.0022, 2.092, 2.942, 1.8291; las sumas de cuadrados para los modelos fueron 48.1089, 52.5376, 79.90938, 23,95009 respectivamente. Así pues, mientras que las desviaciones estandar eran sobreestimadas, los efectos de interacción no fueron siquiera detectados. Cabe aquí hacer notar que las dimensiones de las matrices de datos son pequeñas con sólo 20 elementos por cada matriz. Además puede verse la forma del modelo, que este es lineal en i y también lineal en j . Esto es, se tiene una interacción de tipo lineal, que no es detectada en el modelo.

Algo diferente ocurrió cuando el modelo real era:

$$Y_{ij} = 2(i-1) + 3j + ((i-3)(j-1))^3 + \varepsilon_{ij}$$

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

Donde a σ se le asignaron los valores de .05, .30, .55, .80, 1.05 y en cada caso los resultados obtenidos, aparte de los estimadores, fueron los que en seguida se reportan.

Para el modelo con $\sigma = .05$ el ajuste de un modelo estrictamente aditivo dió como resultado una suma de cuadrados del error de 61147.2634, con un estimador de la desviación estandar $\hat{\sigma} = 71.3$. Las hipótesis rechazadas con los residuos resultantes fueron la de normalidad a los cinco niveles de significancia y la de kurtosis en sus dos niveles.

Al aumentar un término para el efecto de interacción, el estima-

donde θ^2 alcanzó un valor de 61147.243 con lo que la suma de cuadrados se redujo hasta 0.0204, dado un estimador de la desviación estándar $1/\hat{\sigma} = 0.074$, el cuál se encuentra apenas 0.024 arriba del valor real de σ . La única hipótesis rechazada fué la de independencia de los errores al 10% de significancia. Esto indica que con 5% de significancia, no se hubié-
ra rechazado y se tendría un modelo bastante aceptable.

El aumento de un nuevo término para el efecto de interacción redujo aún más la suma de cuadrados siendo el valor de este igual a 0.00486 pero el estimador de la desviación estándar fué de 0.073; es decir la ganancia al introducir el nuevo término fué pequeña. Las hipótesis rechazadas fueron las de homoscedasticidad por renglones, independencia de los residuos en sus dos niveles y kurtosis al 10%.

El estimador de θ^2 para el tercer término del efecto de interacción fué de 0.00485 y para su estimación se hizo uso de 0.66 grados de libertad, con lo que se agotaron los grados de liber-

1/ Los estimadores de σ se calcularon haciendo uso de los valores dados para los grados de libertad en la tabla 7 del anéndice II; esto se hizo de la siguiente manera:

$$\hat{\sigma}_k = \left\{ \frac{SCD_k}{\nu_{k+1} + \dots + \nu_3} \right\}^{1/2}$$

Donde SCD_k es la suma de cuadrados de los residuos obtenidos después de sumentar k términos y $\nu_1 = 8.33$, $\nu_2 = 2.01$ nótese que

$$\nu_1 + \nu_2 + \nu_3 = 12$$

tad y en consecuencia no fué posible dar un estimador de σ para este modelo; sin embargo la suma de cuadrados de los residuos fué de tan sólo 1.457×10^{-17} . Las hipótesis rechazadas fueron las de homoscedasticidad por columnas, normalidad en sus cinco niveles y simetría en sus dos niveles.

Para el caso en el cual la desviación estandar era igual a 0.30, un modelo aditivo dió como resultado una suma de cuadrados de 60920.57 y las hipótesis rechazadas fueron las de normalidad en sus cinco niveles y kurtosis en sus dos niveles.

El estimador de θ^2 para el primer término de la interacción fué de 60920.1628 y los correspondientes grados de libertad de 8.33.

La suma de cuadrados de los residuos se redujo hasta el valor 0.407855 dando un estimador de la desviación estandar de 0.333.

La hipótesis rechazada fué únicamente la de independencia de los residuos en sus dos niveles de significancia.

En vista del tamaño de $\hat{\theta}^2$ y el de la suma de cuadrados la aplicación del criterio de Mandel habría incluido únicamente un término para la explicación del efecto de interacción, sin embargo la inclusión de otro término para la interacción en el modelo, con un estimador de θ^2 igual a 0.3663 y con 3.01 grados de libertad, trajo como resultado una reducción en el valor del estimador de la desviación estandar hasta el valor de 0.25 y el no rechazo de ninguna de las hipótesis, mejorando notablemente el ajuste

de los datos.

Para el valor de $\sigma = 0.55$, la aplicación de ambos criterios para la selección del modelo habrían coincidido ya que la inclusión de un término para el efecto de interacción, con $\hat{\sigma}^2 = 60833.7612$ y 8.33 grados de libertad, trajo como consecuencia una suma de cuadrados de 1.55388, un estimador de la desviación estandar igual a 0.65, estimado con 3.67 grados de libertad, y el no rechazo de ninguna de las hipótesis. Es preciso hacer notar que las hipótesis rechazadas para los residuos obtenidos después de ajustar un modelo aditivo fueron normalidad en sus cinco niveles y kurtosis en sus dos niveles.

Para el mismo modelo pero con $\sigma = 0.80$, el ajuste de un modelo estrictamente aditivo dió como resultado una suma de cuadrados de 61371.5295 y el rechazo de las hipótesis de normalidad en sus cinco niveles y kurtosis en sus dos niveles.

Al incluir un término para la explicación del efecto de interacción, para el cual $\hat{\sigma}^2 = 61363.8595$, el estimador de σ , la desviación estandar, fué de 1.442, es decir, 0.642 arriba del valor real del parámetro. La única hipótesis rechazada en este caso fué la de independencia en sus dos niveles. La inclusión de un segundo término para el efecto de interacción produjo la reducción de la suma de cuadrados hasta el valor 0.466178, la desviación estandar fué estimada por 0.84 y las hipótesis rechazadas

fueron las de homoscedasticidad por renglones y normalidad en sus cuatro niveles de significancia mayores; para este término el estimador de θ^2 fué igual a 7.203824. Para un tercer término para el efecto de interacción se tuvo que el estimador de θ^2 fué 0.466178.

Para este término no restan grados de libertad para la estimación de σ^2 . El programa reporta como suma de cuadrados el valor 1.1778×10^{-16} , aunque este valor puede deberse a errores de aproximación propios del sistema de cómputo. Además se tuvo que las hipótesis rechazadas para los residuos obtenidos fueron las de homoscedasticidad por columnas, normalidad en sus cinco niveles y simetría en sus dos niveles.

Cuando el valor de la desviación estándar cambió a 1.05, los resultados obtenidos al ajustar un modelo aditivo fueron para la suma de cuadrados un valor de 61192.4799, rechazando las hipótesis de normalidad en sus cinco niveles y kurtosis en sus dos niveles. Cuando se incrementó un término para el efecto de interacción, con $\hat{\theta}^2 = 61190.53871$, se obtuvo una suma de cuadrados de 1.94122 de donde el estimador de la desviación estándar alcanzó el valor de 0.72; es decir, el estimador estuvo 0.33 abajo del valor real de el parámetro. En este caso no se rechazó ninguna de las hipótesis, de donde el modelo ajustado fué uno que incluyese un término multiplicativo para el efecto de interacción.

ción.

Cuando el modelo que generó los datos fué cambiado a:

$$Y_{ij} = 2 (i-1) + 3j + ((1-3) (j-1))^5 + \epsilon_{ij}$$

Con $\epsilon_{ij} \sim N(0, \sigma^2)$, $\sigma = 0.05, 0.30, 0.55, 0.80, 1.05$; se tuvo que para los valores de σ iguales a 0.05, 0.80, y 1.05 al menos una de las hipótesis fué rechazada a cada paso, es decir, ninguno de los distintos modelos ajustados satisfizo completamente las hipótesis hechas acerca de la distribución conjunta de los errores.

Los resultados obtenidos para este modelo con $\sigma = 0.05$ se pueden resumir de la siguiente manera. Cuando se ajustó un modelo estrictamente aditivo se obtuvo una suma de cuadrados igual a 33703930.187 y las hipótesis rechazadas fueron homoscedasticidad por renglones y por columnas, normalidad en sus cinco niveles, simetría en sus dos niveles y kurtosis también en sus dos niveles. Cuando se incluyó un término para el efecto de interacción con $\hat{\theta}^2 = 33703925.5996$, la suma de cuadrados de los residuos se redujo hasta alcanzar el valor de 4.587266 dando como resultado $\hat{\sigma} = 1.1178$, nótese que este valor es de casi 22 veces el del valor real de la desviación estándar. Las hipótesis rechazadas fueron la homoscedasticidad por renglones, la de normalidad en sus tres niveles de significancia mayores y la de kurtosis a los dos niveles de significancia usados. El aumento de un término multiplicativo para el efecto de interacción, por el

cual se tuvo que $\hat{\sigma}^2 = 4.5776$ dió como resultado que la suma de cuadrados fuese igual a 0.00968326 y el estimador de la desviación estandar tuviese un valor de 0.1211, el cual es casi 2.5 veces el valor real de la desviación estandar. En términos generales las propiedades distribucionales del error no fueron mucho mejores ya que las hipótesis rechazadas fueron las de normalidad en sus cinco niveles; independencia en los dos niveles y kurtosis en un sólo de los niveles de significancia. Al obtenerse el estimador del tercer término para explicación del efecto de interacción se hizo uso de los grados de libertad restantes con lo cual no es posible tener un estimador de la desviación estandar de los errores, sin embargo la suma de cuadrados, incluye errores de redondeo por la máquina, se redujo hasta 3.3615×10^{-15} .

La única hipótesis rechazada en este caso fué la de normalidad en los cinco niveles usados.

El uso de el criterio de Mandel, dada la magnitud del primer estimador ($\hat{\sigma}^2 = 33703925.5996$), hubiérase aceptado un sólo término para la explicación del efecto de interacción lo que habría dado como resultado un estimado de la varianza de 440 veces el valor real y residuos con propiedades muy nobres.

Cuando se cambió el valor de σ a 0.30 los resultados fueron algo poco común. De el ajuste de un modelo aditivo se obtuvo un va-

lor de la suma de cuadrados de los residuos igual a 84112214.767 y el rechazo de las hipótesis de homoscedasticidad por columnas, normalidad a los cinco niveles usados y kurtosis a los dos niveles utilizados. Al incluir un término multiplicativo al modelo, el ajuste de este modelo trajo como consecuencia que la suma de cuadrados fuera igual a 0.142344261, el estimador de la desviación estandar tomó el valor de 0.1969 y la única hipótesis rechazada fuera la de independencia a los niveles de 0.10 y 0.02. Ahora bien al incluir un nuevo término al efecto de interacción, con $\hat{\theta}^2 = 0.115087$, la suma de cuadrados se redujo hasta el valor 0.027256766, pero el estimador de la desviación estandar "creció" y tuvo el valor de 0.2032 y con este valor para $\hat{\sigma}$ ninguna de las hipótesis fué rechazada. Así pues, el modelo ajustado a este grupo de datos fué uno que incluyó dos términos multiplicativos para explicación del efecto de interacción. La aplicación del criterio de Mandel habría resultado en el ajuste de un modelo con un sólo término para el efecto de interacción. Para el mismo modelo, pero con $\nabla = .55$ el ajuste de un modelo aditivo dió una suma de cuadrados igual a 84132634.0647 y se tuvo que las hipótesis rechazadas fueron las de homoscedasticidad por columnas, normalidad en sus cinco niveles y kurtosis en los dos niveles usados. La inclusión de un término para la explicación del efecto de interacción, con $\hat{\theta}^2 = 84132633.072$ con 8.33

grados de libertad, se tuvo que la suma de cuadrados se redujo hasta 0.993415382 lo que motivó un estimador de la desviación estandar de 0.52 y la única hipótesis rechazada fué la de normalidad en sus tres niveles.

El rechazo de esta hipótesis motivó la inclusión de un nuevo término, para el cual $\hat{\theta}^2 = 0.783317$ con 3.01 grados de libertad, el valor de $\hat{\theta}^2$ hizo que la suma de cuadrados se redujese hasta 0.2100978 y en consecuencia el valor $\hat{\sigma} = 0.564$; esta vez el estimador de la desviación estandar también creció, como en el ejemplo anterior, sin embargo rebasó el valor real del parámetro. Las hipótesis rechazadas fueron las de normalidad en cuatro de sus cinco niveles y la de independencia en los dos niveles utilizados. La consideración de un nuevo término para la explicación del efecto de interacción, para el cual $\hat{\theta}^2 = 0.21009779$, dió una suma de cuadrados igual a $(4.3831) \times 10^{-14}$ y se tuvo el rechazo de las hipótesis de homoscedasticidad por columnas, normalidad en sus cinco niveles y simetría en los dos niveles. La aplicación del criterio de Mandel hubiera ajustado un modelo con un sólo término explicativo para el efecto de interacción, lo que en este caso particular resulta ser lo más adecuado.

Al cambiar el valor de σ a 0.30, el ajuste de un modelo aditivo dió como resultado una suma de cuadrados igual a 84120425.8584 y las hipótesis rechazadas fueron las de homoscedasticidad por

columnas, normalidad en sus cinco niveles y kurtosis en los dos niveles usados. La inclusión de un término multiplicativo para la explicación del efecto de interacción, para el cual $\hat{\theta}^2 = 84120424.4541$, trajo como consecuencia una suma de cuadrados igual a 1.404441, con el estimador de la desviación estándar igual a 0.618; las hipótesis rechazadas fueron normalidad en los tres niveles mayores de significancia, simetría al 10% y kurtosis en los dos niveles de significancia usados, es decir, los estimadores de los errores resultan poseer muchas propiedades no deseables. Cuando se incluyó un nuevo término para el efecto de interacción, para el cual $\hat{\theta}^2 = 1.20003735$, la suma de cuadrados se redujo hasta el valor 0.2044 y en consecuencia $\hat{\sigma} = 0.556$; la única hipótesis rechazada fué la de independencia a ambos niveles de significancia usados. Si se observa que el estimador de la desviación estándar es más pequeño que el valor real de dicho parámetro, se puede pensar que hemos quitado "demasiado" a los errores de modo que estos ya no son independientes. La consideración de un tercer término para explicación del efecto de interacción, con $\hat{\theta}^2 = 0.2044$, se obtuvo una suma de cuadrados igual a $(7.0726) \times 10^{-14}$, que puede deberse a los ya mencionados errores de redondeo: los errores obtenidos no cumplieron con las hipótesis de homocedasticidad por columnas, normalidad en los cinco niveles, simetría y kurtosis en los dos niveles de signifi-

cancia utilizados. En este caso, a través del criterio presentado en este trabajo no es posible encontrar o suponer un modelo óptimo, en todo caso, el modelo "menos malo" sería aquel que incluyese dos términos para explicar el efecto de interacción conclusión que difiere de aquella a la que se llega por el criterio de Mandel.

El comportamiento del modelo para el cual $\sigma = 1.05$ fué aún peor ya que las estimaciones de la desviación estandar fueron muy malas. El ajuste de el modelo aditivo dió como resultado una suma de cuadrados igual a 84132634.0647 resultando rechazadas las hipótesis de homoscedasticidad por columnas, normalidad en sus cinco niveles y kurtosis en sus dos niveles. La inclusión de un término multiplicativo, con $\hat{\sigma}^2 = 84132633.0720$, redujo la suma de cuadrados hasta el valor 0.993415 y dió un estimador de la desviación estandar igual a 0.52, el cual es menor que la mitad del valor real del parámetro; las hipótesis rechazadas fueron las de normalidad en sus tres niveles mayores y simetría al 0.02. Cuando se aumentó un término al efecto de interacción, con $\hat{\sigma}^2 = 0.7833176$, la suma de cuadrados se redujo hasta alcanzar el valor 0.2100978 con un estimador de la desviación estandar igual a 0.464 que sigue siendo significativamente menor que el valor real del parámetro. Las hipótesis que resultaron rechazadas fueron las de normalidad en los cuatro mayores niveles e independencia en los dos niveles de significancia utilizados. Las

propiedades distribucionales de los errores son aún peores al incluir un nuevo término, para el cual $\hat{\sigma}^2 = 0.21009779$, ya que las hipótesis rechazadas fueron las de homoscedasticidad por columnas, normalidad en sus cinco niveles y simetría en sus dos niveles de significancia. Para este término no se cuenta ya con un estimador de la desviación estándar.

Debe hacerse notar en este punto que para los últimos cinco modelos y para los siguientes el efecto de interacción es un polinomio de grado mayor o igual que cinco, en los valores de los renglones y las columnas y no se le ha relacionado con los niveles de los criterios considerados en el desarrollo del experimento, lo que por supuesto es imposible en los ejemplos que se presentan adelante. Lo anterior puede ser la causa del mal comportamiento de los errores y el consecuente mal ajuste de los modelos. En el caso de experimentos reales se espera que los efectos de interacción, si los hay sean de un orden de magnitud pequeño.

El siguiente modelo considerado fué el siguiente:

$$Y_{ij} = 2(i-1) + 3j + ((i-3)(j-1))^7 + \epsilon_{ij}$$

Con los cinco valores, considerados en los modelos anteriores, para la desviación estándar de los errores ϵ_{ij} .

Los resultados obtenidos para el valor de α igual a 0.05 son sorprendentes. Para el primer modelo ajustado, es decir un modelo estrictamente aditivo, dió un valor para la suma de cuadrados igual a $(1.1328) \times 10^{11}$ y las hipótesis rechazadas son las de

homoscedasticidad por columnas, normalidad en sus cinco niveles y kurtosis en sus dos niveles. Ahora bien, a pesar del grado del polinomio en i y en j , la inclusión de el primer término para explicación del efecto de interacción trajo como consecuencia el no rechazo de ninguna de las hipótesis, es decir, se acentó un modelo con un sólo término para la explicación del efecto de interacción. En este caso la suma de cuadrados fué igual a 10.746 dando, para $\hat{\sigma}^2$, el valor de 1.75. La aplicación del criterio de Mandel habría dado precisamente el mismo modelo dado que la magnitud del cuadrado medio para el efecto de interacción es del orden de 10^{11} , mientras que los restantes cuadrados medios sería todos del orden de 10.

Para el caso en que $\sigma = 0.30$ el ajuste de un modelo aditivo dió una suma de cuadrados de los residuos igual a $(1.1328) \times 10^{11}$ con el rechazo de las hipótesis de homoscedasticidad por columnas, normalidad en sus cinco niveles y kurtosis en sus dos niveles. Al incluir un término para la explicación del efecto de interacción resultó, nuevamente, que ninguna de las hipótesis fué rechazada, y como consecuencia que el ajuste de un modelo con un sólo término para la interacción con 6^2 del orden de 10^{10} . La suma de cuadrados alcanzó el valor de 13.19405 y $\hat{\sigma}^2 = 1.89$, valor que no es tan exageradamente mayor como el del ejemplo anterior. Así mismo el modelo que se hubiéramos ajustado haciendo

uso del criterio del Dr. Mandel, sería precisamente el mismo.

En cambio, cuando $G = 0.55$, la decisión acerca de cuál modelo debe ajustarse a los datos no resultó obvia. El ajuste del modelo aditivo dió una suma de cuadrados de $(1.13331) \times 10^{11}$ resultando rechazadas las hipótesis de homoscedasticidad por columnas, normalidad en sus cinco niveles y kurtosis en los dos niveles considerados. La inclusión del primer término multiplicativo para la explicación del efecto de interacción, para el cual $\hat{\theta}^2$ fué de un orden de magnitud de 10^{-11} , dió una suma de cuadrados del error igual a 2.145226, valor que junto con los 3.67 grados de libertad restantes dió un estimador de la desviación estandar de tan solo 0.764; la hipótesis rechazada fué la de normalidad en cuatro de sus niveles únicamente. Al considerar un nuevo término para la interacción, con $\hat{\theta}^2 = 2.13675787$, la suma de cuadrados del error fue igual a 0.008469 lo que dió $\hat{G} = .1132$, valor que es ya cuatro veces menor que el valor real del parámetro; las hipótesis que resultaron rechazadas fueron las de normalidad al 20% de significancia e independencia a los dos niveles considerados. La inclusión de un tercer término multiplicativo para el efecto de interacción, para el cual $\hat{\theta}^2 = 0.0084689$, dió una suma de cuadrados de 2.8896×10^{-11} . En este caso ya no es posible obtener un estimador de la desviación estandar del error. Se rechazaron las hipótesis de normalidad y simetría, ambas a

todos los niveles usados. Obviamente el criterio de Mandel habría considerado el ajuste de un modelo con un sólo término para el efecto de interacción como el adecuado.

En contraste con lo anterior, el caso en que $\sigma = 0.80$ tuvo un comportamiento mejor ya que aunque el ajuste de un modelo aditivo dió una suma de cuadrados de $(1.1333) \times 10^{11}$ y las hipótesis que se rechazaron fueron las de homoscedasticidad por columnas, normalidad y kurtosis, ambas en todos sus niveles, el ajuste de un modelo que incluye un término para la explicación del efecto de interacción, con $\hat{\sigma}^2$ del orden de 10^4 , dió como resultado el que ninguna de las hipótesis fuese rechazada, una suma de cuadrados de tan sólo 2.544475 y un estimador de la desviación estandar de 0.832.

Para los datos con $\sigma = 1.05$ ocurre que la hipótesis de normalidad se rechaza siempre y a todos los niveles para los tres términos aumentados al modelo aditivo. Los estimadores de la desviación estandar fueron 0.844 y 0.935 para los modelos que incluían 1 y 2 términos respectivamente. Otras hipótesis rechazadas fueron las de homoscedasticidad por columnas y kurtosis en sus dos niveles para el modelo aditivo; la de independencia en los dos niveles para el modelo con un término de la interacción; homoscedasticidad por renglones, simetría y kurtosis para los dos niveles en ambas muchas, cuando el modelo incluía dos términos

para explicación del efecto de interacción; finalmente, además de la hipótesis de normalidad, para un modelo con tres términos multiplicativos se rechazó la hipótesis de simetría al 10% y al 2%. En este caso, de acuerdo a los resultados obtenidos, tampoco es posible hablar de un modelo mejor para representar los datos.

Las últimas cinco matrices de datos simulados correspondieron al siguiente modelo:

$$Y_{ij} = 2(i-1) + 3j + ((1-3)(j-1))^9 + \epsilon_{ij}$$

En cada caso, se hizo variar a σ como en todos los datos simulados anteriormente.

Brevemente los resultados obtenidos para cada grupo de datos son:

Para $\sigma = 0.05$ el ajuste de un modelo aditivo, para el cual la suma de cuadrados fué igual a 1.497964×10^{14} , dió como resultado que se rechazaran las hipótesis de homoscedasticidad por columnas, normalidad y simetría, ambas en todos los niveles usados. Para el primer término de la interacción se tuvo una suma de cuadrados de 0.507979325 y un estimador de σ igual a 0.06740; se rechazaron las hipótesis de normalidad a los cinco niveles de significancia, simetría y kurtosis ambas al 10%. Para el segundo término multiplicativo la reducción en suma de cuadrados fué de 0.232297339 ($=\hat{\sigma}^2$) y el valor de $\hat{\sigma}$ fué 0.0592; se rechazaron

las hipótesis de normalidad en los niveles 20% y 15%, independencia en los dos niveles y kurtosis al 10% de significancia, es decir, la hipótesis cuyo rechazo es más grave en este caso es la independencia para el tercer término la suma de cuadrados fué de 0.000000038; no fué posible ya dar un estimador de la desviación estandar y se rechazaron las hipótesis de normalidad y simetría en todos sus niveles. Para $\sigma = 0.30$ cuando se ajustó un modelo lineal se obtuvo una suma de cuadrados de 1.497964×10^{14} y se rechazaron las hipótesis de homoscedasticidad por columnas, normalidad y kurtosis en todos los niveles. Con el primer término de la interacción la suma de cuadrados fué de 0.50797932 y $\hat{\sigma} = 0.37203$; la única hipótesis que se rechazó fué la de normalidad a todos los niveles. Para el segundo término de la interacción se tuvo una suma de cuadrados de 0.12566193 y la estimación para σ fué 0.4363; se rechazaron las hipótesis de normalidad al 20%, 15% y 10% e independencia al 10% y al 2%. Para el mismo modelo, pero con $\sigma = 0.55$, al ajustar un modelo estrictamente aditivo se obtuvo una suma de cuadrados de los residuos igual a $(1.49496487) \times 10^{14}$ y las hipótesis rechazadas con los mismos residuos fueron las de homoscedasticidad por columnas, normalidad en sus cinco niveles y kurtosis en los dos niveles usados. La inclusión de un término multiplicativo para la explicación del efecto de interacción, para el cual $\hat{\theta}^2$ resultó

un valor del orden de 10^{13} , dió como resultado el que la suma de cuadrados de los nuevos residuos fuese igual a 18.639114 lo que daba como resultado un valor para el estimador de la desviación estandar de 2.2536; como hipótesis rechazadas resultaron ser las de homoscedasticidad por renglones, normalidad al 20% y al 15%, simetría y kurtosis, ambas a 10% y 2%. Con la inclusión de un nuevo término para la explicación del efecto de interacción la suma de cuadrados se redujo a 0.450466909 lo que dió como resultado que $\hat{\sigma} = 0.8261$; se rechazaron las hipótesis de normalidad a los tres niveles mayores y la de independencia a los niveles 10% y 2%. Para el tercer término explicativo, la suma de cuadrados se redujo hasta el valor 0.000000042; para este término no se tiene ya un estimador de la desviación estandar; las hipótesis rechazadas fueron las de homoscedasticidad por columnas, normalidad y kurtosis ambas a todos los niveles usados.

Para $\sigma = 0.80$ el ajuste de un modelo estrictamente aditivo dió unos residuos cuya suma de cuadrados fué igual a $(1.497964456) \times 10^{14}$ y con los cuales se rechazaron las hipótesis de homoscedasticidad por columnas, normalidad en sus cinco niveles y kurtosis en los dos niveles usados. En este caso la inclusión de un término para la interacción redujo la suma de cuadrados hasta

10.123865964 con un estimador de σ igual a 1.6609; en este caso ninguna de las hipótesis fue rechazada con lo que se concluyó que un modelo que incluya un término multiplicativo para explicación del efecto de interacción explica satisfactoriamente el comportamiento de los datos.

En contraste con lo anterior para el caso en que $\sigma = 1.05$ los residuos obtenidos después del ajuste de un modelo aditivo tuvieron una suma de cuadrados igual a $(1.49796476) \times 10^{14}$ y con ellos se rechazaron las hipótesis de homoscedasticidad por columnas, normalidad en los cinco niveles utilizados y kurtosis al 10% y 2%. Con la inclusión de un término para el efecto de interacción la suma de cuadrados de los residuos fue igual a 8.6271067 y el estimador de la desviación estándar fue igual a 1.533; se rechazaron las hipótesis de normalidad en cuatro de sus niveles y kurtosis al 10%. Para el segundo término, para explicación del efecto de interacción, para el cual $\hat{\sigma}^2 = 3.31709725$, el valor de la suma de cuadrados descendió hasta 0.310009369 con lo que $\hat{\sigma} = 0.6853$; resultó en rechazo la prueba de independencia usada a los niveles de 10% y 2%. Con la inclusión del tercer término para el efecto de interacción, con $\hat{\sigma}^2 = 0.310009369$, la suma de cuadrados de los residuos resultó igual a 0.000000003, no se tuvo un estimador para la desviación estándar y se rechazaron las hipótesis de homoscedasticidad por columnas, normalidad,

kurtosis y simetría, estas últimas tres lo fueron a todos los niveles a los que se probaron.

Los resultados obtenidos en las últimas matrices de datos simulados pueden deberse a la magnitud de los efectos de interacción, los cuales representaban una gran influencia en los valores originales obtenidos, lo que opacaba los efectos aleatorios (en general las desviaciones estandar resultaban pequeñas) lo que trajo como consecuencia que los criterios usados para las pruebas de las hipótesis resultaran "viciados".

En general es posible observar que en experimentos reales los efectos de interacción no alcanzan valores tan exagerados como los incluidos en las matrices para las cuales $k \geq 5$. El asignar valores tan grandes a los efectos de interacción se hizo con el fin de calibrar las bondades del método para casos extremos.

CAPITULO 7

7.- RESUMEN Y CONCLUSIONES

Se ha propuesto como modelo mas general y más útil para representar algunos fenómenos que son observados en el mundo real, va ser mediante experimentación u observación directa de la naturaleza el siguiente

$$Y_i = \mu + \varepsilon_i$$

con

$$\varepsilon_i \sim N(0, \sigma^2)$$

La suposición de normalidad de los errores o fluctuaciones aleatorias se justifica en algunos casos por la no consideración explícita de aquellos factores o condiciones cuya influencia es nula o pequeña y que individualmente, son causa de una variabilidad pequeña en los resultados observados. El equivalente teórico a esta justificación está dado por el teorema central de límite que plantea más o menos el hecho de que el resultado de sumar variables aleatorias independientes con varianzas finitas, habiendo estandarizado adecuadamente estas variables, y con un número grande de sumandos, es una variable que sigue una distribución normal con media cero y varianza 1.

Como parte importante de estos modelos, se cuentan los modelos lineales de la forma:

$$Y_i = \alpha_0 + \alpha_1 x_{1i} + \dots + \alpha_n x_{ni} + \varepsilon_i$$

con

$$\varepsilon_i \sim N(0, \sigma^2)$$

Casos particulares de estos modelos, en los cuales se ha concen-

trando toda la atención en este trabajo, son los modelos de diseños de experimentos para los cuales X_{ji} toma únicamente los valores cero o uno. En este tipo de modelos se considera que los efectos de los criterios con los cuales se han clasificado las observaciones pueden descomponerse en una parte aditiva, con cada uno de los sumandos compuestos por funciones de una sola de las variables o criterios, y otra parte posiblemente aditiva pero donde los sumandos son funciones de dos o más de las variables. Cuando se consideran grupos de datos con doble criterio de clasificación y una sola observación por cada pareja de niveles de los criterios, la estimación por mínimos cuadrados de los parámetros del modelo trae como consecuencia un ajuste perfecto de los datos por el modelo; razón por la cual se hace imposible el realizar pruebas de hipótesis o el construir intervalos de confianza para los parámetros que forman parte del modelo.

Por lo anterior se han propuesto distintos métodos de estimación de los efectos de interacción, en uno de los cuales, propuesto por el Dr. John Mandel, se basa el desarrollo del presente trabajo. Este método propone que los efectos de interacción pueden descomponerse como una suma de términos de la forma $\theta_{ij} u_i v_j$ donde los valores de θ son estimados por los eigenvalores de la matriz $GD'D'$, siendo la matriz D la matriz de residuos d_{ij} ; los valores de u_i son estimados por las componentes de los vectores propios

de la misma matriz asociados con los estimadores del parámetro θ ; asimismo los parámetros v_j son estimados por los eigenvectores de la matriz $D'D$. Es decir, el modelo que se propone para ser ajustado a las observaciones es el siguiente:

$$y_{ij} = \mu + \alpha_i + \beta_j + \sum_{r=1}^k \theta_r u_{ri} v_{rj} + \varepsilon_{ij}$$

El número de términos multiplicativos que forman la descomposición del efecto de interacción ha de ser determinado mediante criterios adecuados. El criterio propuesto por el Dr. Mandel hace uso de un enfoque parecido al análisis de varianza comparando magnitudes de los cuadrados medios asociados a los términos del modelo. Por otra parte se espera que una vez que los efectos de interacción sean incluidos en el modelo, aquellos factores no controlados producirán variaciones o errores con distribución de probabilidad normal con media cero y desviación estandar común, además de ser independientes.

De acuerdo a este segundo punto de vista se buscó de una manera recursiva eliminar de los residuos tantos términos de la forma $\theta_u v_j$ como fuera necesario para obtener errores independientes y normalmente distribuidos con media cero y varianzas común. Con este punto de vista se ajustaron los modelos que produjeran los errores con las propiedades que se mencionan arriba.

A partir de los resultados obtenidos en los ejemplos de la sección anterior se puede afirmar que el método de descomposición de los

efectos de interacción como una suma de términos, todos ellos múltiplos del producto de funciones de una sola variable, produce errores suficientemente anegados a las hipótesis hechas acerca de la distribución de los errores reales asociados a cada una de las observaciones.

Esto se logra, en ocasiones, haciendo uso de un número distinto de términos a los que se incluirían haciendo uso del enfoque de análisis de la varianza propuesto por el Dr. John Mandel. Lo anterior se observa principalmente en el conjunto de ejemplos reales de la sección anterior, en los cuales las decisiones hechas acerca del número de términos multiplicativos mediante el enfoque de análisis de la varianza produjeron errores no normales (o menos normales), y/o correlacionados y/o heteroscedásticos y/o sobreestimaciones de la varianza real de los errores; algunos de estas desventajas eran eliminadas, parcial o totalmente, con la inclusión de uno o dos términos más aún cuando los cuadrados medios relacionados con estos términos no fueran "significativos" bajo el enfoque de análisis de la varianza; es decir, ir un poco más adelante en la estimación de los efectos de interacción produjo, en ocasiones, notables mejoras en cuanto a los errores obtenidos y a sus propiedades.

Sin embargo para los datos contruados artificialmente, en los casos en que el efecto de interacción era pequeño en magnitud, con

respecto a los valores de los términos que representaban los efectos principales, dichos efectos no fueron detectados ya que los errores obtenidos al ajustar un modelo estrictamente aditivo cumplían satisfactoriamente con las hipótesis planteadas; este hecho propiciaba en la mayoría de los casos que la varianza fuese sobreestimada, en ocasiones era demasiado grande esta diferencia como para ser pasada por alto.

En casos como este el uso previo de una prueba de no aditividad, tal vez la de Tukey, podría proporcionar una mayor confianza en la decisión tomada acerca del modelo que se debe ajustar a los datos, desgraciadamente esta prueba no toma en cuenta las propiedades distribucionales de los errores obtenidos al rechazar la hipótesis de aditividad. Cabe aquí recordar que algo parecido ocurrió con conjuntos de datos reales (ejemplo 3), en los que se sospechaba que ocurría algo parecido, en vista de la magnitud de las sumas de cuadrados por lo que este caso no debe ser pasado por alto.

Por otro lado, también en los datos contruidos artificialmente, es notable la poca eficiencia del método para producir errores satisfactorios en los casos en que el efecto de interacción era tan grande como un polinomio de grado mayor o igual a cinco. Esto es sin embargo una desventaja de tipo técnico únicamente, ya que como se observó en los ejemplos reales, estos casos difícilmente ocurren en la práctica.

Para este tipo de datos sería deseable el contar con información acerca de la distribución de los estimadores de los parámetros incluidos en los términos multiplicativos en los cuales se descompone el efecto de interacción ya que en ese caso sería posible el realizar pruebas estadísticas para niveles de significancia prefijados. Sin embargo, aún cuando se contase con esta información, los errores producidos después de realizar la prueba de significancia adecuada serían poco satisfactorios en cuanto a las propiedades distribucionales de los mismos. La magnitud de los valores originales y las sumas de cuadrados de los residuos, que decrecen rápidamente, pueden servir de indicios de que el efecto de interacción es muy grande.

En todos los casos anteriores es posible observar que los errores obtenidos parecen alcanzar un nivel óptimo para un cierto número de términos del efecto de interacción, a partir del cual los errores que se obtienen aumentando un término empeoran en cuanto a las condiciones que satisfacen. Más aún hay ocasiones en que los errores obtenidos no satisfacen completamente las hipótesis establecidas para ningún número de términos en el efecto de interacción, aunque para un cierto número de ellos las propiedades satisfechas parecen ser mejores que las que se obtuvieron con un término menos y que las que se obtendrían aumentando otro término multiplicativo; si en este caso se audiese pensar que la función "hipótesis no rechazadas" o "propiedades deseables", fun-

ción del número de términos incluidos en el modelo o de los valores de los estimadores de los parámetros, fuese continua, una "interpolación" entre los valores de los estimadores de los parámetros y "mesada" por aquellas hipótesis no satisfechas podría producir un modelo cuyo ajuste diese errores "óptimos" en el sentido que todas las hipótesis fuesen satisfechas.

Es decir, el modelo podría aumentarse con un pequeño conjunto de parámetros (a lo más tres), cuyos valores se determinarían tomando en cuenta las hipótesis rechazadas y los niveles para los cuales tuvo lugar ese rechazo. El modelo aumentado tomaría una forma parecida a la siguiente:

$$Y_{ij} = \mu + \alpha_i + \beta_j + \sum_{r=1}^k \lambda_r \theta_r^u v_{rj} + \varepsilon_{ij}$$

Donde los parámetros λ_r se han aumentado al modelo; los valores posibles para dichos parámetros serían los comprendidos en el intervalo cerrado (0,1).

Para aquellos casos en los que el ajuste de un modelo aditivo diese como resultado el no rechazo de ninguna de las hipótesis pero se sospechase que la varianza ha sido sobreestimada o la hipótesis de aditividad hubiese sido rechazada podría aumentarse un término de la forma $\lambda_1 \theta_1^u v_{1i}$, determinando el tamaño de λ_1 de manera tal que las propiedades de los errores siguieran siendo óptimas. De este modo se obtendría una reducción en sum

de cuadrados y en consecuencia en el valor del estimador de la varianza, logrando además un mejor ajuste del modelo a los datos. Para el caso en que el modelo "óptimo" incluye t términos, aún cuando alguna hipótesis hayan sido rechazadas, los valores de λ_k serían uno para aquellos términos de con $k \leq t-2$, fijándose λ_{t-1} , λ_t , λ_{t+1} de acuerdo a las hipótesis rechazadas y pudiendo alguno de estos parámetros alcanzar el valor cero.

Finalmente, es recomendable como tema de futura investigación el trabajo en la distribución exacta de los parámetros incluidos en el modelo con el objeto de obtener un criterio preciso que nos lleve a determinar aquellos cuadrados medios que resulten realmente significativos. Acerca de esto (la distribución de los eigenvalores de la matriz de varianza y covarianzas), Mandel ha publicado ya un interesante artículo que aparece con el número (15) en la lista de referencias.

APPENDICE I

UNIVERSIDAD NACIONAL AUTONOMA DE MEXICO

CENTRO DE SERVICIOS DE COMPUTO

COMPILADOR ALGOL, NIVEL 2.5.009, BURROUGHS B6700

MARTES 22 DE OCTUBRE DE 1974

06:10 PM.

PROG / PRUEBA

BEGIN

FILE ENT(KIND=READER), SAL(KIND=PRINTER);

ARRAY ALFA[0: 1,15],PSIME,SIME[1:2],KURTO,VON[1:4],COCH[1:2];
INTEGER M,N, PTE;ARRAY JITA[1:3],PNOR[1:2];
REAL PSC,SIGMA,MED,SCDA,X;
LABEL LEE,FIN;

DEFINE ARCONA(A,S)= FOR I=1 STEP 1 UNTIL S DO

SC:=A*(CONT,I)**2;

SC:=SC**(1/2);

FOR I:=1 STEP 1 UNTIL S DO

A*(CONT,I):=-SC #;

PROCEDURE IMPRIMEMATRIZ(SAL,A,M,N);

VALUE M,N;

FILE SAL;

ARRAY A[*,*];

INTEGER M,N;

BEGIN

INTEGER I,J;

FOR I:=1 STEP 1 UNTIL M DO

BEGIN

WRITE(SAL,<X10,"",*(("-----"))>,HIN(N,10));

WRITE(SAL,<X10,"",*(RIS.4,X1,""),(FOR J:=1 STEP 1 UNTIL N DO

A(I,J));

WRITE(SAL,<X10,"",*(("-----"))>,HIN(N,10));

END; END;

BOOLEAN PROCEDURE LEEHATRIZ(FE,FS,A,M,N,IMPRIME);

VALUE M,N,IMPRIME;

FILE FE,FS;

ARRAY A[*,*];

INTEGER M,N;

BOOLEAN IMPRIME;

BEGIN

INTEGER I,J;

LABEL FINDATOS,FORMATO,FIN;

FOR I:=1 STEP 1 UNTIL M DO

B.0000 IS SEGMENT 00003

1

003:0000:0

003:0000:0

DATA IS 0005 LONG

DATA IS 0005 LONG

003:0000:0

003:0007:2

003:0007:2

003:0007:2

003:0007:2

003:0007:2

003:0007:2

003:0007:2

003:0007:2

003:0007:2

003:0007:2

003:0007:2

003:0007:2

003:0007:2

003:0007:2

IMPRIMEMATRIZ IS SEGMENT 00004

2

004:0000:0

004:0004:3

3

004:0004:3

004:0014:1

004:0021:0

004:0029:1

004:0039:1

IMPRIMEMATRIZ(004) IS 0041 LONG

2

003:0007:2

003:0007:2

003:0007:2

003:0007:2

003:0007:2

003:0007:2

003:0007:2

003:0007:2

LEEATRIZ IS SEGMENT 00007

2

007:0000:0

007:0000:0

```

GO TO FIN;
FINDATOS;
WRITE(FS,<<"ERROR, SE TERMINARON LAS TARJETAS EN LA HILEERA ",13>>
1); LEENMATRIZ:=TRUE; GO TO FIN;
FORMATO;
WRITE(FS,<<"ERROR, FORMATO INCOMPATIBLE EN LA HILERA ",13>>,1);
LEENMATRIZ:=TRUE;
FIN;
IF IMPRIME THEN IMPRIMEMATRIZ(FS,A,M,N);
END;

PROCEDURE ESTIMAB(M,N,Y,B);
VALUE M,N; INTEGER M,N; ARRAY Y[*],B[*];
BEGIN
  INTEGER J,K;
  REAL YPP;
  WRITE(SAL,<<"ESTIMA B"/>>);
  FOR J:=1 STEP 1 UNTIL M DO
    FOR K:=1 STEP 1 UNTIL N DO BEGIN
      YPP:= + Y(J,K);
      B(J+1) := + + Y(J,K)/N;
      B(M+K+1) := + + Y(J,K)/M;
    END;
  B(1) := YPP/(M*N);
  FOR J:=1 STEP 1 UNTIL M DO
    B(J+1) := + - B(1);
    FOR K:=1 STEP 1 UNTIL N DO
      B(M+K+1) := + - B(1);
  WRITE(SAL,<<X10,"LA MEDIA GENERAL ES: ",R9.4//X10,"Y LOS EFECTOS POR
  "RENGLON Y COLUMNA ESTAN DADOS POR LOS SIGUIENTES VECTORES:"/>>,B(1)
  );
  WRITE(SAL,<<X10,"(R14.4,">>,M, FOR J:=2 STEP 1 UNTIL M+1 DO B(J));
  WRITE(SAL,<<X10,"(R14.4,">>,N, FOR J:= M+2 STEP 1 UNTIL M+N+1 DO
  B(J));
END;

PROCEDURE RESIDUO1(M,N,Y,D,B,FS);
VALUE M,N; INTEGER M,N; ARRAY Y,D[*],B[*]; FILE FS;
BEGIN
  INTEGER I,J;
  WRITE(FS,<<"RESIDUO 1"/>>);
  FOR I:=1 STEP 1 UNTIL M DO
    FOR J:=1 STEP 1 UNTIL N DO
      D(I,J):=Y(I,J)-B(1)-B(I+1)-B(J+M+1);
  WRITE(FS,<<X10,"LA MATRIZ DE RESIDUOS PRIMERA ES: "/>>);
END RESIDUOS1;

PROCEDURE ORDENA(A,M,N,ADR,PTE);
VALUE M,N; INTEGER M,N; ARRAY A,ADR[*]; INTEGER PTE;
BEGIN
  INTEGER I,J;
  LABEL L;
  ARRAY B,C(10*M*N);
  B(1):=10**5;
  FOR I:=1 STEP 1 UNTIL M DO

```

```

007:0015:0
007:001E:3
007:001E:3
007:0021:2
007:0020:5
007:0020:5
007:003C:1
007:003C:5
007:003C:5
007:0041:0
LEENMATRIZ(007) IS 0049 LONG
2
003:0007:2
003:0007:2
003:0007:2
003:0007:2
ESTIMAB IS SEGMENT 00008
2
008:0000:0
008:0000:0
008:0005:3
008:000A:0
3
008:000E:3
008:0011:3
008:0016:2
008:001B:4
3
008:001C:4
008:001F:2
008:0023:5
008:0027:4
008:002C:1
008:0030:3
008:0032:2
008:003A:1
008:003F:1
008:0055:4
008:0065:2
008:006C:4
ESTIMAB(008) IS 0074 LONG
2
003:0007:2
003:0007:2
003:0007:2
003:0007:2
RESIDUO1 IS SEGMENT 00009
2
009:0000:0
009:0005:4
009:000A:1
009:000E:4
009:0019:1
009:001E:5
RESIDUO1(009) IS 0020 LONG
2
003:0007:2
003:0007:2
003:0007:2
003:0007:2
ORDENA IS SEGMENT 0000A
2
00A:0000:0
00A:0000:0
00A:0004:0
00A:0006:1

```

```

      C[(I-1)*N+J]:=JJ ENDJ
FOR I:=1 STEP 1 UNTIL M*N-1 DO BEGIN
  IF B[I] GTR C[I+1] THEN GO TO L ELSE
  BEGIN
    B[I]:=READLOCK(B[I],B[I+1])
    C[I]:=READLOCK(C[I],C[I+1])
  END
  FOR J:=1 STEP 1 UNTIL 1 DO
  IF B[J] LSS B[J-1] THEN GO TO L ELSE BEGIN
    B[J-1]:=READLOCK(B[J-1],B[J])
    C[J-1]:=READLOCK(C[J-1],C[J]) ENDJ ENDJ
  L: ENDJ
FOR I:=1 STEP 1 UNTIL M DO
FOR J:=1 STEP 1 UNTIL N DO
  ACR[I,J]:=B[(I-1)*N+J]
  PTE:=C[I]
END ORDENA)

PROCEDURE INTEGRAL(P,A,B,INT,SCD)
  VALUE P,A,B: INTEGER P; REAL A,B,INT,SCD;
  BEGIN
    REAL X,CIF,INT1; INTEGER L;REAL VAR;

    LABEL ITERA;
    DEFINE FX(S):=(2+3.1416*S)*((-1/2)*EXP((-X**2)/(2*S)))#;
    PROCEDURE INTEGRAL(P, A,X,B,INT);
      VALUE A,B: REAL A,B,INT,X; INTEGER P;
      BEGIN
        INTEGER I,J; ARRAY UES(0:P+1);

        I:=1; INT:=0;
        FOR X:=A STEP (B-A)/P UNTIL B DO BEGIN
          UES[I+1]:=FX(X,VAR);
          INT:=INT + UES[I];
          FOR J:=0 STEP 1 UNTIL P DO BEGIN
            IF J EQL 0 OR J EQL P THEN
              INT1:=INT + UES[J];
            ELSE IF J MOD(2) EQL 0 THEN
              INT1:=INT1 + 2*UES[J];
            ELSE INT1:=INT1 + 4*UES[J]; ENDJ
            INT:=INT+(B-A)/(P*3);
          ENDJ
        ENDJ

        VAR:=SCD/(M*N-1);
        CIF:=10**(-5);
        ITERA: INT1:=INT;
        L:=L+1; P:=2*P; IF A EQL B THEN INT:=0 ELSE
          INTEGRAL(P,A,X,B,INT);
        IF ABS(INT-INT1) GTR CIF THEN GO TO ITERA;
      ENDJ
    ENDJ

    BOOLEAN PROCEDURE KOLMOGSMIR (M,N,D,SCD,FS,ARR,X,ORD);
      VALUE M,N,SCD: INTEGER M,N; ARRAY D[**],ARR[**];REAL SCD;FILE FS;
      REAL X; ARRAY ORD[**];
      BEGIN
        INTEGER I,J,K,PTE;

        ARRAY DIF,ODIF(1:M,1:N);
        REAL O,OP,COTA,INT;
        DEFINE COMPARE(X)=

```

```

00A:0011:3
00A:0017:1
00A:001C:3
00A:001E:5
00A:0022:0
00A:0025:1
00A:0028:5
00A:002U:1
00A:002E:4
00A:0032:4
00A:0033:1
00A:0037:4
00A:003C:1
00A:0041:4
00A:0043:0
ORDENA(00A) IS 0047 LONG
003:0007:2
003:0007:2
003:0007:2
003:0007:2
003:0007:2
INTEGRA IS SEGMENT 0000H
00B:0000:0
00B:0000:0
00B:0000:0
00B:0000:0
00B:0000:0
00B:0000:0
00B:0000:0
INTEGRAL IS SEGMENT 0000C
00C:0003:2
00C:0005:0
00C:000E:4
00C:0012:3
00C:0019:0
00C:001D:3
00C:001F:1
00C:0020:5
00C:0023:1
00C:0025:1
00C:0029:2
00C:002C:2
INTEGRAL(00C) IS 0031 LONG
00B:0000:0
00B:0002:1
00B:0004:3
00B:0005:2
00B:000A:2
00B:000D:4
00B:0010:2
INTEGRA(00B) IS 0014 LONG
003:0007:2
003:0007:2
003:0007:2
003:0007:2
003:0007:2
003:0007:2
003:0007:2
KOLMOGSMIR IS SEGMENT 0000F
00F:0000:0
00F:0003:5
00F:0003:5

```

```

"CATIVA AL ",F6.3,"%."//X25,"EL VALOR DE LA ESTADISTICA ES"
,X5,R10.8/,>> ARR(0,1),DDIF(1,1))
ELSE BEGIN Q1=Q+1)
WRITE(FS,</"KOLHOGOROV-SHIRNOV"/X20,"LA NORMALIDAD SE ACEPTA CON",
R6.3,"% SIENDO EL VALOR DE LA ESTADISTICA ",X5,R10.8/,>> ARR(0,1),
DDIF(1,1)); END;
INTERPOLA(A,B,C,D,E,F)= F=C+((D-C)/(B-A))*(D-A),
ASINT(X)= ARR(1,X)/((H*N)+.1/2));
ORDEN(A,D,H,N,ORD,PTE)
K1=Q)
FOR J1=N STEP -1 UNTIL 1 DO
FOR J1=N STEP -1 UNTIL 1 DO BEGIN
IF J EQL N AND 1 EQL N THEN
INTEGRAL 6,ORDEN,N1=4*SCD,ORD(H,N),INT,SCD) ELSE IF J EQL N THEN
IF 1 EQL N THEN ELSE INTEGRAL 6,ORD(1+1,1),ORD(1,J),INT,SCD)
ELSE INTEGRAL 6,ORD(1,J+1),ORD(1,J),INT,SCD);
Q1=Q+INT;
K1=K+1;
DIF(1,J)=MAX(ABS(K/(H*N)-Q),ABS((K-1)/(H*N)-Q));
END;
ORDEN(DDIF,H,N,DDIF,PTE);
WRITE(FS,</X10,"LA MATRIZ DE DIFERENCIAS ORDENADAS ES"/>>);
IMPRIMEMATRIZ(FS,DDIF,H,N);
FOR J1=1 STEP 1 UNTIL 5 DO BEGIN
COTA1=ASINT(1);
COMPARA(COTA);END;
IF 5 EQL 5 THEN KOLHOGSMIR:=TRUE ELSE KOLHOGSMIR:=FALSE;
END;

PROCEDURE SUMCUAD(M,N,D,CONT,FCDA,SAL);
VALUE M,N,CONT; INTEGER M,N,CONT; ARRAY D(1..N),REAL [FCDA]FILE SAL;
BEGIN
INTEGER I,J;
REAL SCD,DIF;
FOR I=1 STEP 1 UNTIL N DO
FOR J=1 STEP 1 UNTIL N DO
SCD:= + D(I,J)**2;
WRITE(SAL,</"SUMA DE CUADRADOS"/X15,"SCD= ",R19.9/,>>SCD);
IF CONT GTR 0 THEN BEGIN
SCD:=/PSC; FCDA:=/PSC;
IF ABS(SCD-FCDA) LEQ 0.1 THEN BEGIN
WRITE (SAL,</X15,"LA SUMA DE CUADRADOS DE LOS RESIDUOS CONVERGE PARA"
13," J"/X15,"Y CONVERGE ALREDEDOR DEL VALOR: ",R19.9/,>>CONT,SCD);
END
ELSE BEGIN
DIF:= SCD-FCDA;
WRITE(SAL,</X15,"LA SUMA DE LOS CUADRADOS DE LOS RESIDUOS ACTUALES"
" DIFIERE DE LA DE LOS ANTERIORES EN ",R10.5," J"/X15,"SIENDO LA"
" ACTUAL SUMA DE CUADRADOS IGUAL A ",R19.9," J"/>>DIF,SCD);END;END;
IF CONT GTR 0 THEN FCDA:=SCD+PSC ELSE
FCDA:=SCD;
END;

BOOLEAN PROCEDURE HOMOSCED(D,M,N,SAL,COCH);
VALUE M,N;INTEGER M,N;ARRAY D(1..N),COCH(1..N) FILE SAL;
BEGIN
REAL-SUM,DIVJ

```

```

00F:0003:5
00F:0003:5
00F:0003:5
00F:0003:5
00F:0003:5
00F:0003:5
00F:0003:5
00F:0003:5
00F:0009:0
00F:000C:5
00F:0010:5
00F:0012:4
00F:001C:3
00F:0024:5
00F:002D:2
00F:002E:4
00F:0030:0
00F:0038:3
00F:0039:3
00F:003E:0
00F:0043:4
00F:0047:0
00F:0047:4
00F:004C:4
00F:0077:0
00F:007A:1
KOLHOGSMIR(00F) IS 0084 LONG
2 003:0007:2
003:0007:2
003:0007:2
003:0007:2
SUMCUAD IS SEGMENT 00010
2 010:0000:0
010:0000:0
010:0004:3
010:0009:0
010:000D:2
010:0018:4
010:001C:5
010:001F:4
010:0023:4
010:0025:4
010:0033:4
010:0033:4
010:0034:1
010:0035:3
010:0037:3
010:0037:3
010:0045:4
010:0048:2
010:0049:5
SUMCUAD(010) IS 0051 LONG
2 003:0007:2
003:0007:2
003:0007:2
003:0007:2

```

```

ARRAY RENG,HEDR(11M),COL,MEDC(11N))
DEFINE MAXIMO(A,S)= FOR I=1 STEP 1 UNTIL S-1 DO
  IF A(I) GTR A(I+1) THEN A(I)=READLOCK(A(I),A(I+1)) #,
  SUMA(A,S,K) = SUMI=0)
  FOR I=1 STEP 1 UNTIL S DO SUMI=#+A(I))
  DIVI= ALS)/SUMI)
  IF DIV LEQ COCHK) THEN BEGIN
    WRITE(SAL,<X10,"LA HOMOSCEDASTICIDAD SE ACEPTA AL 1% SIENDO",
    "LA COTA IGUAL A ",R7.4/," Y EL VALOR DE LA ESTADISTICA IGUAL A",
    X1,R7.4/," COCHK(K),DIV)) QI=0+1)
  END ELSE BEGIN
    WRITE(SAL,<X10,"LA HOMOSCEDASTICIDAD SE RECHAZA AL 1% SIENDO LA ",
    "COTA IGUAL A"R10.6" EL VALOR DE LA ESTADISTICA ES"R11.6/,"
    COCHK(K),DIV))
    QI=0) END #)
  FOR JI=1 STEP 1 UNTIL M DO
  FOR JI=1 STEP 1 UNTIL N DO BEGIN
    HEDR(JI)=#+D(I,JI)/N)
    MEDC(JI)=#+D(I,JI)/M)
  END)
  FOR I=1 STEP 1 UNTIL M DO
  FOR JI=1 STEP 1 UNTIL N DO BEGIN
    RENG(I)=#+(D(I,JI)-HEDR(JI))**2)
    COL(JI)=#+(D(I,JI)-MEDC(JI))**2) END)
  MAXIMO(RENG,M) MAXIMO(COL,N)
  SUMA(RENG,M,1) SUMA(COL,N,2))
END) IF Q EQ 2 THEN HOMOSCED =TRUE ELSE HOMOSCED=FALSE)
END)
BOOLEAN PROCEDURE JICUAD(C,M,N,SCD,JITA,PNOR,SAL,ORD))
VALLE M,N,SCD) INTEGER M,N,ORD)REAL SCD) ARRAY D(***),PNOR(1),JITA(1)
FILE SAL) ARRAY ORD(***))
BEGIN
  REAL EXPE/JISQ)
  INTEGER I,J,Q)
  ARRAY OBS(1:6)
  DEFINE COMBOR(COI,VAR,COS,S)=
  IF VAR GTR COI*(SCD/(M*N-1))**((1/2) AND VAR LEQ COS*(SCD/(M*N-1))**
  ((1/2) THEN COS(S)=#+1 ELSE #)
  ORDERACQ,M,N,ORD,PTC)
  EXPE=M*N*0.1666)
  FOR I=1 STEP 1 UNTIL M DO
  FOR JI=1 STEP 1 UNTIL N DO
  IF ORD(I,JI) LEQ -PNOR(1)*(SCD/(M*N-1))**((1/2)-THEN OBS(1)=#+1-ELSE
  COMBOR(-PNOR(1),ORD(I,JI),-PNOR(2),2)
  COMBOR(-PNOR(2),ORD(I,JI),0,3)
  COMBOR(0,ORD(I,JI),PNOR(2),4)
  COMBOR(PNOR(2),ORD(I,JI),PNOR(1),5)
  IF ORD(I,JI) GTR PNOR(1)*(SCD/(M*N-1))**((1/2) THEN OBS(6)=#+1)
  FOR I=1 STEP 1 UNTIL 6 DO
  JISQ=#+((OBS(I)-EXPE)**2)/EXPE)
  FOR JI=1 STEP 1 UNTIL 3 DO
  IF JISQ LSS JITA(JI) THEN BEGIN
    QI=Q+1)
    WRITE(SAL,<"BONDAD DE AJUSTE:"/,"X10,"LA NORMALIDAD SE ACEPTA AL ",
    "12.5% Y EL VALOR DE LA ESTADISTICA ES "R7.4/,"ALFALO,2+JI,JISQ)
    DATA IS OF3 LONG
    HOMOSCED(011) IS OB6 LONG
    JICUAD IS SEGMENT 00013
    013:0000:0
    013:0000:0
    013:0000:0
    013:0000:0
    013:0000:0
    013:0000:0
    013:0004:3
    013:0007:4
    013:0000:1
    013:0010:4
    013:0019:0
    013:0019:3
    013:002A:0
    013:0039:4
    013:0049:1
    013:0064:1
    013:0064:5
    013:006A:3
    013:006B:1
    013:006D:2
    013:006E:4
    013:0070:2
  END)

```

```

WRITE(SAL, <"BONDAD DE AJUSTE"/>, X10, "LA NORMALIDAD NO ES ACEPTADA "
, I2, "X"/>, ALFA(I0, 2+J)) ENDJ
IF Q EQL 3 THEN JICUAD:=TRUE ELSE JICUAD:=FALSEJ
END JICUADJ

PROCEDURE TRANSPON(D, M, N, DP)
VALUE M, N: INTEGER M, N; ARRAY D[*,*], DP[*,*];
BEGIN
  INTEGER I, J;

  FOR I:=1 STEP 1 UNTIL M DO
    FOR J:=1 STEP 1 UNTIL N DO
      DP[J, I] := D[I, J];
    ENDJ
  ENDJ

PROCEDURE MULTIPLICAMATRIZ(A, B, C, M, N, P, FS, FE)
VALUE M, N, P;
FILE FS, FE;
INTEGER M, N, P;
ARRAY A[*,*], B[*,*], C[*,*];
BEGIN
  INTEGER I, J, K;

  LABEL FINJ
  FOR I:=1 STEP 1 UNTIL M DO
    FOR J:=1 STEP 1 UNTIL P DO
      FOR K:=1 STEP 1 UNTIL N DO
        C[I, J] := A[I, K] * B[K, J];
      ENDJ
    ENDJ
  END MULTIPLICAMATRIZJ

BOOLEAN PROCEDURE AUTOCORR(D, M, N, NOR, FS, ARR)
VALUE M, N: INTEGER M, N; ARRAY D[*,*], NOR[*,*]; FILE FS; ARRAY ARR[*,*];
BEGIN
  INTEGER I, J, L;

  ARRAY B[1:1:M];
  REAL ZH, Q, DEN, NUM, DB;
  WRITE(FS, <"AUTOCORRELACION"/>);
  FOR I:=1 STEP 1 UNTIL M DO
    FOR J:=1 STEP 1 UNTIL N DO BEGIN
      B[I-1] := D[I, J] - DB;
      DB := (D[I, J] + DB) / (M+N);
    ENDJ
  ENDJ
  FOR I:=2 STEP 1 UNTIL M+N DO BEGIN
    NUM := ((B[I] - B[I-1]) ** 2);
    DEN := ((B[I] - DB) ** 2);
    DB := ((B[I] - DB) ** 2);
    C := (M * NUM) / ((M+N) * DEN);
    WRITE(FS, <"EL VALOR DE LA RAZON DE VON NEUMANN ES: "/>, R10.6/>, Q);
    FOR I:=1 STEP 1 UNTIL 2 DO
      IF Q GEQ NOR[2*I-1] AND Q LEQ NOR[2*I] THEN BEGIN
        L:=L+1;
        WRITE(FS, <"LA INDEPENDENCIA DE LOS RESIDUOS SE ACEPTA AL "/>,
X1, I2, "X"/>, X15, "SIENDO LOS EXTREMOS DEL INTERVALO IGUALES A ",
2R10.7/>, X1, "X"/>, PSIME[I], NOR[2*I-1], NOR[2*I]) END
      ELSE BEGIN L:=0;
        WRITE(FS, <"LA INDEPENDENCIA DE LOS RESIDUOS NO ES ACEPTADA AL "/>,
X1, I2, "X"/>, X15, "SIENDO LOS EXTREMOS DEL INTERVALO IGUALES A ",
2R10.7/>, X1, "X"/>, PSIME[I], NOR[2*I-1], NOR[2*I]) ENDJ
    ENDJ
  ENDJ

```

```

013:0081:2
013:0083:2
013:0093:3
013:0095:4
JICUAD(013) IS 00A1 LONG
003:0007:2
003:0007:2
003:0007:2
003:0007:2
TRANSPON IS SEGMENT 00014
014:0000:0
014:0004:3
014:0009:0
014:000E:3
TRANSPON(014) IS 0010 LONG
003:0007:2
003:0007:2
003:0007:2
003:0007:2
003:0007:2
003:0007:2
MULTIPLICAMATRIZ IS SEGMENT 00015
015:0000:0
015:0000:0
015:0004:3
015:0009:0
015:000E:3
015:0016:1
MULTIPLICAMATRIZ(015) IS 0016 LONG
003:0007:2
003:0007:2
003:0007:2
003:0007:2
AUTOCORR IS SEGMENT 00016
016:0000:0
016:0003:1
016:0003:1
016:0003:5
016:0008:2
016:0011:5
016:001A:3
016:001C:2
016:0021:3
016:0025:2
016:0025:4
016:0028:0
016:002E:4
016:003C:4
016:003B:2
016:0042:4
016:0043:0
016:0046:0
016:0046:0
016:005A:1
016:0059:2
016:005D:2
016:005D:2

```

END;

```

    BOOLEAN PROCEDURE SIMETRIAC(D,H,N,K2,CONT,SAL,Z,ARR);
    VALUE  M,N,CONT; INTEGER M,N,CONT; ARRAY D[*,*]; REAL Z,K2; FILE SAL;
    ARRAY ARR[*];
    BEGIN
        INTEGER I,J,C;

```

```

        REAL K1,SKEW;
        WRITE(SAL,</"SIMETRIA: ",/,>);
        K2:=0;
        FOR I:=1 STEP 1 UNTIL M DO
            FOR J:=1 STEP 1 UNTIL N DO
                Z1:=D(I,J);
                Z1:=Z1/(CN*M);
                FOR I1:=1 STEP 1 UNTIL M DO
                    FOR J1:=1 STEP 1 UNTIL N DO BEGIN

```

```

                        K11:= + (D(I1,J1)-Z)**3;
                        K21:= + (D(I1,J1)-Z)**2;
                        SKEW:=(K1-(CN*M)**((3/2)/(CN*M*K2**((3/2)))));
                        WRITE(SAL,<X15,"EL COEFICIENTE DE SIMETRIA PARA LA MATRIZ DE "
                        "RESIDUOS DE LA ITERACION ">I3," ES: ">R8.4/,>,CONT,SKEW);
                        FOR I1:=1 STEP 1 UNTIL 2 DO
                            IF ABS(SKEW) LSS SIMET[I] THEN BEGIN C1:=C+1;
                            WRITE(SAL,<X15,"DE DONDE, LOS DATOS TIENEN UNA SIMETRIA MARCADA AL"
                            "I2,">X15,"EL VALOR DE LA COTA ES IGUAL A ">R10.6/,>,PSIMET[I],
                            SIMET[I]); END
                            ELSE BEGIN C1:=0;
                            WRITE(SAL,<X15," DE DONDE LOS RESIDUOS NO SON SIMETRICOS AUN-,">X15,
                            "EL VALOR DE LA COTA ES IGUAL A ">R10.6/,>,SIMET[I]);
                            X15,"EL VALOR DE LA COTA ES IGUAL A ">R10.6/,>,SIMET[I]);
                        END; IF C EQL 2 THEN SIMETRIA := TRUE ELSE SIMETRIA := FALSE; END;

```

```

    BOOLEAN PROCEDURE KURTOSIS(D,H,N,CONT,SAL,Z,K2,ARR);
    VALUE  M,N,CONT; INTEGER M,N,CONT; REAL K2,Z; ARRAY D[*,*],ARR[*];
    FILE SAL;
    BEGIN
        INTEGER I,J; INTEGER C;

```

```

        REAL R1,KURT;
        WRITE(SAL,</"KURTOSIS:">X15/,>);
        FOR I:=1 STEP 1 UNTIL M DO
            FOR J:=1 STEP 1 UNTIL N DO
                R11:= + (D(I,J)-Z)**4;
                KURT:=R1-M*N/(K2**2);
                WRITE(SAL,<X15,"EL COEFICIENTE DE KURTOSIS EN LA ITERACION ">I3,
                " ES IGUAL A ">R10.6/,>,CONT,KURT);
                FOR J1:=1 STEP 1 UNTIL 2 DO
                    IF ABS(KURT) GEQ KURTO(2*J-1) AND ABS(KURT) LEQ KURTO(2*J) THEN
                        BEGIN C1:=C+1;
                        WRITE(SAL,<X15,"POR LO QUE LOS DATOS TIENEN UNA KURTOSIS NORMAL,">X15,
                        "LOS EXTREMOS DEL INTERVALO SON ">X15,">R10.6/,>,KURTO(2*J-1),
                        KURTO(2*J));
                        END
                        ELSE BEGIN C1:=0;
                        WRITE(SAL,<X15,"POR LO QUE LOS DATOS AUN SE ENCUENTRAN LEJOS DE "
                        "TENER UNA KURTOSIS NORMAL,">X15,"LOS EXTREMOS DEL INTERVALO SON"
                        ">X15,">R10.6/,>,KURTO(2*J-1),KURTO(2*J));
                        END;
                IF C EQL 2 THEN KURTOSIS := TRUE ELSE KURTOSIS := FALSE;

```

END;

AUTOCURR(016) IS 0003 LONG

```

2      003:0007:2
      003:0007:2
      003:0007:2
      003:0007:2
      003:0007:2

```

SIMETRIA- IS SEGMENT 00017

```

2      017:0000:0
      017:0000:0
      017:0005:4
      017:0006:3
      017:0008:0
      017:000F:3
      017:0013:4
      017:0015:4
      017:001A:1
      017:001E:4
      017:0022:5
      017:0027:5
      017:0030:0
      017:0032:0
      017:0040:1
      017:0040:5
      017:0044:3
      017:0046:3
      017:004F:3
      017:0056:1
      017:0057:2
      017:0059:2
      017:0066:1

```

SIMETRIAC017) IS 0074 LONG

```

2      003:0007:2
      003:0007:2
      003:0007:2
      003:0007:2
      003:0007:2

```

KURTOSIS- IS SEGMENT 00018

```

2      018:0000:0
      018:0000:0
      018:0005:4
      018:000A:1
      018:000E:4
      018:0013:5
      018:0016:3
      018:0018:3
      018:0026:4
      018:0027:2
      018:002C:3
      018:002E:2
      018:0030:2
      018:003A:1
      018:0041:4
      018:0041:4
      018:0042:5
      018:0044:5
      018:0044:5
      018:0058:3
      018:0058:4

```

SINCLUE "BIBLIO/EIGENVECTOR."

*****WARNING-1-AFTER-THIS-RELEASE-ONLY-DOLLAR-CARDS-WITH-S-IN-COL-2-OR-3-WILL-BE-WRITTEN-ON-THE-NEWTAPE*****

NONSYPEIGENVALUESANDEIGENVECTORS IS SEGMENT 00019

DATA IS 000C LONG

DATA IS 0006 LONG

INNERPRODUCT IS SEGMENT 00019

INNERPRODUCT(018) IS 000A LONG

HESSBERG IS SEGMENT 00010

HESSBERG(010) IS 0076 LONG

TRANSFORM IS SEGMENT 0001E

TRANSFORM(01E) IS 0022 LONG

NORMREAL IS SEGMENT 0001F

NORMREAL(01F) IS 0023 LONG

NORMCOMPLEX IS SEGMENT 00020

NORMCOMPLEX(020) IS 000E LONG

GAUSSM IS SEGMENT 00021

GAUSSM(021) IS 0039 LONG

SOLVE IS SEGMENT 00022

SOLVE(022) IS 0014 LONG

HOR IS SEGMENT 00023

HOR(023) IS 0005 LONG

EIGENVALUES IS SEGMENT 00024

B.0001 IS SEGMENT 00025

B.0001(025) IS 0043 LONG

EIGENVALUES(024) IS 000B LONG

EIGENVECTORS IS SEGMENT 00026

EIGENVECTORS(026) IS 0110 LONG

NONSYPEIGENVALUESANDEIGENVECTORS(019) IS 00AD LONG

PROCEDURE ESTIMAU(CU,TETA,M,D,SAL,CONT,PTE,ENT);

VALUE M;INTEGER M,CONT,PTE;ARRAY U(0:M,1:M),TETA(1:M),D(1:M);FILE SAL,ENT;

BEGIN

REAL SC;

INTEGER I,J;

ARRAY DP(1:M,1:M),UP(1:M,0:M),TETAP,TETA(0:M),TOR(1:M,1:M);

TETAPR(1,1:M);

DEFINE NORMA(A,X)=FOR I=1 STEP 1 UNTIL M DO

FOR J=1 STEP 1 UNTIL M DO

A(I,J):=A(X)/X;

WRITE(SAL, <<"ESTIMA U"/>>);

TRANSFON(D,M,N,DP);

MULTIPLICAMATRIZ(C,DP,T,M,N,M,SAL,ENT);

ORDENACT(M,"TOR,PTE");

IF MAX(ABS(TOR(1,1)),ABS(TOR(M,M))) GTR 1 THEN

IF ABS(TOR(1,1)) LEQ ABS(TOR(M,M)) THEN

NORMA(T,ABS(TOR(M,M)))

ELSE NORMA(T,ABS(TOR(1,1)))

NONSYPEIGENVALUESANDEIGENVECTORS(M,T,TETAP,TETA,UP);

FOR J=1 STEP 1 UNTIL M DO

TETAPR(1,J):=TETAP(J);

ORDENACT(TETAPR,1,M,TETAPR,PTE);

IF MAX(ABS(TOR(1,1)),ABS(TOR(M,M))) GTR 1 THEN

TETA(CONT):=TETAPR(1,1)+MAX(ABS(TOR(1,1)),ABS(TOR(M,M)))-ELSE

TETA(CONT):=TETAPR(1,1);

FOR J=1 STEP 1 UNTIL M DO

U(CONT,J):=UP(J,PTE);

ANDRPA(CU,M);

IMPRIMEMATRIZ(SAL,U,M,N(M,M),M);

WRITE(SAL,<<X15>>"EN LA ITERACION-"I3>> EL EIGENVALOR ELEGIDO RESULTO

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2

2 003:0007:2


```

SINETRIA(D,H,N,SIGMA,CONT,SAL,MED,SIGMA,KURTO) AND
KURTOSIS(D,H,N,CONT,SAL,MED,SIGMA,KURTO) THEN
GO TO YA ELSE GO TO AUNNO
ELSE IF HOMOSCED(D,H,N,SAL,COCH) AND
JICUAD(D,H,N,SCOA,JITA,PNOR,SAL,ORD) AND
SINETRIA(D,H,N,SIGMA,CONT,SAL,MED,SIGMA) AND
AUTOCORR(D,H,N,VON,SAL,ALFA) AND
KURTOSIS(D,H,N,CONT,SAL,MED,SIGMA,KURTO) THEN
GO TO YA ELSE GO TO AUNNO
YA: WRITE(SAL,*,X15,"LOS RESIDUOS TIENEN YA UNA DISTRIBUCION MUY ",
"PROXIMA ALA NORMAL"/>)) ; GO TO LEE
AUNNO: CONT=CONT+1
ESTIMAU(U,TETA,H,D,SAL,CONT,PTE,ENT)
ESTIMAV(CONT,TETA(CONT),U,V,H,N,D,SAL)
RESIDUOS2(H,N,D,U,V,CONT,TETA,SAL)
GO TO REG
END
FIN: END.

```

0201004E14	B,0002(02D)	IS 0092 LONG
0201005315	2	0031008115
0201005415	B,0000(003)	IS 00C4 LONG
0201005911		DATA IS 004B LONG
0201005F13		
0201006510		
0201006A10		
0201006F11		
0201007011		
0201007210		
0201007615		
0201007810		
0201007E10		
0201008415		
0201008B13		
0201008C10		

RESUMEN DE LA COMPILACION EN ALGOL

COMPILACION CORRECTA.

NOMBRE DEL PROGRAMA = PROG/PRUEBA.

TIEMPO EN LA MEZCLA = 51 SEG.
TIEMPO DE COMPILACION = 15 SEG.
TIEMPO DE ENTRADA/SALIDA = 11 SEG.

MEMORIA ESTIMADA = 7630 PALABRAS.
TORRE DE APILAMIENTO = 123 PALABRAS.

NUMERO DE SEGMENTOS = 43.
LONGITUD DE SEGMENTOS = 3684 PALABRAS.
NUMERO DE TARJETAS = 1061.

NUMERO DE ELEMENTOS SINTACTICOS = 9301.
NUMERO DE SEGMENTOS DE DISCO = 277.

```

100 BEGIN
101 FILE WINDOW=LEFOUTE;
102 INTEGER K,L;
103 PROCEDURE IMPRIMERATLIZ(SOL,A,M,N);
104 VALUE M,N; INTEGER M,N; ARRAY A(*,*) FILE SOL;
105 BEGIN INTEGER I,J;
106 WRITE(SOL,<"",<"----->",N);
107 FOR I:=1 STEP 1 UNTIL M DO
108 BEGIN
109 WRITE(SOL,<"",<"----->",N);
110 WRITE(SOL,<"",5*(M-1),<"----->",N);
111 FOR J:=1 STEP 1 UNTIL N DO
112 AT(J);
113 END;
114 END;
115 PROCEDURE NORMAL(DE,EX);
116 VALUE DE,EX; REAL DE,EX;
117 BEGIN
118 ARRAY NOR,UNIT(1:4,1:5);
119 REAL BASE;
120 INTEGER I,J;
121 BASE:=(944153+7)*I;
122 FOR I:=1 STEP 1 UNTIL 4 DO
123 FOR J:=1 STEP 1 UNTIL 5 DO
124 BEGIN
125 UNIT(I,J):=3*I+3*(J-1)+((I-1)*(J-3))*SE+RANDOM(BASE);
126 NOR(I,J):=2*(1-2*RANDOM(BASE))*SINC(6.283185309)*RANDOM(BASE);
127 (14/50)+3*(1+((J-1)+((I-1)*(J-3))*SE);
128 END;
129 IMPRIMERATLIZ(UNIT,NOR,4,5);
130 IMPRIMERATLIZ(LEF,NOR,4,5);
131 AND NOR:
132 FOR K:=1 STEP 5 UNTIL 21 DO
133 FOR L:=1 STEP 2 UNTIL 9 DO
134 FOR M:=1,2,3,4,5,6,7,8,9,10,11,12,13,14,15,16,17,18,19,20,21,22,23,24,25,26,27,28,29,30,31,32,33,34,35,36,37,38,39,40,41,42,43,44,45,46,47,48,49,50,51,52,53,54,55,56,57,58,59,60,61,62,63,64,65,66,67,68,69,70,71,72,73,74,75,76,77,78,79,80,81,82,83,84,85,86,87,88,89,90,91,92,93,94,95,96,97,98,99,100,101,102,103,104,105,106,107,108,109,110,111,112,113,114,115,116,117,118,119,120,121,122,123,124,125,126,127,128,129,130,131,132,133,134,135,136,137,138,139,140,141,142,143,144,145,146,147,148,149,150,151,152,153,154,155,156,157,158,159,160,161,162,163,164,165,166,167,168,169,170,171,172,173,174,175,176,177,178,179,180,181,182,183,184,185,186,187,188,189,190,191,192,193,194,195,196,197,198,199,200,201,202,203,204,205,206,207,208,209,210,211,212,213,214,215,216,217,218,219,220,221,222,223,224,225,226,227,228,229,230,231,232,233,234,235,236,237,238,239,240,241,242,243,244,245,246,247,248,249,250,251,252,253,254,255,256,257,258,259,260,261,262,263,264,265,266,267,268,269,270,271,272,273,274,275,276,277,278,279,280,281,282,283,284,285,286,287,288,289,290,291,292,293,294,295,296,297,298,299,300,301,302,303,304,305,306,307,308,309,310,311,312,313,314,315,316,317,318,319,320,321,322,323,324,325,326,327,328,329,330,331,332,333,334,335,336,337,338,339,340,341,342,343,344,345,346,347,348,349,350,351,352,353,354,355,356,357,358,359,360,361,362,363,364,365,366,367,368,369,370,371,372,373,374,375,376,377,378,379,380,381,382,383,384,385,386,387,388,389,390,391,392,393,394,395,396,397,398,399,400,401,402,403,404,405,406,407,408,409,410,411,412,413,414,415,416,417,418,419,420,421,422,423,424,425,426,427,428,429,430,431,432,433,434,435,436,437,438,439,440,441,442,443,444,445,446,447,448,449,450,451,452,453,454,455,456,457,458,459,460,461,462,463,464,465,466,467,468,469,470,471,472,473,474,475,476,477,478,479,480,481,482,483,484,485,486,487,488,489,490,491,492,493,494,495,496,497,498,499,500,501,502,503,504,505,506,507,508,509,510,511,512,513,514,515,516,517,518,519,520,521,522,523,524,525,526,527,528,529,530,531,532,533,534,535,536,537,538,539,540,541,542,543,544,545,546,547,548,549,550,551,552,553,554,555,556,557,558,559,560,561,562,563,564,565,566,567,568,569,570,571,572,573,574,575,576,577,578,579,580,581,582,583,584,585,586,587,588,589,590,591,592,593,594,595,596,597,598,599,600,601,602,603,604,605,606,607,608,609,610,611,612,613,614,615,616,617,618,619,620,621,622,623,624,625,626,627,628,629,630,631,632,633,634,635,636,637,638,639,640,641,642,643,644,645,646,647,648,649,650,651,652,653,654,655,656,657,658,659,660,661,662,663,664,665,666,667,668,669,670,671,672,673,674,675,676,677,678,679,680,681,682,683,684,685,686,687,688,689,690,691,692,693,694,695,696,697,698,699,700,701,702,703,704,705,706,707,708,709,710,711,712,713,714,715,716,717,718,719,720,721,722,723,724,725,726,727,728,729,730,731,732,733,734,735,736,737,738,739,740,741,742,743,744,745,746,747,748,749,750,751,752,753,754,755,756,757,758,759,760,761,762,763,764,765,766,767,768,769,770,771,772,773,774,775,776,777,778,779,780,781,782,783,784,785,786,787,788,789,790,791,792,793,794,795,796,797,798,799,800,801,802,803,804,805,806,807,808,809,810,811,812,813,814,815,816,817,818,819,820,821,822,823,824,825,826,827,828,829,830,831,832,833,834,835,836,837,838,839,840,841,842,843,844,845,846,847,848,849,850,851,852,853,854,855,856,857,858,859,860,861,862,863,864,865,866,867,868,869,870,871,872,873,874,875,876,877,878,879,880,881,882,883,884,885,886,887,888,889,890,891,892,893,894,895,896,897,898,899,900,901,902,903,904,905,906,907,908,909,910,911,912,913,914,915,916,917,918,919,920,921,922,923,924,925,926,927,928,929,930,931,932,933,934,935,936,937,938,939,940,941,942,943,944,945,946,947,948,949,950,951,952,953,954,955,956,957,958,959,960,961,962,963,964,965,966,967,968,969,970,971,972,973,974,975,976,977,978,979,980,981,982,983,984,985,986,987,988,989,990,991,992,993,994,995,996,997,998,999,1000,1001,1002,1003,1004,1005,1006,1007,1008,1009,1010,1011,1012,1013,1014,1015,1016,1017,1018,1019,1020,1021,1022,1023,1024,1025,1026,1027,1028,1029,1030,1031,1032,1033,1034,1035,1036,1037,1038,1039,1040,1041,1042,1043,1044,1045,1046,1047,1048,1049,1050,1051,1052,1053,1054,1055,1056,1057,1058,1059,1060,1061,1062,1063,1064,1065,1066,1067,1068,1069,1070,1071,1072,1073,1074,1075,1076,1077,1078,1079,1080,1081,1082,1083,1084,1085,1086,1087,1088,1089,1090,1091,1092,1093,1094,1095,1096,1097,1098,1099,1100,1101,1102,1103,1104,1105,1106,1107,1108,1109,1110,1111,1112,1113,1114,1115,1116,1117,1118,1119,1120,1121,1122,1123,1124,1125,1126,1127,1128,1129,1130,1131,1132,1133,1134,1135,1136,1137,1138,1139,1140,1141,1142,1143,1144,1145,1146,1147,1148,1149,1150,1151,1152,1153,1154,1155,1156,1157,1158,1159,1160,1161,1162,1163,1164,1165,1166,1167,1168,1169,1170,1171,1172,1173,1174,1175,1176,1177,1178,1179,1180,1181,1182,1183,1184,1185,1186,1187,1188,1189,1190,1191,1192,1193,1194,1195,1196,1197,1198,1199,1200,1201,1202,1203,1204,1205,1206,1207,1208,1209,1210,1211,1212,1213,1214,1215,1216,1217,1218,1219,1220,1221,1222,1223,1224,1225,1226,1227,1228,1229,1230,1231,1232,1233,1234,1235,1236,1237,1238,1239,1240,1241,1242,1243,1244,1245,1246,1247,1248,1249,1250,1251,1252,1253,1254,1255,1256,1257,1258,1259,1260,1261,1262,1263,1264,1265,1266,1267,1268,1269,1270,1271,1272,1273,1274,1275,1276,1277,1278,1279,1280,1281,1282,1283,1284,1285,1286,1287,1288,1289,1290,1291,1292,1293,1294,1295,1296,1297,1298,1299,1300,1301,1302,1303,1304,1305,1306,1307,1308,1309,1310,1311,1312,1313,1314,1315,1316,1317,1318,1319,1320,1321,1322,1323,1324,1325,1326,1327,1328,1329,1330,1331,1332,1333,1334,1335,1336,1337,1338,1339,1340,1341,1342,1343,1344,1345,1346,1347,1348,1349,1350,1351,1352,1353,1354,1355,1356,1357,1358,1359,1360,1361,1362,1363,1364,1365,1366,1367,1368,1369,1370,1371,1372,1373,1374,1375,1376,1377,1378,1379,1380,1381,1382,1383,1384,1385,1386,1387,1388,1389,1390,1391,1392,1393,1394,1395,1396,1397,1398,1399,1400,1401,1402,1403,1404,1405,1406,1407,1408,1409,1410,1411,1412,1413,1414,1415,1416,1417,1418,1419,1420,1421,1422,1423,1424,1425,1426,1427,1428,1429,1430,1431,1432,1433,1434,1435,1436,1437,1438,1439,1440,1441,1442,1443,1444,1445,1446,1447,1448,1449,1450,1451,1452,1453,1454,1455,1456,1457,1458,1459,1460,1461,1462,1463,1464,1465,1466,1467,1468,1469,1470,1471,1472,1473,1474,1475,1476,1477,1478,1479,1480,1481,1482,1483,1484,1485,1486,1487,1488,1489,1490,1491,1492,1493,1494,1495,1496,1497,1498,1499,1500,1501,1502,1503,1504,1505,1506,1507,1508,1509,1510,1511,1512,1513,1514,1515,1516,1517,1518,1519,1520,1521,1522,1523,1524,1525,1526,1527,1528,1529,1530,1531,1532,1533,1534,1535,1536,1537,1538,1539,1540,1541,1542,1543,1544,1545,1546,1547,1548,1549,1550,1551,1552,1553,1554,1555,1556,1557,1558,1559,1560,1561,1562,1563,1564,1565,1566,1567,1568,1569,1570,1571,1572,1573,1574,1575,1576,1577,1578,1579,1580,1581,1582,1583,1584,1585,1586,1587,1588,1589,1590,1591,1592,1593,1594,1595,1596,1597,1598,1599,1600,1601,1602,1603,1604,1605,1606,1607,1608,1609,1610,1611,1612,1613,1614,1615,1616,1617,1618,1619,1620,1621,1622,1623,1624,1625,1626,1627,1628,1629,1630,1631,1632,1633,1634,1635,1636,1637,1638,1639,1640,1641,1642,1643,1644,1645,1646,1647,1648,1649,1650,1651,1652,1653,1654,1655,1656,1657,1658,1659,1660,1661,1662,1663,1664,1665,1666,1667,1668,1669,1670,1671,1672,1673,1674,1675,1676,1677,1678,1679,1680,1681,1682,1683,1684,1685,1686,1687,1688,1689,1690,1691,1692,1693,1694,1695,1696,1697,1698,1699,1700,1701,1702,1703,1704,1705,1706,1707,1708,1709,1710,1711,1712,1713,1714,1715,1716,1717,1718,1719,1720,1721,1722,1723,1724,1725,1726,1727,1728,1729,1730,1731,1732,1733,1734,1735,1736,1737,1738,1739,1740,1741,1742,1743,1744,1745,1746,1747,1748,1749,1750,1751,1752,1753,1754,1755,1756,1757,1758,1759,1760,1761,1762,1763,1764,1765,1766,1767,1768,1769,1770,1771,1772,1773,1774,1775,1776,1777,1778,1779,1780,1781,1782,1783,1784,1785,1786,1787,1788,1789,1790,1791,1792,1793,1794,1795,1796,1797,1798,1799,1800,1801,1802,1803,1804,1805,1806,1807,1808,1809,1810,1811,1812,1813,1814,1815,1816,1817,1818,1819,1820,1821,1822,1823,1824,1825,1826,1827,1828,1829,1830,1831,1832,1833,1834,1835,1836,1837,1838,1839,1840,1841,1842,1843,1844,1845,1846,1847,1848,1849,1850,1851,1852,1853,1854,1855,1856,1857,1858,1859,1860,1861,1862,1863,1864,1865,1866,1867,1868,1869,1870,1871,1872,1873,1874,1875,1876,1877,1878,1879,1880,1881,1882,1883,1884,1885,1886,1887,1888,1889,1890,1891,1892,1893,1894,1895,1896,1897,1898,1899,1900,1901,1902,1903,1904,1905,1906,1907,1908,1909,1910,1911,1912,1913,1914,1915,1916,1917,1918,1919,1920,1921,1922,1923,1924,1925,1926,1927,1928,1929,1930,1931,1932,1933,1934,1935,1936,1937,1938,1939,1940,1941,1942,1943,1944,1945,1946,1947,1948,1949,1950,1951,1952,1953,1954,1955,1956,1957,1958,1959,1960,1961,1962,1963,1964,1965,1966,1967,1968,1969,1970,1971,1972,1973,1974,1975,1976,1977,1978,1979,1980,1981,1982,1983,1984,1985,1986,1987,1988,1989,1990,1991,1992,1993,1994,1995,1996,1997,1998,1999,2000,2001,2002,2003,2004,2005,2006,2007,2008,2009,2010,2011,2012,2013,2014,2015,2016,2017,2018,2019,2020,2021,2022,2023,2024,2025,2026,2027,2028,2029,2030,2031,2032,2033,2034,2035,2036,2037,2038,2039,2040,2041,2042,2043,2044,2045,2046,2047,2048,2049,2050,2051,2052,2053,2054,2055,2056,2057,2058,2059,2060,2061,2062,2063,2064,2065,2066,2067,2068,2069,2070,2071,2072,2073,2074,2075,2076,2077,2078,2079,2080,2081,2082,2083,2084,2085,2086,2087,2088,2089,2090,2091,2092,2093,2094,2095,2096,2097,2098,2099,2100,2101,2102,2103,2104,2105,2106,2107,2108,2109,2110,2111,2112,2113,2114,2115,2116,2117,2118,2119,2120,2121,2122,2123,2124,2125,2126,2127,2128,2129,2130,2131,2132,2133,2134,2135,2136,2137,2138,2139,2140,2141,2142,2143,2144,2145,2146,2147,2148,2149,2150,2151,2152,2153,2154,2155,2156,2157,2158,2159,2160,2161,2162,2163,2164,2165,2166,2167,2168,2169,2170,2171,2172,2173,2174,2175,2176,2177,2178,2179,2180,2181,2182,2183,2184,2185,2186,2187,2188,2189,2190,2191,2192,2193,2194,2195,2196,2197,2198,2199,2200,2201,2202,2203,2204,2205,2206,2207,2208,2209,2210,2211,2212,2213,2214,2215,2216,2217,2218,2219,2220,2221,2222,2223,2224,2225,2226,2227,2228,2229,2230,2231,2232,2233,2234,2235,2236,2237,2238,2239,2240,2241,2242,2243,2244,2245,2246,2247,2248,2249,2250,2251,2252,2253,2254,2255,2256,2257,2258,2259,2260,2261,2262,2263,2264,2265,2266,2267,2268,2269,2270,2271,2272,2273,2274,2275,2276,2277,2278,2279,2280,2281,2282,2283,2284,2285,2286,2287,2288,2289,2290,2291,2292,2293,2294,2295,2296,2297,2298,2299,2300,2301,2302,2303,2304,2305,2306,2307,2308,2309,2310,2311,2312,2313,2314,2315,2316,2317,2318,2319,2320,2321,2322,2323,2324,2325,2326,2327,2328,2329,2330,2331,2332,2333,2334,2335,2336,2337,2338,2339,2340,2341,2342,2343,2344,2345,2346,2347,2348,2349,2350,2351,2352,2353,2354,2355,2356,2357,2358,2359,2360,2361,2362,2363,2364,2365,2366,2367,2368,2369,2370,2371,2372,2373,2374,2375,2376,2377,2378,2379,2380,2381,2382,2383,2384,2385,2386,2387,2388,2389,2390,2391,2392,2393,2394,2395,2396,2397,2398,2399,2400,2401,2402,2403,2404,2405,2406,2407,2408,2409,2410,2411,2412,2413,2414,2415,2416,2417,2418,2419,2420,2421,2422,2423,2424,2425,2426,2427,2428,2429,2430,2431,2432,2433,2434,2435,2436,2437,2438,2439,2440,2441,2442,2443,2444,2445,2446,2447,2448,2449,2450,2451,2452,2453,2454,2455,2456,2457,2458,2459,2460,2461,2462,2463,2464,2465,2466,2467,2468,2469,2470,2471,2472,2473,2474,2475,2476,2477,2478,2479,2480,2481,2482,2483,2484,2485,2486,2487,2488,2489,2490,2491,2492,2493,2494,2495,2496,2497,2498,2499,2500,2501,2502,2503,2504,2505,2506,2507,2508,2509,2510,2511,2512,2513,2514,2515,2516,2517,2518,2519,2520,2521,2522,2523,2524,2525,2526,2527,2528,2529,2530,2531,2532,2533,2534,2535,2536,2537,2538,2539,2540,2541,2542,2543,2544,2545,2546,2547,2548,2549,2550,2551,2552,2553,2554,2555,2556,2557,2558,2559,2560,2561,2562,2563,2564,2565,2566,2567,2568,2569,2570,2571,2572,2573,2574,2575,2576,2577,2578,2579,2580,2581,2582,2583,2584,2585,2586,2587,2588,2589,2590,2591,2592,259
```

APPENDICE II

TABLA 1VALORES CRITICOS PARA LA PRUEBA DECOCHRAN AL 1%

	4	5	6	7	8	9	10
2	0.864	0.788	0.722	0.664	0.615	0.573	0.536
3	0.781	0.696	0.626	0.568	0.521	0.481	0.447
4	0.721	0.633	0.564	0.508	0.463	0.425	0.393
6	0.641	0.533	0.487	0.435	0.393	0.359	0.331
8	0.590	0.504	0.440	0.391	0.352	0.321	0.294
10	0.554	0.470	0.408	0.362	0.325	0.295	0.270

Tomada de Johnson & Leone, Tabla 13.17, p.54

TABLA 2

VALORES CRITICOS PARA LA PRUEBA DE KOLMOGOROV-SMIRNOV

n	.20	.15	.10	.05	.01
4	.300	.319	.352	.381	.417
5	.285	.299	.315	.337	.405
6	.265	.277	.294	.319	.364
7	.247	.258	.276	.300	.348
8	.233	.244	.261	.285	.331
9	.223	.233	.249	.271	.311
10	.215	.224	.239	.258	.294
11	.206	.217	.230	.249	.284
12	.199	.212	.223	.242	.275
13	.190	.202	.214	.234	.268
14	.183	.194	.207	.227	.261
15	.177	.187	.201	.220	.257
16	.173	.182	.195	.213	.250
17	.169	.177	.189	.206	.245
18	.166	.173	.184	.200	.239
19	.163	.169	.179	.195	.235
20	.160	.166	.174	.190	.231
25	.149	.153	.165	.180	.203
30	.131	.136	.144	.161	.187
	$\frac{.736}{\sqrt{n}}$	$\frac{.768}{\sqrt{n}}$	$\frac{.805}{\sqrt{n}}$	$\frac{.886}{\sqrt{n}}$	$\frac{1.031}{\sqrt{n}}$

TABLA No.3

VALORES PARA LA PRUEBA DEL SESGOPUNTOS PORCENTUALES DE UNA COLA PARA LA DISTRIBUCION DE q .*

TAMAÑO DE MUESTRA n	PUNTOS PORCENTUALES		DESVIACION ESTANDAR
	5%	1%	
25	0.711	1.061	0.4354
30	0.662	0.986	.4052
35	0.621	0.923	.3804
40	0.587	0.870	.3596
45	0.558	0.825	.3418
50	0.534	0.787	.3264
60	0.492	0.723	.3009
70	0.459	0.673	.2806
80	0.432	0.631	.2638
90	0.409	0.596	.2498
100	0.389	0.567	.2377
125	0.350	0.508	.2139
150	0.321	0.464	.1961
175	0.298	0.430	.1820
200	0.280	0.403	.1706
250	0.251	0.360	.1531
300	0.230	0.329	.1400
350	0.213	0.305	.1298
400	0.200	0.285	.1216
450	0.188	0.269	.1147
500	0.179	0.255	.1089

- Como la distribución de q , es simétrica respecto al cero, los puntos porcentuales representan valores del 10% y 2% en pruebas de dos colas.

Tomada de: Snedecor & Cochran (34), Tabla A6 i), del Apéndice.

TABLA 4

VALORES PARA LA PRUEBA DE CURTOSIS (g_2)

TAMAÑO DE MUESTRA n	PUNTOS PORCENTUALES			
	SUPERIOR		INFERIOR	
	1%	5%	1%	5%
50	4.88	3.99	1.95	2.15
75	4.59	3.87	2.08	2.27
100	4.39	3.77	2.18	2.35
125	4.24	3.71	2.24	2.40
150	4.13	3.65	2.29	2.45
200	3.98	3.57	2.37	2.51
250	3.87	3.52	2.42	2.55
300	3.79	3.47	2.46	2.59
350	3.72	3.44	2.50	2.62
400	3.67	3.41	2.52	2.64
450	3.63	3.39	2.55	2.66
500	3.60	3.37	2.57	2.67
550	3.57	3.35	2.58	2.69
600	3.54	3.34	2.60	2.70
650	3.52	3.33	2.61	2.71
700	3.50	3.31	2.62	2.72
750	3.48	3.30	2.64	2.73
800	3.46	3.29	2.65	2.74
850	3.45	3.28	2.66	2.74
900	3.43	3.28	2.66	2.75
950	3.42	3.27	2.67	2.76
1000	3.41	3.26	2.68	2.76
1200	3.37	3.24	2.71	2.78
1400	3.34	3.22	2.72	2.80
1600	3.32	3.21	2.74	2.81
1800	3.30	3.20	2.76	2.82
2000	3.28	3.18	2.77	2.83

Tomada de: SNEDECOR & COCHRAN (34), Tabla A6 ii), del Apéndice.

TABLA 5

VALORES DE V^* (RAZON DE VON NEUMANN)
PARA LOS NIVELES DE SIGNIFICACION DEL 1% Y 5%

n	Valores de K		Valores de K'		n	Valores de K		Valores de K'	
	P=.01	P=.05	P=.95	P=.99		P=.01	P=.05	P=.95	P=.99
4	0.8341	1.0406	4.2927	4.4992	31	1.2469	1.4746	2.6587	2.8864
5	.6724	1.0255	3.9745	4.3276	32	1.2670	1.4817	2.6473	2.8720
6	.6738	1.0682	3.7318	4.1262	33	1.2667	1.4885	2.6365	2.8583
7	.7163	1.0919	3.5748	3.9504	34	1.2761	1.4951	2.6262	2.8451
8	.7575	1.1228	3.4486	3.8139	35	1.2852	1.5014	2.6163	2.8324
9	.7974	1.1524	3.3476	3.7025					
10	.8353	1.1803	3.2642	3.6091	36	1.2940	1.5075	2.6068	2.8202
					37	1.3025	1.5135	2.5977	2.8085
11	.8706	1.2062	3.1938	3.5294	38	1.3108	1.5193	2.5889	2.7973
12	.9033	1.2301	3.1335	3.4603	39	1.3188	1.5249	2.5804	2.7865
13	.9336	1.2521	3.0812	3.3996	40	1.3266	1.5304	2.5722	2.7760
14	.9618	1.2725	3.0352	3.3458					
15	.9880	1.2914	2.9943	3.2977	41	1.3342	1.5357	2.5643	2.7658
					42	1.3415	1.5408	2.5567	2.7560
16	1.0124	1.3090	2.9577	3.2543	43	1.3486	1.5458	2.5494	2.7466
17	1.0352	1.3253	2.9247	3.2148	44	1.3554	1.5506	2.5424	2.7376
18	1.0566	1.3405	2.8948	3.1787	45	1.3620	1.5552	2.5357	2.7289
19	1.0766	1.3547	2.8675	3.1456					
20	1.0954	1.3680	2.8425	3.1151	46	1.3684	1.5596	2.5293	2.7205
					47	1.3745	1.5638	2.5232	2.7125
21	1.1131	1.3805	2.8195	3.0869	48	1.3802	1.5678	2.5173	2.7049
22	1.1298	1.3923	2.7982	3.0607	49	1.3856	1.5716	2.5117	2.6977
23	1.1456	1.4035	2.7784	3.0362	50	1.3907	1.5752	2.5064	2.6908
24	1.1606	1.4141	2.7599	3.0133					
25	1.1748	1.4241	2.7426	2.9919	51	1.3957	1.5787	2.5013	2.6842
					52	1.4007	1.5822	2.4963	2.6777
26	1.1888	1.4336	2.7264	2.9718	53	1.4057	1.5856	2.4914	2.6712
27	1.2012	1.4426	2.7112	2.9528	54	1.4107	1.5890	2.4866	2.6648
28	1.2135	1.4512	2.6969	2.9348	55	1.4156	1.5923	2.4819	2.6585
29	1.2252	1.4594	2.6834	2.9177					
30	1.2363	1.4672	2.6707	2.9016	56	1.4203	1.5955	2.4773	2.6524
					57	1.4249	1.5987	2.4728	2.6465
					58	1.4294	1.6019	2.4684	2.6407
					59	1.4339	1.6051	2.4640	2.6350
					60	1.4384	1.6082	2.4596	2.6294

*Si $V \leq K$ se concluye que existe autocorrelación positiva, y
si $V \geq K'$ se concluye que existe autocorrelación negativa

Tomada de: Yamane T. (39). Tabla 11 de Apéndice.

BIBLIOGRAFIA

- (1) Bancroft, T.A., "Topics in Intermediate Statistical Methods" Ames, Iowa State University, 1968.
- (2) Granér, H., "Metodos Matemáticos de la Estadística" Ed. Aguilar, 1953.
- (3) Hernández, R. y Guerrero, V.N., "El uso de transformaciones en los modelos lineales", Tesis Profesional de Actuario, 1974.
- (4) Hillier, F.S. y Lieberman, G.J.; "Introduction to Operations Research", E. Holden Day, 1972.
- (5) Joanson, M.L. y Leone, F.C., "Statistics and Experimental Design in Engineering and the Physical Sciences", V.II, Ed. John Wiley & Sons, Inc. Segunda Impresión, 1968.
- (6) Johnston, J. "Econometric Methods", New York, McGraw Hill, 1963.
- (7) Illielfors, Journal of the American Statistical Association, Vol.62, 1967.
- (8) Mandel, J., "The Partitioning of Interaction in Analysis of Variance", Journal of Research, National Bureau of Standards, Section B, Math. Sciences, No. 4, 1973.
- (9) Mandel, J., "A new Analysis of Variance Model for Non-Additive Data", Technometrics Vol. 13, No.1, February 1971.
- (10) Mendez, R., I., "Introducción a la Metodología Estadística", Ed. Potens, 1971.
- (11) Yomano, T., "Statistics, an Introductory Analysis", 2 Ed., New York Harper & Row, 1967.
- (12) Graybill, F.A., "Introduction to linear Statistical Models", V.I., McGraw Hill, 1961.
- (12) Rao, B., "Linear Statistical Inference, and its applications", John Wiley & Sons, Inc., 1965.
- (14) Hotelling, H., "Analysis of a Complex of Statistical Variables into Principal Components", Journal of Educational Psychology 24, 417-411; 408-421 (1933).

- (15) Mandel, J., "The Distribution of Eigenvalues of Covariance Matrices of Residuals in Analysis of Variance", Journal of Research, National Bureau of Standards-B. Math. Sciences Vol 74 B, No.3, 1970.
- (16) Jandgren, B.W., "Statistical Theory" Collier MacMillan Ltd. London, 1968.
- (17) Anderson, T.W., "Introduction to Multivariate Statistical Analysis", John Wiley & Sons Ltd.
- (18) Garber, R.J., McIlvane, T.C., and Hoover, M.M. "A Method of Laying out Experiment Plots", Journal of the American Society of Agronomy, 1931.
- (19) Tukey, J.W., "One Degree of Freedom for Non-Additivity" Biometrics 5, 232-242 (1949)
- (20) Panse, V.G. y Sukhatme, P.V., "Metodos Estadísticos para Investigadores Agrícolas", Fondo de Cultura Económica, 2a. Ed., 1963.