

UNIVERSIDAD TÉCNICA DE AMBATO
FACULTAD DE CIENCIAS HUMANAS Y
DE LA EDUCACIÓN

Principios de Estadística con “R”

Una introducción a la aplicación de la
estadística en la investigación científica
utilizando el software libre “R”

Rodrigo Morales Carrasco
2010

El presente curso cubre los principios generales de la Estadística aplicada a la Investigación Científica utilizando el Software libre de estadística “R”. El estudiante aprenderá a utilizar este programa para realizar análisis descriptivo, gráficos, contraste, de hipótesis, ANOVAs, regresión y correlación.

PRINCIPIOS DE ESTADÍSTICA CON “R”

Una introducción a la aplicación de la estadística en la investigación científica utilizando el software libre “R”

Autor:

Rodrigo Morales Carrasco

Diseño Portada:

Edgar Ortiz A.

Diseño, Diagramación e Impresión:

Maxtudio - Ambato

Universidad Técnica de Ambato

Ing. Luis Amoroso
RECTOR

Dr. Galo Naranjo
VICERRECTOR ACADEMICO

Ing. Jorge León
VICERRECTOR ADMINISTRATIVO

Tiraje 1000 ejemplares
Ambato – Ecuador
2010

ÍNDICE DE CONTENIDO

	Pág
Prólogo.....	5
Estadística Descriptiva, Representación gráfica y	
Distribución normal	7
Introducción	7
Estadística descriptiva	8
Representación gráfica	12
Distribución normal	19
Comentarios	23
Referencias	24
 Prueba de hipótesis de tendencia central	
sobre una y dos muestras	25
Introducción	25
Formulación y prueba de hipótesis	25
Pruebas de normalidad de una muestra	29
Pruebas sobre una muestra	31
Pruebas sobre dos muestras	32
Prueba sobre dos muestras apareadas	35
Pruebas sobre dos muestras no normales	38
Comentarios	41
Referencias	42
 Análisis de Varianza (ANOVAs)	43
Introducción	43
ANOVA entr sujetos unifactorial	44
ANOVA factorial entre sujetos con dos factores	53
ANOVA intra-sujetos dew un factor	57
ANOVA factorial intra-sujetos (dos factores “intra”)	60
ANOVA en diseño de cuadrado latino	65
Comentarios	68
Referencias	69
 Correlación y Regresión	71
Introducción	71
Análisis de regresión	75

Correlación y regresión múltiple	81
Comentarios	83
Referencias	84
Glosario	85
Bibliografía	91
Recursos en Internet sobre R	93
Casos prácticos de trabajos de investigación realizados con el Programa Estadístico “R”	97
Creencias sobre la naturaleza de la ciencia, un estudio con Profesores y Estudiantes de la UTA.	99
Estudio comparativo del síndrome de Burnout entre Facultades, de empleados y Profesores de la UTA.	117
Uso de las NTIC’s y su influencia en el rendimiento académico de los estudiantes universitarios	129
Niveles de autoestima en estudiantes universitarios	143
Religión que profesan los estudiantes de la UTA y su relación con el rendimiento académico	147

PROLOGO

Rodrigo Morales Carrasco, profesor titular de la Universidad Técnica de Ambato, es el autor de esta obra que se intitula *“Principios de Estadística con R. Una introducción a la aplicación de la estadística en la investigación científica utilizando el software Ubre R”*. En ella se enfocan de manera compendiada diversos temas de estadística, a saber: (a) “Estadística descriptiva”, páginas 7-24; (b) “Prueba de hipótesis y de tendencia central sobre una y dos muestras”, páginas 25-42; (c) “Análisis de varianza (Anovas)”, páginas 43-69; (d) “Correlación y Regresión”, páginas 71-84.

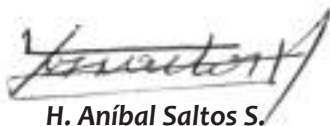
Además, el libro contempla diversos temas aplicados, que son consecuencia de investigaciones realizadas en el ámbito socio educacional en que ha incursionado el citado docente universitario, a saber: “Casos prácticos de trabajos de investigación realizados con el programa estadístico R”, p. 97; “Creencias sobre la naturaleza de la ciencia, un estudio con profesores y estudiantes de la UTA”, p. 99; “Estudio comparativo del síndrome de Brunout entre empleados y profesores de la UT A”, p. 117; “El uso de las NTIC,S y su influencia en el rendimiento académico de los estudiantes universitarios”, p. 129; “Niveles de autoestima en estudiantes universitarios, p. 143 y “Religión que profesan los estudiantes de la UTA y su relación con el rendimiento académico”, p. 147.

Estos temas se desarrollan mediante la aplicación del software libre R, el mismo que puede utilizarse libremente desde las plataformas disponibles de internet. En tal ámbito-conviene destacarse que si bien existen diversos programas utilitarios muy buenos para resolver problemas de índole estadística de diversa complejidad, tales como: SAS, JMP, STAT-GRAPHICS, SPSS, MINITAB, etc, no es menos cierto que su aplicación depende de varios factores entre los que destaco como principales a: (1) la profundidad del conocimiento estadístico matemático del usuario, (2) el ámbito del problema (¿monovariante, multivariante, diseño de experimentos, control de la calidad, econometría, etc? (3) Costo y licen-

cia.

El último factor es determinante para que el desarrollo de software libre y su consiguiente aplicación y uso tenga una tendencia creciente. En consecuencia, la propuesta de Rodrigo Morales es oportuna para introducir en el ámbito de los potenciales usuarios de la UTA las temáticas estadísticas citadas en su trabajo, con una herramienta de cálculo útil y sin costo.

Por lo expuesto, considero meritoria la contribución que hace Rodrigo Morales, la misma que debe merecer la acogida de profesores y estudiantes que consideran a la investigación. científica una actividad relevante.



H. Aníbal Saltos S.

*M.Sc., Estadística Matemática, CIENES-U. de Chile, Santiago de Chile
M.Sc., Economía Agrícola. University of Missouri-Columbia, USA*

ESTADÍSTICA DESCRIPTIVA REPRESENTACIÓN GRÁFICA Y DISTRIBUCIÓN NORMAL

INTRODUCCIÓN

“Es remarcable que una ciencia la cual comenzó con el estudio sobre las chances en juegos de azar, se haya convertido en el objeto más importante del conocimiento humano... las preguntas más importantes sobre la vida son, en su mayor parte, en realidad solo problemas de probabilidad”

Pierre Simon, Marquez de Laplace (1749-1827)

La famosa frase hecha por el Marques de la Place, muestra como a partir del estudio de algo que a simple vista puede ser considerado de una relativa importancia, se puede llegar a algo muy importante, en este caso la estadística es algo realmente importante, cualquiera sea nuestra actividad científica o tecnológica.

La palabra estadística deriva de la palabra italiana statista. Persona que trata asuntos de estado, originalmente se le llamó “aritmética de estado” e involucraba representar con tablas la información relativa a las naciones, especialmente a aquellos datos relacionados con los impuestos y la planificación de la viabilidad de las guerras.

La idea de recolectar estadísticas derivó de requerimientos gubernamentales, pero también de la necesidad en tiempos antiguos de calcular la posibilidad de naufragios y piratería con el propósito de administrar los seguros marítimos para fomentar viajes comerciales y de exploración a sitios lejanos.

El estudio moderno de los índices de mortalidad y seguros de vida se originó en las fosas donde depositaban los cadáveres de las víctimas de la plaga del siglo XVII; allí se contaban los cadáveres de personas muertas en el esplendor de su juventud.

En los comienzos del desarrollo de la estadística en los siglos XVII y XVIII, era usual que se utilizaran los nuevos métodos para probar la existencia de Dios. Por ejemplo, John Arbuthnot descubrió que en Londres, entre los años 1629 y 1790, nacieron mas bebés de sexo masculino que femenino. Mediante lo que se considera el primer caso de utilización de una prueba estadística, probó que el índice de natalidad masculino era mayor de lo que la razón por azar hubiera indicado (asumiendo en este un porcentaje del 50% para cada sexo), llegando a la conclusión de que se estaba cumpliendo un plan determinado para contrarrestar el hecho de que los hombres enfrentaban mayores peligros para obtener el sustento para sus familias y que dicha planificación según él, solo podía haber sido realizado por Dios.

En el año de 1767, John Michell también utilizó la teoría de la probabilidad para demostrar la existencia de Dios, cuando argumentó que las posibilidades de que seis estrellas se ubicaran tan cerca como lo estaban las de la constelación de Pléyades eran de 500000 a 1, y que por ende su ubicación debía ser un acto deliberado del Creador.(tomado del libro Estadística para Psicología de Aron y Aron).

Los ejemplos que se presentaran en este documento fueron implementados con el “lenguaje R”, dicho lenguaje es un entorno con capacidad de programación y graficación, desarrollado originalmente por laboratorios Bell por John Chambers y colaboradores, es fácil de usar y se ha convertido en un proyecto de colaboración entre investigadores a lo largo del mundo, es gratis se lo puede “bajar” por internet en el sitio oficial del proyecto (R project).

Se sugiere al lector remitirse al **glosario** frecuentemente para una correcta utilización de los comandos.

ESTADÍSTICA DESCRIPTIVA

Una definición de estadística descriptiva es describir los datos de una investigación de una manera concisa, la forma más común de describir un conjunto de datos relacionados entre sí es reportar un valor medio y una dispersión alrededor de dicho valor medio.

Para comenzar nuestros ejemplos, necesitamos un conjunto de datos, tenemos las calificaciones finales obtenidas por 80 estudiantes de la UTA:

68,73,61,66,96,79,65,86,84,79,65,78,78,62,80,67,75,88,75,82,89,67,73,73,82,73,87,75,61,97,57,81,68,60,74,94,75,78,88,72,90,93,62,77,95,85,78,63,62,71,95,69,60,76,62,76,88,59,78,74,79,65,76,75,76,85,63,68,83,71,53,85,93,75,72,60,71,75,74,77.

En lenguaje R podemos ingresar este conjunto de datos de la siguiente forma:

```
Calificaciones=c(68,73,61,66,96,79,65,86,84,79,65,78,78,62,80,67,75,88,75,82,89,67,73,73,82,73,87,75,61,97,57,81,68,60,74,94,75,78,88,72,90,93,62,77,95,85,78,63,62,71,95,69,60,76,62,76,88,59,78,74,79,65,76,75,76,85,63,68,83,71,53,85,93,75,72,60,71,75,74,77).
```

Podemos verificar la cantidad de datos ingresados, para esto usamos la función `length` así:

```
length(Calificaciones)
[1] 80
```

La medida básica para describir el valor central de un conjunto de datos es el valor medio o *media* del mismo definido por la ec:

$$\bar{X} = \frac{1}{n} \sum_{i=1}^n X_i$$

En R se calcula simplemente así:

```
mean(Calificaciones)
[1] 75.25
```

Otro promedio utilizados para describir datos es la *mediana* definido como el valor medio cuando los datos son ordenados de menor a mayor, en lenguaje R:

```
median(Calificaciones)
[1] 75
```

Este promedio es utilizado cuando los datos contienen valores extremos.

Para poder describir mejor un conjunto de datos necesitamos una medida de *dispersión* además de un valor central, la más simple es el *rango* el cual muestra los valores mínimo y máximo de los datos, en R:

```
range(Calificaciones)
[1] 53 97
```

La varianza y la desviación estándar son las medidas de dispersión más comunes, la varianza de un conjunto de datos se define como:

$$s^2 = \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2$$

El desvío estándar se define como la raíz cuadrada de la varianza

$$s = \sqrt{\frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2}$$

En lenguaje R se pueden calcular así:

```
var(Calificaciones)
[1] 107.6329
```

```
sd(Calificaciones)
[1] 10.37463
```

Otras medidas de dispersión más sofisticadas pero no menos útiles son los cuantiles o percentilos, por ejemplo la mediana se interpreta como el valor que separa los datos en dos mitades, en otras palabras el 50% de los valores es menor que la mediana y el otro 50% es mayor que la mediana, la mediana puede ser considerado un caso especial de cuantilo. En lenguaje R se calcula así:

```
quantile(Calificaciones,0.5)
50%
75
```

Como podemos apreciar en el párrafo anterior la función de R “quantile” necesita un argumento adicional al conjunto de datos, en este eje.es 0.5 y se denomina probabilidad (debe ser de 0 a 1) . la denominación percentilo se utiliza cuando la probabilidad tiene un valor entre (0% y 100%), siendo los valores mínimo y máximo del conjunto de datos respectivamente. Por ejemplo el percentilo 10 se puede calcular en R así:

```
quantile(Calificaciones,0.1)
10%
61.9
```

El percentilo 10 se interpreta entonces como el valor para el cual el 10% de los datos del conjunto son menores al mismo.

El lenguaje R también provee una función denominada *fivenum* que significa cinco números en castellano, propuesta por el famoso estadístico John W. Tukey, la cual calcula 5 valores que describen concisamente un conjunto de datos, son los valores mínimo, los percentilos 25 50 75 y el valor máximo.

```
fivenum(Calificaciones)
[1] 53.0 67.5 75.0 82.0 97.0
```

Otra medida de dispersión es el *coeficiente de variación* definido por

$$CV = 100\% \frac{s}{x}$$

Esta medida es fácil de interpretar, pero en la mayoría de los casos no es adecuada para comparar dos o más conjuntos de datos. En R tenemos:

```
CV=100*sd(Calificaciones)/mean(Calificaciones)

[1] 13.78688
```

REPRESENTACIÓN GRÁFICA

La estadística descriptiva nos permite caracterizar con números un conjunto de datos, sin embargo en ciertas ocasiones un gráfico nos permite comunicar mejor las características de los datos. El gráfico de caja (box plot en inglés) es la forma gráfica de los cinco números, como podemos ver en la figura 1 la caja muestra los percentilos 25 y 75, la línea en el medio de la caja es la mediana (percentilo 50), los extremos muestran los valores mínimo y máximo.

```
> boxplot(Calificaciones,main="Estudio rendimiento
matemáticas",ylab="Calificaciones")
```

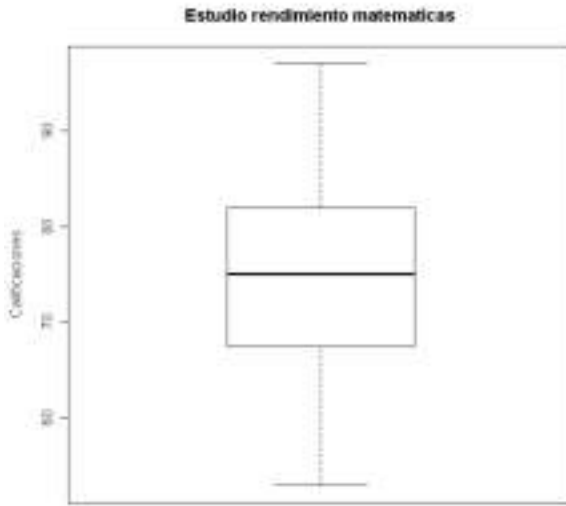


Figura 1 gráfico de caja

Otra opción muy utilizada por estadísticos es el gráfico de *rama y hoja* (stem and leaf plot en inglés), en lenguaje R se puede calcular de esta forma:

```
> stem(Calificaciones)
```

The decimal point is 1 digit(s) to the right of the |

```

5 | 3
5 | 79
6 | 00011222233
6 | 5556778889
7 | 111223333444
7 | 555555566667788888999
8 | 012234
8 | 555678889
9 | 0334
9 | 5567

```

Los números que se encuentran a la izquierda del carácter | son los dígitos más significativos como advierte la leyenda anterior al gráfico el punto decimal está ubicado un dígito a la derecha del carácter | en otras palabras la primera línea 5 | 3 se lee como el primer valor 53 el siguiente es 57, 59, 60, etc.

La representación gráfica más popular de un conjunto de datos es el histograma, el cual representa la frecuencia de aparición de valores dentro del rango del conjunto de datos. La figura 2 muestra el histograma para los datos de la tabla 1

El histograma en Lenguaje R se calcula como se muestra a continuación, note los parámetros adicionales de la función hist para determinar el título principal (main) y los rótulos de cada eje (xlab e ylab):

```
>hist(Calificaciones,main="Estudio rendimiento matematicas", xlab=
"Calificaciones", ylab= "Frecuencia")
```

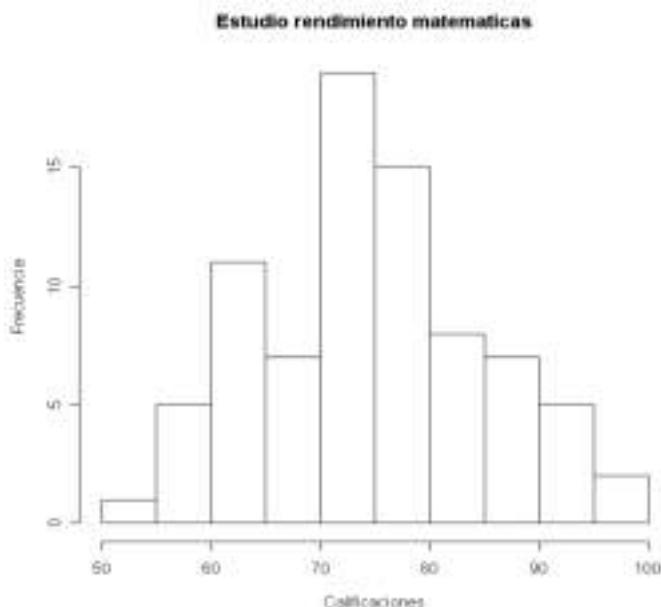


Figura 2: histograma de frecuencias

Como parte de un estudio más amplio sobre el comportamiento social de alumnos universitarios, Tracy McLaughlin-volpe y sus colegas (1998) hicieron que 94 estudiantes del ciclo de Introducción a la Psicología llevaran un diario de sus interacciones sociales durante una semana del semestre. Cada vez que los participantes tuvieran una interacción social de 10 minutos o más, deberían llenar una tarjeta. La tarjeta incluía preguntas tales como quienes eran las otras personas con las que interactuaban, cómo se sintió el alumno durante la interacción y varios aspectos relacionados con la naturaleza de la conversación mantenida. Excluyendo las situaciones familiares y laborales, la cantidad de interacciones sociales fueron las siguientes:

48,15,33,3,21,19,17,16,44,25,30,3,5,9,35,32,26,13,14,14,47,47,29,18,11,5,19,24,17,6,25,8,18,29,1,18,22,3,22,29,2,6,10,29,10,21,38,41,16,17,8,40,8,10,18,7,4,4,8,11,3,23,10,19,21,13,12,10,4,17,11,21,9,8,7,5,3,22,14,25,4,11,10,18,1,28,27,19,24,35,9,30,8,26

Con R tenemos:

```
IS=c(48,15,33,3,21,19,17,16,44,25,30,3,5,9,35,32,26,13,14,14,47,47,29,18,
11,5,19,24,17,6,25,8,18,29,1,18,22,3,22,29,2,6,10,29,10,21,38,41,16,17,8,40,
8,10,18,7,4,4,8,11,3,23,10,19,21,13,12,10,4,17,11,21,9,8,7,5,3,22,14,25,4,11,
10,18,1,28,27,19,24,35,9,30,8,26)
```

```
> length(IS)
[1] 94
```

Los promedios en R son

```
> mean(IS)
[1] 17.39362
> median(IS)
[1] 16.5
```

Igualmente para describir mejor el conjunto de datos necesitamos algunas medidas de dispersión así:

```
> range(IS)
[1] 1 48
> var(IS)
[1] 133.4025
> sd(IS)
[1] 11.55000
> 100*sd(IS)/mean(IS)
[1] 66.40368
```

REPRESENTACIÓN GRÁFICA

```
> stem(IS)
```

The decimal point is 1 digit(s) to the right of the |

```
0 | 112333334444
0 | 5556677888888999
1 | 0000001111233444
1 | 56677778888889999
2 | 1111222344
2 | 55566789999
3 | 0023
3 | 558
4 | 014
4 | 778
```

```
> hist(IS,main="Estudio interacciones Sociales, Univ.
EU",xlab="IS",ylab="frecuencia")
```



Figura 3 Interacciones Sociales EU

Esta investigación se replicó con 33 estudiantes del tercer semestre de la carrera de Psicología de la Universidad Técnica de Ambato obteniéndose los siguientes resultados:

7,11,8,12,7,9,6,4,3,2,15,12,8,7,8,6,6,4,5,6,4,4,2,2,9,6,10,6,8,2,6,4,5

Utilizando R calculamos algunas medidas de tendencia central y de dispersión con algunos gráficos para describir de mejor manera este conjunto de datos y discutirlos con los obtenidos en EU.

```
ISE=c(7,11,8,12,7,9,6,4,3,2,15,12,8,7,8,6,6,4,5,6,4,4,2,2,9,6,10,6,8,2,6,4,5)
> mean(ISE)
[1] 6.484848
> sd(ISE)
[1] 3.153581
> median(ISE)
[1] 6
> range(ISE)
[1] 2 15
```

```
hist(ISE,main="Estudio interacciones Sociales, Univ. Ecuador",xlab="ISE",ylab="frecuencia")
```

Estudio Interacciones Sociales, Univ. Ecuac

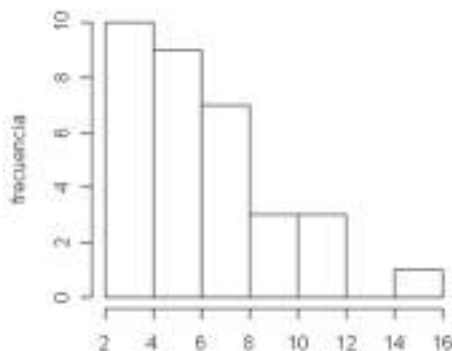


Figura 4 Interacciones Sociales Ecuador

```
boxplot(ISE,main="Estudio interacciones Sociales, Univ. Ecuador",ylab="ISE")
```

Estudio Interacciones Sociales, Univ. Ecuac

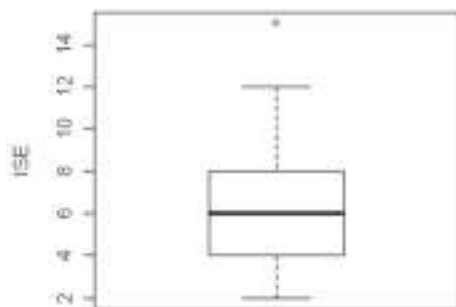


Figura 5 gráfico de caja IS ecuador

En las publicaciones científicas normalmente se hace referencia a la media y al desvío estándar, en algunas oportunidades, la media y el

desvío estándar son incluidos en el texto de una publicación. Por ejemplo, en nuestro caso se puede escribir la cantidad media de interacciones sociales de estudiantes universitarios en EU fue 17 (SD = 11.55) y en Ecuador fue 6 (SD = 3.15).

En las tablas, frecuentemente se hace referencia a la media o al desvío estándar, en especial cuando se involucran varios grupos o cuando los participantes en la investigación son analizados en varias condiciones diferentes. Por ejemplo en nuestro caso se compararon el número de interacciones sociales de universitarios de EU y de Ecuador, el objetivo era probar la teoría de que los universitarios de EU tienen muchas más interacciones sociales que los de Ecuador.

DISTRIBUCIÓN NORMAL

La distribución normal, la cual lo único que tiene de normal es el nombre, fue descrita originalmente por el matemático francés Abraham de Moivre (1667-1754), mas tarde fue utilizada por Pierre Simón Laplace en una variedad de fenómenos de las ciencias naturales y sociales, pero fue el “príncipe de los matemáticos”, el alemán Karl Gauss (1777-1855), quien aplicó la distribución normal al estudio de la forma de la tierra y los movimientos de los planetas, dicho trabajo fue tan influyente que la distribución normal se denomina con mucha frecuencia “Gaussiana”.

La distribución normal se define con su función de densidad de probabilidad, como lo demuestra la siguiente ecuación:

$$f(x) = \frac{1}{\sqrt{2\pi}\sigma} \exp \left[-\frac{1}{2\sigma^2} (x - \mu)^2 \right] \text{ para } -\infty < x < \infty$$

En lenguaje R se puede implementar así:

```
→ mu=0
→ sigma=1
→ x=c(-400:400)/100
```

```
→ fx=(1/sqrt(2*pi*sigma))*exp((x-mu)*(x-mu)/(-2*sigma*sigma))
→ plot(x,fx,main="Distribucion Normal",type="l")
```

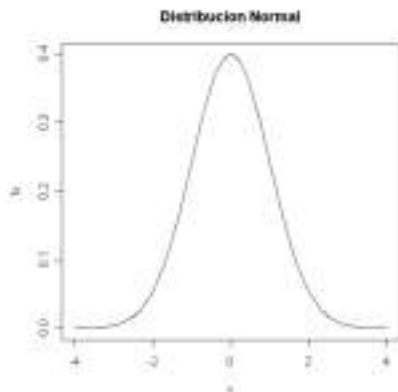


Figura 6 Distribución Normal

Por definición para $-\sigma < x < \sigma$ el área bajo la curva normal es el 68.27% del área total, para $-1.96\sigma < x < 1.96\sigma$ es el 95% del área, y para $-2.576\sigma < x < 2.576\sigma$ el 99% del área total, esto se puede apreciar en la figura 6

La distribución normal es muy importante porque muchas pruebas estadísticas asumen que los datos tienen una distribución normal por lo cual dicho conjunto de datos puede caracterizarse con dos parámetros, la media y el desvío estándar. Una población determinada puede tener una distribución normal, por lo cual dicha población podría ser descrita con sus dos parámetros, una gráfica de la población en cuestión se parecería a la de la figura 6, pero todo esto no significa que una muestra de observaciones a dicha población tenga una distribución normal, esto sucede generalmente cuando la cantidad de observaciones es insuficiente.

El ejemplo que sigue muestra como muestras de diferentes tamaños hechas a una población normal, la función `rnorm` devuelve un vector de `n` cantidad de muestras provenientes de una población de números aleatorios con distribución normal (para más detalles de esta u

otra función de R puede utilizar la función `help()`, para este caso por ejemplo `help(rnorm)`.

Lo citado se puede implementar de la siguiente manera en R:

```
op=par(mfrow=c(2,2))
> ruido10=rnorm(10)
> hist(ruido10,main="A.Histograma ruido10",ylab="Frecuencia")
> ruido50=rnorm(50)
> hist(ruido50,main="A.Histograma ruido50",ylab="Frecuencia")
> ruido500=rnorm(500)
> hist(ruido500,main="A.Histograma ruido500",ylab="Frecuencia")
> ruido1000=rnorm(1000)
> hist(ruido1000,main="A.Histograma ruido1000",ylab="Frecuencia")
```

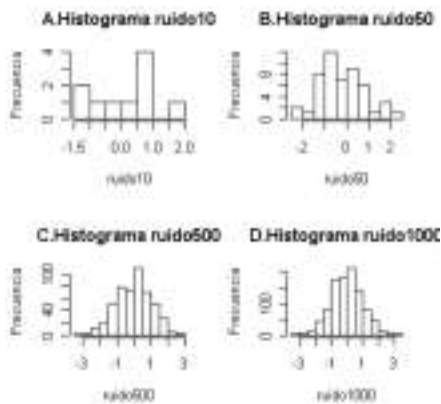


Figura 7 histogramas de 4 muestras con diferente cantidad de observaciones.

En los datos generados con la implementación anterior se puede notar que a medida que aumentamos la cantidad de observaciones el histograma presenta una curva en forma de “campana” (otro de los nombres de la distribución normal), y se parece cada vez más a la distribución de la figura 6.

Como pudimos ver que muestras de una misma población normal

pueden ser diferentes, y por lo tanto la media de las mismas también será diferente, y dichas medias pueden tener también una distribución. Un teorema fundamental de la estadística dice que las medias de muestras aleatorias provenientes de cualquier distribución tiene una distribución normal, dicho teorema se conoce con el nombre de “teorema del límite central”, una consecuencia de este teorema es que cuando trabajamos con muestras de cientos de observaciones podemos olvidarnos de la distribución de la población y asumir que es normal. Una regla práctica muy utilizada dice que muestras con 30 o más observaciones tienen una distribución aceptablemente normal, como lo podemos verificar con nuestro experimento en R.

Una de las aplicaciones más importantes del teorema del límite central es la posibilidad de calcular los denominados *intervalos de confianza* (IC), el más usado es el IC del 95%, por ejemplo si conocemos la media (μ) y el desvío estándar (σ) de una muestra por definición el 95% de los datos se encuentran dentro del intervalo determinado por $1.96\sigma < \mu < 1.96\sigma$.

COMENTARIOS

La estadística descriptiva nos permitió caracterizar con números los datos de las calificaciones, de interacciones sociales de estudiantes universitarios de EU y de Ecuador, calculando la media, el desvío estándar, el rango, los cinco números, pero como se aprecia en los datos de las interacciones sociales figuras 3 y 4 las distribuciones **no son normales**, debido a datos extremos de los estudios en mención.

Por este motivo las muestras de las interacciones sociales no pueden ser descritas solamente con la media y el desvío estándar, es recomendable incluir la mediana y el rango, los cinco números serían inclusive más adecuados en este caso. Como comentario final podemos decir que se deberían calcular y graficar todo lo descrito en este documento cuando se tengan resultados de una investigación, una vez verificado la distribución de la muestra, podemos empezar a decidir qué resultados se reportaran en el trabajo, así como también que métodos estadísticos serán utilizados a posteriori.

REFERENCIAS

Risk MR. Cartas sobre estadística. Revista Argentina de Bioingeniería, 2005

The R Project for Statistical Computing

Aron A, Aron E, Estadística para Psicología, 2001

Carmona F. Un análisis con R Datos Multivariantes, 2005

Tussel F. Lectura manipulación de datos con R, 2005

Murray R. Spiegel Estadística, 1991

PRUEBA DE HIPÓTESIS Y DE TENDENCIA CENTRAL SOBRE UNA Y DOS MUESTRAS

“Una vez uno de mis amigos me dijo que si algunos aseguraban que la tierra rotaba de este a oeste y otros que rotaba de oeste a este, siempre habría un grupo de bien intencionados ciudadanos que sugerirían que quizás haya algo de verdad de ambos lados, y que puede haber un poco de uno y otro poco del otro; o que probablemente la verdad está en los extremos y quizás no rota en absoluto”.

Sir Maurice Kendall(1907-1983), estadístico británico.

INTRODUCCIÓN

En el capítulo anterior revisamos los temas fundamentales sobre la estadística descriptiva, la representación gráfica y la distribución normal; la distribución normal es un tema muy importante en estadística, y aquí vamos a utilizar la misma en forma práctica, como primera prueba antes de cualquier otra prueba estadística. Antes de esto explicaremos la formulación y prueba de hipótesis, luego desarrollaremos como probar la normalidad de una muestra, y pruebas estadísticas sobre una y dos muestras

FORMULACIÓN Y PRUEBA DE HIPÓTESIS

El estadístico británico sir Maurice G. Kendall expresó irónicamente una forma no científica de abordar un problema. El método científico en cambio es un proceso con el cual se investigan en forma sistemática observaciones, se resuelven problemas y se verifican hipótesis. Como parte del método científico la propuesta de una hipótesis, y luego su prueba, son temas muy bien definidos, y a pesar de la incer-

tidumbre asociada al problema es posible cuantificar el error de la conclusión planteada por la hipótesis.

Los pasos del método científico se pueden resumir de la siguiente forma:

- 1) plantear el problema a resolver
- 2) efectuar las observaciones
- 3) formular una o más hipótesis
- 4) probar dichas hipótesis y
- 5) proclamar las conclusiones

La estadística nos puede ayudar en los pasos 2 (diseño de las observaciones) y 4 (prueba de hipótesis).

Una definición de hipótesis es la siguiente: una explicación tentativa que cuenta con un conjunto de hechos y puede ser probada con una investigación posterior. La formulación de una hipótesis se logra examinando cuidadosamente las observaciones para luego proponer un resultado posible.

Por ejemplo un problema a resolver podría ser número de palabras recordadas en función del tipo de texto utilizado, ya tenemos el primer paso del método científico; luego efectuamos observaciones en dos grupos de sujetos uno con texto sencillo y otro con texto complejo; el tamaño de dichas muestras se basa en estudios similares ya publicados y/o experiencia de los investigadores sobre y/o cálculos sobre tamaño de las muestras.

Una vez efectuadas las observaciones en los sujetos de los dos grupos en estudio construimos la tabla:

Texto sencillo	Texto complejo
10	2
5	1
6	7
3	4
9	4

8	5
7	2
5	5
6	3
5	4

Por la observación cuidadosa de la tabla y utilizando métodos de estadística descriptiva, podríamos aventurarnos a decir que el recuerdo se ve afectado por el tipo de texto utilizado y esa sería nuestra hipótesis. La formulación formal de una hipótesis en el método científico se realiza definiendo la hipótesis nula (H_0) y la hipótesis alterna (H_1); generalmente la H_0 establece que no hay diferencias entre los grupos y la H_1 por otra parte, suele indicarse como el complemento de la H_0 , por lo tanto la H_1 establecerá que si hay diferencias entre los grupos en estudio. Por lo tanto la prueba de nuestra hipótesis consistiría en arbitrar los procedimientos necesarios para rebatir H_0 , esto es, rechazarla.

H_0 : recuerdo texto sencillo = recuerdo texto complejo

H_1 : recuerdo texto sencillo $\neq$ recuerdo texto complejo

A la hora de tomar una decisión respecto de la hipótesis nula, surgen situaciones que nos pueden llevar a cometer diferentes errores. Así, una vez arbitradas las técnicas (o realizadas las pruebas) para probar esta hipótesis puede que lleguemos a la conclusión de que el enunciado de nuestra H_0 sea verdadero en tal caso no rechazamos nuestra H_0 o bien que sea falso, en cuyo caso rechazaremos la H_0 . En esta situación puede que hayamos rechazado la H_0 cuando en realidad era cierta, o que la evidencia no haya sido suficiente para rechazarla siendo falsa. Estas diferentes situaciones plantean la existencia de diferentes tipos de errores que se resumen a continuación:

		Decisiones	
		Aceptar H_0	Rechazar H_0
Situaciones	H_0 cierta	<i>Decisión correcta</i>	<i>Error de tipo I</i>
	H_0 Falsa	<i>Error de tipo II</i>	<i>Decisión correcta</i>

El error tipo I, también denominado error α , se produce cuando se rechazó la H_0 y es verdadera. Este, representa la probabilidad de haber cometido este tipo de error. Se establece a priori α como el nivel de significancia o error máximo aceptable para la conclusión. El uso ha impuesto que en estudios de CCNN y/o CCSS este error asuma valores no mayores a 0.01 o 0.05 en estudios experimentales en general. En el caso de que la H_0 sea aceptada siendo falsa se cometerá el error tipo II o β .

El error tipo II está asociado con la **potencia** del método estadístico utilizado para poder detectar diferencias. La potencia de un método estadístico de una determinada situación se calcula como $(1 - \beta)$, lo que se corresponde con la situación de haber rechazado correctamente la H_0 , pues era falsa. Al igual que el valor de significancia α , la potencia del método estadístico se establece por el tamaño de la muestra y la prueba estadística utilizada.

En cualquier caso rechazar la H_0 es lo mismo que aceptar la H_1 y viceversa. El resultado final de un método estadístico para la prueba de una hipótesis es el valor P. que indica la probabilidad de obtener un valor más extremo que el observado si la H_0 es verdadera. Cuando P es menor que α se procede a rechazar la H_0 .

PRUEBAS DE NORMALIDAD DE UNA MUESTRA

Antes de proceder a la prueba de una hipótesis debemos determinar la distribución de las variables consideradas en nuestra muestra. En los métodos convencionales se trabaja con la distribución normal de dichas variables. El paso inicial entonces, es determinar si las variables en estudio pueden ser representadas por una distribución **normal**. En otras palabras, necesitamos verificar esta primera hipótesis. O sea, si las variables medidas en la muestra pueden ser descritas con parámetros de tendencia central y dispersión simétrica alrededor de dichos parámetros y relación media-dispersión conocidas.

La importancia de verificar la normalidad de las muestras en estudio es fundamental en estadística porque si las muestras son normales se pueden aplicar métodos estadísticos paramétricos convencionales, en caso contrario se deben transformar los datos, o bien utilizar métodos como los no paramétricos u otros métodos estadísticos más sofisticados.

Los métodos de estadística descriptiva nos pueden ayudar a verificar la normalidad de las variables, un histograma (Figura 1A) y un gráfico de cajas (Figura 1B) nos muestra en dos formas distintas la distribución de los datos, para el ejemplo propuesto podemos decir por la forma del histograma y por los espacios intercuartiles similares del gráfico de cajas que la muestra parece tener una distribución normal.

El cálculo de los cinco números de tukey > fivenum(Recuerdo)

[1] 3 5 6 8 10 nos muestra numéricamente que no hay evidencia suficiente como para rechazar la distribución normal de la variable recuerdo.

Pruebas de normalidad más formales, no paramétricas, muy recomendables para verificar la normalidad de una variable son las pruebas de Shapiro – Wilk, y de Kolmogorov – Smirnov

Contrariamente a lo que se desea en la mayoría de los casos, en las

pruebas de normalidad se busca aceptar la H_0 , dado que en la mayoría de los métodos estadísticos convencionales es necesaria la distribución normal de la variable de interés, pues siendo así es posible conocer los parámetros que la describen por completo, su media (μ), su desvío (s) y la relación entre ambos y en este sentido estos métodos son más potentes. Un valor $P \geq 0.05$ en los test de normalidad indicaría que no hay prueba suficiente para rechazar la normalidad de la variable.

El siguiente código en lenguaje R fue utilizado para los gráficos y cálculos respectivos

```
>Recuerdo = c(10,5,6,3,9,8,7,5,6,5)
>par(mfrow=c(1,2))
> hist(recuerdo,main="A",xlab="recuerdo",ylab="frecuencias")
> boxplot(Recuerdo,main="B",ylab="Recuerdo")
```

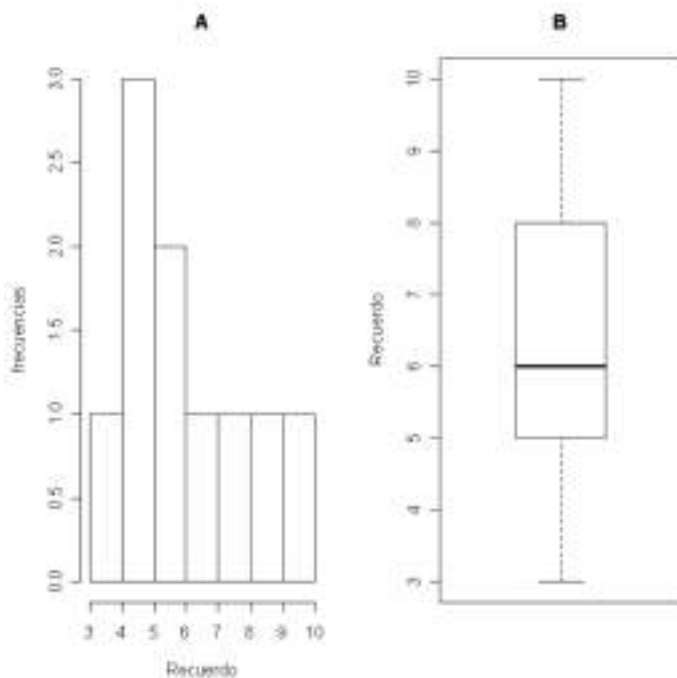


Figura 1: histograma(A) y gráfico de caja (B) para la muestra de recuerdo.

```

> mean(Recuerdo)
[1] 6.4
> sd(Recuerdo)
[1] 2.1187
> fivenum(Recuerdo)
[1] 3 5 6 8 10
> sw=shapiro.test(Recuerdo)
> sw

```

Shapiro-Wilk normality test

```

data: Recuerdo
W = 0.9578, p-value = 0.7604

> mu=mean(Recuerdo);de=sd(Recuerdo)
> ks.test(Recuerdo,"pnorm",mean=mu,sd=de)

```

One-sample Kolmogorov-Smirnov test

```

data: Recuerdo
D = 0.1749, p-value = 0.9198
alternative hypothesis: two.sided

```

Warning message:

cannot compute correct p-values with ties in: ks.test(Recuerdo, "pnorm", mean = mu, sd = de)

Como podemos apreciar en los resultados obtenidos se pudo verificar que la distribución de la variable del ejemplo es normal.

PRUEBAS SOBRE UNA MUESTRA

Una vez verificada la normalidad, podemos realizar la prueba para verificar nuestra H_0 , esto es, que la media del Recuerdo de ambas muestras son iguales. Suponga para ello en primer lugar, que consideramos al grupo 1, los valores para la población normal y los del grupo 2 una muestra. Los parámetros del grupo 1 son considerados

los poblacionales y los del grupo 2 los de la muestra. En este caso existe un estadístico que permite comparar la media muestral con la poblacional.

En ciertas ocasiones se cuenta con una muestra y con la media y el desvío estándar de una población, en dicho caso se utiliza la prueba sobre una muestra, para ello se debe calcular el estadístico t con la siguiente ecuación

$$t = \frac{X - \mu}{\sigma / \sqrt{n - 1}}$$

Para nuestros datos y comparándolos con los poblacionales: $\mu = 3.7$ y $s = 1.77$ tenemos que $t = 0.508$ al ser $P > 0.05$ no podemos rechazar la H_0 y concluimos que la media de la muestra es igual a la media poblacional.

PRUEBAS SOBRE DOS MUESTRAS

En el caso de contar con dos muestras, para nuestro ejemplo el grupo 1 de sujetos con texto simple y el grupo 2 con texto complejo, el test implicado intentará probar si ambas medias no difieren.

El test más difundido es la “ t del estudiante”, publicado por el estadístico británico William Gosset (1876-1937) en 1908 bajo el seudónimo de “estudiante”, según piensan algunos no lo publicó bajo su nombre porque la prueba t fue desarrollada como parte de su trabajo en control de calidad para la cervecería Guinness en el Reino Unido; los resultados de sus estudios probaban que algunos lotes de cerveza no eran de la calidad esperada por Guinness. De esa manera Gosset descubrió la distribución t e inventó la prueba t (la simpleza misma, comparada con la mayoría de los cálculos estadísticos), para aquellas situaciones en las que las muestras son pequeñas y se desconoce la variabilidad de la población que se supone de un tamaño mucho más grande.

La prueba t es la prueba paramétrica más utilizada; la misma está ba-

sada en el cálculo del estadístico t y de los grados de libertad, con estos dos resultados y utilizando o bien una tabla o bien un cálculo de la distribución t se puede calcular el valor de P.

La prueba t se basa en tres supuestos: a) uno es el de la distribución de los errores, b) es la independencia de los mismos y c) es el de la homogeneidad de varianzas, considerando este último como el supuesto más importante.

La siguiente ecuación nos muestra como calcular el estadístico t:

$$t = \frac{X_1 - X_2}{\sqrt{\frac{S_1^2}{n_1} + \frac{S_2^2}{n_2}}}$$

Como dijimos anteriormente se debe verificar la normalidad de las variables, ya lo habíamos hecho para el grupo 1, para el grupo 2 (texto complejo) tenemos lo siguiente: a) Grupo 1 media y (DE) 6.4 (2.1187), b) Grupo 2 media y (DE) 3.7 (1.7669), c) prueba de Shapiro Wilk $P = 0.79$, y d) prueba de Kolmogorov- Smirnov $P = 0.94$; la normalidad de las muestras fueron verificadas con las dos pruebas. La figura 2 muestra gráficamente los datos de nuestro ejemplo.

```
> shapiro.test(Recuerdo)
```

Shapiro-Wilk normality test

```
data: Recuerdo
```

```
W = 0.9604, p-value = 0.7909
```

```
> mu=mean(Recuerdo);de=sd(Recuerdo)
```

```
> ks.test(Recuerdo,"pnorm",mean=mu,sd=de)
```

One-sample Kolmogorov-Smirnov test

```
data: Recuerdo
```

```
D = 0.1674, p-value = 0.942
```

alternative hypothesis: two.sided

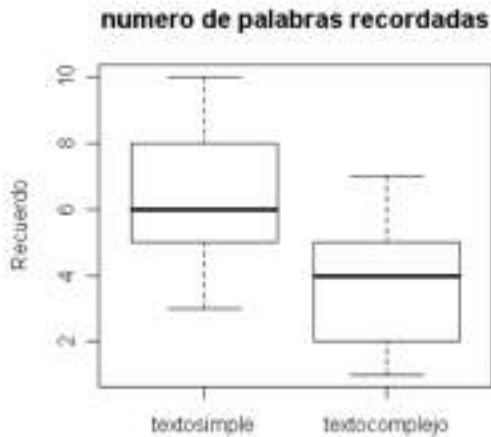


Figura 2 gráfico de cajas para Recuerdo en los dos grupos

Ahora determinamos t, en lenguaje R la prueba t se calcula así:

```
> t.test(textosimple, textocomplejo)
```

Welch Two Sample t-test

data: textosimple and textocomplejo

t = 3.0949, df = 17.438, p-value = 0.006428

alternative hypothesis: true difference in means is not equal to 0

95 percent confidence interval:

0.862875 4.537125

sample estimates:

mean of x mean of y

6.4

3.7

Como $P < 0.05$ se rechaza la H_0 y se concluye que el tipo de texto utilizado afecta el Recuerdo.

PRUEBA SOBRE DOS MUESTRAS APAREADAS

El ejemplo anterior fue sobre dos muestras provenientes de dos grupos distintos de sujetos, en ciertas ocasiones necesitamos trabajar sobre un mismo grupo de sujetos al cual se los observa en forma repetida, por ejemplo antes y después de un tratamiento, en este caso los sujetos son controles de ellos mismos. La prueba t es distinta para poder tener en cuenta que las observaciones son repetidas sobre el mismo grupo de sujetos. La siguiente ecuación muestra como calcular t para el caso de muestras apareadas:

$$t = \frac{\bar{X}_1 - \bar{X}_2}{\sqrt{\frac{\sum d^2 - (\sum d)^2/n}{n(n-1)}}$$

Olthoff (1989) analizó la calidad de comunicación entre parejas comprometidas tres meses antes y tres meses después del matrimonio. Uno de los grupos estudiados estaba conformado por 19 parejas que habían recibido el acostumbrado curso prematrimonial por parte de los ministros que iban a celebrar su matrimonio (para que el ejemplo no se complique, nos concentramos sólo en este grupo, y únicamente en los esposos de este grupo. Los valores de las esposas eran similares aunque un poco más variados).

Los valores de los esposos fueron los siguientes:

Esposo	calidad de comunicación	
	Antes	después
A	126	115
B	133	125
C	126	96
D	115	115
E	108	119
F	109	82
G	124	93
H	98	109

I	95	72
J	120	104
K	118	107
L	126	118
M	121	102
N	116	115
O	94	83
P	105	87
Q	123	121
R	125	100
S	128	118

Las hipótesis serían:

H1: los maridos que asisten al curso prematrimonial (como los maridos que analizo Olthoff) **si** cambian en cuanto a calidad de comunicación antes y después del matrimonio.

H0: los maridos que asisten al curso prematrimonial **no** cambian en cuanto a la calidad de su comunicación antes y después del matrimonio.

Como siempre primero verificamos la normalidad de la variable de interés, los resultados de las pruebas Shapiro-Wilk y Kolmogorov-Smirnov fueron a) antes: $P=0.073$ y $P=0.785$ y b) después $P=0.21$ y $P=0.567$, la normalidad de las muestras es verificada.

El código en lenguaje R para calcular la t para dos muestras apareadas es el siguiente:

```
>antes=c(126,133,126,115,108,109,124,98,95,120,118,126,121,116,94,105,123,125,128)
```

```
>despues=c(115,125,96,115,119,82,93,109,72,104,107,118,102,115,83,87,121,100,118)
```

```
> t.test(antes,despues,paired=T)
```

Paired t-test

data: antes and despues

$t = 4.2404$, $df = 18$, $p\text{-value} = 0.000492$

alternative hypothesis: true difference in means is not equal to 0
95 percent confidence interval:

6.081135 18.024128

sample estimates:

mean of the differences

12.05263

```
> mean(antes)
```

```
[1] 116.3158
```

```
> mean(después)
```

```
[1] 104.2632
```

Con estos resultados se rechaza la H_0 y se concluye que existe una disminución significativa de la calidad de comunicación, decreciendo de 116.32 antes del matrimonio a 104.26 después del matrimonio $t(18) = 4.24$, $P < 0.05$

PRUEBAS SOBRE DOS MUESTRAS NO NORMALES

Hasta el momento todos los ejemplos tuvieron distribuciones normales, por lo cual la aplicación de pruebas paramétricas normales es totalmente válido; pero si estamos ante muestras no normales nos olvidamos de las pruebas paramétricas y, se pueden **transformar** (procedimiento matemático) los resultados de una muestra para que sea de distribución normal y así aplicar los métodos clásicos, o buscamos la equivalente no paramétrica, una de las más utilizadas en CCSS es la Chi cuadrado. Existen muchas: para **una variable** se utiliza la prueba de bondad de ajuste que consiste en determinar si los datos de cierta muestra corresponden a cierta distribución poblacional, y para **dos variables** se utilizan la prueba de homogeneidad que consiste en comprobar si varias muestras de un carácter cualitativo proceden de una misma población y la prueba de independencia que consiste en comprobar si dos características cualitativas están relacionadas entre sí.

Por ejemplo, para este último caso, se trata de averiguar qué relación existe entre la organización del hogar y el rendimiento instructivo de los alumnos del segundo curso del colegio Técnico Yaruquí (Quito-Ecuador), obteniéndose los siguientes resultados

→ rendimiento

	def.	reg.	bueno	M.bueno
HO	1	3	5	6
HD	2	4	3	6

Donde HO = hogares organizados y HD = hogares desorganizados.

Ho: la organización del hogar no se relaciona con el rendimiento

H1: la organización del hogar se relaciona con el rendimiento

Para probar esta hipótesis en lenguaje R tenemos

```

> rendimiento=matrix(c(1,3,5,6,2,4,3,6),2,4,byrow=T)
> dimnames(rendimiento)
NULL
> dim(rendimiento)
[1] 2 4
> organizacion=c("HO","HD")
> calificaciones=c("def.,""reg.,""bueno","M.bueno")
> dimnames(rendimiento)=list ( organización, calificaciones)
> rendimiento

```

	def.	reg.	bueno	M.bueno
HO	1	3	5	6
HD	2	4	3	6

```

> chisq.test(rendimiento)

```

Pearson's Chi-squared test

data: rendimiento

X-squared = 0.9762, df = 3, p-value = 0.807

La organización del hogar no tiene relación con el rendimiento de los alumnos. La organización del hogar es el producto de la estructura social donde los alumnos de hogares desorganizados hacen múltiples esfuerzos para ponerse al nivel de los alumnos de hogares organizados.

En el año 1994, Riehl realizó un estudio para analizar la experiencia universitaria de alumnos de primer año que eran la primera generación de la familia en asistir a la universidad. Los alumnos fueron comparados con otros alumnos que no eran la primera generación de la familia que asistía a la universidad (todos los alumnos pertenecían a la Universidad de Indiana). Una de las variables que midió Riehl fue si los alumnos abandonaban o no los estudios durante el primer semestre.

Generación que asiste a la universidad

	Primera	Otras
Abandono	73	89
Continuidad	657	1226

En lenguaje R

```
> experienciaU=matrix(c(73,89,657,1226),2,2,byrow=T)
> dimnames(experienciaU)
NULL
> dim(experienciaU)
[1] 2 2
> estudios=c("abandono", "continuidad")
> generación=c("primera", "otras")
> dimnames(experiencia U)=list(estudios, generación)
> experiencia U
```

	Primera	Otras
Abandono	73	89
Continuidad	657	1226

```
> chisq.test(experiencia U)
```

Pearson's Chi-squared test with Yates' continuity correction

data: experienciaU

X-squared = 6.2863, df = 1, p-value = 0.01217

Se rechaza la hipótesis de independencia y se acepta la relación entre las dos variables estudiadas.

COMENTARIOS

En este capítulo se definieron conceptos fundamentales sobre la formulación y la prueba de hipótesis, así como también ejemplos prácticos para pruebas de normalidad, pruebas con una muestra, pruebas de dos muestras apareadas y no apareadas. Finalmente se describió la prueba no paramétrica chi cuadrado de independencia, cuyo inventor fue Karl Pearson quien vio en las ideas de Galton, una forma de convertir la Psicología, la Antropología y la Sociología en campos tan científicos como lo eran la Física y la Química. Esperaba evitar la cuestión de la causalidad a través de la utilización de esta categoría más amplia de correlación, asociación, o contingencia (con un rango de 0, independencia, a 1, "unidad de causalidad". "Ningún fenómeno es causal"-expresó. "Todos los fenómenos son contingentes, y el problema que enfrentamos es el de medir el grado de contingencia".

La formulación de la hipótesis se hizo en cada ejemplo con el criterio más conservador, planteando la hipótesis nula como la igualdad entre las muestras en estudio, por lo cual la hipótesis alternativa siempre fue la desigualdad entre las mismas. Esta forma se apoya en el concepto de las "dos colas" y no asume a priori que la diferencia es por menor o mayor, es aconsejable plantear las hipótesis siempre de esta forma, a menos que el conocimiento sobre el problema en estudio amerite lo contrario.

Un comentario especial merecen las pruebas de normalidad, a veces omitidas por algunos investigadores, pero se consideran fundamentales para poder verificar la normalidad de las muestras, y de esta forma poder aplicar apropiadamente las pruebas estadísticas paramétricas. La prueba de normalidad de Shapiro Wilk está considerada como la más poderosa para verificar la normalidad de una muestra, por lo cual algunos estadísticos consideran que por sí sola es suficiente

REFERENCIAS

Risk MR Cartas sobre estadística 1.Revista Argentina de Bioingeniería.
En prensa

Kanji GK. 100 Statistical test . sage Publications.1999

Venables WN. Smith DM. An Introduction to R, Version 1.5.1. R development Core Team, 2002

Conover WJ. Practical non Parametric statistics. New York: John Wiley & Sons. 1971

Aron A, Aron E. Estadística para Psicología , segunda edición. 2001

Flores, L. Diseños de Investigación Educativa, Quito Ecuador 1982

Greene,J. Oliveira, M. Learning to use Statistical Test in Psychology. A student's Guide

ANÁLISIS DE VARIANZA (ANOVAs)

INTRODUCCIÓN

El análisis de varianza es una maravillosa idea básica que vale la pena analizarlo, por dos motivos: primero, porque en la medida que el alumno lea o realice investigaciones irá progresivamente adoptando esa manera de razonar y, segundo, porque de hecho es la forma en la que ya piensa.

Al realizar cualquier investigación lo que nos interesa saber es si determinada variable realmente origina alguna diferencia. Organizamos dos (o más) grupos para poder demostrar que cualquier diferencia en los resultados existe puramente porque un grupo recibió la influencia de la variable y el otro no.

El análisis de varianza es similar a la forma en que siempre hemos pensado. Kelley (1971) sugirió que, en el fondo, todos somos científicos, puesto que todos formulamos hipótesis y las sometemos a prueba; y el método que utilizamos para distinguir y tomar decisiones acerca de la causalidad aplica el razonamiento del análisis de varianza.

Si el alumno está familiarizado con el trabajo del psicólogo Jean Piaget, especialista en el campo del desarrollo, reconocerá que el tipo de razonamiento del análisis de varianza es parte de lo que él llama "operaciones formales", el estilo de pensamiento adquirido alrededor de los 14 años; por lo tanto no deberíamos tener inconvenientes en comprender el análisis de varianza, ¡inocentemente lo hemos estado utilizando durante años!.

En este capítulo analizaremos los ANOVAs de un factor, de dos factores entre sujetos e intra sujetos, el diseño factorial y el diseño de cuadrado latino.

ANOVA ENTRE SUJETOS UNIFACTORIAL

Los supuestos del análisis de varianza son básicamente los mismos que los de la prueba *t* para medias independientes. Es decir, obtenemos resultados estrictamente precisos sólo cuando las poblaciones siguen una distribución normal y tienen la misma varianza. Además, al igual que con la prueba *t*, en la práctica obtenemos resultados bastante aceptables aún cuando las poblaciones son moderadamente distintas de lo normal y tienen diferencias moderadas en cuanto a las varianzas.

La hipótesis nula en un análisis de varianza establece que las diversas poblaciones que se comparan tienen la misma media. **Por ejemplo**, en el estudio sobre el recuerdo citado anteriormente, a tres grupos diferentes de 6 sujetos cada uno se les dio una lista de 10 palabras para aprender. Al primer grupo se le dieron las palabras lentamente a una velocidad de una palabra cada cinco segundos, al segundo grupo se le dieron a una velocidad media de una palabra cada 2 segundos y al tercer grupo a una velocidad alta de una palabra por segundo.

Se predijo que los puntajes de recuerdo se verían afectados por la velocidad de presentación.

Los resultados se presentan a continuación:

Grupo 1 (v baja)	Grupo 2 (v media)	Grupo 3 (v alta)
8	7	4
7	8	5
9	5	3
5	4	6
6	6	2
8	7	4

El juego de hipótesis será:

$$H_0: \bar{X}_1 = \bar{X}_2 = \bar{X}_3$$

$$H_1: \bar{X}_1 \neq \bar{X}_2 \neq \bar{X}_3$$

Para resolver este problema en lenguaje R procedemos así:

```
> palabrasrecordadas=c(8,7,9,5,6,8,7,8,5,4,6,7,4,5,3,6,2,4)
> velocidad=factor(rep(1:3,each=6),labels=c("v1","v2","v3"))
> recuerdo=data.frame(velocidad,palabrasrecordadas)
```

```
> recuerdo
```

	Velocidad	palabras recordadas
1	v1	8
2	v1	7
3	v1	9
4	v1	5
5	v1	6
6	v1	8
7	v2	7
8	v2	8
9	v2	5
10	v2	4
11	v2	6
12	v2	7
13	v3	4
14	v3	5
15	v3	3
16	v3	6
17	v3	2
18	v3	4

Hacemos el **análisis de varianza** (aov) en R

```
> a=aov(palabrasrecordadas~velocidad)
> summary(a)
```

	Df	Sum Sq	Mean Sq	F value	Pr(>F)
velocidad	2	31.444	15.722	7.4474	0.005672 **
Residuals	15	31.667	2.111		

Signif. codes: 0 ‘***’ 0.001 ‘**’ 0.01 ‘*’ 0.05 ‘.’ 0.1 ‘.’ 1

En este caso hemos encontrado un efecto significativo de la variable **velocidad** (el valor de P ha sido menor de 0.05).

Naturalmente es necesario indicar las medias por condición. El comando que se aplica es **tapply**:

```
> tapply(velocidad,palabrasrecordadas,mean)
```

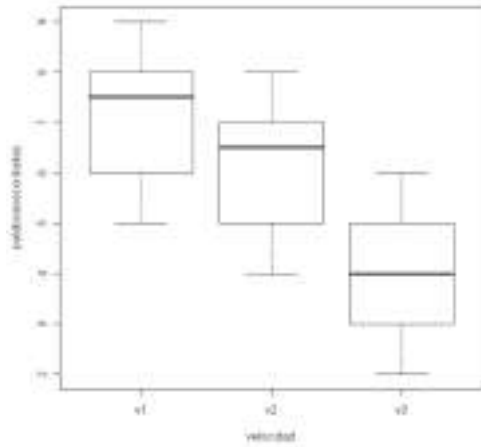
v1	v2	v3
7.166667	6.166667	4.000000

Un aspecto muy importante en cualquier investigación es observar los datos en cada grupo, para lo cual podemos tener el diagrama de caja y bigotes para cada grupo. El comando que se emplea es **plot**:

```
> plot(palabrasrecordadas~velocidad)
```

En la fórmula se indica la variable dependiente seguida (tras el signo ~) de la variable independiente.

En el gráfico se puede observar una diferencia clara entre el grupo v1 con el de v3. Naturalmente, viendo el gráfico y la existencia de un efecto significativo de *velocidad* en el ANOVA, lo que hemos de hacer ahora son comparaciones múltiples. (En caso de que el ANOVA no



hubiera sido significativo, no se debe proceder a efectuar estas pruebas). De esta manera podemos determinar entre qué condiciones experimentales hay diferencias significativas. Para ello, empleamos el método de Tukey, empleando la función **TukeyHSD**

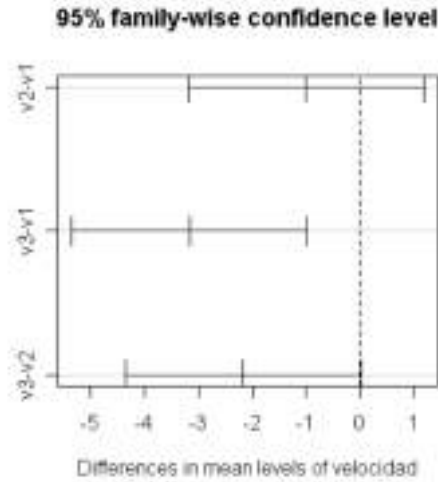
```
> aTukey=TukeyHSD(a,"velocidad")
> aTukey
Tukey multiple comparisons of means
95% family-wise confidence level
```

```
Fit: aov(formula = palabrasrecordadas ~ velocidad)
```

\$velocidad	diff	lwr	upr	p adj
v2-v1	-1.000000	-3.178941	1.17894113	0.4756623
v3-v1	-3.166667	-5.345608	-0.98772554	0.0049063*
v3-v2	-2.166667	-4.345608	0.01227446	0.0514005

```
> plot(aTukey)
```

Las diferencias entre medias en las que el intervalo de confianza que engloba los límites inferior y superior no contienen el valor 0, son estadísticamente significativas con el método de Tukey. En nuestro



caso, al comparar las medias $v2-v1$ y $v3-v2$ no se encuentran diferencias significativas, pero al comparar las medias $v1-v3$ si se encuentran diferencias significativas, es decir, el número de palabras recordadas es diferente en estas condiciones, o, las variaciones encontradas no se deben al azar sino a la velocidad de presentación de las palabras. Esto puede verse también gráficamente en los trazos de los intervalos de confianza, se puede apreciar que sólo el intervalo de confianza de la diferencia entre $v1$ y $v3$ no toca 0, es decir, que la diferencia entre estas dos medias es estadísticamente significativa.

Tenemos otro ejemplo

Un psicólogo especializado en asuntos empresariales estaba interesado en averiguar si los individuos que trabajaban en diferentes sectores de la empresa tenían diferentes actitudes hacia la misma. Los resultados correspondientes a las tres personas entrevistadas del área de ingeniería fueron 10,12, y 11; los resultados de los tres del área de comercialización 6,6 y 8; los resultados de los tres miembros de contaduría 7,4 y 4; y los resultados de los tres de producción 14,16 y 13 (los números más altos indican actitudes más positivas).

La hipótesis sería:

$$H_0: \bar{X}_1 = \bar{X}_2 = \bar{X}_3 = \bar{X}_4$$

$$H_1: \bar{X}_1 \neq \bar{X}_2 \neq \bar{X}_3 \neq \bar{X}_4$$

En lenguaje R lo resolvemos de la siguiente manera:

```
> Actitud=c(10,12,11,6,6,8,7,4,4,14,16,13)
> sector=factor(rep(1:4,each=3),labels=c("Ing.", "Comer.", "Conta.",
"Produc."))
```

```
> a=data.frame(sector,Actitud)
```

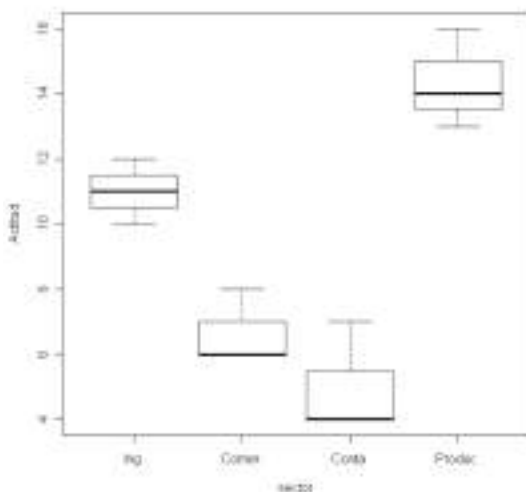
```
> a
```

	sector	Actitud
1	Ing.	10
2	Ing.	12
3	Ing.	11
4	Comer.	6
5	Comer.	6
6	Comer.	8
7	Conta.	7
8	Conta.	4
9	Conta.	4
10	Produc.	14
11	Produc.	16
12	Produc.	13

```
> tapply(Actitud,sector,mean)
```

Ing.	Comer.	Conta.	Produc.
11.000000	6.666667	5.000000	14.333333

```
> plot(Actitud~sector)
```



Hacemos el **ANOVA**

```
> b=aov(Actitud~sector)
```

```
> summary(b)
```

	Df	Sum Sq	Mean Sq	F value	Pr(>F)
sector	3	160.917	53.639	27.985	0.0001360 ***
Residuals	8	15.333	1.917		

Signif. codes: 0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1

Se observa que el sector donde trabajan los individuos tiene mucho que ver en la actitud hacia la empresa, es decir, se encuentran dife-

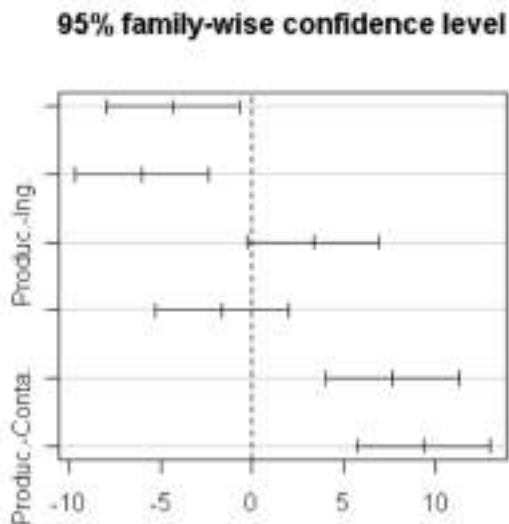
rencias significativas en las actitudes hacia la empresa. Para averiguar si difieren estadísticamente los sectores en la actitud se utiliza la prueba de Tukey en R así:

```
> bTukey=TukeyHSD(b,"sector")
> bTukey
Tukey multiple comparisons of means
95% family-wise confidence level
```

```
Fit: aov(formula = Actitud ~ sector)
```

\$sector	diff	lwr	upr	p adj
Comer.-Ing.	-4.333333	-7.953235	-0.713432	0.0208506*
Conta.-Ing.	-6.000000	-9.619901	-2.380099	0.0031963*
Produc.-Ing.	3.333333	-0.286568	6.953235	0.0715001
Conta.-Comer.	-1.666667	-5.286568	1.953235	0.4933347
Produc.-Comer.	7.666667	4.046765	11.286568	0.0006387*
Produc.-Conta.	9.333333	5.713432	12.953235	0.0001605*

```
> plot(bTukey)
```



Al comparar las medias de actitudes de los sectores comercialización-ingeniería, contabilidad-ingeniería, producción-comercialización, y producción-contabilidad se encuentran diferencias significativas, es decir, son diferentes las actitudes hacia la empresa en estos sectores, teniendo las mejores actitudes hacia la empresa el sector de producción (14.33), y el de menor actitud el departamento de contabilidad (5).

No se encuentran diferencias en las medias de actitud de producción-ingeniería, y contabilidad-comercialización, es decir, a pesar de tener diferentes valores de actitud estadísticamente no son diferentes.

ANOVA FACTORIAL ENTRE SUJETOS CON DOS FACTORES

El procedimiento estadístico para analizar los resultados de un experimento factorial en dos sentidos se denomina ANOVA de dos criterios. La lógica básica es la misma que venimos utilizando. Este análisis se utiliza cuando se estudian dos variables (factores) y se usan **diferentes sujetos** para cada condición. Tenemos un **ejemplo**

Wong y Csikszentmihalyi (1991) realizaron un estudio en el cual, durante una semana, 170 alumnos portaron equipos buscapersonas y se los llamaba a intervalos aleatorios (aproximadamente cada dos horas durante el día). Cada vez que recibían una llamada, los alumnos debían llenar un formulario indicando qué estaban haciendo en ese momento. El estudio era un diseño factorial 2x2 que analizaba el efecto del sexo y el nivel de los alumnos obtenido en una prueba acerca del deseo de relacionarse. La variable medida era la cantidad de ocasiones, durante la semana, en las que el alumno estaba realizando actividades sociales cuando se lo llamaba. (También había otras variables pero nos concentraremos sólo en éstas). Para que el ejemplo sea simple a los efectos de aprendizaje, se incluyen sólo 10 participantes por casilla.

ACTIVIDADES SOCIALES

BAJO NIVEL DE DESEO DE RELACIONARSE		ALTO NIVEL DE DESEO DE RELACIONARSE	
Hombres	Mujeres	Hombres	Mujeres
12.1	17.4	11.1	22
11.4	17.1	10.4	20.5
11.2	16.8	10.2	19.9
10.9	16.7	9.8	19.1
10.3	15.5	9.2	18.5
9.8	15.3	9.1	17.4
9.7	15	8.9	17
9.5	15.4	8.7	17.1
9.3	14.3	8.2	17.1
8.8	14	6.6	16.5

```
> actsocial=c(12.1,11.4,11.2,10.9,10.3,9.8,9.7,9.5,9.3,8.8,17.4,17.1,16.8,16.
7,15.5,15.3,15,15.4,14.3,14,11.1,10.4,10.2,9.8,9.2,9.1,8.9,8.7,8.2,6.6,22,20.
5,19.9,19.1,18.5,17.4,17,17.1,17.1,16.5)
```

```
> sexo=factor(rep(rep(1:2,each=10),2),labels=c("H","M"))
> deseorel=factor(rep(1:2,each=20),labels=c("bajo","alto"))
> ejem=data.frame(deseorel,sexo,actsocial)
```

ejem

	deseorel	sexo	asocial		deseorel	sexo	asocial
1	bajo	H	12.1	21	alto	H	11.1
2	bajo	H	11.4	22	alto	H	10.4
3	bajo	H	11.2	23	alto	H	10.2
4	bajo	H	10.9	24	alto	H	9.8
5	bajo	H	10.3	25	alto	H	9.2
6	bajo	H	9.8	26	alto	H	9.1
7	bajo	H	9.7	27	alto	H	8.9
8	bajo	H	9.5	28	alto	H	8.7
9	bajo	H	9.3	29	alto	H	8.2
10	bajo	H	8.8	30	alto	H	6.6
11	bajo	M	17.4	31	alto	M	22.0
12	bajo	M	17.1	32	alto	M	20.5
13	bajo	M	16.8	33	alto	M	19.9
14	bajo	M	16.7	34	alto	M	19.1
15	bajo	M	15.5	35	alto	M	18.5
16	bajo	M	15.3	36	alto	M	17.4
17	bajo	M	15.0	37	alto	M	17.0
18	bajo	M	15.4	38	alto	M	17.1
19	bajo	M	14.3	39	alto	M	17.1
20	bajo	M	14.0	40	alto	M	16.5

Se predice que se observará una interacción significativa entre las dos variables y que habrá un deseo de relacionarse alto y más en mujeres.

Hacemos el ANOVA

```
> a=aov(actsocial~deseorel*sexo)
> summary(a)
```

	Df	Sum Sq	Mean Sq	F value	Pr(>F)
deseorel	1	7.06	7.06	3.7772	0.0598
sexo	1	543.17	543.17	290.7670	< 2.2e-16 ***
deseorel:sexo	1	36.86	36.86	19.7339	8.14e-05 ***
Residuals	36	67.25	1.87		

Signif. codes: 0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1

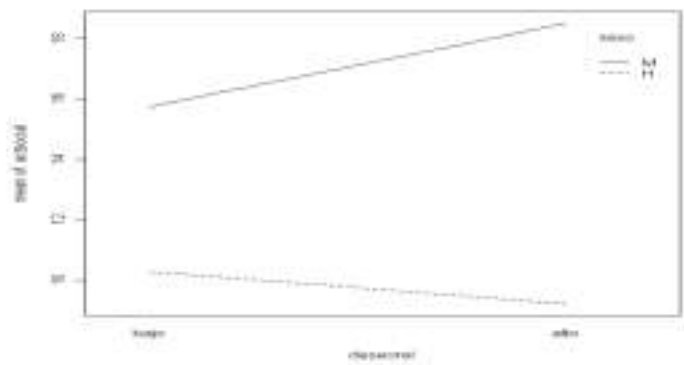
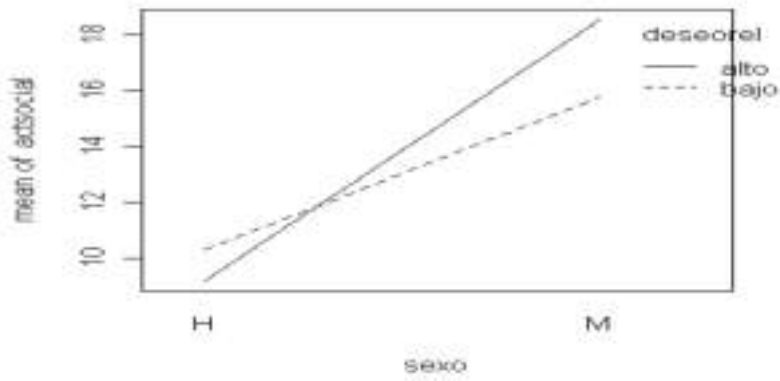
```
tapply(actsocial,list(deseorel,sexo),mean)
```

	H	M
Bajo	10.30	15.75
alto	9.22	18.51

La interacción deseo de relacionarse-sexo y la variable sexo son altamente significativas, mientras que la variable deseo de relacionarse no es significativa con relación a la actividad social, es decir, las mujeres tienen alto deseo de relacionarse (18.51).

En diseños factoriales es de interés tener un gráfico para apreciar los efectos principales y la interacción. El comando a emplear es **interaction.plot**

```
→ interaction.plot(sexo,deseorel,actsocial)
→ interaction.plot(deseorel,sexo,actsocial)
```



ANOVA INTRA – SUJETOS DE UN FACTOR

Este análisis se usa cuando se estudia una variable en tres o más condiciones y se usan **los mismos sujetos** (o sujetos **igualados**), para todas las condiciones experimentales.

Utilizaremos los mismos resultados del primer ejercicio de ANOVA. sin embargo, esta vez supondremos que **los mismos** seis sujetos se desempeñaron en las tres condiciones experimentales, aprendiendo listas diferentes de palabras a baja, media y alta velocidad de presentación.

Suj.	Grupo 1 (v baja)	Grupo 2 (v media)	Grupo 3 (v alta)
1	8	7	4
2	7	8	5
3	9	5	3
4	5	4	6
5	6	6	2
6	8	7	4

La variable velocidad de presentación entre grupos representa las diferencias predichas en los puntajes de recuerdo entre las condiciones experimentales. Sin embargo en esta ocasión no es lo mismo entre condiciones que entre sujetos. Dado que todos los sujetos se desempeñan en las tres condiciones, es posible observar el desempeño global de cada uno de los seis sujetos en las tres condiciones. Esto significa que las diferencias en los puntajes debidas a los sujetos individuales pueden considerarse como una fuente diferente de varianza.

La varianza debida al error representa las diferencias individuales **entre** los sujetos **dentro** de cada condición como consecuencia de

variables irrelevantes que afectan su desempeño. Es esta varianza debida al error la que se compara con la varianza entre condiciones en la razón F. Entonces las fuentes de variación serán: velocidad de presentación (**entre**), sujetos (**intra**) y el error.

Introducimos los datos en R así:

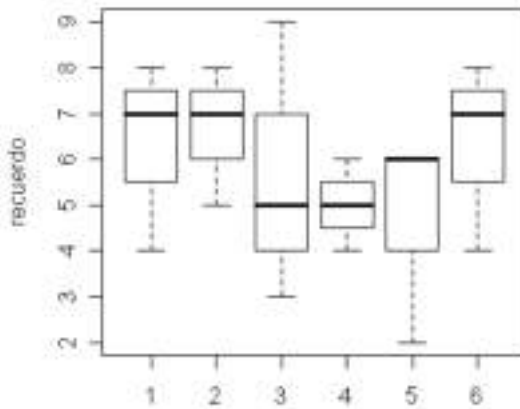
```
> recuerdo=c(8,7,9,5,6,8,7,8,5,4,6,7,4,5,3,6,2,4)
> velocidad=factor(rep(1:3,each=6),labels=c("v1","v2","v3"))
> sujetos=factor(rep(1:6,3))
> p=data.frame(sujetos,velocidad,recuerdo)
> p
```

	sujetos	velocidad	recuerdo
1	1	v1	8
2	2	v1	7
3	3	v1	9
4	4	v1	5
5	5	v1	6
6	6	v1	8
7	1	v2	7
8	2	v2	8
9	3	v2	5
10	4	v2	4
11	5	v2	6
12	6	v2	7
13	1	v3	4
14	2	v3	5
15	3	v3	3
16	4	v3	6
17	5	v3	2
18	6	v3	4

```
> tapply(recuerdo,sujetos,mean)
```

1	2	3	4	5	6
6.333333	6.666667	5.666667	5.000000	4.666667	6.333333

Gráficamente tenemos:



Se puede apreciar que las diferencias son mínimas

Ahora hacemos el ANOVA en R

```
> a=aov(recuerdo~velocidad+sujetos)
```

> summary(a)	Df	Sum Sq	Mean Sq	F value	Pr(>F)
velocidad	2	31.4444	15.72	227.1827	0.01164 *
sujetos	5	9.7778	1.9556	0.8934	0.52078
Residuals	10	21.8889	2.1889		

Signif. codes: 0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1

De acuerdo con estos resultados podemos concluir que existen diferencias significativas entre condiciones pero no entre sujetos, es decir, la velocidad de presentación incide sobre los puntajes de recuerdo pero no los sujetos.

ANOVA FACTORIAL INTRA – SUJETOS (DOS FACTORES “INTRA”)

Este análisis se usa cuando se estudian dos o más variables con dos o más condiciones para cada variable y se usan los mismos sujetos en todas las condiciones experimentales.

Ejemplo

Se dio una prueba a cuatro sujetos bajo cuatro condiciones, que representan una combinación de dos niveles de la variable A (longitud de las palabras) y dos niveles de la variable B (velocidad de presentación).

Puntajes de recuerdo

A ₁ (Palabras cortas)			A ₂ (Palabras largas)	
Sujeto	B ₁ (v. alta)	B ₂ (v. baja)	B ₁ (v. alta)	B ₂ (v. baja)
1	7	7	3	5
2	5	6	1	3
3	6	8	2	5
4	4	9	2	4

Predecimos que las variables A y B tendrán un efecto significativo sobre la variable, puntajes de recuerdo, que los sujetos podrán recordar más palabras cortas que palabras largas y que los puntajes serán significativamente más altos con una baja velocidad de presentación, también predecimos que no habrá interacción entre las dos variables A y B.

Introducimos los datos en R

```
> recuerdo=c(7,5,6,4,7,6,8,9,3,1,2,2,5,3,5,4)
```

```
> velocidad=factor(rep(rep(1:2,each=4),2),labels=c("v1","v2"))
```

```
> longitud=factor(rep(1:2,each=8),labels=c("l1","l2"))
```

```
> sujetos=factor(rep(rep(1:4,each=1),4),labels=c("1","2","3","4"))
```

```
> ejem=data.frame(sujetos,longitud,velocidad,recuerdo)
```

> ejem	sujetos	longitud	velocidad	recuerdo
1	1	l1	v1	7
2	2	l1	v1	5
3	3	l1	v1	6
4	4	l1	v1	4
5	1	l1	v2	7
6	2	l1	v2	6
7	3	l1	v2	8
8	4	l1	v2	9
9	1	l2	v1	3
10	2	l2	v1	1
11	3	l2	v1	2
12	4	l2	v1	2
13	1	l2	v2	5
14	2	l2	v2	3
15	3	l2	v2	5
16	4	l2	v2	4

Como con el ANOVA intra sujetos de un factor, las diferencias globales entre los sujetos en cada condición pueden considerarse como una fuente adicional de varianza. Por eso, para un diseño relacionado 2 x 2 la tabla de fuentes de varianza debe incluir no sólo las dos varia-

bles A y B, sino también una tercera variable, los sujetos (S). El efecto de esta variable adicional es bastante dramático puesto que ahora tenemos que tener en cuenta no sólo la interacción entre las dos variables A x B sino también la varianza debida al error de cada una de las variables por separado AS, BS, ABS, que representan en cada caso las diferencias individuales **entre** sujetos dentro de cada una de estas condiciones. Por lo tanto en la tabla de ANOVA se verá que el valor de F para A, B, y A x B, incluyen términos de error diferentes. En resumen este diseño puede considerarse como un diseño 2 x 2 pero con una **tercera** variable consistente en todas las diferencias en el desempeño de los sujetos en las diferentes condiciones. Esto sucede porque los mismos sujetos están desempeñándose en todas las condiciones, mientras que en el diseño factorial entre sujetos con dos factores cada sujeto se desempeñó en una sola condición.

Hacemos el ANOVA en R

```
>a=aov(recuerdo~(longitud*velocidad)+Error(sujetos/(longitud*velocidad)))
> summary(a)
```

Error: sujetos

	Df	Sum Sq	Mean Sq	F value	Pr(>F)
Residuals	3	7.1875	2.3958		

Error: sujetos:longitud

	Df	Sum Sq	Mean Sq	F value	Pr(>F)
longitud	1	45.562	45.562	729	0.0001115 ***
Residuals	3	0.188	0.063		

Signif. codes: 0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1

Error: sujetos:velocidad

	Df	Sum Sq	Mean Sq	F value	Pr(>F)
velocidad	1	18.0625	18.0625	14.695	0.03129 *
Residuals	3	3.6875	1.2292		

Signif. codes: 0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1

Error: sujetos:longitud:velocidad

	Df	Sum Sq	Mean Sq	F value	Pr(>F)
longitud:velocidad	1	0.0625	0.0625	0.0508	0.836
Residuals	3	3.6875	1.2292		

El efecto predicho de la longitud de palabras sobre los puntajes de recuerdo es altamente significativo, es decir, se recuerdan más palabras cortas.

El efecto de la velocidad de presentación de palabras sobre el puntaje de recuerdo es significativo, se recuerdan más palabras a velocidad baja.

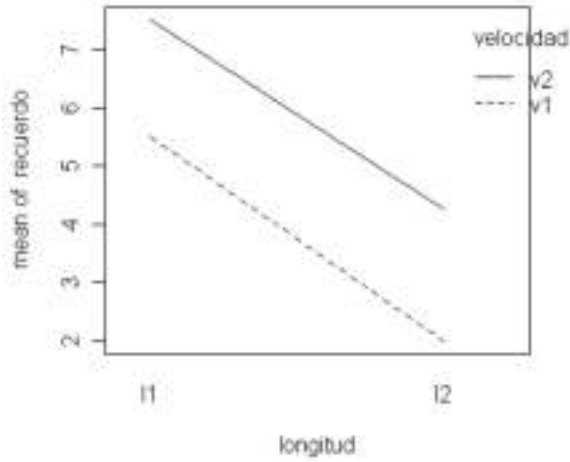
Sin embargo, la interacción no incide sobre la variable puntajes de recuerdo. Tampoco los sujetos afectan el recuerdo.

La tabla de medias la tenemos con **tapply**

```
> tapply(recuerdo,list(longitud,velocidad),mean)
  v1 v2
l1 5.5 7.50
l2 2.0 4.25.
```

Para el gráfico con **interaction. plot**

```
> interaction.plot(longitud,velocidad,recuerdo)
```



Es decir se ratifica lo dicho anteriormente

ANOVA EN DISEÑO DE CUADRADO LATINO

Tenemos un **ejemplo**:

Se desea descubrir qué efecto produce el tipo de música sobre el rendimiento, para lo cual ponemos en comparación cuatro programas de diferentes tipos de música que llamamos A, B, C, D, y un testigo sin música con la letra E. Cada grupo de cinco alumnos recibirá el estímulo un día laborable de la semana y cada día con un programa distinto, rotando los programas de una hora a otra, de tal manera que al final de los cinco días cada programa habrá ocupado una hora distinta del día.

```

B C D E A
C A E B D
D E A C B
E D B A C
A B C D E

```

Las condiciones experimentales son:

A = san juanitos, B = pasillos, C = vales criollos, D = vales clásicos, y
E = sin música

Resultados:

Lunes					Martes					Miércoles					Jueves					Viernes				
B	C	D	E	A	C	A	E	D	B	D	E	A	B	C	E	B	C	A	D	A	D	B	C	E
9	8	6	5	10	7	6	10	7	8	10	8	9	8	7	4	10	4	5	6	7	8	9	7	5
8	7	4	10	6	6	8	7	5	9	7	6	7	9	5	7	8	6	10	4	6	6	8	6	10
10	9	8	9	9	9	7	8	4	10	8	7	4	5	8	6	5	8	7	8	4	7	6	4	7
7	10	9	6	8	4	10	6	8	7	5	4	8	7	9	5	7	9	4	7	8	8	5	8	6
4	8	10	7	9	8	5	9	7	10	9	8	10	6	10	8	6	10	7	5	9	10	7	9	8

B	C	D	E	A
7.6	6.8	7.8	6.0	6.8
C	A	E	B	D
8.4	7.2	6.6	7.2	7.8
D	E	A	C	B
7.4	8.0	7.6	7.4	7.0
E	D	B	A	C
7.4	6.2	7.0	6.6	6.8
A	B	C	D	E
8.4	8.8	7.8	6.0	7.2

Hipótesis:

$$\mu_A = \mu_B = \mu_C = \mu_D = \mu_E$$

$$\mu_A \neq \mu_B \neq \mu_C \neq \mu_D \neq \mu_E$$

Introducimos los resultados en R

```
>rendimiento=c(7.6,8.4,7.4,7.4,8.4,6.8,7.2,8,6.2,8.8,7.8,6.6,7.6,7,7.8,6,7.2,7.4,6.6,6,6.8,7.8,7,6.8,7.2)
```

```
>fila=factor(rep(1:5,5))
```

```
>columna=factor(rep(1:5,each=5))
```

```
>musica=c("B","C","D","E","A","C","A","E","D","B","D","E","A",  
"B","C","E","B","C","A","D","A","D","B","C","E")
```

```
>p=data.frame(fila,columna,música,rendimiento)
```

>p	fila	columna	música	rendimiento
1	1	1	B	7.6
2	2	1	C	8.4
3	3	1	D	7.4
4	4	1	E	7.4
5	5	1	A	8.4
6	1	2	C	6.8
7	2	2	A	7.2
8	3	2	E	8.0
9	4	2	D	6.2
10	5	2	B	8.8
11	1	3	D	7.8
12	2	3	E	6.6
13	3	3	A	7.6
14	4	3	B	7.0
15	5	3	C	7.8
16	1	4	E	6.0
17	2	4	B	7.2
18	3	4	C	7.4
19	4	4	A	6.6
20	5	4	D	6.0
21	1	5	A	6.8
22	2	5	D	7.8
23	3	5	B	7.0
24	4	5	C	6.8
25	5	5	E	7.2

```
> a=aov(rendimiento~fila+columna+música,data=p)
```

```
> summary(a)
```

	Df	Sum Sq	Mean Sq	F value	Pr(>F)
fila	4	2.5184	0.6296	1.4344	0.2820
columna	4	3.8464	0.9616	2.1908	0.1316
música	4	0.9984	0.2496	0.5687	0.6904
Residuals	12	5.2672	0.4389		

No se encuentran diferencias significativas de los tratamientos con relación al rendimiento, tampoco existe incidencia significativa del día sobre el rendimiento.

COMENTARIOS

El análisis de varianza es una técnica extremadamente versátil. Con el análisis de varianza de un criterio o de dos criterios, tal como lo hemos descrito aquí se pueden analizar muchas situaciones de investigación. En un diseño de investigación factorial los participantes son divididos en grupos según las combinaciones de las variables cuyos efectos están siendo analizados. A través de los diseños factoriales podemos analizar los efectos de dos o más variables sin necesidad de convocar el doble de participantes. Además, estos diseños hacen posible el análisis de efectos interactivos, es decir, los efectos de las combinaciones de las dos variables. Específicamente, un efecto interactivo ocurre cuando el efecto de una variable depende del nivel de la otra variable. Un efecto principal es el efecto promedio general de una variable, ignorando el efecto de la otra variable. Los efectos principales e interactivos pueden describirse verbal, numérica y gráficamente.

REFERENCIAS

F. Carmona, Modelos lineales, publicaciones UB, 2005

J.J. Faraway, Linear Models with R, Chapman & Hall/CRC, 2004

J. Versani, Using R for Introductory Statistics. Chapman & Hall/CRC, 2004

C.M. Cuadras, Problemas de Probabilidades y Estadística. Vol.2: Inferencia Estadística. EUB, 2000

M. Perea, ANOVAs con el Programa Estadístico R. U de Valencia, 2008

R. Morales, Diseño Experimental. U. T. Ambato, 1996

A. Aron, E. Aron, Estadística para Psicología, 2001

CORRELACIÓN Y REGRESIÓN

“Una respuesta apropiada para un problema bien formulado es mucho mejor que una respuesta exacta para un problema aproximado”

**John Wilder Tukey (1915-2000),
estadístico estadounidense.**

INTRODUCCIÓN

La frase formulada por John W. Tukey, nos deja como enseñanza la importancia de formular correctamente un problema, nuestra primera preocupación debe ser esa, una vez formulado el problema en forma correcta podemos elegir el método más apropiado para resolverlo, una respuesta apropiada puede no ser exacta, como es el caso de pruebas estadísticas.

El concepto de correlación no fue inventado en realidad por los especialistas en estadística. Es uno de los procesos mentales más básicos. Los primeros humanos deben haber pensado en términos de correlaciones todo el tiempo, al menos aquellos que sobrevivieron. “Cada vez que nieva, los animales que cazamos huyen. La nieve es sinónimo de ausencia de animales. Cuando vuelva a nevar tendremos que seguir a los animales para no morir de hambre”.

De hecho, la correlación es un proceso mental tan típicamente humano que parecíamos tener una organización psicológica tal que nos lleva a encontrar un grado de correlación mayor que el que en realidad existe, como ocurría con los aztecas, quienes pensaban que las buenas cosechas estaban correlacionadas con los sacrificios humanos.

En este capítulo revisaremos la correlación, la regresión (predicción), regresión múltiple.

Ejemplo

Supongamos que una empresa esta pensando aumentar la cantidad de personal bajo el mando de cada uno de sus gerentes de piso. Sin embargo, la empresa esta preocupada por el estrés que esto podría provocar a sus gerentes. La empresa supone que cuantas mas personas supervise un gerente, mayor será el estrés sufrido por él. Para analizar la situación un psicólogo laboral sugiere estudiar a cinco gerentes seleccionados al azar de entre todos los gerentes de piso de la empresa. (En la práctica debería utilizarse un grupo mucho mayor, pero aquí utilizaremos sólo cinco casos para simplificar el ejemplo). Se entrega a cada uno de los cinco gerentes un cuestionario de medición de estrés en el cual los posibles registros van de 0 (estrés nulo) a 10 (estrés extremo). Los resultados son los siguientes:

Empleados supervisados y nivel de estrés (datos ficticios)

Empleados supervisados	nivel de estrés
6	7
8	8
3	1
10	8
8	6

Ingresamos los datos en R con la funcion **data.frame** que concatena todas las variables en un sólo conjunto de datos:

```
> supervisados=c(6,8,3,10,8)
> estrés =c(7,8,1,8,6)
> c=data.frame(supervisados, estrés)
> c
```

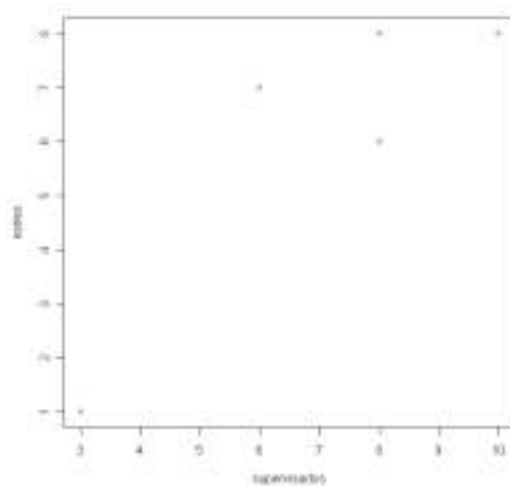
	supervisados	estrés
1	6	7
2	8	8
3	3	1
4	10	8
5	8	6

Podemos verificar los datos de este conjunto utilizando la función **summary**, la cual nos brinda para cada variable dentro del conjunto, los valores mínimos, mediana, media, máximo y los cuantiles primero y tercero, todo esto para el caso de las variables en escala intervalar:

> summary(c)

supervisados	estrés	supervisados	estrés
Min. :3	Min. :1	Mean :7	Mean :6
1st Qu.:6	1st Qu.:6	3rd Qu :8	3rd Qu :8
Median: 8	Median: 7	Max :10	Max : 8

Con la función **plot** podemos verificar gráficamente los datos
> plot(supervisados ,estrés)



El cálculo de la correlación se utiliza para cuantificar el grado de asociación de dos variables. El coeficiente de correlación r , también denominado de Pearson, se define con la siguiente ecuación:

$$r = \frac{\sum (x_1 - \bar{x})(y_1 - \bar{y})}{\sqrt{\sum (x_1 - \bar{x})^2 \sum (y_1 - \bar{y})^2}}$$

El coeficiente de correlación puede tener valores de -1 a $+1$, el signo indica la dirección de la asociación, por lo tanto las asociaciones más fuertes son -1 y $+1$, en el centro $r = 0$ indica falta de asociación. Otra medida derivada de r es el coeficiente de determinación, el cual simplemente se calcula como el cuadrado de r (r^2), el coeficiente de determinación multiplicado por 100 se pueden interpretar como el porcentaje de pares de puntos que se pueden explicar con la recta de regresión. Con R tenemos:

```
> cor(supervisados, estrés)
```

```
[1] 0.875075
```

Por lo tanto $r = 0.765756$ es decir se explica el 76.57 % de los casos.

Es muy importante y muchas veces omitido, el cálculo de significancia estadística del valor de r , para el caso de una distribución normal se puede calcular el estadístico t para un r dado con la siguiente ecuación:

$$t = \frac{r \sqrt{n-2}}{\sqrt{1-r^2}}$$

Para $gl = n-2$. por lo tanto la hipótesis nula es que $r = 0$. La función **cor.test** calcula el coeficiente de correlación, el estadístico t y la P correspondiente:

```
> cor.test(supervisados, estrés)
Pearson's product-moment correlation
data: supervisados and estrés
t = 3.1316, df = 3, p-value = 0.052
alternative hypothesis: true correlation is not equal to 0
95 percent confidence interval:
-0.03154805 0.99170008
sample estimates:
cor
0.875075
```

ANÁLISIS DE REGRESIÓN

El análisis de regresión fue introducido por Sir Francis Galton (1822-1911), científico británico de la época victoriana, el cual contribuyó al conocimiento de la antropometría, la psicología diferencial, la geografía y la estadística entre otras; además fue primo de Sir Charles Darwin. Uno de sus trabajos más influyentes es “Regresión Towards Mediocrity in Hereditary Stature”, publicado en 1886 en el Journal of the Anthropological Institute; en dicho trabajo Francis Galton analizó gráficamente la relación entre la altura de padres e hijos, concluyendo que la altura media de hijos nacidos de padres de una dada altura tienden a valores de la media de la población, luego Galton explicó esto diciendo que fue una **regresión** a los valores medios de la población.

El análisis de regresión se puede utilizar para describir la **relación**, su extensión, dirección e intensidad, entre una o varias variables independientes con escala intervalar y una variable dependiente también intervalar. Cabe destacar que **no** debe utilizarse el análisis de regresión o correlación como prueba de **causalidad**, la dirección de causalidad (justamente, qué es la causa de qué) no puede determinarse solamente a partir de la correlación. Tomemos el ejemplo del estrés de los gerentes. El estudio comenzó con la noción implícita de que supervisar un mayor número de personas (x) causa un aumento del nivel de estrés (y). El resultado del estudio fue una marcada correlación positiva entre x e y, que ciertamente coincide con la idea de que

x es la causa de y. Sin embargo, también coincide de la misma forma con la idea de que Y es la causa de X. (Tal vez los gerentes que parecen sufrir de estrés sean considerados muy trabajadores y ese sea el motivo por el cual sus superiores asignen mayor cantidad de personas a su cargo). También es posible que la correlación sea el resultado de algún tercer factor que cause que X e Y se desarrollen de manera conjunta. Por ejemplo, algunos sectores de la fábrica podrían necesitar más personal y también generar más estrés. Es decir, determinado sector de la fábrica causa estrés y requiere de muchos empleados para supervisar.

Existe bastante confusión acerca de este asunto de la correlación y la causalidad. El tema se complica al existir dos usos de la palabra **correlación**. Algunas veces se utiliza para describir un procedimiento estadístico (como lo hemos hecho en este capítulo), y otras veces se utiliza para describir un tipo de **diseño de investigación** en el que se miden dos variables en un grupo de personas, sin realizar una asignación aleatoria de sujetos a determinados valores de una de las variables. Comúnmente los **diseños de investigación** correlacionales son analizados estadísticamente utilizando el coeficiente de correlación, y los diseños de investigación experimentales se analizan utilizando pruebas estadísticas como las vistas en los anteriores capítulos.

En el análisis de regresión moderno, en su forma más simple es decir lineal y con una sola variable independiente (VI), la variable independiente X se relaciona con la variable dependiente (VD) Y, de acuerdo con la siguiente ecuación:

$$Y = \beta_0 + \beta_1 X + \varepsilon$$

Donde ε son los residuos, es decir la diferencia entre la estimación y los valores reales para cada par de puntos XY. Los coeficientes de la ecuación son: β_0 denominado intersección (cuando $X=0$) y β_1 pendiente.

En el ejemplo podemos apreciar el uso de la función **lm (modelo lineal)** para calcular la regresión lineal entre el número de personas su-

pervisadas como VI y el nivel de estrés como la VD:

```
> a=lm(estrés~supervisados)
```

```
> Summary (a)
```

```
Call:
```

```
lm(formula = estrés ~ supervisados)
```

Residuals:

1	2	3	4	5
1.9643	1.0357	-1.1429	-0.8929	-0.9643

Coefficients:

	Estimate	Std.Error	t value	Pr (> t)
(Intercept)	-0.7500	2.2753	-0.330	0.763
supervisados	0.9643	0.3079	3.132	0.052 .

Signif. codes: 0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1

Residual standard error: 1.629 on 3 degrees of freedom

Multiple R-Squared: 0.7658, Adjusted R-squared: 0.6877

F-statistic: 9.807 on 1 and 3 DF, p-value: 0.052

El resultado del análisis de regresión con la función **lm** muestra primero las variables utilizadas en la fórmula de regresión y el conjunto de datos (data) de donde provienen dichas variables. En una segunda parte muestra los valores mínimos, mediana, máximo y cuartiles primero y tercero de los residuos, esto es muy importante para poder verificar la distribución de los residuos. Finalmente muestra para cada coeficiente el valor estimado, el error estándar de la estimación, el estadístico t y la probabilidad. El estadístico t se utiliza para rechazar la hipótesis nula que dice que el coeficiente es igual a cero, en el caso de que no pueda rechazarse la H_0 se dice que no hay asociación entre la VI y la VD.

Reemplazando valores en la ecuación tendremos:

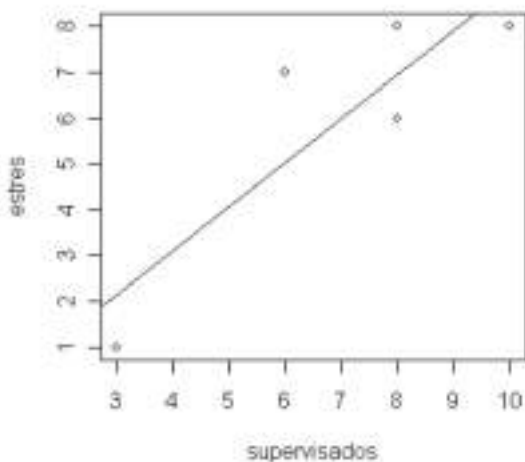
$$Y = 0.75 + 0.9643 X +$$

Con lo cual podemos hacer predicciones

Con la función plot se puede observar gráficamente la relación entre el número de supervisados y el nivel de estrés, junto con la recta de regresión, note en el código R la función **abline** para graficar dicha recta.

```
> plot(supervisados, estrés)
```

```
> abline(a)
```



¿Qué conclusiones podemos sacar de nuestro ejemplo? la primera es que la estimación de la pendiente no es estadísticamente significativa ($P=0.052$), está cerca con lo cual podemos concluir que existe una asociación pobre ($P>5\%$) entre el número de supervisados y el nivel de estrés, el signo positivo de la estimación (nro. de supervisados), nos indica que a medida que se incrementa, aumenta el nivel de estrés a un ritmo de 0.9643, el coeficiente de intersección no es estadísticamente significativo, con una estimación de 0.75, dicho valor

toma el nivel de estrés cuando el número de supervisados es = 0.

Sin embargo, el modelo no es lo suficientemente bueno, así es que probamos el mismo pero sin término independiente:

```
a=lm(estrés~supervisados-1)
> summary(a)
Call:
lm(formula = estrés ~ supervisados - 1)
```

Residuals:

1	2	3	4	5
1.7912	1.0549	-1.6044	-0.6813	-0.9451

Coefficients:

Estimate Std. Error t value Pr(>|t|)

Supervisados 0.86813 0.08693 9.986 0.000565 ***

Signif. codes: 0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1

Residual standard error: 1.436 on 4 degrees of freedom

Multiple R-Squared: 0.9614, Adjusted R-squared: 0.9518

F-statistic: 99.72 on 1 and 4 DF, p-value: 0.000565

Por lo tanto el modelo quedaría: **estrés = 0.868 supervisados**, que es mejor ya que los coeficientes de determinación son cercanos a 1, que sería el mejor modelo, además el p-value es significativo.

CORRELACIÓN Y REGRESIÓN MÚLTIPLE

Hasta aquí hemos aprendido a predecir el valor de una persona en la variable dependiente utilizando el valor de esa misma persona en una sola variable predictora. Es decir, se predice una variable dependiente (como puede ser el nivel de estrés) sobre la base de una variable predictora (como la cantidad de personal supervisado). Qué sucedería si se pudieran utilizar variables predictoras adicionales? Por ejemplo al predecir el nivel de estrés de los gerentes, supongamos que además de la cantidad de personal supervisado, también se conociera el nivel de ruido. Con esta información adicional, se podría realizar una predicción del nivel de estrés mucho mas acertada.

La asociación de una variable dependiente y dos o más variables independientes se denomina **correlación múltiple**. Realizar predicciones en la situación anteriormente descrita se denomina **regresión múltiple**. Con este método también es posible estudiar la interacción de dos variables calculando un coeficiente extra, por ejemplo:

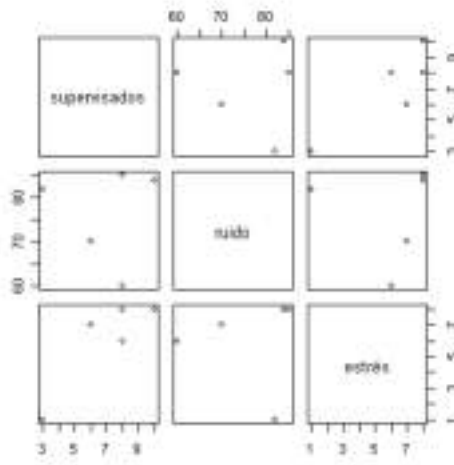
$$Y = \beta_0 + \beta_1 \text{supervisados} + \beta_2 \text{ruido} + \beta_3 \text{supervisadosruido} + \varepsilon$$

Introducimos los datos en R

```
> supervisados=c(6,8,3,10,8)
> ruido=c(70,85,82,84,60)
> estrés=c(7,8,1,8,6)
> c=data.frame(supervisados,ruido,estrés)
> c
```

	supervisados	ruido	estrés
1	6	70	7
2	8	85	8
3	3	82	1
4	10	84	8
5	8	60	6

Para captar visualmente las posibles relaciones lineales entre las variables podemos representar una matriz de diagramas de dispersión: **pairs(c)**



Numéricamente podemos cuantificar el grado de relación lineal mediante la matriz de coeficientes de correlación:

>cor(c)

	supervisados	ruido	estrés
supervisados	1.00000000	-0.00869125	0.87507503
ruido	-0.00869125	1.00000000	-0.01577436
estrés	0.87507503	-0.01577436	1.00000000

En lenguaje **R** la ecuación anterior se implementa así:

> a=lm(estrés~supervisados+ruido+supervisados*ruido)

> summary(a)

Call:

lm(formula = estrés ~ supervisados + ruido + supervisados * ruido)
Residuals:

1	2	3	4	5
0.17583	1.37571	-0.49233	-1.05499	-0.00422

Coefficients:

	Estimate	Std. Error	t value	Pr(> t)
(Intercept)	67.35234	57.25624	1.176	0.449
supervisados	-7.85454	7.38707	-1.063	0.480
ruido	-0.84018	0.70612	-1.190	0.445

supervisados:ruido 0.10812 0.09047 1.195 0.444

Residual standard error: 1.811 on 1 degrees of freedom

Multiple R-Squared: 0.9036, Adjusted R-squared: 0.6142

F-statistic: 3.123 on 3 and 1 DF, p-value: 0.3889

El resultado nos muestra que el estrés se decremento a un ritmo de -7.85 y -0.84 para el numero de supervisados y el ruido, sin embargo, los coeficientes estimados no son significativos.

Como en todos los casos anteriores podemos verificar la normalidad de los residuos. El coeficiente de determinación múltiple, $r = 0.90$ en nuestro ejemplo, cuantifica la cercanía de los puntos al plano de regresión; por otro lado el coeficiente de determinación ajustado tiene en cuenta la cantidad de coeficientes de la regresión múltiple, por lo tanto su valor es menor al r , para nuestro ejemplo $r = 0.61$. Ambos coeficientes de determinación, junto con el estadístico F y el P asociado nos indican la validez de la regresión, en la mayoría de los casos de la regresión múltiple no es posible una representación gráfica, en algunos casos se puede graficar en forma parcial las asociaciones.

COMENTARIOS

El análisis de regresión y la correlación son métodos utilizados para verificar y cuantificar la asociación entre dos o más variables. Es muy importante tener en cuenta que el análisis de regresión plantea un **modelo estadístico**, por lo tanto no es posible verificar **causalidad**, solamente podemos verificar **asociación**. Para poder verificar causalidad debemos utilizar **modelos determinísticos**; no existen métodos estadísticos para analizar causa – efecto. Como dijo John w. Tukey con su famosa frase el primer paso es formular correctamente el problema, en nuestro ejemplo el nivel de estrés es la variable dependiente y el número de supervisados es la variable independiente, formular de esta forma el problema nos muestra como el estrés aumenta con el número de supervisados, si cambiamos el orden diríamos que el número de supervisados aumenta con el estrés, esto no es correcto porque el número de supervisados no depende de el estrés ! La verificación de las condiciones anteriormente descritas para poder hacer el análisis de regresión es muy importante y muchas veces no tenida en cuenta.

Existe una gran ventaja en utilizar la correlación (o la regresión / correlación múltiples si es necesario) en lugar de la prueba **t** o el **ANOVA**. El método correlacional proporciona información directa acerca del grado de relación entre la(s) variable(s) de predicción y la variable dependiente, a la vez que permite realizar una prueba de significación. La prueba **t** y el ANOVA sólo brindan la significación estadística. Generalmente en los reportes científicos no se muestran los resultados del análisis de regresión en forma gráfica, a menos que muestre algo muy interesante, pero es muy aconsejable hacerlo para verificación personal de los investigadores.

REFERENCIAS

Faraway j. Practical Regression and Analysis of variance, 1990

Aron, A. Aron, E. Estadística para Psicología 2001

Downie, N. Métodos Estadísticos Aplicados 1986

Ryan TP. Modern Regresión Análisis, John Wiley & sons 1997

Risk, M. Cartas sobre Estadística. Argentina, 2005

GLOSARIO

Datos:

read.table("D:\\... dirección\\archivo", header=FALSE, sep=" ", dec=".",...) ### Lee archivo y crea un data frame (casos en filas y variables en columnas).

***data.frame* ###** crea un conjunto de datos (vectores, matrices, listas...) en formato data.frame.

Vectores y Matrices:

***c(...)* ###** concatena elementos en un vector: *c(1:10)*, *c("a","b")*.

***seq(from=a, to=b, by=0.1)* ###** crea una sucesión de números (*by*= o *length=...*)

***rep(x, r)* ###** repite los elementos del vector *x*, *r* veces.

***matrix(x, nrow=..., ncol=..., byrow=FALSE)* ###** crea una matriz con el vector *x*.

***cbind(x, y, z, ...)* ###** pega en una matriz con columnas *x*, *y*, *z*,...

***rbind(x, y, z, ...)* ###** pega en una matriz con filas *x*, *y*, *z*,...

***length(x)* ###** longitud del objeto *x*.

***dim(m)* ###** dimensión o dimensiones del objeto *m*.

***names(M)* ###** nombres de los elementos de *M*.

***colnames(M) <- c("col1", "col2",...)* ###** asigna nombres a las columnas de *M*.

***rownames* ###** como *colnames*, pero para las filas.

***apply(M, 1, funcion)* ###** aplica la función a las filas o a las columnas de *M*.

***tapply(x, factor, funcion)* ###** aplica la función a cada grupo de valores del vector *x* que ocupan la misma posición que los diferentes niveles del factor. Si factor no lo es se "coerce" a factor.

***diag(M)* ###** vector con los elementos diagonales de la matriz *M*.

`list( ... )` ### crea una lista de objetos.

Creación de Funciones:

`function( x,... ) { ... } ###` crea una función de argumentos x,...

Bucles:

`for( i in 1:100 ) { ... } ###` genera un bucle.

`if( cond ) expr ###` condicional.

Distribuciones:

`pnorm( q, m, s ) ###` calcula la función de distribución de una $N(m,s)$ en q.

`qnorm( p, m, s ) ###` cuantil de p en una $N(m,s)$.

`dnorm( x, m, s ) ###` calcula la densidad de una $N(m,s)$ en x.

`rnorm( n, m, s ) ###` calcula n valores pseudo aleatorios de una.

$N(m,s)$. ### Estas funciones se tienen para las distribuciones más habituales.

`pchisq( q, r ) ###` función de distribución en q de una χ^2 con r gdl.

`qchisq( p, r ) ###` función cuantil en p de una χ^2 con r gdl.

`set.seed( semilla ) ###` fija la “semilla” del generador de números pseudoaleatorios.

Muestreo aleatorio:

`sample(x, n, replace = FALSE, prob = NULL) ###` extrae una m.a. de x de tamaño n, sin o con reemplazamiento y prob especificadas.

Factores:

`factor( x ) ###` convierte el vector x en factor.

`gl( n, k, length=n*k, labels=...) ###` genera factores: n=n° de niveles; k=n° de réplicas; length=longitud total del factor.

Conversión a vector, matriz, data.frame,...

`as.vector`

as.matrix
as.data.frame
as.numeric
as.factor
as.table

Gráficos de dispersión:

x11() ### abre nueva ventana para un nuevo plot (sin anular el anterior).

plot(x, y,...) ### gráfico con pares de puntos (x,y).

lines(x, y, type="l") ### añade línea que une los puntos (x,y).

points(x, y, pch=1) ### añade los puntos(x,y).

abline(.....) ### añade a un plot rectas: *abline(a,b)* añade una línea de ordenada en el origen a y pendiente b; *abline(h=a, v=b)* añade una horizontal y=a y una vertical x=b.

rug(x) #### marca sobre un eje del plot los valores de x.

segments(x0, y0, x1, y1) ### añade segmentos que unen el extremo inicial (x0,y0) con el extremo final en (x1,y1).

curve(ff(x), add=T) #### añade la curva ff(x) al plot.

text(x, y, labels="el texto") ### añade el texto en las coordenadas (x,y).

mtext(text, side = 3, line = 0) ### añade texto en alguno de los 4 ejes.

legend(x, y, legend="...", bty="o") ### añade una leyenda al plot.

axis(side=1, at=, labels=...) ### añade un eje al plot.

matplot(x, y) ### plot de las columnas de la matriz x frente a las de la matriz y.

Histogramas y diagramas:

hist(x, probability=T,...) ### histograma de frecuencias de x.

barplot ### diagrama de barras.

boxplot ### diagrama de cajas.

Funciones de optimización:

uniroot(f, lower=a, upper=b) ### busca una raíz de la función f en el intervalo (a,b).

***nlm(f, p)* ###** minimiza la función *f*, con algoritmo de Newton y parámetros iniciales *p*. Puede calcular el hessiano (*hessian=T*).

Modelos lineales y no lineales (mínimos cuadrados):

***lm(y ~ x1+x2+x1*x3)* ###** ajusta un modelo lineal general.

***nls* ###** ajusta un modelo no lineal por mínimos cuadrados.

Tablas de Contingencia:

***xtabs(y ~ a + b)* ####** genera tabla con las frecuencias en *y*, cruzando los factores *a* y *b*.

***summary(x)* ###** da un resumen del objeto *x*. Depende del tipo de objeto.

***table* ###** tabla cruzada.

***as.table(M)* ###** coerce a *table* la matriz *M*.

***addmargins(tabla)* ###** da la tabla con las marginales.

***margins.table(tabla, margin=1)* ###** da una las marginales de la tabla.

***prop.table(tabla, margin=1)* ###** calcula las proporciones en la tabla.

***fable(tabla)* ###** da la tabla en formato más “elegante”.

***mosaicplot(table)* ###** crea un plot mosaico de una tabla de contingencia.

Tests sobre proporciones:

***binom.test(x, n, p=po)* ###** test exacto de $p=p_0$, cuando observamos *x* de una $b(n,p)$.

***prop.test(x, n, p)* ###** test χ^2 de las proporciones *p*.

Tests de no asociación y de simetría:

***chisq.test(tabla)* ###** test χ^2 de no asociación en la tabla *I* x *J*.

***\$expected* ###** valores esperados estimados bajo no asociación.

***\$residuals* ###** residuos de pearson.

***fisher.test(tabla)* ###** test exacto de Fisher de no asociación.

***mcnemar.test(tabla)* ###** test de simetría de McNemar.

Test CMH:

mantelhaen.test(tabla2xk) ### test CMH de independencia condicional en una tabla 2x2xK.

library(vcd)

woolf_test ### test de homogeneidad (OR iguales) en estratos de una tabla 2x2xK.

Modelos log-lineales, regresión logística y de poisson:

glm(y ~ x, family="binomial") ### ajusta un modelo lineal generalizado.

summary(mglm) ### resumen del objeto mglm.

anova.glm(mglm) ### análisis secuencial de la deviance del modelo mglm.

anova(modelo, test="Chisq") ### con test chi2 de efectos.

drop1(mglm) ### deviance al excluir cada término del modelo mglm.

anova(mo, m1) ### comparación de los modelos mo y m1 mediante deviance.

model.matrix(mglm) ### matriz de diseño del modelo mglm, incluidos factores.

glm values:

\$coef ### coeficientes del modelo.

\$fitted.values ### probabilidades o medias ajustadas.

\$linear.predictors ### predictores lineales del glm: xbeta's.

\$residuals ### working residuals.

\$df.residual ### grados de libertad residuales.

coef ### *\$coef*.

fitted ### *\$fitted.values*.

residuals.glm ### residuos tipo deviance, pearson, ...

predict ### predictor lineal.

deviance ### deviance (residual) del modelo ajustado.

AIC ### AIC del modelo ajustado.

Vcov ### matriz de covarianzas de los estimadores de los betas.

Algunas “librerías” (packages) útiles para ajustar diferentes modelos

Estos paquetes deberían estar instalados en R, antes de cargarlos en la sesión en la que se vayan a utilizar sus funciones.

`library( MASS )`

`corresp()` *### análisis de correspondencias de una tabla de contingencia. (svd)*

`glm.nb()` *### ajuste de un glm con la binomial negativa.*

`polr()` *### para ajustar modelos de riesgos proporcionales en situaciones de respuesta multinomial ordinal.*

`glmmPQL()` *### ajuste de modelos lineales generalizados mixtos.*

`library( nnet )`

`multino()` *### para ajustar modelos logit de respuesta multinomial con categorías nominales.*

`library( vcd )`

`woolf_test()` *### test de homogeneidad (OR iguales) en estratos de una tabla 2x2xK.*

`library( brglm )`

`brglm()` *### regresión logística con reducción de sesgo.*

`library( lme4 )`

`lmer()` *### ajuste de modelos lineales mixtos.*

`library( gee )`

`gee()` *### ajuste de GEE, ecuaciones estimadoras generalizadas.*

Ayuda en la consola de R: `help(nombre)` o `?nombre`, proporciona información sobre la función de R denominada nombre.

El lenguaje de R diferencia entre letras mayúsculas y minúsculas.

Tipos de objetos de R: data frame, vector, matrix, list, table, factor, lm, glm,...

La instalación básica de R y de los paquetes de interés puede hacerse desde su página web:

<http://www.r-project.org/>

BIBLIOGRAFÍA

John Fox (2002). *An R and S-Plus Companion to Applied Regression*. Sage Publications, Thousand Oaks, CA, USA.

<http://socserv.socsci.mcmaster.ca/jfox/Books/Companion/index.html>

Peter Dalgaard (2002). *Introductory Statistics with R*. Springer.

<http://www.biostat.ku.dk/~pd/ISwR.html>

William N. Venables and Brian D. Ripley (2002). *Modern Applied Statistics with S. Fourth Edition*. Springer, New York.

<http://www.stats.ox.ac.uk/pub/MASS4>

John Maindonald and John Braun (2003). *Data Analysis and Graphics Using R*. Cambridge University Press, Cambridge.

<http://www.maths.anu.edu.au/~johnm/r-book.html>.

Julian J. Faraway (2004). *Linear Models with R*. Chapman & Hall/CRC, Boca Raton, FL.

<http://www.maths.bath.ac.uk/~jjf23/LMR/>.

John Verzani (2005). *Using R for Introductory Statistics*. Chapman & Hall/CRC, Boca Raton, FL.

<http://wiener.math.csi.cuny.edu/UsingR/>

Paul Murrell (2006). *R Graphics*. Chapman & Hall/CRC, Boca Raton,

<http://www.stat.auckland.ac.nz/~paul/RGraphics/rgraphics.html>.

Michael J. Crawley (2005). *Statistics: An Introduction using R*. Wiley.

<http://www.bio.ic.ac.uk/research/crawley/statistics/>

Brian S. Everitt (2005). *An R and S-Plus Companion to Multivariate Analysis*. Springer.

<http://biostatistics.iop.kcl.ac.uk/publications/everitt/>.

Brian Everitt and Torsten Hothorn (2006). *A Handbook of Statistical*

Analyses Using R. Chapman & Hall/CRC, Boca Raton, FL.

<http://www.cran.r-project.org/src/contrib/Descriptions/HSAUR.html>

Julian J. Faraway (2006). *Extending Linear Models with R: Generalized Linear, Mixed Effects and Nonparametric Regression Models*. Chapman & Hall/CRC, Boca Raton, FL.

<http://www.maths.bath.ac.uk/~jjf23/ELM/>.

Deepayan Sarkar (2007). *Lattice Multivariate Data Visualization with R*. Springer, New York.

Jonathan D. Cryer and Kung-Sik Chan (2008). *Time Series Analysis with Applications in R*. Springer, New York.

W. John Braun and Duncan J. Murdoch (2007). *A First Course in Statistical Programming with R*. Cambridge University Press, Cambridge.

<http://www.stats.uwo.ca/faculty/braun/statprog/>

Robert H. Shumway and David S. Stoffer (2006). *Time Series Analysis and Its Applications With R Examples*. Springer, New York.

Fionn Murtagh (2005). *Correspondence Analysis and Data Coding with JAVA and R*. Chapman & Hall/CRC, Boca Raton, FL.

<http://www.cs.rhul.ac.uk/home/fionn/>]

RECURSOS EN INTERNET SOBRE R

The R Project for Statistical Computing:

<http://www.r-project.org/>

Wikipedia. R Proyect (español):

<http://es.wikipedia.org/wiki/R-project>

Wikipedia. R Proyect (Inglés):

http://en.wikipedia.org/wiki/R_%28programming_language%29

Trellis Graphics User's Manual:

<http://www.stat.auckland.ac.nz/~ihaka/787/trellis.user.pdf>

R para Principiantes:

http://cran.r-project.org/doc/contrib/rdebuts_es.pdf

Introducción a R:

<http://cran.r-project.org/doc/contrib/R-intro-1.1.0-espanol.1.pdf>

Cartas sobre Estadística de la Revista Argentina de Bioingeniería:

<http://cran.r-project.org/doc/contrib/Risk-Cartas-sobre-Estadistica.pdf>

Regresión Lineal con R:

<http://math.uprag.edu/Rsampleregression.pdf>

Curso introductorio de R:

http://www.es.geocities.com/r_vaquerizo/Manual_R_menu.htm

Using R for Introductory Statistics:

<http://www.cran.r-project.org/doc/contrib/Verzani-SimpleR.pdf>

Página de Charles J. Geyer:

<http://www.stat.umn.edu/geyer/5102/examp/>

Análisis Exploratorio y Confirmatorio de Datos de Experimentos de Microarrays:

http://www.dm.uba.ar/materias/analisis_expl_y_conf_de_datos_de

[_exp_de_marrays_Mae/2006/1/practicas.html](#)

Tutoriales sobre R:

<http://www.stat.auckland.ac.nz/~ihaka/120/software.html>

Universidad de Minnesota:

<http://www.stat.umn.edu/geyer/5102/examp/reg.html>

Curso básico de R:

<http://www.ub.es/stat/docencia/EADB/Curso%20basico%20de%20R.htm>

Curso básico de R:

<http://www.ub.es/stat/docencia/EADB/Curso%20basico%20de%20R-bn.pdf>

Paul Murrell's Home Page:

<http://www.stat.auckland.ac.nz/~paul/>

R Graphics:

<http://www.stat.auckland.ac.nz/~paul/RGraphics/rgraphics.html>

R graphics: overview, desiderata:

<http://biosun1.harvard.edu/~carey/CompMeth/StatVis/dem.pdf>

Un análisis con R. Datos Multivariantes:

<http://www.ub.es/stat/docencia/EADB/Ejemplo.pdf>

Practical Regression and Anova using R:

<http://cran.r-project.org/doc/contrib/Faraway-PRA.pdf>

Gráficos Estadísticos con R:

<http://cran.r-project.org/doc/contrib/grafi3.pdf>

Graphics with R:

<http://csg.sph.umich.edu/docs/R/graphics-1.pdf>

Graphics with R :

<http://www.stat.auckland.ac.nz/~ihaka/120/Notes/cho3.pdf>

An Introduction to R:

<http://www.stat.auckland.ac.nz/~ihaka/120/Notes/cho2.pdf>

Página de Ross Ihaka:

<http://www.stat.auckland.ac.nz/~ihaka/>

Página de John Maindonald:

<http://www.maths.anu.edu.au/~johnm/Personal.html>

Página de Julian Faraway:

<http://www.maths.bath.ac.uk/~jjf23/>

R Reference Card;

<http://cran.r-project.org/doc/contrib/Short-refcard.pdf>

Análisis exploratorio de datos:

<http://www.ciberconta.unizar.es/LECCION/aed/ead.pdf>

Exploratory Data Analysis:

<http://www.itl.nist.gov/div898/handbook/index.htm>

Exploratory data analysis (Wikipedia):

http://en.wikipedia.org/wiki/Exploratory_data_analysis

Análisis exploratorio de datos:

<http://pavlov.psicol.unam.mx:8080/CIM2000/Estadistica/Exploratorio/1exploratorio.htm>

Notes on the use of R for psychology experiments and questionnaires:

<http://cran.r-project.org/doc/contrib/Baron-rpsych.pdf>

Taller de introducción al uso de R para análisis estadísticos:

http://www.cricyt.edu.ar/interactivo/cursos/r_intro/

Programación en R:

http://www.faculty.ucr.edu/%7Eetgirke/Documents/R_BioCond/R_Programming.html

Lenguaje R:

<http://www.et.bs.ehu.es/~etptupaf/r.php>

Página de F. Tusell:

<http://www.et.bs.ehu.es/~etptupaf/nuevo/es/index.php>

Statistics and Statistical Graphics Resources:

<http://www.math.yorku.ca/SCS/StatResource.html#DataVis>

Tutorial R:

<http://www.traba.org/wikitraba/index.php/R>

Apuntes en R:

<http://lbe.uab.es/vm/salud/r/docs/apuntesr.pdf>

Curso de R:

<http://lbe.uab.es/vm/salud/r/docs/curso-R-xsjv.pdf>

Urso Diaz-Uriarte:

<http://cran.r-project.org/doc/contrib/curso-R.Diaz-Uriarte.pdf>

Casos prácticos de trabajos de investigación realizados con el programa estadístico “R”

Introducción

La misión de la Universidad es promover el desarrollo del país, formando cuadros profesionales preparados con rigor científico, capaces de “ver” la realidad, y de transformarla, usando la investigación científica. En este contexto, se presentan trabajos de investigación bajo el tema: “Quienes somos los universitarios”, realizado por estudiantes de las carreras de Psicología, Informática y Computación del sistema presencial y de Educación Básica del sistema semipresencial de la Facultad de Ciencias Humanas y de la Educación, en diferentes variables, tales como: Creencia sobre la naturaleza de la ciencia, uso de las NTICs, síndrome de “burnout”, religión, autoestima, rendimiento; que nos darán pautas para desarrollar proyectos de mitigación de impacto, en una demostración de que la investigación no es privativa de centros de investigación en donde la élite científica se reúne, sino que a lo interno de la Universidad con estudiantes y profesores, es posible generar nuevos conocimientos.

TEMA:

CREENCIAS SOBRE LA NATURALEZA DE LA CIENCIA UN ESTUDIO CON PROFESORES Y ESTUDIANTES DE LA UTA.

Bustos, M¹ y Morales, R.²

1,2 Profesores UTA

RESUMEN

Este estudio muestra las creencias sobre la naturaleza de la ciencia de una muestra de profesores científicos y no científicos y estudiantes de los últimos semestres de la UTA en base a las respuestas a algunas cuestiones del cuestionario de opiniones sobre ciencia y sociedad (COCS), de Acevedo(1994), un cuestionario de opción múltiple diseñado para evaluar opiniones sobre temas de ciencia tecnología y sociedad, adaptado del “*Views on Science-Technology-Society*” (VOSTS), preparado por Aikenhead, Fleming y Ryan (1987) y modificado poco después (Aikenhead y Ryan, 1992; Aikenhead, Ryan y Fleming; 1989. Los tópicos sobre la naturaleza de la ciencia estudiados aquí son los siguientes: Supuestos de la ciencia, elegancia de las teorías y leyes, el papel de los errores en la ciencia, el estatus epistemológico del conocimiento científico (realismo vs instrumentalismo), la coherencia de conceptos entre distintos paradigmas y cuestiones axiológicas. En general las creencias del profesorado son eclécticas, pero más bien inclinadas hacia posiciones positivistas; y que sus concepciones sobre el tema influyen significativamente en las creencias

acerca de la naturaleza de la ciencia de sus estudiantes. El contraste estadístico se hizo mediante la prueba no paramétrica “chi-cuadrado” ($p\text{-value} = 0.005737$) con un nivel de exigencia estadística mínimo de $p = 0.05$, para poder rechazar la hipótesis de nulidad. Esto evidencia la necesidad de incluir la naturaleza de la ciencia en la formación de los estudiantes como instrumento para mejorar la enseñanza y aprendizaje de la naturaleza de la ciencia en la UTA. Todo esto supone descartar los enfoques formativos reduccionistas, sesgados hacia el estudio de una sola corriente de pensamiento como sumo paradigma capaz de explicar los planteamientos sociales o filosóficos de la ciencia.

1.- INTRODUCCIÓN

Bajo la denominación de naturaleza de la ciencia se engloban todos aquellos aspectos que configuran la ciencia como una manera especial de llegar al conocimiento, es decir, los valores y suposiciones propias del desarrollo del conocimiento científico y que constituyen lo que se denomina el método científico (Aikenhead, 1979). Lejos de considerar este método como una vía única constituida por una serie de etapas o recetas algoritmizadas que seguidas mecánicamente permiten llegar a resultados seguros, el método científico se entiende hoy como un conjunto de supuestos no escritos y valores aceptados por la comunidad científica que sirven para avalar una racionalidad común. Así, la fundamentación en el cuerpo de conocimientos, la emisión y contrastación de hipótesis, la predecibilidad, la coherencia y la referencia empírica de las teorías y modelos constituyen lugares comunes habituales de esta metodología, cuyas exigencias necesarias e ineludibles son la comunicabilidad (publicidad) y la replicabilidad (Manassero, M. 2001)

La filosofía y la sociología de la ciencia han tenido un extraordinario desarrollo a lo largo de este siglo, analizando los valores y supuestos que caracterizan la actividad científica. Así aparecen diversas interpretaciones y enfoques, que además de superar los viejos planteamientos positivistas excesivamente ligados al empirismo lógico, representan análisis y críticas que han contri-

buido a acotar y precisar aspectos esenciales de la ciencia y la metodología científica, aumentando la profundidad y precisión de su conocimiento, y en consecuencia, lo que debe ser una correcta comprensión de la naturaleza de la ciencia, e incluso a exagerar ciertas críticas, subrayando nuevas perspectivas que hacen más complejo el panorama conceptual sobre la naturaleza de la ciencia (Popper, 1977; Kuhn, 1962; Lakatos, 1983; Feyerabend, 1982; Latour, y Woolgar, 1996; Laudan, 1986; Toulmin, 1977; y Bunge, 1980).

2.- REVISIÓN DE LITERATURA

Aunque la preocupación por la comprensión de la naturaleza de la ciencia como objetivo de la educación viene desde comienzo de siglo (Lederman, 1992; Matthews, 1998), la realidad es que ha estado poco presente en los currículos escolares durante muchos años, y en consecuencia, también de la formación general del profesorado, hasta que en la segunda mitad de este siglo ha comenzado a enfatizarse este objetivo reiterado especialmente en documentos más recientes, como el proyecto de ciencia para todos (AAAS, 1990) y números especiales de revistas como *Science Education* (1991, 75), coincidente con las propuestas en el mismo sentido de otros autores (Hazen y Trefil, 1991; Ruthford y Ahlgren, 1990) e instituciones (UNESCO, 1994). En el marco de la investigación en Didáctica de la Ciencia, de la conceptualización de la alfabetización científica y los movimientos Ciencia para todos y Ciencia – Tecnología – Sociedad han contribuido decisivamente a la extensión del conocimiento de la naturaleza de la ciencia como objetivo central de la enseñanza de la ciencia en muchos países (Stinner y Williams, 1998). Como consecuencia, las reformas curriculares emprendidas por la mayoría de estos países dentro de los sistemas educativos han introducido este objetivo en las materias de ciencias de los diversos niveles escolares, la UTA tiene en su currículo actual algunos niveles de investigación.

Siguiendo la estela de la epistemología y la sociología de la ciencia, investigación didáctica relacionada con la naturaleza de la

ciencia se ha tenido tres orientaciones principales:

1. Historia de la ciencia: reanálisis de casos históricos desde una perspectiva didáctica y educativa, para mejorar y renovar la enseñanza de la ciencia (Matthews, 1994; Solbes y Traver, 1996; Stinner y Williams, 1998).
2. Filosofía de la ciencia: análisis didáctico de los fundamentos de los distintos autores y corrientes en filosofía de la ciencia, como una parte importante para la fundamentación epistemológica de la ciencia escolar (Aliberas, Gutiérrez e Izquierdo, 1989; López 1990, 1995; Niaz, 1993).
- 3.- Analogías y relaciones entre naturaleza de la ciencia y teorías de la enseñanza y el aprendizaje de la ciencia: (Brickhose, 1990; Burbules y Linn, 1991; Mellado y Carracedo, 1993; Porlán, 1995).

Las concepciones de los estudiantes sobre la naturaleza de la ciencia han sido objeto de estudio temprano (Wilson, 1994). La mayoría de los estudios concluyen que los estudiantes tienen concepciones inadecuadas y tradicionales, ancladas en el positivismo e ignorantes de las principales aportaciones realizadas por la filosofía y sociología de la ciencia (Lederman, 1992; Désautels, y Laroche, 1998), aunque esta afirmación general tiene, no obstante, muchos matices diferentes según los estudios. Por otro lado estas investigaciones han dado pie a críticas metodológicas importantes tales como la ausencia de una definición clara del marco epistemológico empleado por los investigadores, que no refleja suficientemente la naturaleza diversa y conflictiva de la naturaleza de la ciencia y dificulta la comparación entre diferentes estudios e incluso puede dar lugar a contradicciones internas de los instrumentos que podrían afectar a la validez de algunas conclusiones (Gardner, 1996). Además, la costumbre de resumir en una única puntuación asuntos tan complejos y variados como el método, el status del conocimiento, los patrones de cambio, los criterios de demarcación, y diversas dualidades enfrentadas – objetivismo/subjetivismo, realismo/instrumenta-

lismo, idealismo/ objetivismo, inducción/invencción, definitivo/provisional, etc. - también puede llevar a conclusiones sesgadas (Koulaidis y Ogborn, 1995).

A pesar de estas objeciones, se acepta que los resultados negativos en las concepciones de los estudiantes exigen la necesidad de mejorar la enseñanza y la comprensión de la naturaleza de la ciencia en los diferentes niveles educativos. Para ello se han sugerido dos caminos diferentes, pero relacionados: por un lado, la mejora de los currículos de ciencias, introduciendo la naturaleza de la ciencia como un contenido de la enseñanza de la ciencia cuando está ausente o actualizando su presentación didáctica; por otro lado, la mejora de las concepciones del profesorado sobre naturaleza de la ciencia, en el supuesto implícito que esta mejora tendrá una inmediata traducción en la mejora de la enseñanza de la ciencia a través de la acción del profesorado.

Además de los inconvenientes derivados de los defectos metodológicos de los instrumentos empleados, las contradicciones personales observadas en las conclusiones de diversos estudios hacen muy compleja la situación global de la investigación del pensamiento epistemológico del profesorado (Acevedo, 1994; Lakin y Wellington, 1994). Una misma persona puede sostener, simultáneamente, ideas positivistas en determinados aspectos conviviendo con otras de corte más constructivista y sociológico, que resultan parcialmente contradictorias con las primeras. Aunque esta contradicción resulta llamativa para los investigadores, no reviste tanta importancia para el profesorado, ya que éste carece de los fundamentos necesarios y la reflexión epistemológica suficiente sobre la naturaleza de la ciencia para construir unas teorías implícitas personales coherentes. (Lederman y O'Malley, 1990)

Los **propósitos** del presente trabajo fueron:

Determinar si las creencias que tienen los profesores de la Universidad Técnica de Ambato sobre la naturaleza de la ciencia, influyen sobre las creencias que sus estudiantes tienen sobre el

tema.

Evaluar las concepciones o creencias que tienen los profesores de la UTA sobre la Naturaleza de la Ciencia.

Evaluar las concepciones o creencias que tienen los estudiantes de la UTA sobre la Naturaleza de la Ciencia.

Relacionar las concepciones del profesorado sobre la NdC, y las concepciones del alumnado sobre el tema.

3. METODOLOGÍA

Se realizaron encuestas a 125 estudiantes de los últimos semestres de las facultades de Ciencias Humanas y de la Educación, Ingeniería en Alimentos, Ingeniería Agronómica, Ciencias de la Salud y Ciencias Administrativas, mediante un muestreo intencional y por estratos se encuestaron a 35 docentes de estas mismas facultades.

En primer lugar los encuestados contestaron a las preguntas del Cuestionario de Opiniones sobre Ciencia y Sociedad COCS que consta de 20 enunciados expresados unos en términos positivos y otros en términos negativos simplificados en seis grandes tópicos derivados de la sociología y la epistemología de la ciencia, con los cuales se establecieron categorías (A,B,C,D) a fin de darles un significado que facilitara su interpretación. En segundo lugar, una vez analizados e interpretados los datos, se compararon los resultados de las respuestas dadas por docentes y estudiantes.

4. RESULTADOS OBTENIDOS

TABLA 1 (Docentes)

Cuestiones	Acuerdo	Dudoso	Desacuerdo	Valoración grupal
Los modelos teóricos elaborados por los científicos, por ejemplo los modelos atómicos o del ADN, pretenden describir los mas exactamente posible la realidad	25	7	3	Bastante de acuerdo
Los mejores científicos son los que siguen en sus investigaciones las etapas del método científico los más escrupulosamente posible.	17	11	7	De acuerdo
En general los científicos son más objetivos e imparciales en sus investigaciones que la mayoría de los demás ciudadanos en sus trabajos.	20	13	2	De acuerdo
Los contactos sociales de los científicos no influyen en su trabajo profesional, ni en el contenido del conocimiento científico de sus descubrimientos.	12	11	12	Algo de acuerdo
La política de un país tiene poca influencia sobre el trabajo de sus científicos, porque sus preocupaciones investigadoras se en-	14	11	10	Algo de acuerdo

cuentran en general al margen de la política.				
Quando las investigaciones científicas son correctas el conocimiento que se deriva de ellas no cambia.	14	9	12	Algo de acuerdo

TABLA 2 (Estudiantes)

Cuestiones	Acuerdo	Dudoso	Desacuerdo	Valoración grupal
Los modelos teóricos elaborados por los científicos, por ejemplo los modelos atómicos o del ADN, pretenden describir los mas exactamente posible la realidad	81	36	8	Bastante de acuerdo
Los mejores científicos son los que siguen en sus investigaciones las etapas del método científico los más escrupulosamente posible.	52	49	24	De acuerdo
En general los científicos son más objetivos e imparciales en sus investigaciones que la mayoría de los demás ciudadanos en sus trabajos.	74	27	24	Bastante de acuerdo
Los contactos sociales de los científicos no influyen en su trabajo profesional, ni en el contenido del co-	48	52	25	Algo de acuerdo

nocimiento científico de sus descubrimientos.				
La política de un país tiene poca influencia sobre el trabajo de sus científicos, porque sus preocupaciones investigadoras se encuentran en general al margen de la política.	46	49	30	Algo en contra
Cuando las investigaciones científicas son correctas el conocimiento que se deriva de ellas no cambia.	62	30	33	Bastante de acuerdo

Se observa en las tablas anteriores (1 y 2) que la tendencia de profesores y estudiantes de la UTA, en general van hacia el Realismo Ontológico, predomina el absolutismo empirista, basado en la excelencia del método científico que posee un status jerárquicamente superior desde un punto epistemológico, sobre el pluralismo metodológico. Es claramente mayoritaria la visión Objetivista, en comparación con la subjetivista. La mayoría muestra posiciones contextualistas, y una importante mayoría tiene un punto de vista dinámico del conocimiento científico.

Esto se confirma con el siguiente cuadro:

	Docentes	Alumnos
A	3	1
B	4	22
B – C	1	9
C	17	79
D	10	14

En donde la mayoría está en la categoría C.

Principales características y descripciones de la tipología tomadas del documento de Acevedo (2001).

Tipos	Rasgos	Descripciones
A	Idealistas ontológicos Relativistas epistemológicos: subjetivistas por el contexto	La construcción del conocimiento científico depende del contexto sociopolítico. La práctica científica no garantiza la objetividad de la ciencia. No es posible describir una realidad única, puesto que no existe; con nuestras teorías tan solo podemos hacer interpretaciones. Así, el conocimiento científico se modifica porque en ocasiones cambian de manera ontológica las perspectivas conceptuales con las que se interpreta el mundo.
B	Realistas ontológicos Relativistas epistemológicos subjetivistas por el contexto	La construcción del conocimiento científico depende del contexto sociopolítico. El modo de trabajo científico no garantiza la objetividad de la ciencia. Es posible hacer una descripción de la realidad, pero siempre desde una determinada perspectiva. De esta manera, el conocimiento científico puede cambiar aunque proceda de investigaciones correctas.
C	Realistas ontológicos Empiristas contextualistas. Objetivistas y positivistas.	Aunque el contexto sociopolítico puede influir más o menos, hay una realidad única que es posible describir con objetividad accediendo a ella empíricamente, preferentemente por inducción, mediante la utilización sistemática y rigurosa del método científico. El contexto puede facilitar o dificultar esta labor, esto es, algunos contextos sociales, políticos y culturales favorecen el acceso al conocimiento científico, mientras que otros nos alejan de él. El

		conocimiento científico suficientemente probado por investigaciones correctas no cambia básicamente, cuando se modifica no es tanto por un cambio de perspectiva en la forma de ver el mundo, sino por una ampliación acumulativa del dominio de aplicación de la teoría elaborada.
D	Realistas ontológicos Empiristas radicales. Objetivistas y positivistas	El contexto sociopolítico no influye en el conocimiento científico correcto, porque este es universal y se encuentra libre de la carga de la subjetividad que conlleva tal influencia. Existe una realidad única que se puede describir con objetividad accediendo a ella empíricamente, preferentemente por inducción, mediante la utilización sistemática del método científico. El conocimiento científico suficientemente probado por investigaciones correctas no cambia básicamente, cuando se modifica no es por cambio de perspectiva en la forma de ver el mundo, sino por una ampliación acumulativa del dominio de aplicación de la teoría elaborada.

5. CONTRASTE DE HIPOTESIS

Ho: independencia

H1: relación entre profesores y estudiantes

El contraste estadístico se hizo con la prueba chi cuadrado, en el software estadístico libre r.

```
>creencias=matrix(c(3,4,1,17,10,1,22,9,79,14),2,5,byrow=T)
```

```
> dimnames(creencias)
```

```
NULL
```

```
> dim(creencias)
```

[1] 2 5

```
> uta=c("Profesores","Alumnos")
> tendencias=c("A","B","BC","C","D")
> dimnames(creencias)=list(uta,tendencias)
> creencias
```

	A	B	BC	C	D
Profesores	3	4	1	17	10
Alumnos	1	22	9	79	14

```
> chisq.test(creencias)
```

Pearson's Chi-squared test

data: creencias

X-squared = 14.5479, df = 4, p-value = 0.005737

Es decir, existe relación entre las creencias que tienen los profesores con las que tienen los estudiantes acerca de la naturaleza de la ciencia.

6. DISCUSIÓN DE RESULTADOS

La mayoría de los profesores (49%) y estudiantes (69%) cree en la objetividad de la ciencia y de los científicos en su trabajo, asociando su objetivismo a la existencia de un potente, riguroso y universal método científico, que permite acceder al conocimiento de manera empírica y preferentemente inductiva, de acuerdo con la clásica secuencia Observación – Hipótesis – Experimentación – Teoría (O-H-E-T). así pues, para nuestros investigados, es mas bien la metodología de las ciencias experimentales la que las hace mas objetivas para elaborar un conocimiento valido y fiable, ra-

dicando su eficiencia en la posibilidad de someter los descubrimientos al control de la comunidad científica mediante la replica imparcial de las investigaciones. Por lo tanto por esta cuestión podría decirse sin ambigüedades que la mayoría de los investigados manifiesta una posición epistemológica, bastante coherente, acorde con uno de los constructos del pensamiento de Gallagher 1991, Abell y Smith 1994 y Porlan y Rivero 1998.

En este estudio la mayoría de los investigados, además del objetivismo, tienen puntos de vista epistemológicos próximos al empirismo y al positivismo, aunque son pocos profesores (28%) y estudiantes (12%) los que se manifiestan empiristas radicales (D), si se comparan con los que admiten cierta clase de contextualismo, pero con un sentido bastante diferente al de los relativistas epistemológicos; es decir, opinan que el conocimiento científico se ve intrínsecamente afectado por factores culturales, sociales, históricos y políticos. En relación con esto, Koulaidis y Ogborn (1989) señalaron que los profesores de ciencias suelen asumir posiciones eclécticas sobre la naturaleza de la ciencia, más próximas al contextualismo de Kuhn que al empirismo. En nuestro trabajo al igual que ocurre en otros recientes desarrollados en España (por ejemplo Rebollo, 1998) también se encuentra un elevado porcentaje de referencias al contexto en nuestros investigados. En efecto, la mayor parte de los que admiten la influencia de los aspectos sociales, políticos etc. Lo hacen asignándole un significado diferente, consideran que en algunos contextos socioculturales, como el de la civilización occidental, se apoya públicamente la ciencia creando el marco adecuado, aportando fondos y subvenciones, y dando reconocimiento a los científicos más destacados, y premiándolos con honores y recompensas. De esta manera se protege y promueve, en general, el conocimiento científico.

Otra cuestión importante abordada en este trabajo es la visión del cambio del conocimiento científico que tienen los profesores (55%) y estudiantes (50%) de la UTA. En relación con esto, Rebollo (1998) ha proporcionado resultados que muestran que prácticamente la totalidad de los investigados, licenciados en Biología o en Química de la Universidad de Málaga, admite el cambio de

conceptos y teorías científicas, concediendo un estatus temporal al conocimiento científico. Visión que es compartida por estudiantes portugueses de profesorado de educación primaria (Thomaz, Cruz, Martin y Cachapus, 1996) y esto lo hacen desde una perspectiva positivista y acumulativa del conocimiento, que es otro de los rasgos señalados por Porlán (1994) de acuerdo con el principio de simplicidad y máxima economía del conocimiento, que es una forma de “reduccionismo epistemológico”.

7. CONCLUSIONES IMPLICACIONES EDUCATIVAS

Existe una estrecha relación entre las creencias de los profesores y estudiantes acerca de la naturaleza de la ciencia

Predominan algunas creencias, en profesores y estudiantes, ya mostradas en otros trabajos, tales como, realismo, objetivismo, estatus jerárquicamente superior del método científico, empirismo, visión acumulativa del conocimiento científico, positivismo.

Se encuentra un panorama muy rico y complejo sobre algunos puntos de vista epistemológicos, contrarios a la mayoría como: pluralismo metodológico, subjetivismo, relativismo, visión cambiante del conocimiento científico y la posible influencia de factores sociales, culturales y políticos en la ciencia y los conocimientos que esta elabora.

Nuestra posición es opuesta al adoctrinamiento de quienes pretenden imponer una determinada perspectiva de la naturaleza de la ciencia a profesores y estudiantes, presentándola como si fuera la mejor o inmutable, por el contrario, aquí se defiende la necesidad de mostrar a profesores y alumnos, diversos puntos de vista sobre el tema, dándoles a conocer las distintas formas de entender la naturaleza de la ciencia para que puedan comprenderla mejor, valorarla críticamente, y sobre todo, adquirir la idea clave de que las conceptualizaciones sobre la naturaleza de la ciencia también cambian, tal y como ocurren con los propios conocimientos científicos (Vásquez y Manassero, 1995). Todo esto supone descartar

los enfoques formativos reduccionistas, sesgados hacia el estudio de una sola corriente de pensamiento como sumo paradigma capaz de explicar los planteamientos sociales o filosóficos de la ciencia; por el contrario es necesario presentar a la comunidad universitaria una variedad de pensamientos, opiniones y enfoques para que puedan someterlos a un reflexivo análisis crítico (Vázquez, Acevedo, Manassero y Acevedo, 2001). Se cumpliría así el objetivo de que la introducción en la formación de profesionales de la reflexión epistemológica sobre la ciencia conduzca a que se adquiriera una visión más plural, evitando en lo posible posturas más o menos dogmáticas. (Jiménez, 1995).

8. BIBLIOGRAFÍA

ACEVEDO, J.A. (1992): "Cuestiones de sociología y epistemología de la ciencia. La opinión de los estudiantes", en: *Revista de Educación de la Universidad de Granada*, 6, 167-182.

ACEVEDO, J.A. (1993): "¿Qué piensan los estudiantes sobre la ciencia? Un enfoque CTS", en: *Enseñanza de las Ciencias*, n.º extra (IV Congreso), 11-12.

ACEVEDO, J.A. (1994): "Los futuros profesores de enseñanza secundaria ante la sociología y la epistemología de las ciencias", en: *Revista Interuniversitaria de Formación del Profesorado*, 19, 111-125. Versión electrónica corregida y actualizada en *Sala de Lecturas CTS+I de la OEI*

<<http://www.campusoei.org.salactsi/acevedo8.htm>>, 2001.

ACEVEDO, J.A. (1996): "La formación del profesorado de enseñanza secundaria y la educación CTS. Una cuestión problemática", en: *Revista Interuniversitaria de Formación del Profesorado*, 26, 131-144. Versión electrónica corregida y actualizada en *Sala de Lecturas CTS+I de la OEI*

<<http://www.campusoei.org.salactsi/acevedo9.htm>>, 2001.

Vázquez, A., Acevedo, J. A., Manassero, M.A., y Acevedo, P. (2006) Creencias ingenuas sobre naturaleza de la ciencia: consensos en

sociología interna de ciencia y tecnología. Actas del IV Seminario Ibérico de CTS en la educación científica.

MANASSERO, M., VASQUEZ, A. “Creencias del Profesorado sobre la Naturaleza de la Ciencia”. (2001).

TEMA

ESTUDIO COMPARATIVO DEL SINDROME DE BURNOUT DE EMPLEADOS Y PROFESORES DE LA UTA

Villagómez, S¹ y Morales, R.²

¹Estudiante Psicología

² Profesor UTA

INTRODUCCION

El estrés laboral es el segundo problema de salud más frecuente en la UE. Para la Unión Europea, el estrés laboral está considerado como el segundo problema de salud más frecuente por detrás los trastornos musculoesqueléticos. Su coste anual en Europa se ha llegado a cifrar en 20.000 millones de euros.

En los últimos 40 años hubo un incremento muy importante de los problemas relacionados con el estrés laboral, por lo que es necesario analizar el síndrome de Burnout que fue utilizado por primera vez por el psicólogo clínico Herbert Freudenberger para definir el desgaste extremo de un empleado, al ser este estado causa del mencionado "estrés laboral". Fue adoptado por los sindicatos y abogados como elemento de ayuda para mencionar los problemas físicos generados por un grado de agotamiento excesivo. En la actualidad es una de las causas más importante de incapacidad laboral.

El presente proyecto se lo realizó con el propósito determinar los niveles de este síndrome en empleados y profesores de la uta.

METODOLOGIA

La investigación se la realizó en Ecuador, en la provincia de Tungurahua, en su ciudad capital Ambato, en la Universidad Técnica de Ambato en las facultades de Ciencias Humanas, de Sistemas, de la Salud y en el departamento de Administración Central. En una muestra de 66 personas.

Para evaluar el Burnout, el M.B.I. (Maslach Burnout Inventory), es el instrumento usado con mayor frecuencia y fue elaborado por Maslach y Jackson (1981-1986). En esta ocasión, fue el instrumento que se aplicó a todos los integrantes de la muestra.

Consiste en un inventario auto administrado, creado para evaluar tres aspectos fundamentales del síndrome:

- El agotamiento o cansancio emocional
- La despersonalización
- La falta de realización personal.

Las características de cada subescala que integra el Inventario de Burnout de Maslach son:

1. Subescala de agotamiento emocional. Valora la vivencia de estar exhausto emocionalmente por las demandas del trabajo. Puntuación máxima 54.
2. Subescala de despersonalización. Valora el grado en que cada uno reconoce actitudes de frialdad y distanciamiento. Puntuación máxima 30.
3. Subescala de realización personal. Evalúa los sentimientos de auto eficacia y realización personal en el trabajo. Puntuación máxima 48.

Se consideran que las puntuaciones son bajas entre 1 y 33. Puntuaciones altas en los dos primeros y baja en el tercero definen el síndrome.

Se aplicó del test Maslach Burnout Inventory en la muestra y tomando como uno de los mayores indicadores el área Agotamiento Emocional como la más significativa para determinar el grado de burnout, la escala es:

BAJO	(= 0 <18)	Igual o menor a 18 un nivel de burnout bajo.
MEDIO	(19 - 26)	De 19 a 26 un nivel de burnout medio.
ALTO	(= 0 >27)	Igual o mayor a 27 un nivel de burnout alto.

RESULTADOS

FACULTAD DE CIENCIAS HUMANAS

En la Facultad de Ciencias Humanas se tomó una muestra de 23 personas distribuidas entre los sectores administrativos, docentes y empleados de la facultad.

Estos son los datos obtenidos en este estudio.

```
> burnout=c(1,3,3,3,6,3,6,6,5,5,8,5,9,3,4,2,5,8,6,4,6,5,7)
> mean(burnout)
[1] 4.913043
> median(burnout)
[1] 5
> range(burnout)
[1] 1 9
> var(burnout)
[1] 4.083004
> sd(burnout)
[1] 2.020644
```

Teniendo un rango de entre 1 – 9 es un nivel de burnout sumamente bajo de acuerdo a los parámetros que indica la tabla de valores del test.

Para la determinación del grado de burnout que esta facultad pre-

senta, hemos tomado en cuenta la median (5) que la representamos en este boxplot que al revisar la tabla de valores nos indica un grado muy bajo de agotamiento emocional en la Facultad de Ciencias Humanas.

FACULTAD DE INGENIERÍA EN SISTEMAS

En la Facultad de Ingeniería en Sistemas se tomó una muestra de 11 personas distribuidas entre los sectores administrativos, docentes y empleados.

Estos son los datos obtenidos en este estudio.

```
burnout=c(5,8,5,5,5,6,3,11,5,7,4)
```

```
> mean(burnout)
```

```
[1] 5.818182
```

```
> median(burnout)
```

```
[1] 5
```

```
> range(burnout)
```

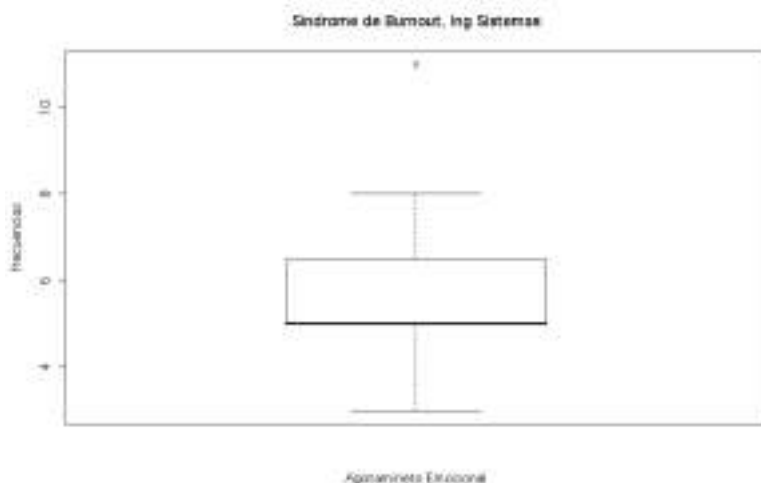
```
[1] 3 11
```

```
> var(burnout)
```

```
[1] 4.763636
```

```
> sd(burnout)
```

```
[1] 2.182576
```



Teniendo un rango de entre 3 - 11 que sigue siendo un nivel de burnout sumamente bajo de acuerdo a los parámetros que indica la tabla de valores del test.

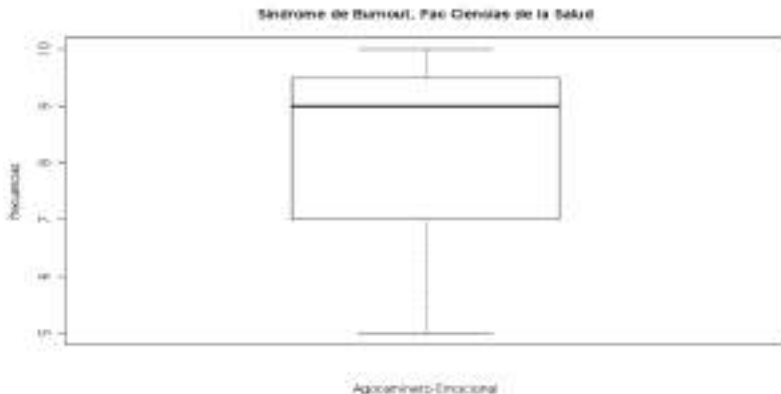
Para la determinación del grado de burnout que esta facultad presenta, hemos tomado en cuenta la median (5) que la representamos en este boxplot que al revisar la tabla de valores nos indica un grado muy bajo de agotamiento emocional en la Facultad de Ingeniería en Sistemas.

CIENCIAS DE LA SALUD

En la Facultad de Ciencias de la Salud se tomó una muestra de 4 personas distribuidas entre los sectores administrativos, docentes y empleados de la facultad.

Estos son los datos obtenidos en este estudio.

```
> burnout=c(5,10,9,9)
> mean(burnout)
[1] 8.25
> median(burnout)
[1] 9
> range(burnout)
[1] 5 10
```



```
> var(burnout)
[1] 4.916667
> sd(burnout)
[1] 2.217356
```

Teniendo un rango de entre 5 - 10 que sigue siendo un nivel de burnout sumamente bajo de acuerdo a los parámetros que indica la tabla de valores del test.

Para la determinación del grado de burnout que esta facultad presenta, hemos tomado en cuenta la median (9) que la representamos en este boxplot que al revisar la tabla de valores nos indica un grado muy bajo de agotamiento emocional en la Facultad de Ciencias de la Salud.

ADMINISTRACION CENTRAL

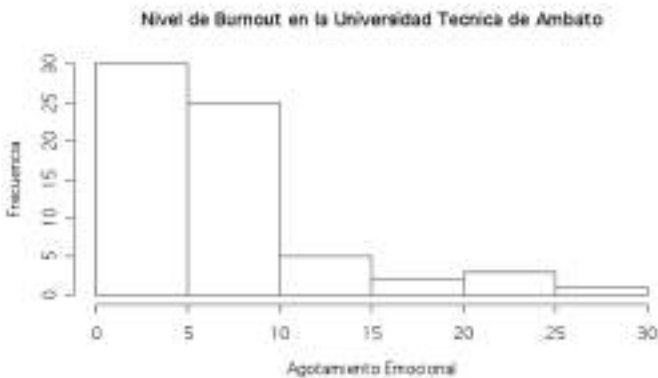
En Administración Central se tomó una muestra de 28 personas distribuidas entre los sectores administrativos. Estos son los datos obtenidos en este estudio.

```
> burnout=c(6,7,9,7,0,10,5,2,4,6,25,6,12,7,6,4,13,18,17,24,21,15,
26,9,12,5,4,2)
> mean(burnout)
[1] 10.07143
> median(burnout)
[1] 7
> range(burnout)
[1] 0 26
> var(burnout)
[1] 53.03175
> sd(burnout)
[1] 7.28229
```

Teniendo un rango de entre 0 - 26 que en este caso se puede observar que existen niveles muy separados y que se encuentran en los niveles bajo y medio de acuerdo a los parámetros que indica la tabla de valores del test, es decir, es una muestra heterogénea.



Para la determinación del grado de burnout que la Administración Central presenta, hemos tomado en cuenta la median (7) que la representamos en este boxplot que al revisar la tabla de valores nos indica un grado muy bajo de agotamiento emocional en este grupo humano.



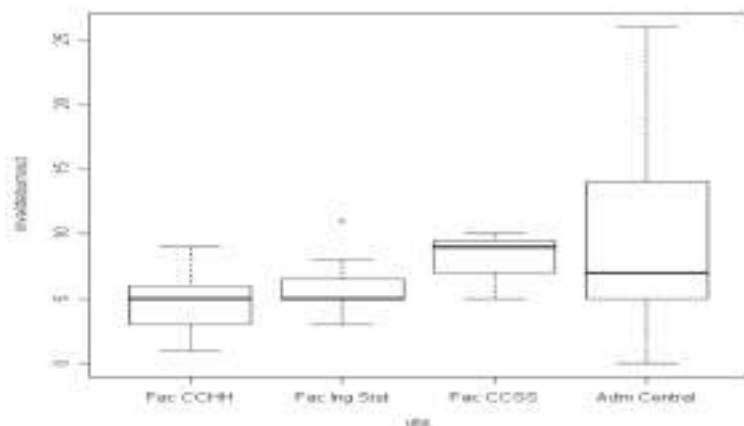
Este histograma nos da una clara visión del nivel generalizado del síndrome de burnout que presentan al nivel de los sectores administrativos, docentes y empleados de la Universidad con la aplicación del test Maslach Burnout Inventory como uno de los instrumentos mejor elaborados para este síndrome.

El nivel de la universidad viene a ser bajo casi muy insignificante para ser tomado en cuenta como una aparición del síndrome de burnout.

La pregunta que surge es ¿existen diferencias de burnout entre las facultades y departamento investigados? Y esto lo resolvemos con un anova. Pero primero observamos el plot así:

```
>niveldeburnout=c(1,3,3,3,6,3,6,6,5,5,8,5,9,3,4,2,5,8,6,4,6,5,7,
NA,NA,NA,NA,NA,5,8,5,5,5,6,3,11,5,7,4,NA,NA,NA,NA,NA,NA,
NA,NA,NA,NA,NA,NA,NA,NA,NA,NA,NA,NA,5,10,9,9,NA,NA,NA,N
A,NA,NA,NA,NA,NA,NA,NA,NA,NA,NA,NA,NA,NA,NA,NA,NA,NA,
NA,NA,NA,NA,6,7,9,7,0,10,5,2,4,6,25,6,12,7,6,4,13,18,17,24,21,1
5,26,9,12,5,4,2)
```

```
> uta=factor(rep(1:4,each=28),labels=c("Fac CCHH","Fac Ing
Sist","Fac CCSS","Adm Central"))
```



```
> burnout=data.frame(uta,niveldeburnout)
> plot(niveldeburnout~uta)
> burnout=aov(niveldeburnout~uta)
> summary(burnout)
```

	Df	Sum	Sq Mean Sq	F value	Pr(>F)
uta	3	372.29	124.10	4.8572	0.004242 **
Residuals	62	1584.07	25.55		
Signif. codes: 0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1					

La Facultad de la Salud tiene el más alto índice de burnout, para buscar el origen de la significación usamos la prueba de Tukey así:

```
burnoutTukey=TukeyHSD(burnout,"uta")
> burnoutTukey
Tukey multiple comparisons of means
 95% family-wise confidence level
Fit: aov(formula = niveldeburnout ~ uta)
$uta
```

	diff	lwr	upr	p adj
Fac Ing Sist-Fac CCHH	0.9051383	-3.9869167	5.797193	0.9613833
Fac CCSS-Fac CCHH	3.3369565	-3.8924076	10.566321	0.6175109
Adm Central-Fac CCHH	5.1583851	1.4030016	8.913769	0.0031866*
Fac CCSS-Fac Ing Sist	2.4318182	-5.3598675	10.223504	0.8428971
Adm Central-Fac Ing Sist	4.2532468	-0.4953914	9.001885	0.0947773
Adm Central-Fac CCSS	1.8214286	-5.3116681	8.954525	0.9064806

Al comparar las medias de agotamiento emocional entre (CCHH-IngSiste)(CCHH-CCSS)(IngSiste-CCSS)(IngSiste-AdCent)y(CCSS-AdCent) no presenta diferencia significativa, pero

al comparar las medias (CCHH-AdCent) si presenta diferencia significativa es decir que el nivel de burnout aunque ambas tienen niveles bajos la Facultad de Ciencias Humanas tiene un nivel mucho más bajo que la Administración Central que tiene un nivel mayor con relación a los puntajes de la Facultad de Ciencias Humanas.

CONCLUSIONES

Los empleados y profesores de la uta no presentan el síndrome de burnout

Existen diferencias entre las medias de burnout de la facultad de CH Y EE (4.91) y Administración Central (10.07), en las demás facultades no se encuentran diferencias significativas.

BIBLIOGRAFÍA

Burchell B & alters: Job insecurity and work intensification (2002) london

Casas i Hilari, M: Cuando querer no es poder: el síndrome del burnout Cyclops 2002 N°46

http://www.europa.eu.int/comm/employment_social/health_safety/publicat/stress_es.pdf

<http://www.ondasalud.com/edicion/noticia/o,2458,235850,00.html>

<http://www.ecofield.com.ar/notas/n-135.htm>

http://geosalud.com/Salud%20Ocupacional/estres_laboral.htm

http://agency.osha.eul.int/publications/magazine/5/es/MAGAZINES_ES.pdf

<http://www.agency.osha.eu.int/publications/reports/>

http://www.ucm.es/info/seas/estrés_lab.pdf

TEMA:

El uso de las NTIC'S y su influencia en el rendimiento académico de las estudiantes Universitarios

Curillo, M,; Morales, R

UNIVERSIDAD TÉCNICA DE AMBATO

RESUMEN

Las NTIC'S (Nuevas Tecnologías de la Información y la Comunicación). Consideradas, como un conjunto de procesos y productos derivados de las nuevas herramientas utilizadas en la enseñanza, cumplen un papel vital en el entorno educativo y de preparación del estudiante, que han hecho de ésta una forma de vida y estudio.

Con el propósito de conocer el uso de las Ntic's en los estudiantes universitarios y de determinar su incidencia en el rendimiento académico, se trabajó en las Facultades Ciencias Humanas, Ciencias Administrativas y de Contabilidad y Auditoría, con una muestra representativa de 148 estudiantes, con un margen de error del 8% ;, se utilizó un cuestionario de opción múltiple bajado del internet (<http://proyectoeducativo.iespana.es/Actividades/EncuestaTICs.html>). Los resultados muestran que existen diferencias significativas en el uso de las Ntic's entre facultades, al comparar F1 (Auditoría y Contabilidad) – F2 (Ciencias Humanas), se encuentran diferencias significativas con la prueba de Tukey al 5% siendo más alto el uso de las Ntic's en la Facultad de Contabilidad y Auditoría (24,66) contra (3.66) de la Facultad de Ciencias Humanas. Sin embargo, no se pudo detectar diferencias significativas en el rendimiento académico de los estudiantes, por el uso de las Ntic's.

INTRODUCCIÓN

En lo que se refiere al tema, existen investigaciones realizadas anteriormente, en los cuales se muestra diferentes enfoques generales sobre las NTIC'S y su impacto en el estudiante.

Sin embargo considerando que las NTIC'S cumplen un papel vital en el entorno educativo y de preparación del estudiante, existe material bibliográfico que trata por un lado la evolución de la tecnología en las aulas y por otro la utilización de las NTIC'S por parte de los estudiantes que han hecho de ésta una forma de vida y estudio.

Las tecnologías de la información y la comunicación (TICs) ejercen actualmente una influencia cada vez mayor en la educación científica, tanto en la enseñanza secundaria como en la universitaria, no sólo en lo que respecta a la mejora del aprendizaje de la ciencia por parte de los alumnos de tales niveles, sino que también desempeñan un papel creciente en la formación inicial y permanente del profesorado. Sobre esta temática hemos elaborado un trabajo de revisión, que por su extensión se ha desglosado en dos partes. En este primer se realiza un análisis panorámico de tales aplicaciones, abordando las posibles funciones educativas y los tipos de recursos informáticos que pueden utilizar los profesores de ciencias

Todos estos temas constituyen una panorámica suficientemente amplia como para propiciar el debate y la reflexión entre los profesionales de la enseñanza en los comienzos del siglo XXI, que será probablemente un periodo de grandes cambios en la educación a consecuencia de la incorporación de las TICs al mundo de la enseñanza. En este trabajo no pretendemos dar respuestas exhaustivas a todas estas cuestiones pero podemos realizar una revisión de esta problemática y apuntar algunas ideas o sugerencias que permitan seguir avanzado en el desarrollo de la informática educativa aplicada a la enseñanza de las ciencias.

Objetivos educativos	Funciones a desarrollar
Conceptuales	- Facilitar el acceso a la información - Favorecer el aprendizaje de conceptos
Procedimentales	- Aprender procedimientos científicos - Desarrollar destrezas intelectuales
Actitudinales	- Motivación y desarrollo de actitudes favorables al aprendizaje de la ciencia

Tabla 1.- Fines y funciones de las TICs en la formación de estudiantes.

Por último hay que indicar que el uso educativo de las TICs fomenta el desarrollo de actitudes favorables al aprendizaje de la ciencia y la tecnología. Como han puesto de manifiesto diversos trabajos sobre el tema (Jegede, 1991; Yalcinalp et al., 1995; Escalada y Zollman, 1997), el uso de programas interactivos y la búsqueda de información científica en Internet ayuda a fomentar la actividad de los alumnos durante el proceso educativo, favoreciendo el intercambio de ideas, la motivación y el interés de los alumnos por el aprendizaje de las ciencias. Muchos alumnos también participan en foros de debate sobre temas científicos o llegan a elaborar sus propias páginas webs y pequeños programas de simulación.

En un informe del programa EURYDICE de la Comisión Europea para el desarrollo de la Educación y la Cultura (Pépin, 2001), sobre los indicadores básicos que describen la incorporación de las TICs en los sistemas educativos de los diversos países europeos, se resaltan estas funciones de las TICs en la formación del profesorado, aunque se advierte que en muchos casos los profesores han ido adquiriendo formación docente sobre el uso de las TICs de forma autónoma, por interés personal o por la necesidad de ponerse al día en estos temas, ya que no existe una planificación general en todos los países sobre la forma adecuada en que debe llevarse a cabo la formación inicial y permanente del profesorado en relación al uso docente de las TICs.

OBJETIVOS:

General:

Determinar el uso de las NTIC'S y su influencia en el rendimiento académico de los estudiantes de los séptimos y octavos semestres de las Facultades de: Ciencias de la Educación: las carreras de Docencia en Informática e Idiomas, Ciencias Administrativas: Carreras: Economía, Administración de Empresas, Facultad de Jurisprudencia y Ciencias Sociales: Carreras: Comunicación Social y Derecho durante el período 2007 – 2008.

Específicos:

- a) Conocer el nivel de rendimiento académico de los estudiantes universitarios.
- b) Determinar el uso de los implementos y herramientas para la aplicación de las NTIC'S en el aula.
- c) Relacionar las NTIC'S con el rendimiento académico de los estudiantes.

METODOLOGÍA

LUGAR

Con el fin de determinar el uso de las NTIC'S y su influencia en el rendimiento académico de los estudiantes desarrollamos el proceso de investigación en la F.C.H.E., carreras de Cultura física y Educación Básica, ubicado en el sector Ingahurco, Av Colombia y Chile; de igual manera en la Facultad de Ciencias Administrativas en las carreras de Mercadotecnia, Administración de Empresas, y finalmente en la Facultad de Contabilidad y Auditoria en las carreras de Contabilidad y Economía las dos últimas ubicadas en Huachi Chico del cantón Ambato, provincia de Tungurahua.

POBLACIÓN Y MUESTRA

La población estudiada estuvo conformada de la siguiente manera:

FACULTAD DE CONTABILIDAD Y AUDITORIA

Auditoria =74

Economía=13

FACULTAD DE CIENCIAS ADMINISTRATIVAS

Org. Empresas=34

Mercadotecnia=11

FACULTAD DE CIENCIAS HUMANAS Y DE LA EDUCACIÓN

Cul. Física =5

Educ. Básica = 11

TOTAL= 148

NIVEL DE ERROR=8%

TÉCNICAS E INSTRUMENTOS DE MEDICIÓN

Con el propósito de llevar a cabo la investigación y obtener resultados se utilizó y aplicó técnicas e instrumentos de medición tales como: ficha de observación y encuestas como recursos básicos para el cumplimiento de nuestro objetivo.

PROCEDIMIENTO

Para el cumplimiento del objetivo específico, **Determinar el nivel de rendimiento académico de los estudiantes universitarios.** Concurrimos a las secretarías de cada una de las carreras y obtuvimos los promedios globales de los alumnos, permitiéndonos así determinar el rendimiento académico de los estudiantes de cada una de las carreras.

Se utilizó una encuesta bajada del internet

(<http://proyectoeducativo.iespana.es/Actividades/Encuesta-TICs.html>), con el fin de **Determinar el uso de los implementos y herramientas para la aplicación de las NTIC'S en el aula.**

Con el objeto de **Relacionar las NTIC'S con el rendimiento académico de los estudiantes**. Se tomó los resultados obtenidos mediante la media aritmética del rendimiento de los estudiantes y los datos tabulados, ordenados y categorizados que se obtuvo en la encuesta, a más de eso nos basamos en los datos de la ficha de observación con todos estos instrumentos llegamos al cumplimiento de nuestro objetivo.

RESULTADOS

Resultado del Uso de las NTICS

FACULTAD DE CONTABILIDAD Y AUDITORÍA

AUDITORIA

ALTO	MEDIO	BAJO
52	14	8
45	5	24
39	17	18
14	28	32
50	0	24
20	12	42
43	12	19
47	9	18
38,75	12,125	23,125

TOTAL 24,6666667

ECONOMÍA

ALTO	MEDIO	BAJO
7	4	2
6	3	4
7	4	2
1	8	4
9	0	4

9	3	1
8	2	3
8	3	2
6,875	3,375	2,75

TOTAL 4,33333333

MEDIA DE FACULTAD 1 14,5



FACULTAD DE CIENCIAS ADMINISTRATIVAS

ORGANIZACIÓN DE EMPRESAS

ALTO	MEDIO	BAJO
11	19	4
21	10	3
25	5	4
7	23	4
7	0	27
9	21	4
25	7	2
5	6	23
13,75	11,375	8,875

TOTAL 11,33333333

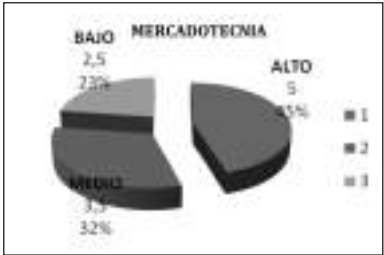
MERCADOTECNIA

ALTO	MEDIO	BAJO
7	3	1
4	5	2
7	3	1
6	4	1

2	0	9
4	5	2
8	2	1
2	6	3
5	3,5	2,5

TOTAL 3,66666667

MEDIA FACULTAD 2 7,5



FACULTAD DE CIENCIAS HUMANAS
Y DE LA EDUCACIÓN

EDUCACION BÁSICA

ALTO	MEDIO	BAJO
1	6	4
1	7	3
2	5	4
6	4	1
4	0	7
5	4	2
8	2	1
2	6	3
3,625	4,25	3,125

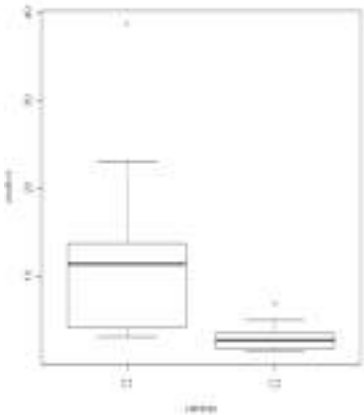
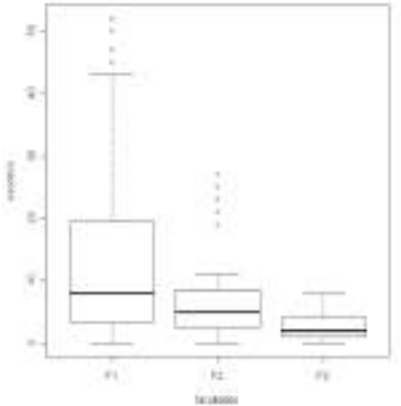
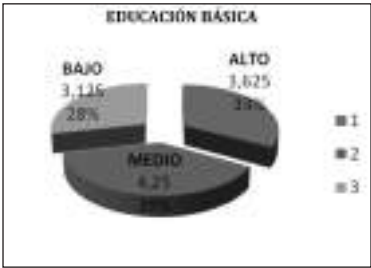
TOTAL 3,66666667

CULTURA FÍSICA

ALTO	MEDIO	BAJO
1	3	1
3	1	1
2	2	1
1	1	3
2	0	3
1	4	0
4	1	0
0	2	3
1,75	1,75	1,5

TOTAL 1,66666667

MEDIA FACULTAD 32,66666667

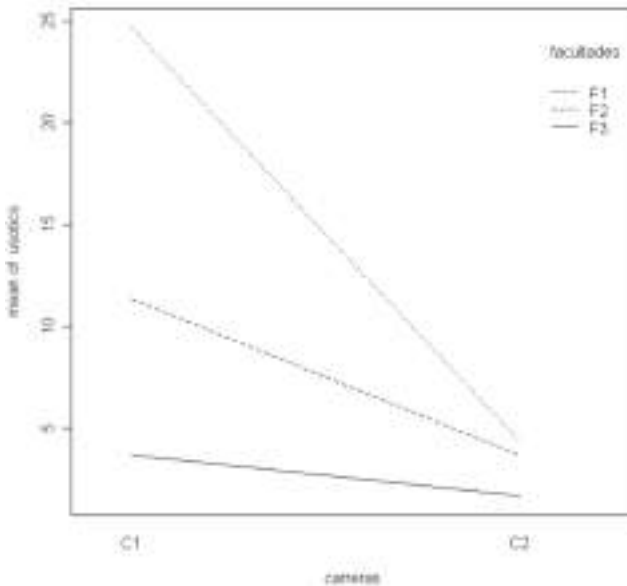


Con el comando tapply en “R” se pudo establecer los siguientes resultados:

Los valores presentados muestran las medias obtenidas en cada carrera y facultad en cuanto al uso de las Ntic’s

	F1	F2	F3
C1	24.66333	11.330000	3.663333
C2	4.33000	3.666667	1.666667

Mediante un interaction.plot se establece que la facultad de auditoria tiene un nivel más alto que las demás facultades:



CONTRASTE DE HIPÓTESIS: ANOVA:

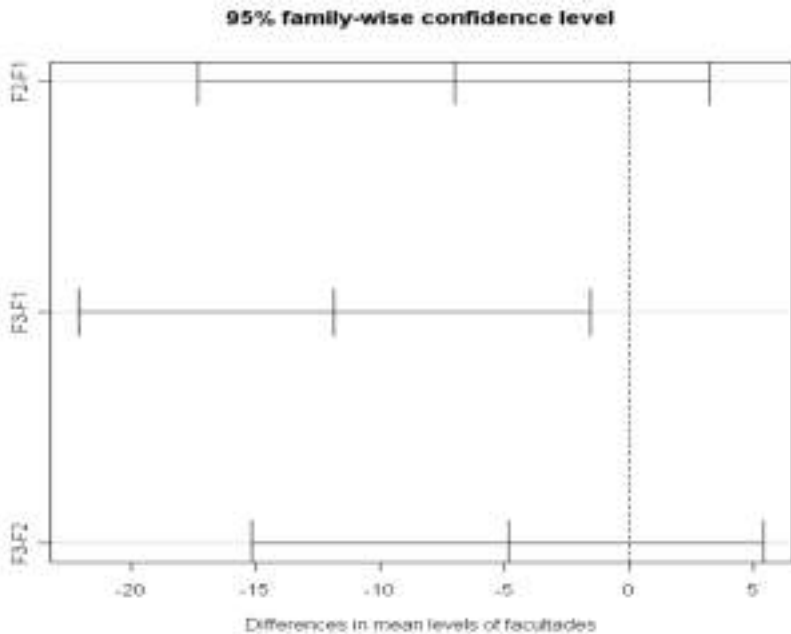
	Df	Sum	Sq Mean	Sq F value	Pr(>F)
carreras	1	449.80	449.80	9.7147	0.007575 **
facultades	2	424.65	212.33	4.5858	0.029391 *
Residuals	14	648.22	46.30		

Existen diferencias significativas entre Facultades y también entre carreras en cuanto al uso de las Tic's.

Buscamos el origen de las significaciones con la prueba de Tukey así:

PRUEBA DE TUKEY

	diff	lwr	upr	p adj
F2-F1	-6.998333	-17.28053	3.283860	0.2115261
F3-F1	-11.831667	-22.11386	-1.549474	0.0237765
F3-F2	-4.833333	-15.11553	5.448860	0.4556564



La facultad 1 con la facultad 3 denotan diferencias en cuanto al uso de las Ntic's.

RENDIMIENTO ACADÉMICO:

Se establecieron los 3 grupos:

G1= Alto

G2=Medio

G3=Bajo

Se obtuvo los siguientes resultados en medias de:

G1

8.0075

G2

8.11450

G3

8.0075

Se obtuvo los siguientes resultados en desviaciones en los tres grupos de:

G1

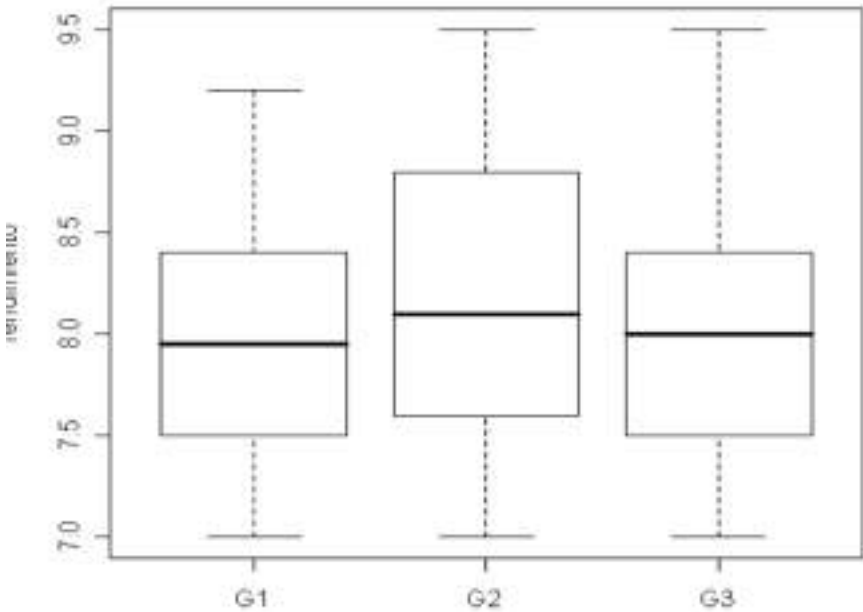
0.6513841

G2

0.7056366

G3

0.6149369



**MEDIAS OBTENIDAS EN EL RENDIMIENTO
ACADÉMICO:**

Se observan que aparentemente no existen diferencias entre las medias de rendimientos comparadas. Los cual se ratifica con el siguiente anova.

	Df	Sum	Sq Mean	Sq F value	Pr(>F)
Grupos	2	0.504	0.252	0.5816	0.5606

No se encuentran diferencias significativas en el rendimiento de los grupos estudiados.

CONCLUSIONES:

De los resultados obtenidos se deduce que si existe una diferencia significativa en la facultad de Auditoria y Contabilidad, es decir que en esta facultad es mayor el uso de las NTICS. De la misma manera existen diferencias entre las Carreras.

Al comparar las medias obtenidas entre facultades, pues existió diferencias, mediante el método de Tukey se deduce que:

No se encuentran diferencias en las medias entre la **F1** (Auditoria y Contabilidad) y **F2** (Administración), así mismo entre la Facultad 2 con la Facultad 3.

Existen diferencias significativas entre la Facultad 1 (Auditoria y Contabilidad) y la Facultad 3 (Ciencias de la Educación).

No se encuentran diferencias significativas en los rendimientos de los estudiantes, es decir que al comparar los niveles de NTICS obtenidos anteriormente, con el rendimiento, se concluye que no influye el uso y conocimiento de las NTIC'S en el rendimiento académico.

Sin embargo es necesario destacar que los estudiantes que tienen un nivel medio en conocimiento y NTIC'S, son las personas que tiene un promedio un poco mayor que las anteriores,

RECOMENDACIONES

Implementar en la malla curricular la asignatura como tal, con el fin de que los estudiantes conozcan y relacionen las NTICS con su carrera y puedan hacer uso de éste en bien de su desarrollo profesional.

Se recomienda implementar herramientas actualizadas y que

estén acorde con la tecnología actual y al alcance de los estudiantes.

Brindar charlas y conferencias que permitan a los estudiantes conocer sobre las NTICS, dotándoles de libros, textos o folletos que les permitan auto educarse.

REFERENCIAS

<http://www.scribd.com/doc/95709/INFORME-NTICS>

<http://educacionyntics.ning.com/>

<http://pad-ntics.blogspot.com/>

<http://ntics-nuevastecnologias.blogspot.com/>

<http://www.jonhy.com/aula/>

TEMA:

Niveles de autoestima en estudiantes universitarios

Un estudio del grado de influencia del género y la carrera en estudiantes del Seminario de graduación de la Facultad de Ciencias Humanas de la UTA.

La autoestima es un elemento importante en el desarrollo interno del ser humano ya que determina como nos vemos y sentimos, como la imagen que proyectamos a los demás. Nos ayuda a ver qué esta mal en qué estamos fallando y por ende corregir la falla. Las personas de baja autoestima es decir, que tienen deteriorada la imagen de si mismos no poseen una imagen clara, son poco o nada exitosas se autocritican en exceso piensan que son una carga para los demás, van a donde otros quieren que él vaya, no tienen un norte fijado, de la misma manera las personas con alta autoestima se creen superiores a los demás, creen que siempre tienen la razón.

Así encontramos un punto de inicio hacia la lucha contra este mal social que afecta no sólo a los estudiantes sino a nivel general.

PROPÓSITO

Determinar el nivel de autoestima de los estudiantes de las diferentes carreras del seminario de graduación de la Facultad de Ciencias Humanas.

Comprobar si existen diferencias de autoestima entre hombres y mujeres

METODOLOGÍA

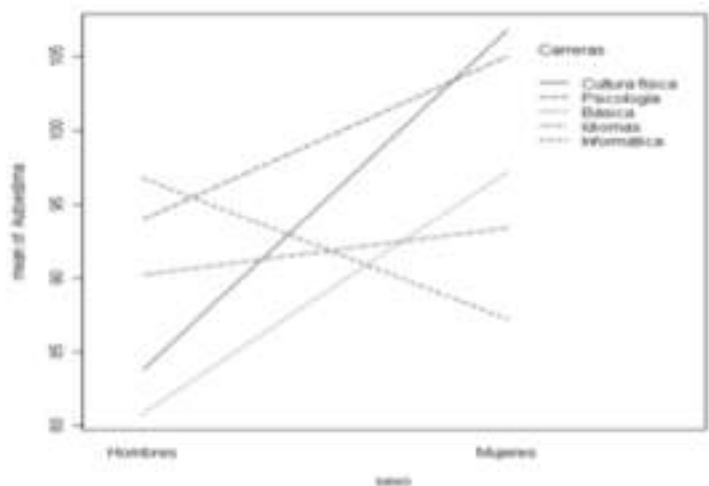
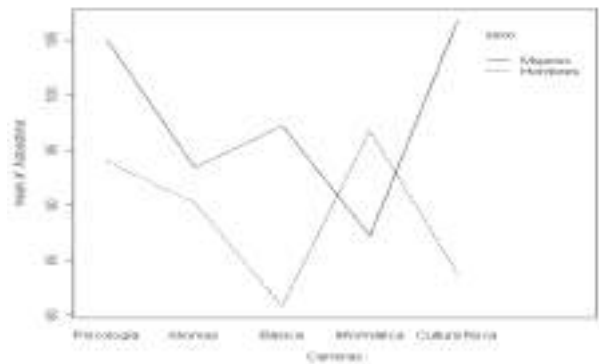
La muestra esta integrada por 50 estudiantes de las carreras de Informática, Idiomas, Educación Básica, Cultura Física, y Psicología.

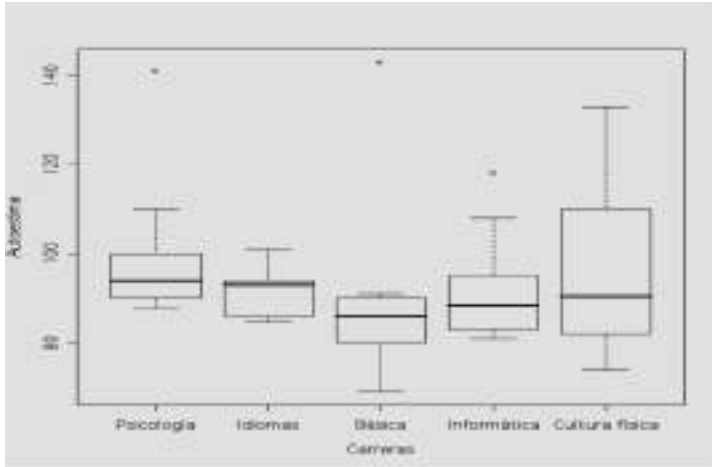
Para obtener la información necesaria se utilizó el test de Rosenberg que tiene las siguientes valoraciones: autoestima baja negativa de 40-73, autoestima baja positiva de 74-83, autoestima alta positiva de 84-103, y autoestima alta negativa 104-160.

Se utilizó para el análisis el software estadístico R.

RESULTADOS

AUTOESTIMA ENTRE SEXO Y CARRERAS





Al observar los gráficos podemos decir que la autoestima de las mujeres es más alta que la de los hombres en la mayoría de las carreras investigadas y que en la carrera de Educación Básica la autoestima es más baja. Es decir, tenemos que resolver las sts. Hipótesis:

- 1.- H1: autoestima de mujeres > que autoestima de hombres
Ho: autoestima de mujeres = autoestima de hombres
- 2.- H1: la autoestima de las diferentes carreras es diferente
Ho: la autoestima de las diferentes carreras es la misma

Esto lo comprobamos con un análisis de varianza

ANOVA AUTOESTIMA SEXO Y CARRERAS

```
> a=aov(Autoestima~sexo+carreras)
```

```
> summary(a)
```

	Df	Sum	Sq Mean	Sq F value	Pr(>F)
sexo	1	968.0	968.0	4.5127	0.0393 *
Carreras	4	646.3	161.6	0.7532	0.5613
Residuals	44	9438.2	214.5		
Signif codes: 0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1					

CONCLUSIONES

La autoestima de la mayoría de los investigados esta entre baja positiva y alta positiva (74-103)

Se encuentran diferencias significativas en la autoestima entre hombres y mujeres en la mayoría de carreras estudiadas siendo más alta la autoestima en las mujeres excepto en Informática, en donde la autoestima de las mujeres es más baja (84) contra la de los hombres (93).

No se encuentran diferencias significativas entre carreras.

REFERENCIAS:

http://www.vida7.cl/articulos/psicologia/archivos/200703/psi_.htm

http://www.geomundos.com/mujeres/gema/la_autoestima_doc_1861_html

http://www.renuevodeplenitd.com/excelencia_en_la_autoestima.html

TEMA:

Religión que profesan los estudiantes de la UTA y su relación con el rendimiento académico

Acosta, G. ¹ y Morales, R. ²

^{1,2} Universidad Técnica de Ambato

RESUMEN

Varias Religiones se han creado y otras han desaparecido a lo largo de la historia, sus seguidores han adoptado nuevos ritmos de vida, los cuales van ligados directamente a la religión que profesan. Se ha observado conductas que marcan la diferencia a la hora de reconocer a una persona de diferente religión, y siendo estas personas parte de la sociedad, y miembros de instituciones educativas de tercer nivel o universitarios, nos hemos propuesto investigar a que religiones pertenecen los estudiantes de la UTA y la actitud de los estudiantes de diferentes religiones con respecto a su desenvolvimiento académico, pues desde niveles inferiores hemos podido observar que gran parte de los estudiantes que pertenecen a una religión ajena a la Católica, presentan un mayor nivel de rendimiento académico que los que no pertenecen a esta.

Con este propósito se trabajó en las facultades de: Ingeniería Civil, Contabilidad y Auditoría, Ingeniería en Sistemas con una muestra representativa de 97 estudiantes. Para obtener información acerca de la religiosidad de los investigados se procedió a adaptar por el grupo investigador un test de varios bajados de la Internet:
(<http://www.fungamma.org/religion1.htm>
<http://www.google.com.ec/translate?hl=es&sl=en&u>
http://www.gotoquiz.com/the_religious_identity_test&ei=WwIn

SsPPE5ik8AToxMyBDw&sa=X&oi=translate&resnum=6&ct=result&prev=/search%3Fq%3Dtest%252Breligion%26hl%3Des). Además se obtuvo los rendimientos académicos de los estudiantes en el UTAMÁTICO.

Los resultados obtenidos ratifican que la mayoría de los estudiantes (73.33%) son Católicos, seguidos de un 20% que tienen otras religiones (Evangélicos, Testigos de Jehová, Mormones, etc) y un 6.67% que no profesan ninguna religión. Las diferencias de religiosidad entre facultades son mínimas, y entre hombre y mujeres se nota que las mujeres son más religiosas que los hombres. Al comparar los grupos religiosos (Grupo 1: Católico, Grupo 2: Ninguna religión, Grupo 3: Otras religiones) con el rendimiento académico se encontró con un Anova que existen diferencias significativas. Al comparar los grupos 1-2, 1-3, y 2-3, con la prueba de Tukey, se encuentran diferencias significativas y se estableció que los rendimientos más bajos están en el segundo grupo y el más alto rendimiento en el grupo 3.

INTRODUCCIÓN

La investigación constituye un elemento fundamental e imprescindible para lograr un proceso de enseñanza – aprendizaje de excelencia, motivo por el cual es necesario conocer todos los factores que diferencian a quienes son parte activa de la institución educativa universitaria, y por supuesto uno de ellos es la RELIGIÓN, sobre la cual se ha realizado el presente proyecto, pretendiendo conocer los aspectos de esta que influyen en el comportamiento humano.

Religión (del latín *religio*: cuidado, en el sentido de prácticas de culto) a veces usada como sinónimo de fe o sistema de creencias, se define comúnmente como creencia.

Religiones

Hay diferentes clasificaciones de las religiones, por ejemplo:

Por concepción teológica

Teísmo: es la creencia en una o más deidades. Dentro del teísmo existe el monoteísmo, panteísmo, cosmoteísmo y el deísmo.

Monoteístas: aquellas religiones que afirman la existencia de un solo Dios, que a menudo es creador del universo. Las religiones monoteístas más numerosas son el cristianismo y el islam. Otras más minoritarias son el judaísmo o la bahai.

Henoteístas: las que reconocen la existencia de varios dioses, pero sólo uno de ellos es suficientemente digno de adoración. El zoroastrismo, algunas escuelas hindúas y los primitivos hebreos son considerados henoteístas.

Politeístas: creen en la existencia de diversos dioses organizados en una jerarquía o panteón, como ocurre en el hinduismo, el shinto japonés, o las antiguas religiones de la humanidad como la griega, la romana o la egipcia.

No teístas: Las religiones no-teístas como el budismo y el taoísmo no mantienen la existencia de dioses absolutos o creadores universales. En ocasiones, existen deidades que son vistas como recursos metafóricos utilizados para referirse a fenómenos naturales o estados de la mente.

Por revelación

Otra división que se utiliza consiste en hablar de religiones reveladas o no reveladas.

Las religiones reveladas se fundamentan en una verdad revelada de carácter sobrenatural desde una deidad o ámbito trascendente y que indica a menudo cuáles son los dogmas en los que se debe creer y las normas y ritos que se deben seguir.

Las religiones no reveladas no definen su origen según un mensaje

dado por deidades o mensajeros de ellas, aunque pueden contener sistemas elaborados de organización de deidades reconociendo la existencia de éstas deidades y espíritus en las manifestaciones de la naturaleza.

Por origen

Otra clasificación de las religiones es por origen o familia. Las religiones se agrupan en troncos de donde derivan, por ejemplo: Usualmente se acepta que las principales familias de religiones son las siguientes:

Familia de religiones abrahámicas o semíticas.

Familia de religiones dhármicas o índicas.

Familia de religiones iranianas.

Familia de religiones neopaganas.

Familia de religiones tradicionales.^[cita requerida]

Sectas o Nuevos Movimientos Religiosos

Algunas religiones, de reciente creación, tienen un estatus complejo ya que no son reconocidas como religiones de manera universal. Una **secta** o Nuevo Movimiento Religioso, según la antropología y la sociología, es, desde el punto de vista sociológico, un grupo de personas con afinidades comunes: culturales, religiosas, políticas, esotéricas, etc. Habitualmente es un término peyorativo, frente al que ha surgido el eufemismo «nuevos movimientos religiosos».

Aunque el vocablo «secta» esté relacionado con grupos que posean una misma afinidad, con el paso de los años ha adquirido una connotación relacionada con grupos de carácter religioso, a los que se califica como «secta destructiva». Estos grupos pueden tener un historial judicial en uno o varios países, por manipulación mental o por ser grupos de carácter destructivo. En algunos países, algunas de estas no están reconocidas o autorizadas. A menudo una secta está centrada en el culto personal al profeta o líder, del grupo. La palabra secta se ha concebido derivada, principalmente, del latín *sequi*: ‘seguir’.

Las religiones en cifras

La mayoría de las diversas religiones gozan de buena salud en número de seguidores y su número ha aumentado en todo el mundo. En los países con anteriores regímenes comunistas, la religión se ha revitalizado a una velocidad sorprendente como muestran los casos de Rusia y China.

Ganesh, la popular deidad hindú destructora de obstáculos y patrón de las artes, las ciencias y la sabiduría.

Estadística de religiones en el mundo (Estos datos son antiguos, aproximativos, provisionales y discutidos)

Cristianismo:	2000 millones
Islam:	1500 millones
Secularismo/agnosticismo/ateísmo:	1100 millones
Hinduismo:	900 millones
Budismo:	entre los 1691 millones y los 230 millones.
Religión tradicional china:	394 millones
Religiones indígenas:	300 millones
Religiones afroamericanas:	100 millones
Sijismo:	23 millones
Mormonismo:	20 millones
Espiritismo:	15 millones
Judaísmo:	14 millones
Baha'i:	7 millones
Testigos de Jehová:	7,1 millones
Jainismo:	4 millones
Shintoísmo:	4 millones
Caodaísmo:	4 millones
Zoroastrismo:	2 millones
Tenrikyo:	2 millones
Neopaganismo:	1 millón
Unitarismo/universalismo:	0,8 millones
Rastafarianismo:	0,6 millones
Otras:	0,4 millones

Religiones en el mundo

Lista de las principales religiones actualmente practicadas en el mundo, por orden alfabético.

Budismo: fundada por Siddhrtha Gautama (Buda Gautama o El Buda) en el siglo VI a. C. Actualmente extendida por todo el mundo a excepción de la mayoría de países africanos.

- **Theravada:** rama más antigua del budismo surgida alrededor de la primera compilación budista escrita. Asentada originalmente en India y el Sudeste asiático
- **Mahayana:** movimiento de reforma surgido en el siglo I. Es el más numeroso actualmente. Asentada originalmente en China, Japón y el Sudeste asiático.
- **Vashraiana:** parte del mahayana pero definido propiamente por su influencia del tantrismo hindú. Asentada originalmente en la región de los Himalayas, Kalmukia, Japón y Mongolia.

Confucianismo: sistema ético y moral que rige la sociedad china. No es propiamente una religión, si bien esta denominación es discutida.

Cristianismo: centrada en la figura de Jesús de Nazaret (siglo I). Presente en casi todo el mundo, excepto el norte de África y gran parte de Asia (presente en Rusia, antiguos países soviéticos asiáticos y Filipinas)

- **Iglesia católica:** iglesia proveniente del cristianismo en Europa Occidental. Principalmente en América Latina, buena parte de Europa occidental (excepto Reino Unido, norte de Alemania y países nórdicos), Filipinas y Guinea Ecuatorial
- **Iglesia ortodoxa:** iglesia proveniente del cristianismo en Europa Oriental y Asia Menor. Está presente principalmente en Rusia, Grecia y varios países de la Europa del Este, y actualmente se ha expandido alrededor del mundo principalmente gracias a emigrantes de esos territorios
- **Iglesia copta:** surgida en el siglo II d. C. En origen son los cristianos nativos de Egipto (coptos), de teología no calcedoniana. Principalmente en Egipto
- **Iglesias orientales católicas:** agrupa a 22 iglesias que aceptan la autoridad del papa católico romano pero mantienen ritos independientes.
- **Anglicanos** (episcopalianos): surgida por la escisión creada por Enrique VIII (1491-1547) de la iglesia católica romana.
- **Protestantismo:** conjunto de iglesias cristianas aparecidas desde el

siglo XVI tras la reforma de Martín Lutero y que pretendían volver a los fundamentos de la iglesia cristiana del siglo I. Actualmente en países desarrollados, Nigeria, Brasil, República Democrática del Congo, Kenia, Indonesia y Sudáfrica.

- ✚ **Luteranismo:** fundado por Martín Lutero (1483-1546) rechazando la autoridad del papa católico romano.
- ✚ **Baptista:** surgido en el siglo XVII desde el protestantismo.
- ✚ **Evangelismo:** agrupa diferentes iglesias cristianas protestantes.
- ✚ **Metodismo:** movimiento surgido desde el protestantismo en Gran Bretaña, en el siglo XVIII. Extendida por EE. UU..
- ✚ **Pentecostalismo:** movimiento impulsado en 1901 por Charles Fox Parham, predicador metodista de EE. UU.
- ✚ **Cristianos reformados:** profesan el espíritu de Calvino (1509-1564). Actualmente agrupa a numerosas iglesias protestantes de Australia y EE. UU.
- ✚ **Cuáqueros:** movimiento protestante fundado en el siglo XVII en Inglaterra, rechaza la jerarquización del protestantismo y se centra en la «luz interior» o chispa divina en cada ser humano.
- ✚ **Unitarios:** nace a partir del pensamiento desarrollado principalmente por Miguel Servet y Fausto Socino en el siglo XVI, niega la Santísima Trinidad y afirma el uso de la razón en la religión.
- ✚ **Universalistas:** surge del metodismo inglés aunque arraiga principalmente en EE. UU., afirma la salvación universal y la inexistencia del infierno.
- ✚ **Iglesia Unificada de Cristo:** formada en 1957, agrupa a iglesias reformadas, evangélicas y congregacionales de EE. UU.
- ✚ **Adventistas:** familia de iglesias protestantes de carácter conservador o literalista, la mayoría originadas en EE. UU.
 - ✓ **Adventistas cristianos:** fundada en 1860.
 - ✓ **Davidianos:** fundada en el siglo XX.
 - ✓ **Cristadelfianos:** fundada en 1844, son evangélicos de teología unitarista.
 - ✓ **Conferencia General de Dios:** fundada en 1921.
 - ✓ **Iglesia Adventista del Séptimo Día:** fundada en 1863.
 - ✓ **Iglesia de Dios y los Santos de Cristo:** fundada en 1896.
 - ✓ **Adventistas del Séptimo Día:** fundada en 1845.
- ✚ **Testigos de Jehová:** fundada en 1870 y conocidos como «los estudiantes de la Biblia» hasta 1931. Presentes en 236 países.
- ✚ **Espiritismo:** Fundado en Francia en 1857. Escuela científico-fi-

losófica, y religiosa. Basado en la codificación de Allan Kardec.

✚ **Iglesia Mundial de Dios:** fundada en 1933.

✚ **Mormonismo:** fundada el 6 de abril de 1830 por José Smith hijo. Su nombre oficial es: «La Iglesia de Jesucristo de los Santos de los Últimos Días»

Bahaísmo: fundada por Bahá'u'lláh (1817-1892), considerado por sus partidarios como el prometido de todas las religiones. Su enseñanza central es la unidad de la humanidad.

Hinduismo: originada en India. Agrupa distintas creencias alrededor de las Escrituras védicas (aprox. del siglo X a. C., la cultura de textos y religión de la India.

- **Shivaísmo:** se centra en el dios Shiva; sus seguidores se llaman shivaístas. El texto más antiguo es del siglo V a. C. aprox.
- **Vaishnavismo:** se centra en la deidad Vishnú.
- **Advaita Vedanta:** basada en la doctrina *vedanta* y el *prasthan trayi* (los tres textos canónicos de la doctrinas hinduistas).
- **Australianas:** practicadas por los aborígenes de Australia, suelen usar la interpretación de sueños.
- **Americanas:** realizan un culto a la naturaleza y pueden utilizar plantas y elementos psicoactivos como el peyote.
 - ✚ **Andinas:** recogen elementos de la mitología inca y de otras antiguas, realizando un sincretismo chamánico.
 - ✚ **Mexicanas:** recogen elementos de la mitología azteca y maya realizando un sincretismo chamánico.
- **Africanas:** agrupan multitud de creencias transmitidas oralmente.
 - ✚ **Yoruba** (yorubá): de ella se derivan multitud de sincretismos en toda América.
 - ✚ **Vudú:** originada en África Occidental y asentada en el Caribe y sur de EE. UU.
 - ✚ **Santería:** originada desde un sincretismo entre el animismo y las creencias cristianas.
 - ✚ **Candomblé:** de origen totémico, es un sincretismo de religiones afrobrasileñas.
 - ✚ **Kimbanda:** originada en Brasil por el sincretismo del cristianismo con religiones africanas y creencias cristianas.
 - ✚ **Umbanda:** originada desde un sincretismo entre candomblé, el kardecismo espiritualista y las creencias cristianas.
- **Asia:** que incluyen los cultos animistas y chamánicos de: Bön: religión tradicional de Tíbet.

- ✦ **Chamanismo:** extendido por toda Asia en poblaciones tribales.
- ✦ **Chondogyo** de Corea.
- ✦ **La religión tradicional china.**

Islam: basada en las enseñanzas del *Corán*, transmitido por el profeta Mahoma (nacido en el 570 d. C.).

- **Chiismo** (*shii*): siguen el Ahl al-Bayt o autoridad de la familia de Muhammad y sus descendientes. Es la segunda afiliación más grande al islam.
- **Sunismo** (*sunni*): a diferencia de los chiíes, los suníes aceptan el califato de Abu Bakr (573-634). Es la rama más grande del islam.
- **Sufismo:** el sufismo no es propiamente una rama del islam sino una tradición mística que aparece tanto con seguidores chiíes como suníes.

Jainismo: fundado en la India en el siglo VI a. C. por Mahavira.

Judaísmo: basado en las enseñanzas de la *Torá*. Principalmente en Israel, pero después de la diáspora están extendidos en el mundo.

- **Conservador:** llamado *Maserti*. Señalan la importancia del movimiento sionista en el judaísmo.
- **Secular:** el judaísmo secular es aquel que se ve independiente de organizaciones.
- **Ortodoxo:** llamado *Haredi*. Es la línea teológica más conservadora del judaísmo.

Shinto: religión nativa de Japón, en su origen chamánica y animista. Es seguida por muchos japoneses.

Sijismo: fundada por Gurú Nanak en el siglo XV en la India, en la región del Panyab.

Mandeísmo: una religión muy antigua que parece ser descendiente del *gnosticismo* antiguo y rinde culto a Juan el Bautista. Probablemente son los sabeos mencionados en el *Corán*. Cuenta con 38.000 seguidores, casi todos en Iraq.

Samaritanismo: una rama disidente del judaísmo, muy antigua, con sede en Samaria (Israel), que es pretalmúdica y de hecho, no reconoce al Talmud.

Taoísmo: conjunto de enseñanzas filosóficas y religiosas originadas en China partir de Lao-Tse (Laozi) en el siglo VI a. C.

Yazidismo: una religión autóctona de Kurdistán de influencias islámicas y zoroástricas seguida por alrededor de 200.000 kurdos. Profesan culto a los ángeles y arcángeles de las religiones abrahámicas dándoles una explicación propia.

Zoroastrismo: de orígenes inciertos, aparece como religión alrededor del siglo V a. C. Sus enseñanzas se basan en el profeta y poeta Zoroastro del antiguo Imperio Persa.

OBJETIVOS

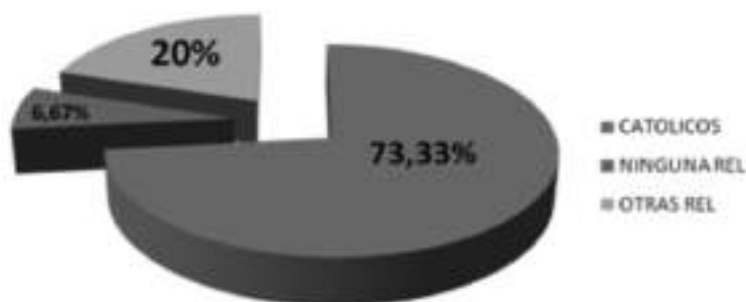
- Determinar si la religión influye en el rendimiento académico universitario, favoreciéndolo notablemente
- Determinar porcentajes y grupos religiosos dentro de la Universidad Técnica de Ambato
- Determinar diferencias existentes entre las relaciones: religión – sexo, religión – rendimiento.

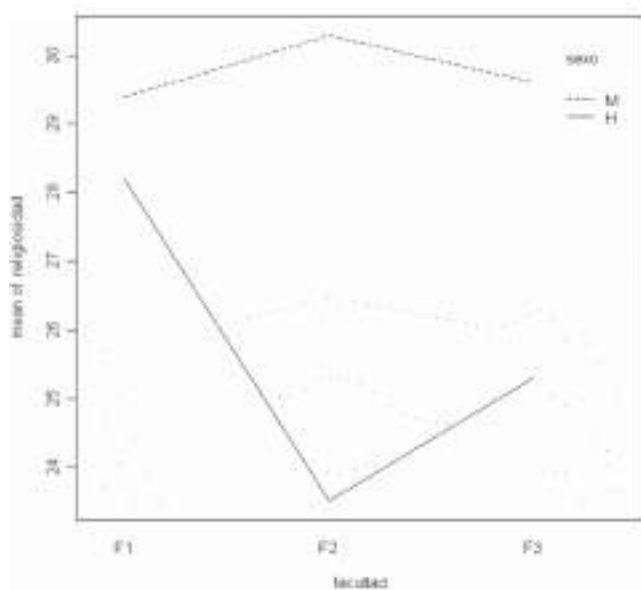
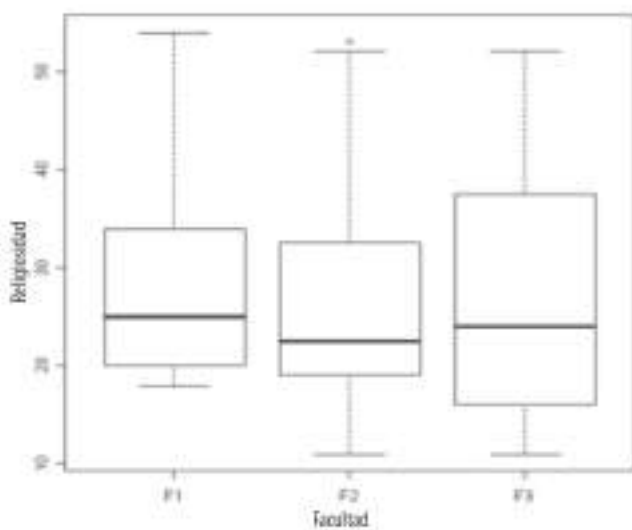
METODOLOGÍA:

Para obtener información a cerca de la religiosidad de los investigados, se procedió a la aplicación de un test de conocimiento religioso, el cual contenía preguntas estratégicas de medición de religiosidad y obtención del culto profesado, existía un punto el cual separaba grupos religiosos de no religiosos por medio de puntuaciones que fluctuaban entre: <30 no religiosos, >29 religiosos, con una pregunta directa sobre el credo profesado que contenía varios grupos a escoger

Continuando con la investigación se procedió a medir el rendimiento académico de los estudiantes anteriormente investigados siguiendo este procedimiento: Solicitar el número de cédula de cada estudiante, pudiendo así verificar sus calificaciones en el sistema de consignación de notas UTAMÁTICO, obteniendo el promedio respectivo por estudiante.

RESULTADOS





ANOVA

Rendimiento estudiantes - uta

	Sum of Squares	df	Mean Square	f	Sig
Between Groups	44,365	2	22,183	33,155	.000
Within Groups	33,136	57	,669		
Total	82,502	59			

Descriptivos

Rendimiento estudiantes - uta

	N	Mean	Std Deviation	Std Error	95% Confidence Interval for Mean		Minimum	Maximum
					Lower Bound	Upper Bound		
1,00	44	7,4773	,86880	,13098	7,2131	7,7414	5,60	8,80
2,00	4	5,8500	,85440	,42720	4,4905	7,2095	5,20	7,00
3,00	12	9,2417	,56320	,16258	8,8838	9,5995	8,00	9,90
Total	60	7,7217	1,18251	,15266	7,4162	8,0271	5,20	9,90

Multiple Comparisons

Dependent Variable: rendimiento estudiantes - uta

	(I) religiosidad uta	(J) religiosidad uta	Mean Difference (I-J)	Std Error	Sig	95% Confidence Interval	
						Lower Bound	Upper Bound
Tukey HSD	1,00	2,00	1,62727(*)	,42717	,001	,5993	2,6552
		3,00	-1,76439(*)	,26638	,000	-2,4054	-1,1234
	2,00	1,00	-1,62727(*)	,42717	,001	-2,6552	-,5993
		3,00	-3,39167(*)	,47225	,000	-4,5281	-2,2552
	3,00	1,00	1,76439(*)	,26638	,000	2,2552	4,5281
		2,00	3,39167(*)	,47225	,000	2,2552	4,5281
LSD	1,00	2,00	1,62727(*)	,42717	,000	,7719	2,4827
		3,00	-1,76439(*)	,26638	,000	-2,2978	-1,2310
	2,00	1,00	-1,62727(*)	,42717	,000	-2,4827	-,7719
		3,00	-3,39167(*)	,47225	,000	-4,3373	-2,4460
	3,00	1,00	1,76439(*)	,26638	,000	1,2310	2,2978
		2,00	3,39167(*)	,47225	,000	2,4460	4,3373

Tabla. 1, 2,3. Tablas de resultados de datos obtenidos.

DISCUSIÓN DE RESULTADOS

No se encuentran diferencias significativas de religiosidad entre facultades, tampoco en sexo.

Existen diferencias significativas en los rendimientos de los grupos estudiados (católicos, ninguna religión, otras religiones).

Al comparar los grupos 1 – 2, 1 – 3 y 2 – 3 también se encuentran diferencias significativas encontrándose que los rendimientos más bajos están en el segundo grupo (ninguna religión) y el más alto rendimiento en el grupo 3 (Otras religiones entre las que constan cristianos evangélicos, testigos de Jehová y otros). Lo que concuerda con los resultados obtenidos por otros investigadores especialmente en Europa y en Estados Unidos, (http://www.oei.es/noticias/spip.php?article4839&debut_5ultimasOEI=50), que manifiestan que los estudiantes de otras religiones en efecto son mejores estudiantes que aquellos que pertenecen a la religión católica y/u otros grupos religiosos

CONCLUSIONES

- La religión predominante en la Universidad Técnica de Ambato es la católica.
- No se encuentran diferencias significativas en las relaciones: Religión – Sexo, Religión – Facultad, es decir, las mujeres no son más religiosas que los hombres.
- La religión influye en el rendimiento académico de los estudiantes universitarios.

REFERENCIAS

1. <http://www.fungamma.org/religion1.htm>
2. http://www.google.com.ec/translate?hl=es&sl=en&u=http://www.gotoquiz.com/the_religious_identity_test&ei=WwInSsPPE5ik8AToxMyBDw&sa=X&oi=translate&resnum=6&ct=result&prev=/search%3Fq%3Dtest%252Breligion%26hl%3Des
3. <http://es.wikipedia.org/wiki/Religi%C3%B3n>

NOTAS

[illegible]

NOTAS

This image shows a single sheet of white paper with horizontal ruling lines. The lines are evenly spaced and run across the width of the page. There are no margins, text, or other markings on the paper.

NOTAS

[illegible]

NOTAS

[illegible]