

PROFESOR:

INGENIERO INDUSTRIAL UNEXPO

INTRODUCCION

El siguiente material es un bosquejo de los diferentes apuntes tomados de diversas fuentes bibliográficas específicas de la asignatura con el objeto de preparar las clases a dictar en la unidad curricular CALCULO DE PROBABILIDADES, correspondiente la especialidad de Producción Industrial y obviamente sigue las pautas dictadas por el programa correspondiente a la unidad curricular.

Los métodos estadísticos de acuerdo a su función, se dividen en métodos estadísticos descriptivos y métodos estadísticos inductivos; los primeros son aquellos que tratan de condensar o resumir todos los datos o características de una serie de valores para de esta forma describir varios aspectos de la serie. Los inductivos tratan de estimar las características de la población, universo o colectivo; a través del estudio de una o varias partes de esta población llamadas muestras.

GLOSARIO BASICO

- **Dato estadístico:** es toda información de carácter cuantitativo o cualitativo que permite obtener una idea del estado del fenómeno en estudio.
- **Población o Universo:** Es un conjunto finito o infinito de elementos que tiene características comunes.
- **Censo:** Significa abarcar todos los elementos integrantes de una población para definir las características estudiadas.
- **Parámetro:** Medida de resumen que describe una característica de un universo.
- **Muestra:** Parte de una población o subconjunto de un conjunto de elementos que resulta de la aplicación de algún proceso, generalmente selección deliberada, con el objeto de investigar las propiedades de la población o conjunto de los cuales proviene.
- **Muestreo:** Proceso mediante el cual se obtienen una o más muestras representativas de un universo. El muestreo lleva implícito las siguientes ventajas: Economía en la realización de la investigación y la rapidez en la obtención de resultados.
- **Muestra aleatoria:** Muestra tomada al azar. Donde todos los elementos de la población tienen la misma posibilidad de ser seleccionados.
- **Teoría del muestreo.** Es el estudio de la relación existente entre una población y las muestras tomadas de ella. Comprende aspectos como la representatividad de la muestra, el tamaño y los tipos y conveniencia del muestreo.
- **Estadístico:** Medida que resume y describe una característica de una muestra.
- **Variable:** Es un símbolo, generalmente una letra, que puede tomar un conjunto de valores prefijados llamado dominio de esa variable.
- **Función:** Es una regla o relación que asigna a cada valor de una variable independiente X un valor por, correspondencia, de una variable dependiente Y . Se escribe $Y=f(x)$, se lee Y es función de X .
- **Variable independiente:** Es aquella a la cual se le puede asignar un valor cualquiera dentro de su dominio. Es decir si queremos establecer una relación entre el peso y la edad de los niños menores de 6 años se puede decir entonces que la edad de los niños es la variable independiente y puede tomar valores entre 0 y 6.

- **Variable dependiente:** Es aquella cuyo valor depende del valor asignado a la variable independiente. En el caso anterior el peso de los niños es la variable dependiente.

INSTITUTO UNIVERSITARIO DE TECNOLOGIA INDUSTRIAL

CURSO INTENSIVO FEB. 99. ASIGNATURA: CALCULO DE PROBABILIDADES

PROFESOR: NESTOR GUERRA

INGENIERO INDUSTRIAL UNEXPO LUIS CABALLERO MEJIAS 1989

DESCRIPCION DE DATOS ESTADISTICOS

Medidas de Ubicación en los Conjuntos de Datos:

Una medida de ubicación es un valor que se calcula para un grupo de datos y que se utiliza para describir los datos en alguna forma. Generalmente se busca que el valor sea representativo de todos los valores del grupo y por lo tanto se desea un estadístico de tendencia central. Existen diversos estadísticos de tendencia central pero para el alcance del curso estudiaremos cuatro de ellos: La Media Simple, La Media Ponderada, La Mediana y La Moda.

La Media Aritmética: También llamada promedio simple se define matemáticamente como el cociente entre la suma de una serie de valores y el número de valores de la serie.

$\bar{X} = \sum X_i / N$. Siendo N el número de datos. Este estadístico nos permite conocer el valor alrededor del cual se presentan los valores de una serie.

Para datos agrupados en distribuciones de frecuencia se utiliza el punto medio de cada clase para el cálculo de la media aritmética. La cual se obtiene a través de esta expresión.

$\bar{X} = \sum (F_i X_i) / \sum F_i$, en donde X_i es el punto medio de clase y F_i es la frecuencia absoluta de cada clase. En este caso se puede decir que $\sum F_i = N$, siendo N el número de datos.

La Media Ponderada: También conocido como promedio ponderado es una media aritmética en la cual cada valor se pondera de acuerdo a su importancia en el grupo total. Esta importancia se determina a través de un valor k que puede estar sujeto a consideraciones subjetivas del investigador. La expresión matemática para el cálculo de la media ponderada es: $\bar{X}_p = \sum (k_i X_i) / \sum k_i$, siendo k_i el valor de ponderación correspondiente al dato X_i .

La Mediana: La mediana de un grupo de valores es el valor del ítem medio cuando todos los ítems del grupo se han dispuestos en orden ascendente o descendente, en términos de valor. Para un grupo con un número par de elementos se supone que la mediana está en la posición intermedia entre dos valores adyacentes al medio. Cuando trabajamos con datos agrupados en distribuciones de frecuencia, debemos determinar primero la clase que contiene el valor de la mediana (aquella cuya frecuencia absoluta acumulada iguala o excede a la mitad del número total de observaciones). Una vez identificada esta clase se procede a interpolar a través de la fórmula:

$$\text{Med} = CL + \{ (N/2 - \sum f_i - 1) / f_i \} I$$

Siendo CL = frontera inferior de la clase que contiene a la mediana

N = Número total de datos u observaciones en la distribución de frecuencia

$\sum f_i - 1$ = Frecuencia acumulada en la clase precedente a la clase que contiene la mediana

f_i = Frecuencia absoluta en la clase que contiene la mediana

I = Tamaño del intervalo de clase

La Moda: La moda es el valor que ocurre mas frecuentemente en un conjunto de valores. Dicha distribución se describe como unimodal. Para los conjuntos pequeños de valores en los cuales no se repiten valores medidos, no hay moda. Puede darse el caso de que una distribución de valores tenga mas de una moda. En este caso hablamos de distribuciones multimodales.

Para los datos agrupados en una distribución de frecuencia, con intervalos de clase iguales, se determina primero la clase modal (aquella que contiene el valor de la moda), identificada con el numero mayor de observaciones (mayor frecuencia absoluta). Después interpolamos a través de la expresión:

$$Mo = CL + \left\{ \frac{d_1}{d_1 + d_2} \right\} I$$

En donde

CL = frontera inferior de la clase que contiene a la moda

d_1 = diferencia entre la frecuencia de la clase modal y la frecuencia de la clase precedente

d_2 = diferencia entre la frecuencia de la clase modal y la frecuencia de la clase siguiente

I = tamaño del intervalo de clase.

Cuartiles, Deciles y Percentiles: Los cuartiles, los deciles y los percentiles se asemejan mucha a la mediana porque también subdividen una distribución de mediciones de acuerdo con la proporción de frecuencias observadas. Mientras la mediana divide una distribución en dos partes iguales, los cuartiles la dividen en cuatro cuartos, los deciles la dividen en diez decimos y los percentiles la dividen en cien partes.

Para los datos agrupados antes de usar la formula se debe determinar primero la clase apropiada que contenga el punto de interés luego se hace la interpolación:

$$\text{Cuartil } w = CL + \left\{ \frac{(wN/4 - \sum f_i - 1)}{f_i} \right\} I$$

$$\text{Decil } w = CL + \left\{ \frac{(wN/10 - \sum f_i - 1)}{f_i} \right\} I$$

$$\text{Percentil } w = CL + \left\{ \frac{(wN/100 - \sum f_i - 1)}{f_i} \right\} I$$

Donde

w es el numero de cuartil ($0 \leq w \leq 4$), decil ($0 \leq w \leq 10$) o percentil ($0 \leq w \leq 100$) que se quiere calcular

el resto de los símbolos tienen el mismo significado que en la formula de la mediana.

Medida de Dispersión en los Conjuntos de Datos

Las medidas de dispersión describen un grupo de valores en función de la variación o dispersión de los ítems incluidos dentro de ese grupo. Existen varias técnicas para medir el grado de dispersión de un grupo de datos en este curso incluiremos El Rango, La Desviación Promedio, La Desviación estándar y El Coeficiente de

Variación.

El Rango: Es la diferencia entre el valor mas alto (VM) y el mas bajo (Vm) de los valores de una serie que no se han agrupado en una distribución de frecuencia de esta manera:

$$R = VM - Vm$$

Para los datos agrupados en una distribución de frecuencia, el rango se define como la diferencia entre el limite superior de la clase mas alta o ultima clase (VM) y el limite inferior de la clase mas baja o primera clase (Vm).

Desviación Promedio: Se basa en la diferencia entre cada valor del conjunto de datos y la media del grupo. Es la media de estas desviaciones la que se calcula. Se calcula la media de las sumas de los valores absolutos de las diferencias.

$$Dp = \sum |xi - X| / N$$

Para xi = valor de la serie

X = media aritmética de la serie

N = numero de datos.

Para datos agrupados en distribución de frecuencia la desviación promedio se calcula a partir de los puntos medio de clases (xi) y las frecuencias absolutas de clases (fi).

$$Dp = \sum (fi |xi - X|) / N$$

La Varianza y La desviación Estándar: La varianza es similar a la desviación promedio en cuanto a la base en la diferencia entre cada valor del conjunto de datos y la media del grupo, difiere de ella porque esas diferencias se elevan al cuadrado antes de sumarse. Para la varianza de la población se utiliza la letra griega sigma 2 la formula es.

$$\sigma^2 = \sum (xi - X)^2 / N$$

A diferencia de lo que sucede con otras muestras estadísticas que hemos analizado, la varianza para una muestra no es exactamente equivalente a la varianza de una población, en lo que al calculo se refiere. Mas bien, el denominador de la formula de varianza de la muestra es ligeramente diferente. En esencia en esta formula se incluye un factor de corrección, de manera que la varianza de la muestra es un estimador no sesgado (un estimador no sesgado es un estadístico de muestra que tiene un valor esperado igual al parámetro que va a ser estimado) de la varianza de la población y su formula es

$$s^2 = \sum (xi - X)^2 / (n-1)$$

Para datos agrupados en distribución de frecuencia las expresiones son:

$$\sigma^2 = \sum fi (xi - X)^2 / N$$

$$s^2 = \sum fi (xi - X)^2 / (n-1)$$

La desviación estándar es la raíz cuadrada del valor de la varianza.

Coeficiente de Variación: El coeficiente de variación V, indica la magnitud relativa de la desviación estándar comparada con la media de la distribución de mediciones:

$$V = s/X$$

El coeficiente de variación es útil cuando tenemos que comparar la variabilidad de dos conjuntos de datos en relación con el nivel general de valores en cada conjunto.

Ejercicios

Una muestra de 20 obreros de producción de una pequeña compañía gana los siguientes salarios durante una semana cualquiera: 140, 165, 240, 140, 140, 155, 140, 155, 140, 165, 140, 180, 180, 140, 190, 200, 140, 205, 225, 230.

Determine: La distribución de frecuencia adecuada, la media, la moda, la mediana, la desviación estándar, la desviación promedia, ¿Cuánto gana el 37% de los obreros?

Si comparamos los datos anteriores con los de otro departamento donde elegimos a 30 obreros con los siguientes salarios. 123, 234, 142, 165, 133, 133, 112, 200, 203, 205, 213, 222, 345, 134, 145, 123, 156, 267, 289, 234, 123, 245, 300, 305, 234 256, 278, 256, 241, 245. Determine además de todo lo anterior ¿Cuál de los dos departamentos tiene un nivel de salario mas uniforme?

INSTITUTO UNIVERSITARIO DE TECNOLOGIA INDUSTRIAL

CURSO INTENSIVO FEB. 99. ASIGNATURA: CALCULO DE PROBABILIDADES

PROFESOR: NESTOR GUERRA

INGENIERO INDUSTRIAL UNEXPO LUIS CABALLERO MEJIAS 1989

INTRODUCCION A LA TEORIA DE PROBABILIDADES

El concepto de probabilidad nace con el deseo del hombre de conocer con certeza los eventos futuros. Es por ello que el estudio de probabilidades surge como una herramienta utilizada por los nobles para ganar en los juegos y pasatiempos de la época. El desarrollo de estas herramientas fue asignado a los matemáticos de la corte. Con el tiempo estas técnicas matemáticas se perfeccionaron y encontraron otros usos muy diferentes para la que fueron creadas. Actualmente se continuo con el estudio de nuevas metodologías que permitan maximizar el uso de la computación en el estudio de las probabilidades disminuyendo, de este modo, los márgenes de error en los cálculos

A través de la historia se han desarrollado tres enfoques conceptuales diferentes para definir la probabilidad y determinar los valores de probabilidad:

El enfoque clásico: Dice que si hay x posibles resultados favorables a la ocurrencia de un evento A y z posibles resultados desfavorables a la ocurrencia de A, y todos los resultados son igualmente posibles y mutuamente excluyente (no pueden ocurrir los dos al mismo tiempo), entonces la probabilidad de que ocurra A es:

$$P(A) = \frac{x}{x+z}$$

$$(x+z)$$

El enfoque clásico de la probabilidad se basa en la suposición de que cada resultado sea igualmente posible. Este enfoque es llamado enfoque a priori porque permite, (en caso de que pueda aplicarse) calcular el valor de probabilidad antes de observar cualquier evento de muestra.

Ejemplo: Si tenemos en una caja 15 piedras verdes y 9 piedras rojas. La probabilidad de sacar una piedra roja en un intento es:

$$P(A) = \frac{9}{9+15} = 0.375 \text{ o } 37.5\%$$

9+15

El enfoque de frecuencia relativa: También llamado Enfoque Empírico, determina la probabilidad sobre la base de la proporción de veces que ocurre un evento favorable en un número de observaciones. En este enfoque no se utiliza la suposición previa de aleatoriedad. Porque la determinación de los valores de probabilidad se basa en la observación y recopilación de datos.

Ejemplo: Se ha observado que 9 de cada 50 vehículos que pasan por una esquina no tienen cinturón de seguridad. Si un vigilante de tránsito se para en esa misma esquina una vez cualquiera ¿Cuál será la probabilidad de que detenga un vehículo sin cinturón de seguridad?

$$P(A) = \frac{9}{50} = 0.18 \text{ o } 18\%$$

50

Tanto el enfoque clásico como el enfoque empírico conducen a valores objetivos de probabilidad, en el sentido de que los valores de probabilidad indican al largo plazo la tasa relativa de ocurrencia del evento.

El enfoque subjetivo: Dice que la probabilidad de ocurrencia de un evento es el grado de creencia por parte de un individuo de que un evento ocurra, basado en toda la evidencia a su disposición. Bajo esta premisa se puede decir que este enfoque es adecuado cuando solo hay una oportunidad de ocurrencia del evento. Es decir, que el evento ocurrirá o no ocurrirá esa sola vez. El valor de probabilidad bajo este enfoque es un juicio personal.

El valor de la probabilidad: El valor más pequeño que puede tener la probabilidad de ocurrencia de un evento es igual a 0, el cual indica que el evento es imposible, y el valor mayor es 1, que indica que el evento ciertamente ocurrirá. Entonces si decimos que $P(A)$ es la probabilidad de ocurrencia de un evento A y $P(A')$ la probabilidad de no-ocurrencia de A, tenemos que:

$$0 \leq P(A) \leq 1$$

$$P(A) + P(A') = 1$$

Eventos mutuamente excluyentes y eventos no excluyentes: Dos o más eventos son mutuamente excluyentes o disjuntos, si no pueden ocurrir simultáneamente. Es decir, la ocurrencia de un evento impide automáticamente la ocurrencia del otro evento (o eventos).

Ejemplo: Al lanzar una moneda solo puede ocurrir que salga cara o sello pero no los dos a la vez, esto quiere decir que estos eventos son excluyentes.

Dos o más eventos son no excluyentes, o conjuntos, cuando es posible que ocurran ambos. Esto no indica que necesariamente deban ocurrir estos eventos en forma simultánea.

Ejemplo: Si consideramos en un juego de domino sacar al menos un blanco y un seis, estos eventos son no excluyentes porque puede ocurrir que salga el seis blanco.

Reglas de la Adición

La Regla de la Adición expresa que: la probabilidad de ocurrencia de al menos dos sucesos A y B es igual a:

$$P(A \text{ o } B) = P(A) \cup P(B) = P(A) + P(B) \text{ si A y B son mutuamente excluyentes}$$

$$P(A \text{ o } B) = P(A) + P(B) - P(A \text{ y } B) \text{ si A y B son no excluyentes}$$

Siendo: $P(A)$ = probabilidad de ocurrencia del evento A

$P(B)$ = probabilidad de ocurrencia del evento B

$P(A \text{ y } B)$ = probabilidad de ocurrencia simultanea de los eventos A y B

Eventos Independientes: Dos o más eventos son independientes cuando la ocurrencia o no-ocurrencia de un evento no tiene efecto sobre la probabilidad de ocurrencia del otro evento (o eventos). Un caso típico de eventos independiente es el muestreo con reposición, es decir, una vez tomada la muestra se regresa de nuevo a la población donde se obtuvo.

Ejemplo: lanzar al aire dos veces una moneda son eventos independientes por que el resultado del primer evento no afecta sobre las probabilidades efectivas de que ocurra cara o sello, en el segundo lanzamiento.

Eventos dependientes: Dos o más eventos serán dependientes cuando la ocurrencia o no-ocurrencia de uno de ellos afecta la probabilidad de ocurrencia del otro (o otros). Cuando tenemos este caso, empleamos entonces, el concepto de probabilidad condicional para denominar la probabilidad del evento relacionado. La expresión $P(A|B)$ indica la probabilidad de ocurrencia del evento A sí el evento B ya ocurrió. Se debe tener claro que $A|B$ no es una fracción.

$$P(A|B) = P(A \text{ y } B)/P(B) \text{ o } P(B|A) = P(A \text{ y } B)/P(A)$$

Reglas de Multiplicación

Se relacionan con la determinación de la ocurrencia de conjunta de dos o más eventos. Es decir la intersección entre los conjuntos de los posibles valores de A y los valores de B, esto quiere decir que la probabilidad de que ocurran conjuntamente los eventos A y B es:

$$P(A \text{ y } B) = P(A \text{ B}) = P(A)P(B) \text{ si A y B son independientes}$$

$$P(A \text{ y } B) = P(A \text{ B}) = P(A)P(B|A) \text{ si A y B son dependientes}$$

$$P(A \text{ y } B) = P(A \text{ B}) = P(B)P(A|B) \text{ si A y B son dependientes}$$

INSTITUTO UNIVERSITARIO DE TECNOLOGIA INDUSTRIAL

CURSO INTENSIVO FEB. 99. ASIGNATURA: CALCULO DE PROBABILIDADES

PROFESOR: NESTOR GUERRA

INGENIERO INDUSTRIAL UNEXPO LUIS CABALLERO MEJIAS 1989

DISTRIBUCIONES DE VARIABLES ALEATORIAS DISCRETAS

Variable aleatoria:

Hasta ahora hemos visto el desarrollo de una idea de un intento que resulta en la aparición aleatoria de un estado del conjunto $E_1, E_2, E_3, E_4, \dots, E_n$, y se introdujo la noción de espacio de muestra o espacio de muestra como un modelo conveniente de los resultados. Aquí introducimos el concepto de **variable aleatoria**, la cual es simplemente un conjunto de números $X_1, X_2, X_3, \dots, X_n$, uno para cada estado de manera que el resultado de un intento no es solamente el estado E_i , sino también el número X_i de interés.

Ejemplo: si el intento es el lanzamiento de dos dados, una variable aleatoria puede definirse como la suma de los puntos obtenidos en los dados. Si denotamos esta variable aleatoria por z , entonces z tendrá el valor 2 para el estado $\{1,1\}$, 3 para el estado $\{1,2\}$ y así sucesivamente. Observe que a pesar de haber 36 puntos en el espacio de muestra, solamente hay 11 valores posibles para z , estos son con su respectiva función de frecuencia (relativa respecto al espacio de muestra):

Variable aleatoria z Función de frecuencia $P(z)$

- $1/36$
- $2/36$
- $3/36$
- $4/36$
- $5/36$
- $6/36$
- $5/36$
- $4/36$
- $3/36$
- $2/36$
- $1/36$

Total $36/36$

Se puede observar que la suma de las probabilidades individuales en cualquier función de frecuencia es 1, ello resulta del hecho de que uno y solo uno de los resultados posibles se materializa como resultado de un intento.

Pueden obtenerse muchas variables aleatorias en el mismo espacio de muestra. Si el intento es el lanzamiento de dos dados, podemos definir también la variable aleatoria z como el número menor de los dos que aparecen en el lanzamiento. En este caso z tendría los valores 1, 2, 3, 4, 5, 6. El evento " $z = 6$ " se presenta solamente en el punto de muestra (6,6) y tiene probabilidad $1/36$. El evento " $z = 4$ " se observa en los puntos (4,4), (4,5), (5,4), (4,6), (6,4) y tiene una probabilidad de $5/36$ y así sucesivamente. La función de frecuencia completa para esta variable aleatoria es:

Variable aleatoria z Función de frecuencia $P(z)$

- $11/36$
- $9/36$
- $7/36$
- $5/36$
- $3/36$

Total $36/36$

Generalmente una distribución de frecuencia de una variable aleatoria se caracteriza por dos estadísticos derivados: *su media y su varianza*. La descripción mediante estos números es una función de frecuencia de probabilidad, aunque incompleta, es valiosa en muchas de las aplicaciones.

Sea $x_i \{1, 2, 3, \dots, n\}$ los diversos valores posibles que puede tomar una variable aleatoria, y sea $P(x)$ la probabilidad de que la variable toma el valor x_i , entonces su media será:

$$\bar{X} = \sum x_i P(x_i)$$

$$s^2 = \sum x_i^2 P(x_i) - \bar{X}^2$$

La media de x es un valor que puede esperarse que tome x en un intento, y la varianza es una medida de la dispersión esperada de los valores que alcanza x , alrededor del valor esperado

Al igual que en la aplicación de las formulas para la probabilidad de eventos compuestos, el espacio de muestra fundamental no necesita estar en forma explícita para que se utilicen los conceptos de variable aleatoria y de distribución de probabilidad. La función de frecuencia de la variable aleatoria es entonces un conjunto de números no negativos.

$P(x_i), P(x_{ii}), P(x_{iii}), \dots, P(x_n)$ uno para cada x_i , tal que $\sum P(x_i) = 1$

Bajo esta premisa, como sucede en cualquier tratamiento de espacios de muestras, el numero de valores posible de la variable aleatoria x no tiene que ser finito. Es por ello que existen funciones de distribución de probabilidad para variables aleatorias discretas y para variables aleatorias continuas. En esta unidad estudiaremos las distribuciones de probabilidad para variables discretas o distribuciones de probabilidad discretas.

Distribuciones de Probabilidad Discretas: Son funciones de probabilidad en las cuales la variable aleatoria toma valores discretos, entre las más importantes tenemos: La Distribución de Bernoulli o Distribución Binomial, la Distribución Hipergeométrica, Distribución de Poisson

Distribución Binomial: La distribución binomial es una distribución de probabilidad discreta, aplicable cada vez que se suponga que un proceso de muestreo conforma un proceso de Bernoulli. Es decir que ocurra un proceso de muestreo en el cual:

1. – Hay dos resultados posibles mutuamente excluyentes en cada ensayo u observación.
2. – La serie de ensayos u observaciones constituyen eventos independientes.
3. – La probabilidad de éxito permanece constantede ensayo a ensayo, es decir el proceso es estacionario.

Para aplicar esta distribución al calculo de la probabilidad de obtener un numero dado de éxitos en una serie de experimentos en un proceso de Bernoulli, se requieren tres valores: el numero designado de éxitos (m), el numero de ensayos y observaciones (n); y la probabilidad de éxito en cada ensayo (p). Entonces la probabilidad de que ocurran m éxitos en un experimento de n ensayos es:

$$P(x = m) = {}^nC_m p^m (1-p)^{n-m}$$

Siendo nC_m el numero total de combinaciones posibles de m elementos en un conjunto de n elementos. En otras palabras $P(x = m) = \frac{n!}{m!(n-m)!} p^m (1-p)^{n-m}$

Ejemplo. La probabilidad de que un alumno apruebe la asignatura Calculo de Probabilidades es de 0,15. Si en un semestre intensivo se inscriben 15 alumnos ¿Cuál es la probabilidad de que aprueben 10 de ellos?

$$P(x = 10) = {}^{15}C_{10}(0,15)^{10}(0,85)^5 = \frac{10!}{10!(15-10)!}(0,15)^{10}(0,85)^5 = 7,68 * 10^{-6}$$

Generalmente existe un interés en la probabilidad acumulada de "**m** o más " éxitos o "**m** o menos" éxitos en **n** ensayos. En tal caso debemos tomar en cuenta que:

$$P(x < m) = P(x = 1) + P(x = 2) + P(x = 3) + \dots + P(x = m-1)$$

$$P(x > m) = P(x = m+1) + P(x = m+2) + P(x = m+3) + \dots + P(x = n)$$

$$P(x \leq m) = P(x = 1) + P(x = 2) + P(x = 3) + \dots + P(x = m)$$

$$P(x \geq m) = P(x = m) + P(x = m+1) + P(x = m+2) + \dots + P(x = n)$$

Supongamos que del ejemplo anterior se desea saber la probabilidad de que aprueben:

a.- al menos 5

b.- mas de 12

a.- la probabilidad de que aprueben al menos 5 es $P(x \leq 5)$ es decir que

$$P(x \leq 5) = P(x = 1) + P(x = 2) + P(x = 3) + P(x = 4) + P(x = 5)$$

$$P(x \leq 5) = 0,2312 + 0,2856 + 0,2184 + 0,1156 + 0,045 = 0,8958$$

b.- la probabilidad de que aprueben mas de 12 es $P(x > 12)$ es decir que

$$P(x > 12) = P(x = 13) + P(x = 14) + P(x = 15)$$

$$P(x > 12) = 1,47 * 10^{-9} + 3,722 * 10^{-11} + 4,38 * 10^{-13} = 1,507 * 10^{-9}$$

La esperanza matemática en una distribución binomial puede expresarse como

$$E(x) = \sum np$$

Y la varianza del numero esperado de éxitos se puede calcular directamente:

$$\text{Var}(x) = np(1-p)$$

Distribución Binomial expresada en Proporciones: En lugar de expresar la variable aleatoria como el numero de éxitos X , podemos designarla en términos de la proporción de éxitos, p , que es la relación entre el numero de éxitos y el numero de ensayos:

$$p = \frac{X}{n}$$

n

En tales casos la formula se modifica solo respecto de la definición de la proporción:

$$P(p = P/n) = nCpx(1-p)^{n-x}$$

Ejemplo: La probabilidad de que Juan pueda conquistar una chica es de 0,20. Si se seleccionan 5 chicas al azar, que se encontraran con Juan, ¿Cuál es la probabilidad la proporción de chicas interesadas en Juan sea exactamente 0,2?

$$P(p = 0,2 = 1/5) = 5C1(0,20)^1(0,80)^4 = 0,4096$$

Cuando la variable binomial se expresa como una proporción, la distribución es aun discreta y no continua. Solo pueden ocurrir las proporciones para las que el numero de éxitos X es un numero entero. El valor esperado para una distribución de probabilidad binomial expresada por proporciones es igual a la proporción de la población:

$$E(p) = p$$

La varianza de una proporción de éxitos para una distribución de probabilidad binomial es:

$$\text{Var}(p) = \frac{p(1-p)}{n}$$

N

Distribución Hipergeometrica: Cuando el muestreo se hace **sin reemplazo** de cada articulo muestreado tomado de una población finita de artículos, no se aplica el proceso de Bernoulli porque hay un cambio sistemático en la probabilidad de éxitos a medida que se retiran los ítems de la población. Es por ello que se utiliza la distribución de probabilidad Hipergeometrica por ser la mas apropiada.

Si X es el numero designado de éxitos, N es el numero total de ítems en la población, XT es el numero total de éxitos incluidos en la población y n es el numero de ítems de la muestra, la formula para determinar la probabilidad hipergeometrica es:

$$N - X_T$$

$$n - X$$

$$P(X|N, X_T, n) = \frac{(N - X_T)! (X_T)!}{N!} \cdot \frac{n!}{(n - X)! X!}$$

N

.n

Ejemplo: De seis estudiantes de Cálculos de Probabilidades, tres han cursado la materia tres veces o más. Si se escoge cuatro estudiantes del grupo de seis ¿cuál es la probabilidad de que dos hayan cursado la materia en mas de una oportunidad?

$$X_T = 3, X = 2, N = 6, n = 4$$

$$6 - 3 = 3$$

$$4 - 2 = 2$$

$$P(X=2|6,3,4) = \frac{(6-3)! (3)!}{6!} \cdot \frac{4!}{(4-2)! 2!} = 0,60$$

Observe que en esta distribución el valor de probabilidad se calcula determinando el número de combinaciones diferentes que incluirían dos alumnos con mayor índice de repitencia y dos nuevos con una relación total de combinaciones de cuatro alumnos de los seis. Cuando la población es grande y la muestra es relativamente pequeña, el hecho de que el muestreo se efectúe sin reemplazo tiene poco efecto sobre la probabilidad de éxito de cada ensayo.

Distribución de Poisson: Se utiliza para determinar la probabilidad de un número designado de éxitos cuando los eventos ocurren en un espectro continuo de tiempo y espacio. Tal proceso se denomina Proceso de Poisson, es semejante al proceso de Bernoulli excepto que los eventos ocurren en un espectro continuo en vez de ocurrir en ensayos u observaciones fijas. Por ejemplo la entrada de materiales a una celda de producción, la llegada de clientes a un servidor cualquiera, etc.

Solo se requiere un valor para determinar la probabilidad de un número designado de éxitos en un proceso de Poisson: el número promedio de éxitos para la dimensión específica de tiempo o espacio de interés. Este número promedio se representa generalmente por λ o $\bar{x}$. La expresión matemática de la distribución de Poisson es.

$$P(x|\lambda) = \frac{\lambda^x e^{-\lambda}}{x!}$$

Ejemplo: Un puesto de trabajo en una línea recibe un promedio de 4 productos por hora. ¿Cuál es la probabilidad de que reciba al menos 2 productos?

$$\lambda = 5, x \leq 2$$

$$P(x \leq 2 | \lambda = 5) = \frac{5^1 e^{-5}}{1!} + \frac{5^2 e^{-5}}{2!} = 0,1179$$

Puesto que el proceso de Poisson es estacionario, se concluye que la media del proceso es siempre proporcional a la longitud del espectro continuo de tiempo o espacio.

Aproximación de Poisson de Probabilidades Binomiales: Cuando el número de observaciones o ensayos n en un proceso de Bernoulli es muy grande, los cálculos son bastante tediosos. Mas aun, en general no se encuentran tablas de probabilidad con valores muy pequeños de p . En estos casos la distribución de Poisson es conveniente como una aproximación de probabilidades binomiales cuando n es grande y p o $(1-p)$ es pequeño. Empíricamente esta aproximación se puede hacer cuando $n \geq 30$, y $np < 5$. La media de la distribución de Poisson, utilizada para aproximar probabilidades binomiales es.

$$\lambda = np$$

INSTITUTO UNIVERSITARIO DE TECNOLOGIA INDUSTRIAL

CURSO INTENSIVO FEB. 99. ASIGNATURA: CALCULO DE PROBABILIDADES

PROFESOR: NESTOR GUERRA

INGENIERO INDUSTRIAL UNEXPO LUIS CABALLERO MEJIAS 1989

DISTRIBUCIONES DE VARIABLES ALEATORIAS CONTINUAS

Variables aleatorias continuas: A diferencia de una variable aleatoria discreta, una *variable aleatoria continua* es aquella que puede tener cualquier valor fraccional dentro de un rango definido de valores. Es decir, este tipo de variable se define en espacios de muestra con un numero de puntos infinitos no denumerable. Debido a las dificultades matemáticas inherentes, no desarrollaremos el concepto en detalle, pero retendremos la noción de un "intento" como conjunto de operaciones que dan por resultado algún valor definido pero aleatorio de x . Entonces si a es algún valor particular en el dominio de definición de x , el evento "el valor de x que resulta de un intento es menor o igual que a " tiene una probabilidad, definida $P\{x \leq a$, para cada valor a en el dominio de x . De esta manera, para una distribuciones de probabilidad, no se pueden enumerar todos los valores posibles para una variable aleatoria continua x junto con un valor de probabilidad correspondiente. En este caso el enfoque más conveniente es elaborar una función de densidad de probabilidad, basada en la función matemática correspondiente. La proporción del arrea incluida entre dos puntos cualesquiera por debajo de la curva de probabilidad identifica la probabilidad de que una variable continua aleatoriamente seleccionada tenga un valor entre esos puntos.

Ejemplo:

Supongamos que $U(a) = P\{x \leq a\}$, decimos que $U(a)$ es una función de distribución acumulativa de x , puesto que $P\{x \leq a\}$ esta definida para cada valor de a en el dominio de x . Si a y b son dos valores en el recorrido de x , con $a < b$, podemos denotar mediante $P\{a < x \leq b\}$ la probabilidad del evento "el valor de x que resulta de un intento es mayor que a , pero menor o igual que b . Podemos expresar $P\{a < x \leq b\}$ en términos de la función de distribución acumulativa, si vemos primero que

$$P\{x \leq b\} = P\{x \leq a\} + P\{a < x \leq b\}$$

Entonces

$$P\{a < x \leq b\} = U(b) - U(a)$$

oo b

$$U(a) = \int_{-\infty}^a f(x) dx; P\{a < x \leq b\} = \int_a^b f(x) dx$$

$-\infty$ a

Siendo $f(x)$ una función de distribución acumulativa diferenciable en todos los puntos de recorrido de x . Estas funciones así definidas se denominan también "función de densidad de probabilidad de x ".

Estas funciones deben cumplir lo siguiente:

•
oo

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

$-\infty$

- Si a es un valor fijo de x $P\{x = a\} = 0$
- $F(x) dx$, puede considerarse como la probabilidad de que la variable aleatoria tome algún valor entre x y $x + dx$.

Como en el caso discreto, la distribución de probabilidad de una variable aleatoria con recorrido sobre un continuo puede describirse (de manera incompleta, desde luego) mediante su media y su varianza:

00

$$\text{Media} = \int_{-\infty}^{\infty} xf(x) dx$$

−∞

∞

$$\text{varianza} = \int_{-\infty}^{\infty} (x - \mu)^2 f(x) dx, \text{ siendo } \mu \text{ igual a la media}$$

−∞

Distribución de probabilidad normal: Es una distribución de probabilidad continua que es tanto simétrica como mesocurtica. La curva que representa la distribución de probabilidad normal se describe generalmente como en forma de campana. Esta distribución es importante en inferencia estadística por tres razones diferentes:

- Se sabe que las medidas producidas en muchos procesos aleatorios siguen esta distribución.
- Las probabilidades normales pueden utilizarse generalmente para aproximar otras distribuciones de probabilidad, tales como las distribuciones binomial y de Poisson.
- Las distribuciones estadísticas tales como la media de la muestra y la proporción de la muestra, siguen a menudo la distribución normal, sin tener en cuenta la distribución de la población

Los valores de los parámetros de la distribución de probabilidad normal son $\mu = 0$ y $\sigma = 1$. Cualquier conjunto de valores X normalmente distribuido pueden convertirse en valores normales estándar z por medio de la formula:

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

haciendo posible el uso de la tabla de proporciones de área y hace innecesario el uso de la ecuación de la función de densidad de cualquier distribución normal dada.

Para aproximar las distribuciones discretas binomial y de Poisson se debe hacer:

Binomial	$\mu = np$	$\sigma = \sqrt{np(1-p)}$	Si $n \geq 30$ $np \geq 5$ $n(1-p) \geq 5$
Poisson	$\mu = \lambda$	$\sigma = \sqrt{\lambda}$	$\lambda \geq 10$

Distribución de probabilidad exponencial. Si en el contexto de un proceso de Poisson ocurren eventos o éxitos en un espectro continuo de tiempo y espacio. Entonces la longitud del espacio o tiempo entre eventos sucesivos sigue una distribución de probabilidad exponencial. Puesto que el tiempo y el espacio son un espectro continuo, esta es una distribución continua. En caso de este tipo de distribución no vale la pena preguntarse ¿cuál es la probabilidad de que el primer pedido de servicio se haga exactamente de aquí a un minuto?. Mas bien debemos asignar un intervalo dentro del cual el evento puede ocurrir, preguntándonos, ¿cuál es la probabilidad de que el primer pedido se produzca en el próximo minuto?.

Dado que el proceso de Poisson es estacionario, la distribución exponencial se aplica ya sea cuando estamos interesados en el tiempo (o espacio) hasta el primer evento, el tiempo entre dos eventos sucesivos, o el tiempo hasta que ocurra el primer evento después de cualquier punto aleatoriamente seleccionado.

Donde δ es la cifra media de ocurrencias para el intervalo de interés, la probabilidad exponencial de que el primer evento ocurra dentro del intervalo designado de tiempo o espacio es.

$$P(T \leq t) = 1 - e^{-\delta}$$

De manera que la probabilidad exponencial de que el primer evento no ocurra dentro del intervalo designado de tiempo o espacio es:

$$P(T > t) = e^{-\delta}$$

Ejemplo δ Un departamento de mantenimiento recibe un promedio de 5 llamadas por hora. Comenzando en un momento aleatoriamente seleccionado, la probabilidad de que una llamada llegue dentro de media hora es:

Promedio 5 por hora, como el intervalo es media hora tenemos que $\delta = 2,5/\text{media hora}$.

$$P(T \leq 30 \text{ min.}) = 1 - e^{-5} = 1 - 0,08208 = 0,91792$$