

ÍNDICE :

- **Introducción**
- **Número factorial**
- **Variaciones**
- **Permutaciones**
- **Combinaciones**
- **Números combinatorios**
- **Triángulo de Tartaglia**
- **Binomio de Newton**

Introducción : La combinatoria estudia las diferentes formas en que se pueden realizar la ordenación o agrupamiento de unos cuantos objetos siguiendo unas determinadas condiciones o reglas . Una forma de hacer estos recuentos es utilizar los **diagramas en árbol** . Estos recuentos están intimamente relacionados con la **probabilidad** .

Número factorial : es el producto de los consecutivos naturales

$$n! = (n) \cdot (n-1) \cdot (n-2) \cdot \dots \cdot 3 \cdot 2 \cdot 1$$

Todo producto tiene al menos dos factores , luego debemos admitir que $0! = 1$ y que $1! = 1$

Variaciones ordinarias :

Se llama variaciones ordinarias de m elementos tomados de n en n (n en m) a los distintos grupos formados por n elementos de forma que :

- los n elementos que forman el grupo son distintos (no se repiten)
- Dos grupos son distintos si se diferencian en algún elemento o en el orden en que están colocados (influye el orden) .

$$V_{m,n} = \frac{m!}{(m-n)!}$$

Variaciones con repetición : se llama variaciones con repetición de m elementos tomados de n en n a los distintos grupos formados por n elementos de manera que :

- los elementos que forman cada grupo pueden estar repetidos
- Dos grupos son distintos si se diferencian en algún elemento o en el orden en que estos están colocados (influye el orden) .

$$VR_{m,n} = mn$$

Permutaciones ordinarias : se llama permutaciones de m elementos a las diferentes agrupaciones de esos m elementos de forma que :

- en cada grupo intervienen los m elementos sin repetirse ninguno (intervienen todos los elementos)
- dos grupos son diferentes si el orden de colocación de alguno de esos m elementos es distinto (influye el orden) .

$$P_m = m!$$

Permutaciones con repetición : se llama permutaciones con repetición de m elementos donde el primer elemento se repite a veces , el segundo b veces , el tercer c a los distintos grupos que pueden formarse con esos m elementos de forma que :

- intervienen todos los elementos
- dos grupos se diferencian en el orden de colocación de alguno de sus elementos .

$$PR_{m,a,b,c,\dots} = \frac{m!}{a!b!c!\dots}$$

Combinaciones : se llama combinaciones de m elementos tomados de n en n (n m) a todas las agrupaciones posibles que pueden hacerse con los m elementos de forma que :

- cada agrupación está formada por n elementos distintos entre sí
- dos agrupaciones distintas se diferencian al menos en un elemento , sin tener en cuenta el orden .

$$C_{m,n} = \frac{m!}{n!(m-n)!}$$

$$= \frac{m}{n}$$

= número combinatorio

Combinaciones con repetición : se llama combinaciones con repetición de m elementos tomados de n en n , a los distintos grupos formados por n elementos de manera que :

- los elementos que forman cada grupo pueden estar repetidos
- dos agrupaciones distintas se diferencian al menos en un elemento , sin tener en cuenta el orden .

$$CR_{m,n} = \frac{(m+n-1)!}{n!(m-1)!} = \frac{m+n-1}{n}$$

Por ejemplo las combinaciones con repetición de los elementos (a,b,c,d) tomados de dos en dos son :

aa ab ac ad

bb bc bd

cc cd

dd

Otro ejemplo : en una bodega hay 12 botellas de ron , 12 de ginebra y 12 de anís .Un cliente compró 8 botellas en total . ¿Cuántas posibilidades hay ?

$$CR_{8,3} = 120$$

Resumen :

Intervienen todos los elementos Permutaciones

Influye el orden Variaciones

No intervienen todos los elementos

No influye el orden Combinaciones

Números combinatorios : se llama número combinatorio de índice m y orden n al número de combinaciones de m elementos tomados de n en n tales que $n \leq m$.

$$= \frac{m!}{n!(m-n)!}$$

Propiedades :

$$\begin{aligned} & \bullet \quad \frac{m}{0} \\ &= \frac{m}{n} \\ &= 1 \end{aligned}$$

$$\begin{aligned} & \bullet \quad \frac{m}{n} \\ &= \frac{m}{m-n} \end{aligned}$$

$$\begin{aligned} & \bullet \quad \frac{m}{n} \\ &+ \frac{m}{n+1} \\ &= \frac{m+1}{n+1} \end{aligned}$$

$$\begin{aligned} & \bullet \quad \frac{n}{n} \\ &+ \frac{n}{n+1} \\ &+ \dots + \frac{n}{m} \end{aligned}$$

$$= \frac{m+1}{n+1}$$

Triángulo de Tartaglia o Pascal :

0

0

1

0

1

1

2

0

2

1

2

2

3

0

3

1

3

2

3

3

4

0

4

1

4

2

4

3

4

4

1 1

1 2 1

1 3 3 1

1 4 6 4 1

Binomio de Newton :

$$(a + b) = a + b$$

$$(a + b)^2 = a^2 + 2ab + b^2$$

$$(a + b)^3 = a^3 + 3a^2b + 3ab^2 + b^3$$

$$(a + b)^4 = a^4 + 4a^3b + 6a^2b^2 + 4ab^3 + b^4$$

.....

Si nos fijamos atentamente , los coeficientes coinciden con los del triángulo de Pascal , los exponentes de a van disminuyendo desde n hasta 0 y los de b van aumentando desde 0 hasta n , y en cada término la suma de los exponentes de a y b es igual a n .

Generalizando :

$$(a + b)^n = \begin{matrix} n \\ 0 \end{matrix}$$

$$a^nb^0 + \begin{matrix} n \\ 1 \end{matrix}$$

$$a^{n-1}b^1 + \dots + \begin{matrix} n \\ n-1 \end{matrix}$$

$$a^1b^{n-1} + \begin{matrix} n \\ n \end{matrix}$$

$$a^0b^n$$

ÍNDICE

- Definición de Estadística
- Conceptos generales
- Tratamiento de la información
- Representación de los datos
- Medidas de centralización
- Medidas de dispersión
- Estadística bidimensional
- Correlación
- Regresión

Definición de Estadística : la palabra estadística procede del vocablo "estado" pues era función principal de los gobiernos de los estados establecer registros de población , nacimientos , defunciones , etc . Hoy en día la

mayoría de las personas entienden por estadística al conjunto de datos , tablas , gráficos , que se suelen publicar en los periodicos .

En la actualidad se entiende por estadística como un método para tomar decisiones , de ahí que se emplee en multitud de estudios científicos .

La estadística se puede dividir en dos partes :

- **Estadística descriptiva o deductiva** , que trata del recuento , ordenación y clasificación de los datos obtenidos por las observaciones . Se construyen tablas y se representan gráficos , se calculan parámetros estadísticos que caracterizan la distribución , etc.
- **Estadística inferencial o inductiva** , que establece previsiones y conclusiones sobre una población a partir de los resultados obtenidos de una muestra . Se apoya fuertemente en el cálculo de probabilidades .

Población : es el conjunto de todos los elementos que cumplen una determinada característica . Ejemplo : alumnos matriculados en COU en toda España .

Muestra : cualquier subconjunto de la población . Ejemplo : alumnos de COU del Sotomayor .

Carácter estadístico : es la propiedad que permite clasificar a los individuos , puede haber de dos tipos :

- **Cuantitativos** : son aquellos que se pueden medir . Ejemplo : nº de hijos , altura , temperatura .
- **Cualitativos** : son aquellos que no se pueden medir . Ejemplo : profesión , color de ojos , estado civil .

Variable estadística : es el conjunto de valores que puede tomar el carácter estadístico cuantitativo (pues el cualitativo tiene "modalidades") . Puede ser de dos tipos :

- **Discreta** : si puede tomar un número finito de valores . Ejemplo : nº de hijos
- **Continua** : si puede tomar todos los valores posibles dentro de un intervalo . Ejemplo : temperatura , altura .

Frecuencia absoluta f_i : (de un determinado valor x_i) al número de veces que se repite dicho valor .

Frecuencia absoluta acumulada F_i : (de un determinado valor x_i) a su frecuencia absoluta más la suma de las frecuencias absolutas de todos los valores anteriores .

Frecuencia relativa h_i : es el cociente f_i/N , donde N es el número total de datos .

Frecuencia relativa acumulada H_i : es el cociente F_i/N

Si las frecuencias relativas las multiplicamos por 100 obtenemos los % .

Tratamiento de la información : se deben de seguir los siguientes pasos :

- recogida de datos
- ordenación de los datos
- recuento de frecuencias
- agrupación de los datos , en caso de que sea una variable aleatoria continua o bien discreta pero con un número de datos muy grande se agrupan en clases .

$$N^{\circ} \text{ de clases} = \sqrt{N}$$

Los puntos medios de cada clase se llaman marcas de clase .

Además se debe adoptar el criterio de que los intervalos sean cerrados por la izquierda y abiertos por la derecha .

- construcción de la tabla estadística que incluirá , clases , marca de clase , f_i , F_i , h_i , H_i .

Ejemplo : Las notas de Matemáticas de una clase han sido las siguientes :

5 3 4 1 2 8 9 8 7 6 6 7 9 8 7 7 1 0 1 5 9 9 8 0 8 8 8 9 5 7

Construir una tabla :

x_i	f_i	F_i	h_i	H_i
0	2	2	2/30	2/30
1	3	5	3/30	5/30
2	1	6	1/30	6/30
3	1	7	1/30	7/30
4	1	8	1/30	8/30
5	3	11	3/30	11/30
6	2	13	2/30	13/30
7	5	18	5/30	18/30
8	7	25	7/30	25/30
9	5	30	5/30	30/30
	30		1	

Representaciones gráficas : para hacer más clara y evidente la información que nos dan las tablas se utilizan los gráficos , que pueden ser :

- Diagramas de barras (datos cualitativos y cuantitativos de tipo discreto) . En el eje y se pueden representar frecuencias absolutas o relativas .
- Histogramas (datos cuantitativos de tipo continuo o discreto con un gran número de datos) . El histograma consiste en levantar sobre cada intervalo un rectángulo cuyo área sea igual a su frecuencia absoluta

$$\text{área} = \text{base} \cdot \text{altura } f_i = \Delta x_i \cdot n_i$$

luego la altura de cada rectángulo vendrá dada por n_i que se llama **función de**

densidad . Si por ejemplo un intervalo es doble de ancho que los demás su altura n_i debe ser la mitad de la frecuencia absoluta y así no se puede inducir a errores . Normalmente la amplitud de los intervalos es cte por lo que n_i será

proporcional a f_i y por tanto podemos tomar f_i como la altura ni ya que la forma del gráfico será la misma , aunque ahora el área del rectángulo ya no sea exactamente la frecuencia absoluta (a no ser que la amplitud del intervalo sea igual a 1) .

- Polígono de frecuencias
- Diagrama de sectores
- Cartogramas
- Pirámides de población
- Diagramas lineales
- Pictogramas

CÁLCULO DE PARÁMETROS :

Medidas de centralización :

- **Media aritmética :**

$$\bar{x} = \frac{x_1 + x_2 + \dots}{N}$$

$$= \frac{x_i}{N}$$

si son pocos datos

$$\bar{x} = \frac{x_1 f_1 + x_2 f_2 + \dots}{f_1 + f_2 + \dots}$$

$$= \frac{x_i f_i}{N}$$

si son muchos valores pero se repiten mucho

En el caso de que los datos estén agrupados en clases , se tomará la marca de clase

como x_i .

No siempre se puede calcular la media aritmética como por ejemplo cuando los

datos son cualitativos o los datos están agrupados en clases abiertas .

Ejemplo : hacer los cálculos para el ejercicio de las notas

- **Moda :** es el valor de la variable que presenta mayor frecuencia absoluta . Puede haber más de una . Cuando los datos están agrupados en clases se puede tomar la marca de clase o utilizar la fórmula :

$$M_0 = \text{Linf} + \Delta \frac{d_1}{d_1 + d_2}$$

donde : Linf = límite inferior de la clase modal , Δ
=amplitud

del intervalo , d_1 = diferencia entre la f_i de la clase modal y la f_i de la clase anterior y

d_2 = diferencia entre la f_i de la clase modal y la f_i de la clase posterior .

También se puede hacer gráficamente :

La moda si sirve para datos cualitativos , pero no tiene por qué situarse en la zona central del gráfico .

Ejemplo : en el ejercicio de las notas la moda sería $x=8$

- **Mediana** : es el valor de la variable tal que el número de observaciones menores que él es igual al número de observaciones mayores que él . Si el número de datos es par , se puede tomar la media aritmética de los dos valores centrales .

Cuando los datos están agrupados la mediana viene dada por el primer valor de la variable cuya F_i excede a la mitad del número de datos . Si la mitad del número de datos coincide con F_i se tomará la semisuma ente este valor y el siguiente .

Cuando los datos estén agrupados en clases se puede utilizar reglas de tres o la fórmula :

$$M = \text{Linf} + \Delta \frac{\frac{N}{2} - F_{i-1}}{f_i}$$

Gráficamente se hace a partir del polígono de frecuencias acumuladas .

Ejemplo : En el caso de las notas podrías ordenar de menor a mayor los datos y obtendríamos : 0 0 1 1 1 2 3 4 5 5 5 6 6 7 7 7 7 7 8 8 8 8 8 8 8 9 9 9 9 9

dato número 15–16 (por ser par)

luego la mediana sería 7

También se podría observar las F_i y ver que en el 7 se excede a la mitad del n° de datos , es decir , sobrepasa el 15 .

- **Cuantiles** : son parámetros que dividen la distribución en partes iguales , así por ejemplo la mediana los divide en dos partes iguales , los **cuartiles** son tres valores que dividen a la serie de datos en cuatro partes iguales , los **quintiles** son cuatro valores que lo dividen en 5 partes , los **deciles** en 10 y los **percentiles** en 100 . Se calculan de la misma manera que la mediana .

También se puede utilizar la fórmula : $C_n = \text{Linf} + \Delta \frac{n \frac{N}{100} - F_{i-1}}{f_i}$

donde n es el valor que deja el $n\%$ de valores por debajo de él .

Medidas de dispersión :

- **Rango o recorrido** : es la diferencia entre el mayor valor y el menor . Depende mucho de los valores extremos por que se suele utilizar el rango intercuartílico =

$Q3 - Q1$ o el rango entre percentiles = $P90 - P10$

Ejemplo : Para el caso de las notas sería $9 - 0 = 9$

- **Varianza s²** : es la media aritmética de los cuadrados de las desviaciones respecto a la media (desviación respecto a la media $d = x_i - \bar{x}$).

$$s^2 = \frac{(x_1 - \bar{x})^2 + (x_2 - \bar{x})^2 + \dots}{N}$$

$$= \frac{(x_i - \bar{x})^2}{N}$$

$$s^2 = \frac{f_1(x_1 - \bar{x})^2 + f_2(x_2 - \bar{x})^2 + \dots}{f_1 + f_2 + \dots}$$

$$= \frac{f_i(x_i - \bar{x})^2}{N}$$

Al igual que la media en el caso de que los datos estén agrupados en clases , se tomará la marca de clase como x_i .

Otra forma de calcular s² es :

$$s^2 = \frac{f_i(x_i - \bar{x})^2}{N}$$

$$= \frac{f_i(x_i^2 + \bar{x}^2 - 2x_i\bar{x})}{N} =$$

$$\frac{f_i x_i^2}{N} + \bar{x}^2 - 2\bar{x}^2$$

$$= \frac{f_i x_i^2}{N} - \bar{x}^2$$

Se llama **desviación típica s** a la raíz cuadrada de la varianza . Es más útil que la varianza ya que tiene las mismas dimensiones que la media

Ejemplo : Hacer los cálculos para el ejercicio de las notas

- **Coefficiente de variación** : es el cociente entre la desviación típica y la media aritmética . Valores muy bajos indican muestras muy concentradas .

$$C.V. = \frac{\sigma}{\bar{x}}$$

DISTRIBUCIONES BIDIMENSIONALES :

Variables estadísticas bidimensionales : es cuando al estudiar un fenómeno obtenemos dos medidas x e y , en vez de una como hemos hecho hasta ahora .

Ejemplo : pulso y tª de los enfermos de un hospital , ingresos y gastos de las familias de los trabajadores de una empresa , edad y nº de días que faltan al trabajo los productores de una fábrica .

Tipos de distribuciones bidimensionales :

- cualitativa – cualitativa
- cualitativa – cuantitativa (discreta o continua)
- cuantitativa (discreta o continua) – cuantitativa (discreta o continua)

Tipos de tablas :

- Tabla de dos columnas x_i, y_i (pocos datos)
- Tabla de tres columnas x_i, y_i, f_i (muchos datos y pocos valores posibles)
- Tablas de doble entrada (muchos datos y muchos valores posibles)

	x_1	x_2		x_n	f^*_j
y_1	f_{11}	f_{21}		f_{n1}	f^*_1
y_2	f_{12}	f_{22}		f_{n2}	f^*_2
.....					
y_m	f_{1m}	f_{2m}		f_{nm}	f^*_m
f_i^*	f_{1^*}	f_{2^*}		f_{n^*}	$f^{**}=N$

Diagramas de dispersión :

Si hay pocos datos (tabla de dos columnas), se representan las variables en los ejes x e y .

Si hay muchos datos pero muy agrupados (tabla de tres columnas y tablas de doble entrada), se hace igual pero con los puntos más gordos según la f_i ,o se pintan muchos puntos juntos , o se pinta en tres dimensiones x , y , f_i , con lo que obtendríamos un diagrama de barras en tres dimensiones .

Si hay muchos datos y muchos valores posibles , se pueden agrupar en clases , y se utilizan los **estereogramas** (3 dimensiones) en los que el volumen de cada prisma es proporcional a la frecuencia . También se puede tomar la marca de clase de los intervalos y tratar la variable continua como si fuese discreta .

Cálculo de parámetros :

- Cuando hay pocos datos o están muy agrupados (tablas de 2 o 3 columnas)

$$\bar{x} = \frac{\sum x_i f_i}{N}$$

$$\bar{y} = \frac{\sum y_i f_i}{N}$$

$$s_x^2 = \frac{\sum f_i (x_i - \bar{x})^2}{N}$$

$$s_y^2 = \frac{\sum f_i (y_i - \bar{y})^2}{N}$$

Aparece un parámetro nuevo que es la **covarianza** que es la media aritmética de las desviaciones de cada una de las variables respecto a sus medias respectivas .

$$s_{xy} = \frac{f_i(x_i - \bar{x})(y_i - \bar{y})}{N}$$

$$= \frac{f_i x_i y_i}{N} - \bar{x} \bar{y}$$

- Cuando hay muchos datos (tablas de doble entrada)

$$\bar{x} = \frac{x_i f_{i*}}{N} = \frac{x_i f_{ij}}{N}$$

$$\bar{y} = \frac{y_j f_{*j}}{N} = \frac{y_j f_{ij}}{N}$$

$$s_x^2 = \frac{f_{i*}(x_i - \bar{x})^2}{N} = \frac{f_{ij}(x_i - \bar{x})^2}{N} = f_{ij}x_i^2 - \bar{x}^2$$

$$s_y^2 = \frac{f_{*j}(y_j - \bar{y})^2}{N} = \frac{f_{ij}(y_j - \bar{y})^2}{N} = f_{ij}y_j^2 - \bar{y}^2$$

$$s_{xy} = \frac{f_{ij}(x_i - \bar{x})(y_j - \bar{y})}{N}$$

$$= \frac{f_{ij}x_i y_j}{N} - \bar{x} \bar{y}$$

Correlación o dependencia : es la teoría que trata de estudiar la relación o dependencia entre las dos variables que intervienen en una distribución bidimensional , según sean los diagramas de dispersión podemos establecer los siguientes casos :

- **Independencia funcional o correlación nula :** cuando no existe ninguna relación entre las variables .($r = 0$)
- **Dependencia funcional o correlación funcional :** cuando existe una función tal que todos los valores de la variable la satisfacen (a cada valor de x le corresponde uno solo de y o a la inversa) ($r = 1$)
- **Dependencia aleatoria o correlación curvilínea (ó lineal):** cuando los puntos del diagrama se ajustan a una línea recta o a una curva , puede ser positiva o directa , o negativa o inversa ($-1 < r < 0$ ó $0 < r < 1$)

Ejemplo : a 12 alumnos de COU se les toma las notas de los últimos exámenes de Matemáticas , Física y Filosofía :

Matemáticas	Física	Filosofía
2	1	2
3	3	5
4	2	7
4	4	8
5	4	5
6	4	3
6	6	4

7	4	6
7	6	7
8	7	5
10	9	5
10	10	9

Si representamos las variables matemáticas– física en un diagrama y matemáticas–filosofía en otro vemos que la correlación es mucho más fuerte en el primero que en el segundo ya que los valores están más alineados .

Coefficiente de correlación lineal : es una forma de cuantificar de forma más precisa el tipo de correlación que hay entre las dos variables .

$$r = \frac{s_{xy}}{s_x s_y}$$

Regresión : consiste en ajustar lo más posible la nube de puntos de un diagrama de dispersión a una curva . Cuando esta es una recta obtenemos la recta de regresión lineal , cuando es una parábola , regresión parabólica , cuando es una exponencial , regresión exponencial , etc . (logicamente r debe ser distinto de 0 en todos los casos) .

La **recta de regresión de y sobre x** es :
$$y - \bar{y} = \frac{s_{xy}}{s_x^2} (x - \bar{x})$$

en la cual se hace mínima la distancia entre los valores yj obtenidos experimentalmente y los valores teóricos de y.

A valor
$$\frac{s_{xy}}{s_x^2}$$

se le llama **coeficiente de regresión de y sobre x** (nos da la pendiente de la recta de regresión).

La **recta de regresión de x sobre y** es :
$$x - \bar{x} = \frac{s_{xy}}{s_y^2} (y - \bar{y})$$

en la cual se hace mínima la distancia entre los valores xi obtenidos experimentalmente y los valores teóricos de x.

A valor
$$\frac{s_{xy}}{s_y^2}$$

se le llama **coeficiente de regresión de x sobre y** (su inversa nos da la otra pendiente) .

ÍNDICE

- Métodos de muestreo
- Distribución del muestreo
- Intervalos de confianza
- Contraste de hipótesis

Métodos de muestreo : para no tener que trabajar con toda la población se utiliza el muestreo . Puede ser :

- **Muestreo no probabilístico** : no se usa el azar , sino el criterio del investigador , suele presentar grandes sesgos y es poco fiable .
- **Muestreo probabilístico** : se utilizan las leyes del azar . Puede ser :

- **Muestreo aleatorio simple** (es el más importante) : cada elemento de la población tiene la misma probabilidad de ser elegido , las observaciones se realizan con reemplazamiento , de manera que la población es idéntica en todas las extracciones , o sea , que la selección de un individuo no debe afectar a la probabilidad de que sea seleccionado otro cualquiera aunque ello comporte que algún individuo pueda ser elegido más de una vez . ("se hacen tantas papeletas numeradas como individuos hay , se coge una y se devuelve , se vuelve a coger otra y se devuelve , etc")
- **Muestreo sistemático** : es cuando los elementos de la población están ordenados por listas . Se elige un individuo al azar y a continuación a intervalos constantes se eligen todos los demás hasta completar la muestra . Si el orden de los elementos es tal que los individuos próximos tienden a ser más semejantes que los alejados , el muestreo sistemático tiende a ser más preciso que el aleatorio simple , al cubrir más homogéneamente toda la población .
- **Muestreo estratificado** : es cuando nos interesa que la muestra tenga la misma composición a la de la población la cual se divide en clases o estratos . Si por ejemplo en la población el 20% son mujeres y el 80% hombres , se mantendrá la misma proporción en la muestra .

Distribuciones de muestreo : al obtener conclusiones de la muestra y las comparamos con las de la población puede que se aproximen o no . No obstante , las medias muestrales se comportan estadísticamente bien y siguen leyes perfectamente previsibles , esto nos permitirá hacer inferencias precisas a partir de ellas , incluso determinar el riesgo que asumimos al hacerlas .

Si una población está formada por N elementos , el nº de muestras diferentes de tamaño n que se pueden obtener , si se pueden repetir los elementos (m.a.s.) sería :

$$VRN, n = N^n$$

Distribución de medias muestrales : aunque al tomar una muestra no podemos estar seguros de que los parámetros obtenidos sean buenos estimadores de los parámetros poblacionales si se puede afirmar que :

- La media de las medias muestrales es igual a la media real de la población es decir :

$$\bar{X} = \frac{\bar{X}_1 + \bar{X}_2 + \dots + \bar{X}_n}{n \text{ de muestras posibles}} = \mu$$

- La desviación típica de las medias muestrales vale :

$$\sigma_{\bar{x}} = \frac{\sigma}{\sqrt{n}}$$

Esto significa que la distribución de medias muestrales de tamaño n extraídas de una población (normal o no normal) se distribuye según una $N(\mu, \frac{\sigma}{\sqrt{n}})$

Ejemplo : Supongamos que tenemos los elementos 2,4,6,8 .

$$\mu = \frac{2+4+6+8}{4} = 5$$

$$\sigma = \sqrt{\frac{2^2 + 4^2 + 6^2 + 8^2}{4} - 5^2} = 2.23$$

En esta población vamos a tomar todas las muestras posibles de tamaño 2 :

Elementos e1 e2	Media de la muestra
2 2	2
2 4	3
2 6	4
2 8	5
4 2	3
4 4	4
4 6	5
4 8	6
6 2	4
6 4	5
6 6	6
6 8	7
8 2	5
8 4	6
8 6	7
8 8	8

La media de las medias muestrales será : $\bar{X} = \frac{2+3+4+5+3+....}{16} = 5$

La varianza de las medias muestrales será : $\sigma_{\bar{x}} = \sqrt{\frac{(2-5)^2 + (3-5)^2 + ...}{16}} = 1.58$

Por lo tanto $\bar{X} = \mu$

y $\sigma_{\bar{x}} = \frac{\sigma}{\sqrt{n}}$

Ejemplo : el peso de los recién nacidos en una maternidad se ha distribuido según una ley normal de media $\mu = 3100$ gr y de desviación típica $\sigma = 150$ gr ¿Cuál será la probabilidad de que la media de una muestra de 100 recién nacidos se superior a 3130gr ?

La distribución muestral sigue una $N(3100,15)$ por lo que $p(\bar{x} > 3130) =$

$$p\left(Z > \frac{3130 - 3100}{15}\right) = p(Z > 2)$$

$$= 1 - 0.9772 = 0.0228$$

por lo tanto solo un 2'28% de las muestras tendrá una media por encima de los 3130gr

Intervalos de probabilidad : inferencia estadística ; Como la distribución de medias muestrales es

$$N \quad \mu, \frac{\sigma}{\sqrt{n}}$$

se tendrá por ejemplo que :

$$P \left(\mu - \frac{\sigma}{\sqrt{n}} \leq \bar{X} \leq \mu + \frac{\sigma}{\sqrt{n}} \right) = 0.6826$$

$$P \left(\mu - 2 \frac{\sigma}{\sqrt{n}} \leq \bar{X} \leq \mu + 2 \frac{\sigma}{\sqrt{n}} \right) = 0.9544$$

$$P \left(\mu - 3 \frac{\sigma}{\sqrt{n}} \leq \bar{X} \leq \mu + 3 \frac{\sigma}{\sqrt{n}} \right) = 0.9974$$

Esto significa que por ejemplo el 68'26% de las muestras de tamaño n extraídas de una población de media μ y desviación σ

tendrán una media perteneciente al intervalo $(\mu - \sigma, \mu + \sigma)$

En general el $100 \cdot (1 - \alpha)$

% de las muestras de tamaño n tendrán una media comprendida entre :

$$\mu - z_{\alpha/2} \frac{\sigma}{\sqrt{n}}, \mu + z_{\alpha/2} \frac{\sigma}{\sqrt{n}}$$

siendo $\alpha/2$

el valor de la probabilidad que queda a cada lado del intervalo . O lo que es lo mismo :

$$P \left(\mu - z_{\alpha/2} \frac{\sigma}{\sqrt{n}} \leq \bar{X} \leq \mu + z_{\alpha/2} \frac{\sigma}{\sqrt{n}} \right) = 1 - \alpha$$

(nivel de confianza)

Así por ejemplo si α

= 0.05 entonces el 95% de las muestras tendrán una media comprendida entre

$$\mu - z_{0.025} \frac{\sigma}{\sqrt{n}}, \mu + z_{0.025} \frac{\sigma}{\sqrt{n}}$$

$$= \mu - 1.96 \frac{\sigma}{\sqrt{n}}, \mu + 1.96 \frac{\sigma}{\sqrt{n}}$$

μ

$1 - \alpha$

$\alpha/2$

$\alpha/2$

$$\mu + 1.96 \frac{\sigma}{\sqrt{n}}$$

$$\mu - 196 \frac{\sigma}{\sqrt{n}}$$

Sin embargo lo normal será que se desconozca la media y la desviación típica de la población y que mediante técnicas de muestreo se busque estimarlas con la fiabilidad necesaria .

Por lo tanto si nos hacen la pregunta de otra forma (¿ cuál es la probabilidad de que la media poblacional se encuentre entre ...?) podremos transformar la desigualdad obteniendo :

$$\bar{X} - Z_{\alpha/2} \frac{\sigma}{\sqrt{n}} \leq \mu \leq \bar{X} + Z_{\alpha/2} \frac{\sigma}{\sqrt{n}}$$

$$P \left(\bar{X} - Z_{\alpha/2} \frac{\sigma}{\sqrt{n}} \leq \mu \leq \bar{X} + Z_{\alpha/2} \frac{\sigma}{\sqrt{n}} \right) = 1 - \alpha$$

(nivel de confianza)

A este intervalo se le llama **intervalo de confianza** para la media poblacional . A lo que está fuera del intervalo se le llama **región crítica** .

Al valor $1 - \alpha$
se le llama **nivel de confianza** .

Al valor α
se le llama **nivel de significación** .

$\bar{X}$

Por lo tanto el nuevo dibujo sería :

$$\bar{X} - 196 \frac{\sigma}{\sqrt{n}}$$

$$\bar{X} + 196 \frac{\sigma}{\sqrt{n}}$$

Por lo tanto podemos afirmar que en ese intervalo tenemos una probabilidad del 95

% de que está la media poblacional .

Como ya hemos dicho lo normal será que se desconozca la desviación , por lo que debemos sustituir σ por $s_{n-1} = \sqrt{\frac{\sum (x_i - \bar{X})^2}{n-1}}$

donde s_{n-1} es la cuasivarianza muestral .

La relación entre la varianza muestral y la cuasivarianza muestral es :

$$s_{n-1}^2 = \frac{n}{n-1} s^2$$

Aunque para valores grandes de n (mayores de 30) coinciden aproximadamente la cuasivarianza y la varianza por lo que se puede sustituir σ por s .

Por lo tanto :

$$P \left(\bar{X} - z_{\alpha/2} \frac{s_{n-1}}{\sqrt{n}} \leq \mu \leq \bar{X} + z_{\alpha/2} \frac{s_{n-1}}{\sqrt{n}} \right) = 1 - \alpha$$

Para n grandes :

$$P \left(\bar{X} - z_{\alpha/2} \frac{s}{\sqrt{n}} \leq \mu \leq \bar{X} + z_{\alpha/2} \frac{s}{\sqrt{n}} \right) = 1 - \alpha$$

Distribución para proporciones : cuando se trata de determinar la proporción de una población que posee un cierto atributo (hombre/mujer , video/no video , éxito/fracaso , etc) su estudio es equiparable al de una distribución binomial . Así pues si tomamos muestras aleatorias de tamaño n , la media y la desviación típica de las medias muestrales será :

$$\bar{p} = p \quad \sigma_{\bar{p}} = \sqrt{\frac{pq}{n}}$$

Esta distribución es aproximadamente normal para valores grandes de n (mayor de 30) en consecuencia puede estudiarse como una N

$$p, \sqrt{\frac{pq}{n}}$$

Si hablamos de intervalos de probabilidad entonces :

$$P \left(p - z_{\alpha/2} \sqrt{\frac{pq}{n}} \leq \bar{p} \leq p + z_{\alpha/2} \sqrt{\frac{pq}{n}} \right) = 1 - \alpha$$

(nivel de confianza)

Como lo que no se suele saber es la media y la varianza podemos hacer :

$$P \left(\bar{p} - z_{\alpha/2} \sqrt{\frac{\bar{p}q}{n}} \leq \bar{p} \leq \bar{p} + z_{\alpha/2} \sqrt{\frac{\bar{p}q}{n}} \right) = 1 - \alpha$$

Error admitido y tamaño de la muestra :

cuando decimos que $\bar{X} - z_{\alpha/2} \frac{\sigma}{\sqrt{n}} \leq \mu \leq \bar{X} + z_{\alpha/2} \frac{\sigma}{\sqrt{n}}$

estamos admitiendo un error máximo de $z_{\alpha/2} \frac{\sigma}{\sqrt{n}}$

esto es : la diferencia máxima entre la media poblacional y la media muestral debe ser menor que este valor . Como se puede observar de este valor se puede controlar dos parámetros , n y z .

El tamaño mínimo de una encuesta depende de la confianza que se desee para los resultados y del error máximo que se esté dispuesto a asumir :

$$E = z_{\alpha/2} \frac{\sigma}{\sqrt{n}}$$

despejando

$$n = z_{\alpha/2}^2 \frac{\sigma^2}{E^2}$$

Analogamente se puede hacer para la distribución de proporciones .

Ejemplo : se desea realizar una investigación para estimar el peso medio de los hijos de madres fumadoras . Se admite un error máximo de 50 gr , con una confianza del 95% . Si por estudios se sabe que la desviación típica es de 400 gr ¿ Qué tamaño mínimo de muestra se necesita en la investigación ?

El tamaño mínimo de la muestra debe ser $n = 196 \frac{400^2}{50^2} = 24586$

= 246

Contraste de hipótesis sobre la media poblacional : La media muestral puede ser diferente de la media poblacional . Lo normal es que estas diferencias sean pequeñas y estén justificadas por el azar , pero podría suceder que no fuesen debidas al azar sino a que los parámetros poblacionales sean otros , que por los motivos que sea , han cambiado .

El contraste de hipótesis es el instrumento que permite decidir si esas diferencias pueden interpretarse como fluctuaciones del azar (hipótesis nula) o bien , son de tal importancia que requieren una explicación distinta (hipótesis alternativa). Como en los intervalos de confianza las conclusiones se formularán en términos de probabilidad .

Comparando la media poblacional y la media muestral ¿ Podemos asegurar que esa muestra procede de una población de media μ

0 ? La respuesta será no cuando μ

0 no pertenezca al intervalo de confianza de $\bar{x}$

, para el nivel de significación prefijado , por el contrario la respuesta será sí cuando sí pertenezca a tal intervalo .

Sí pertenece a la población..... $\mu_0 \quad \bar{x} - z_{\alpha/2} \frac{\sigma}{\sqrt{n}}, \bar{x} + z_{\alpha/2} \frac{\sigma}{\sqrt{n}}$

se acepta la hipótesis nula . Otra forma de verlo es que : $\left| \frac{\bar{x} - \mu_0}{\sigma / \sqrt{n}} \right| < z_{\alpha/2}$

No pertenece a la población $\mu_0 \quad \bar{x} - z_{\alpha/2} \frac{\sigma}{\sqrt{n}}, \bar{x} + z_{\alpha/2} \frac{\sigma}{\sqrt{n}}$

se rechaza la hipótesis nula . Otra forma de verlo es que : $\left| \frac{\bar{x} - \mu_0}{\sigma / \sqrt{n}} \right| > z_{\alpha/2}$

Error de tipo I : es el que cometemos cuando rechazamos la hipótesis nula siendo verdadera .

Error de tipo II : es el que cometemos cuando aceptamos la hipótesis nula siendo falsa .

Podemos hacer todavía dos preguntas :

¿ La muestra procede de una población con media mayor que la supuesta ?

Se acepta que la media poblacional es mayor que la supuesta μ_0

cuando : $\frac{\bar{x} - \mu_0}{\sigma / \sqrt{n}} \geq z_{\alpha}$

desarrollando la igualdad obtenemos que : $\bar{x} - z_{\alpha/2} \frac{\sigma}{\sqrt{n}} \geq \mu_0$

μ_0

$\bar{x}$

$$\bar{x} - 196 \frac{\sigma}{\sqrt{n}}$$

La media poblacional debe de estar por encima de $\bar{x} - z_{\alpha/2} \frac{\sigma}{\sqrt{n}}$

y por lo tanto por encima de μ_0

La rechazamos en caso contrario .

¿ La muestra procede de una población con media menor que la supuesta ?

Se acepta que la media es menor que la supuesta cuando : $\frac{\bar{x} - \mu_0}{\sigma / \sqrt{n}} \leq -z_{\alpha}$

desarrollando la igualdad obtenemos que : $\bar{x} + z_{\alpha/2} \frac{\sigma}{\sqrt{n}} \leq \mu_0$

μ_0

$\bar{x}$

$$\bar{x} + 196 \frac{\sigma}{\sqrt{n}}$$

La media poblacional debe de estar por debajo de $\bar{x} + z_{\alpha/2} \frac{\sigma}{\sqrt{n}}$

y por lo tanto por debajo de μ_0

La rechazamos en caso contrario .

Nota : No olvidemos que en todas las ecuaciones anteriores si se desconoce la desviación típica de la población debemos sustituirla por la cuasivarianza de la muestra.

Contraste de hipótesis sobre la proporción p : por analogía con el apartado anterior par responder a la pregunta : ¿ Puede asegurarse que esa muestra de proporción $\bar{p}$ procede de una población con proporción p_0 ?

La respuesta será sí cuando : $\left| \frac{\bar{p} - p_0}{\sqrt{pq/n}} \right| \geq z_{\alpha/2}$
 con una probabilidad de $1 - \alpha$

Se admite que la media poblacional es mayor que un valor p_0 si :

$$\frac{\bar{p} - p_0}{\sqrt{pq/n}} \geq z_{\alpha}$$

Se admite que la media poblacional es menor que un valor p_0 si :

$$\frac{\bar{p} - p_0}{\sqrt{pq/n}} \leq -z_{\alpha}$$

ÍNDICE

- Sucesos aleatorios
- Definición de probabilidad
- Probabilidad condicionada
- Teorema de la Bayes
- Variable aleatoria
- Función de probabilidad
- Función de distribución
- Media y varianza
- Distribución binomial
- Distribución normal

SUCESOS ALEATORIOS

Experimento aleatorio : es aquel que se caracteriza porque al repetirlo bajo análogas condiciones jamás se puede predecir el resultado que se va a obtener . En caso contrario se llama experimento determinista .

Espacio muestral E : (de un experimento aleatorio) es el conjunto de todos los resultados posibles del experimento .

Suceso de un experimento aleatorio : es un subconjunto del espacio muestral . Puede haber los siguientes tipos :

- suceso elemental
- suceso compuesto (de varios sucesos elementales)
- suceso seguro
- suceso imposible
- suceso contrario

Operaciones con sucesos :

- **Unión de sucesos** : la unión de dos sucesos A y B es el suceso que se realiza cuando se realiza A ó B
- **Intersección de sucesos** : la intersección de A y B es el suceso que se realiza cuando se realizan simultáneamente los sucesos A y B . Cuando es imposible que los sucesos se realicen

simultaneamente se dice que son **incompatibles** . Si $A \cap B = \Phi$ entonces son incompatibles . En caso contrario se dice que son compatibles .

Propiedades :

	Unión	Intersección
Asociativa	$(A \cup B) \cup C = A \cup (B \cup C)$	$(A \cap B) \cap C = A \cap (B \cap C)$
Conmutativa	$A \cup B = B \cup A$	$A \cap B = B \cap A$
Idempotente	$A \cup A = A$	$A \cap A = A$
Simplificativa	$A \cup (A \cap B) = A$ $A \cap (A \cup B) = A$	
Distributiva	$A \cup (B \cap C) = (A \cup B) \cap (A \cup C)$ $A \cap (B \cup C) = (A \cap B) \cup (A \cap C)$	
Suceso contrario	$\overline{\overline{A}} = A$ $\overline{A \cup B} = \overline{A} \cap \overline{B}$ $\overline{A \cap B} = \overline{A} \cup \overline{B}$	$\overline{\overline{A}} = A$

Sistema completo de sucesos : Se dice que un conjunto de sucesos $A_1, A_2, \dots$ constituyen un sistema completo cuando se verifica :

- $A_1 \cup A_2 \cup \dots = E$
- $A_1, A_2, \dots$ son incompatibles 2 a 2 .

$A_1, A_2, \dots, A_n$

PROBABILIDAD

Ley de los grandes números : La frecuencia relativa de un suceso tiende a estabilizarse en torno a un número , a medida que el número de pruebas del experimento crece indefinidamente . Este número lo llamaremos **probabilidad de un suceso** .

Definición clásica de probabilidad : (regla de Laplace)

$$p(A) = \frac{n \text{ de casos favorables}}{n \text{ de casos posibles}}$$

(para aplicar esta definición se supone que los sucesos elementales son equiprobables)

Definición axiomática de probabilidad : (Kolmogorov) Se llama probabilidad a una ley que asocia a cada suceso A un número real que cumple los siguientes axiomas :

- La probabilidad de un suceso cualquiera del espacio de sucesos siempre es positiva , es decir **$p(A) \geq 0$**
- La probabilidad del suceso seguro es 1 , es decir , **$p(E) = 1$**
- La probabilidad de la unión de sucesos incompatibles es igual a la suma de probabilidades de cada uno de ellos , o sea , **$p(A \cup B) = p(A) + p(B)$**

Consecuencias de los axiomas :

- **$p(\overline{A}) = 1 - p(A)$**
- **$p(\emptyset) = 0$**
- **$0 \leq p(A) \leq 1$**
- Si $A \subset B$ **$p(A) \leq p(B)$**
- Si los sucesos son compatibles : **$p(A \cup B) = p(A) + p(B) - p(A \cap B)$**

Para el caso de tres sucesos compatibles sería :

$$p(A \cup B \cup C) = p(A) + p(B) + p(C) - p(A \cap B) - p(A \cap C) - p(B \cap C) + p(A \cap B \cap C)$$

Probabilidad condicionada $p(A/B)$: Se llama probabilidad del suceso A condicionado por B a la probabilidad de que se cumpla A una vez que se ha verificado el B .

$$p(A/B) = \frac{p(A \cap B)}{p(B)}$$

A B

a b c

$$p(A|B) = \frac{b}{a+b+c}$$

$$p(B) = \frac{b+c}{a+b+c}$$

$$p(A/B) = \frac{b}{b+c}$$

Otra forma de ver la fórmula es :

$$p(A|B) = p(B) \cdot p(A/B) = p(A) \cdot p(B/A) = p(B|A)$$

Generalizando : $p(A|B|C) = p(A) \cdot p(B/A) \cdot p(C/A|B)$

Ejemplo :

	Hombres	Mujeres	
Fuman	70	40	110
No Fuman	20	30	50
	90	70	160

$$p(H) = 90/160 \quad p(M) = 70/160 \quad p(F) = 110/160 \quad p(NF) = 50/160$$

$$p(H/NF) = 20/50 \quad p(H/F) = 70/110 \quad p(M/NF) = 30/50 \quad p(M/F) = 40/110$$

$$p(H|F) = 70/160 = p(F) \cdot p(H/F) = (110/160) \cdot (70/110)$$

Lo mismo se podría hacer con color de ojos (marrones y azules) y color de pelo (rubio y castaño) .

Sucesos independientes : dos sucesos A y B se dice que son independientes si

$p(A) = p(A|B)$. En caso contrario , $p(A|B) \neq p(A)$, se dice que son dependientes .

Probabilidad de la intersección o probabilidad compuesta :

- Si los sucesos son dependientes $p(A|B) = p(A) \cdot p(B/A) = p(B) \cdot p(A/B)$
- Si los sucesos son independientes $p(A|B) = p(A) \cdot p(B)$

Ejemplo : si al extraer dos cartas de una baraja lo hacemos con devolución tendremos dos sucesos independientes , $p(A|B) = p(A) \cdot p(B)$

$B) = p(A) \cdot p(B)$ pero si lo hacemos sin devolución ahora si son dependientes $p(A$

$B) = p(A) \cdot p(B/A)$.

Teorema de la probabilidad total : sea un sistema completo de sucesos y sea un suceso B tal que $p(B/A_i)$ son conocidas , entonces :

$$p(B) = p(B/A_1) + p(B/A_2) + \dots = p(B/A_i)$$

$A_1 A_2 A_3 A_4$

B

B

Teorema de Bayes : sea un sistema completo de sucesos y sea un suceso B tal que $p(B/A_i)$ son conocidas , entonces :

$$p(A_i/B) = \frac{p(A_i) p(B/A_i)}{p(A_i) p(B/A_i)}$$

Ejemplo importante : Se va a realizar el siguiente experimento , se tira una moneda , si sale cara se saca una bola de una urna en la que hay 4 bolas negras , 3 turquesa y 3 amarillas , si sale cruz se saca una bola de otra urna en la que hay 5 bolas negras , 2 turquesa y 3 amarillas .

NNNN

TTT

AAA

Cara -----

NNNNN

TT

AAA

Cruz -----

N 4/10 $p(\text{Cara})$
 $N) = 1/2 \cdot 4/10 = 4/20$

Cara 1/2 T 3/10 $p(\text{Cara})$
 $T) = 1/2 \cdot 3/10 = 3/20$

A 3/10 $p(\text{Cara})$
 $A) = 1/2 \cdot 3/10 = 3/20$

$$N \text{ 5/10 } p(\text{Cruz}) \\ N) = 1/2 \cdot 5/10 = 5/20$$

$$\text{Cruz } 1/2 \text{ T } 2/10 \text{ } p(\text{Cruz}) \\ T) = 1/2 \cdot 2/10 = 2/20$$

$$A \text{ 3/10 } p(\text{Cruz}) \\ A) = 1/2 \cdot 3/10 = 3/20$$

$$\text{Tª de la probabilidad total : } p(N) = p(\text{Cara} \\ N) + p(\text{Cruz} \\ N) = 4/20 + 5/20 = 9/20$$

$$\text{Tª de Bayes : } p(\text{Cara}/N) = \frac{p(\text{Cara } N)}{p(\text{Cara } N) + p(\text{Cruz } N)}$$

que no es ni más ni menos que casos favorables entre casos posibles .

DISTRIBUCIONES DISCRETAS : DISTRIBUCIÓN BINOMIAL

Variable aleatoria X : es toda ley que asocia a cada elemento del espacio muestral un número real . Esto permite sustituir los resultados de una prueba o experimento por números y los sucesos por partes del conjunto de los números reales .

Las variables aleatorias pueden ser discretas o continuas .

Por ejemplo en el experimento aleatorio de lanzar tres monedas el espacio muestral es $E = [CCC , CCX , CXC , XCC , CXX , XCX , XXC , CCC]$. Supongamos que a cada suceso le asignamos un número real igual al número de caras obtenidas . Esta ley o función que acabamos de construir la llamamos variable aleatoria (discreta) que representa el nº de caras obtenidas en el lanzamiento de tres monedas .

Consideremos el experimento que consiste en elegir al azar 100 judías de una plantación y medimos su longitud . La ley que asocia a cada judía su longitud es una variable aleatoria (continua) .

Por ejemplo al lanzar un dado podemos tener la variable aleatoria x_i que asocia a cada suceso el nº que tiene en la parte de arriba .

Por ejemplo al lanzar dos dados podemos tener la variable aleatoria x_i que asocia a cada suceso el producto de los dos números que tiene en la parte de arriba .

Función de probabilidad : (de una variable aleatoria) es la ley que asocia a cada valor de la variable aleatoria x_i su probabilidad $p_i = p(X = x_i)$.

Función de distribución $F(x)$: (de una variable aleatoria) es la ley que asocia a cada valor de la variable aleatoria , la probabilidad acumulada de este valor .

$$F(x) = p (X \\ x)$$

$$\text{Media de una variable aleatoria discreta : } \mu = \sum x_i p_i$$

$$\text{Varianza de una variable aleatoria discreta : } \sigma^2$$

$$= (x_i - \bar{x})^2 p_i$$

Ejemplo : en una bolsa hay bolas numeradas : 9 bolas con un 1 , 5 con un 2 y 6 con un 3 . Sacamos una bola y vemos que número tienen .

La función de probabilidad es :

x_i	1	2	3
p_i	9/20	5/20	6/20

La función de distribución es :

x_i	1	2	3
p_i	9/20	14/20	20/20

La media es $1 \cdot (9/20) + 2 \cdot (5/20) + 3 \cdot (6/20) = 1'85$

La varianza es $(1-1'85)^2 \cdot 9/20 + (2-1'85)^2 \cdot 5/20 + (3-1'85)^2 \cdot 6/20 = 0'72$

Distribución binomial : Una variable aleatoria es binomial si cumple las siguientes características :

- Los elementos de la población se clasifican en dos categorías , éxito o fracaso .
- El resultado obtenido en cada prueba es independiente de los resultados anteriores
- La probabilidad de éxito y fracaso es siempre constante

Ejemplos : fumadores de una población , nº de aprobados de la clase , días de lluvia a lo largo de un año , nº de caras al tirar una moneda , etc .

• **Función de probabilidad** $p(X = r) = \binom{n}{r} p^r q^{n-r}$

pr q^{n-r} donde p es la probabilidad de éxito , q la probabilidad de fracaso , n el numero total de pruebas y r el número de éxitos .

• **Función de distribución** $P(X \leq x) = \sum_{r=0}^x \binom{n}{r} p^r q^{n-r}$

$$\sum_{r=0}^x \binom{n}{r} p^r q^{n-r}$$

$$r$$

$$p^r q^{n-r}$$

• **Media** $\mu = n \cdot p$

• **Varianza** $\sigma^2 = n \cdot p \cdot q$

Ejemplo : Se lanza una moneda 11 veces :

¿Cuál es la probabilidad de obtener 5 caras ?

¿Cuál es la probabilidad de obtener 5 o menos caras ?

¿ Cuántas caras se obtienen por término medio ?

¿Cuál es la desviación típica ?

DISTRIBUCIONES CONTINUAS : DISTRIBUCIÓN NORMAL

Función de densidad $f(x)$: cuando en un histograma de frecuencias relativas de una variable continua aumentamos el nº de clases y por lo tanto su amplitud es más pequeña vemos que el polígono de frecuencias relativas se acerca a una función $f(x)$ que llamaremos función de densidad que cumple las siguientes propiedades :

- $f(x) \geq 0$

- $\int_{-\infty}^{+\infty} f(x)dx = 1$

-

el área encerrada bajo la curva de la función es igual a la unidad .

- $\int_a^b f(x)dx = p(a \leq X \leq b)$

-

área bajo la curva correspondiente a ese intervalo .

Función de distribución $F(x) = p(X \leq x)$

$F(x)$: cuando en un histograma de frecuencias relativas acumuladas de una variable continua aumentamos el nº de clases y por lo tanto su amplitud es más pequeña vemos que el polígono de frecuencias relativas acumuladas se acerca a una función $F(x)$ que llamaremos función de distribución que cumple las siguientes propiedades :

- $F(a) = \int_{-\infty}^a f(x)dx$

-

$$= p(X \leq a)$$

a) por lo tanto :

$$p(a < X \leq b) =$$

$$\int_a^b f(x)dx$$

$$= F(b) - F(a)$$

- $F(x)$ es nula para todo valor de x anterior al menor valor de la variable aleatoria y es igual a la unidad para todo valor posterior al mayor valor de la variable aleatoria . Si es continua se dice que $F(-\infty)=0$ y $F(+\infty)=1$

- Por ser una probabilidad $0 \leq F(x) \leq 1$

.

- Es una función creciente .

Media de una variable aleatoria continua :

$$\mu = \int_a^b x f(x)dx$$

Varianza de una variable aleatoria continua : σ^2

$$= \int_a^b (x - \bar{x})^2 f(x) dx$$

Distribución normal : una variable aleatoria es normal si se rige según las leyes del azar . La mayoría de las distribuciones más importantes son normales . Por ejemplo la distribución de los pesos de los individuos de cualquier especie , la estatura de una población , Tª del mes de agosto a lo largo de 100 años , la longitud de los tornillos que salen de una fábrica , etc .

No todas las distribuciones son normales por ejemplo si clasificamos según el nivel de renta a los ciudadanos españoles son muy pocos los que poseen niveles de rentas altas y en cambio son muchos los que poseen niveles de rentas bajas , por tanto la distribución no sería simétrica y en consecuencia no se adapta al modelo normal .

Función de densidad : una variable continua X sigue una distribución normal de media μ y desviación típica σ , y se designa por $N(\mu, \sigma)$, si cumple que

$$f(x) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{1}{2} \frac{x-\mu}{\sigma}^2}$$

Podríamos comprobar que :

$$\begin{aligned} & \int_{-\infty}^{+\infty} x \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{1}{2} \frac{x-\mu}{\sigma}^2} dx \\ &= \mu \\ & \int_{-\infty}^{+\infty} (x-\mu)^2 \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{1}{2} \frac{x-\mu}{\sigma}^2} dx \\ &= \sigma^2 \end{aligned}$$

Para calcular los máximos y mínimos deberíamos hacer :

$$f(x) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{1}{2} \frac{x-\mu}{\sigma}^2}$$

$$f'(x) = -\frac{x-\mu}{\sigma}$$

$f(x)$, puesto que $f(x)$ nunca puede valer 0 entonces , si $x = \mu$
 $f'(x) = 0$

por lo que será un posible máximo o mínimo .

$$f''(x) = -\frac{1}{\sigma^2} \left(1 - \frac{(x-\mu)^2}{\sigma^2}\right) f(x)$$

luego $f''(\mu) < 0$ por lo que es hay un máximo en el punto $(\mu, \frac{1}{\sigma\sqrt{2\pi}})$

Conviene observar que cuando la desviación típica es elevada aumenta la dispersión y se hace menos puntiaguda la función ya que disminuye la altura del máximo . Por el contrario para valores pequeños de σ obtenemos una gráfica menos abierta y más alta .

Cuando $\mu = 0$ y $\sigma = 1$, $N(0,1)$ se dice que tenemos una distribución normal reducida , estandar o simplificada .

Función de distribución :
$$F(x) = \int_{-\infty}^x \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{1}{2} \frac{(x-\mu)^2}{\sigma^2}} dx$$

$= p(X \leq x)$

Distribución Normal Estándar $N(0,1)$: La distribución $N(0,1)$ se encuentra tabulada , lo cual permite un cálculo rápido de las probabilidades asociadas a esta distribución . Pero en general la media no suele ser 0 , ni la varianza 1 , por lo que se hace una transformación que se llama **tipificación de la variable** , que consiste en hacer el siguiente cambio de variable :

$$Z = \frac{x - \mu}{\sigma}$$

a partir del cual obtenemos una variable Z que si es $N(0,1)$ y que por lo tanto podemos calcular sus probabilidades .

$$F(x) = \int_{-\infty}^x \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{1}{2} \frac{(x-\mu)^2}{\sigma^2}} \sigma dz$$

Ejemplo : si tenemos $N(2,4)$ y queremos calcular $p(x < 7)$ entonces :

$$p(x < 7) = p\left(\frac{x-2}{2} < \frac{7-2}{2}\right)$$

$$= p(z < 5/4) = 0'1056$$

Manejo de tablas : pueden presentarse los siguientes casos :

$$p(z < 1'45) = 0'9265$$

$$p(z < -1'45) = 0'0735$$

$$p(1'25 < z < 2'57) = 0'1005$$

$$p(-2'57 < z < -1'25) = 0'1005$$

$$p(-0'53 < z < 2'46) = 0'695$$

Utilización conjunta de μ y σ
:

En $(\mu \pm \sigma)$
está el 68'26% de los datos ya que :

$$p\left(\mu - \sigma < X < \mu + \sigma\right) = p\left(\frac{\mu - \sigma - \mu}{\sigma} < Z < \frac{\mu + \sigma - \mu}{\sigma}\right) = p(-1 < Z < 1) = 0.6826$$

Análogamente se puede comprobar que en $(\mu \pm 2\sigma)$
está el 95'4% de los datos y en $(\mu \pm 3\sigma)$
está el 99'7% .

Ejemplo : El C.I. de los 5600 alumnos de una provincia se distribuyen $N(112,6)$. Calcular aproximadamente cuántos de ellos tienen :

- más de 1122800 alumnos.....la mitad de los alumnos
- entre 106 y 1183823 alumnoseste es el caso : $(\mu \pm \sigma)$
- entre 106 y 1121911 alumnos
- menos de 100128 alumnos
- más de 1307 alumnos
- entre 118 y 124761 alumnos

(ojo hay que multiplicar % obtenido en la tabla por 5600/100 , que sale de una regla de tres)

Aproximación normal para la binomial :

Cuando los valores a calcular para la binomial superan a los de las tablas para obtener un resultado aproximado se utiliza la distribución normal , es decir , la variable $y = \frac{x - np}{\sqrt{npq}}$

obedece a una distribución $N(0,1)$

El resultado es tanto más fiable cuanto mayor es el tamaño de la muestra n y cuanto más cerca está p de 0'5 .

Ejemplo : Se ha comprobado que la probabilidad de tener un individuo los ojos marrones es 0'6 . Sea X la variable aleatoria que representa el nº de individuos que tienen los ojos marrones de un grupo de 1100 . Calcular $p(X > 680)$ y $p(X = 680)$

$$p(X > 680) = 1 - p(X < 680) = 1 - p\left(Y < \frac{680 - 110 \cdot 0'6}{\sqrt{1100 \cdot 0'6 \cdot 0'4}}\right) = 1 - p(Y < 1'23) = 0'1093$$

$p(X = 680) = p(679.5 < X < 680.5)$ se debe hacer así puesto que en una variable continua no tiene sentido calcular probabilidades de valores puntuales .

COMBINATORIA

