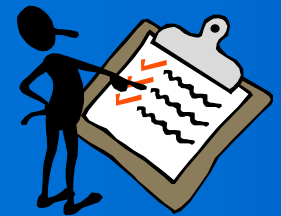

Inteligencia Artificial

Técnicas de clasificación

Prof. Dra. Silvia Schiaffino
ISISTAN - CONICET

Clasificación: Agenda

- Concepto
 - Clasificación
 - Predicción
 - Evaluación
 - Árboles de Decisión
 - Construcción
 - Uso
 - Poda
 - Clasificador Bayesiano
 - Ejemplos
-



Supervisado vs. No Supervisado

- Aprendizaje Supervisado: **Clasificación**
 - Conocemos las clases y el número de clases
- Aprendizaje No Supervisado: **Clustering**
 - No conocemos las clases y podemos no conocer el número de clases

- El objetivo de la clasificación de datos es organizar y categorizar los datos en clases diferentes
 - Se **crea un modelo** basándose en la distribución de los datos
 - El **modelo es luego usado** para clasificar nuevos datos
 - Dado el modelo, se puede predecir la clase de un nuevo dato
- Clasificación = predicción para valores discretos y nominales



- El objetivo de la predicción es deducir el valor de un atributo basándose en los valores de otros atributos
 - Se crea un modelo basándose en la distribución de los datos
 - El modelo se usa para predecir valores futuros o desconocidos

Si se deduce un valor discreto → Clasificación

Si se deduce un valor continuo → Predicción
(Regresión)



Definición de clasificación

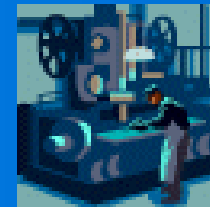
Dada una base de datos $D = \{t_1, t_2, \dots, t_n\}$ de tuplas o registros (individuos) y un conjunto de clases $C = \{C_1, C_2, \dots, C_m\}$, el problema de la clasificación es encontrar una función

$f: D \rightarrow C$ tal que cada t_i es asignada una clase C_j .

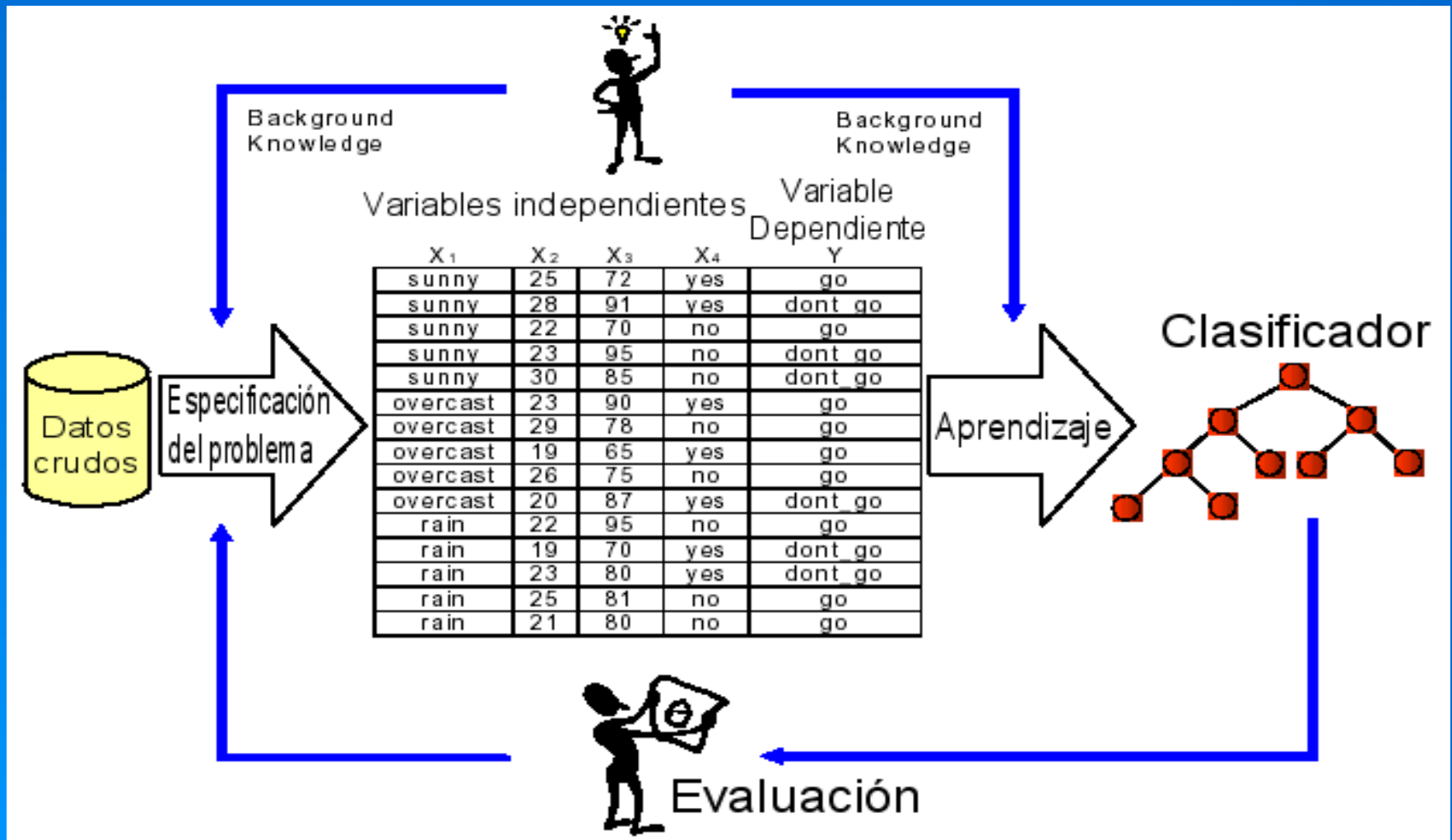
$f: D \rightarrow C$ podría ser una Red Neuronal, un Árbol de Decisión, un clasificador Bayesiano...

Aplicaciones

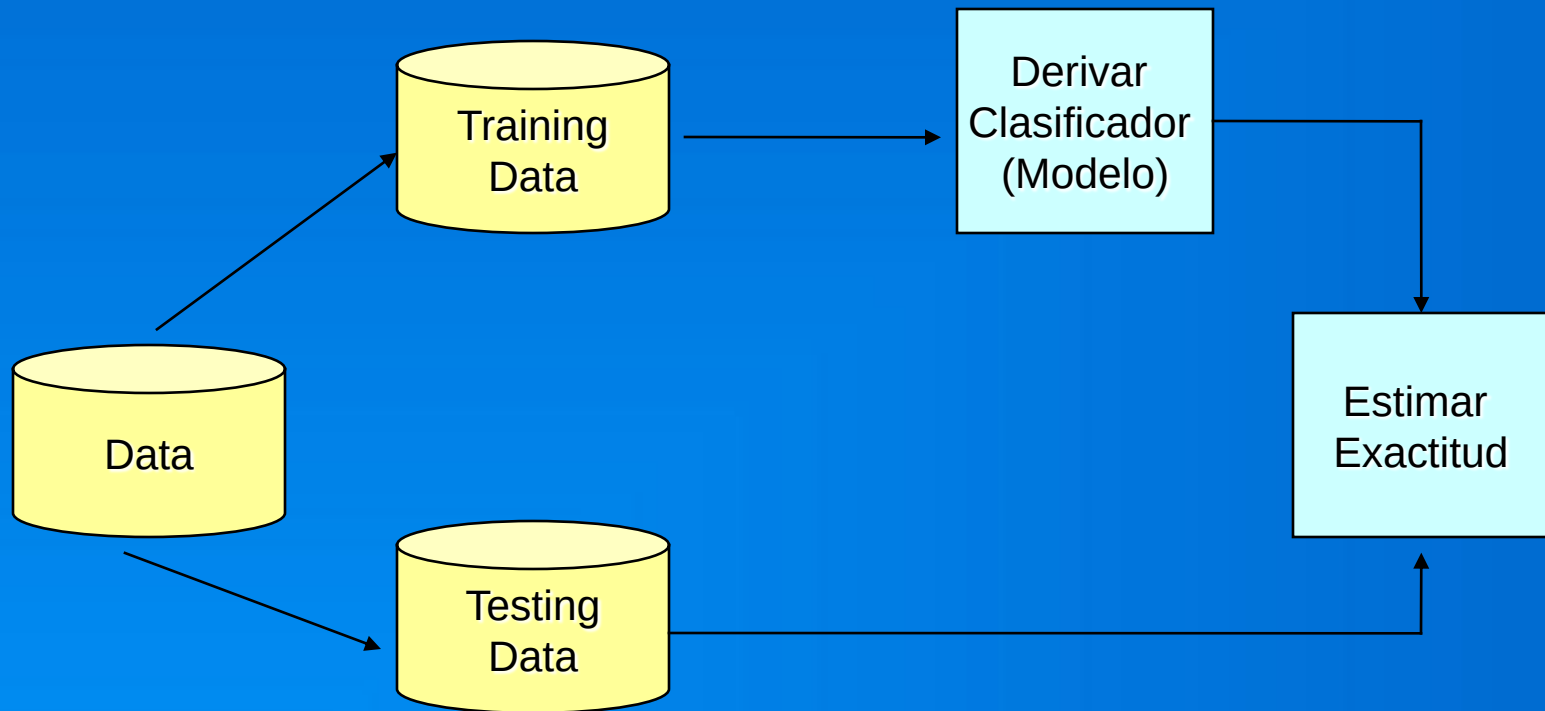
- Aprobación de créditos
- Marketing
- Diagnóstico médico
- Identificación de partes defectuosas en manufactura
- Zonificación (crímenes)
- Clasificación de e-mail
- Clasificación de documentos
- ...



Proceso de clasificación



Etapas: Terminología



Etapas: Construcción del modelo - Aprendizaje

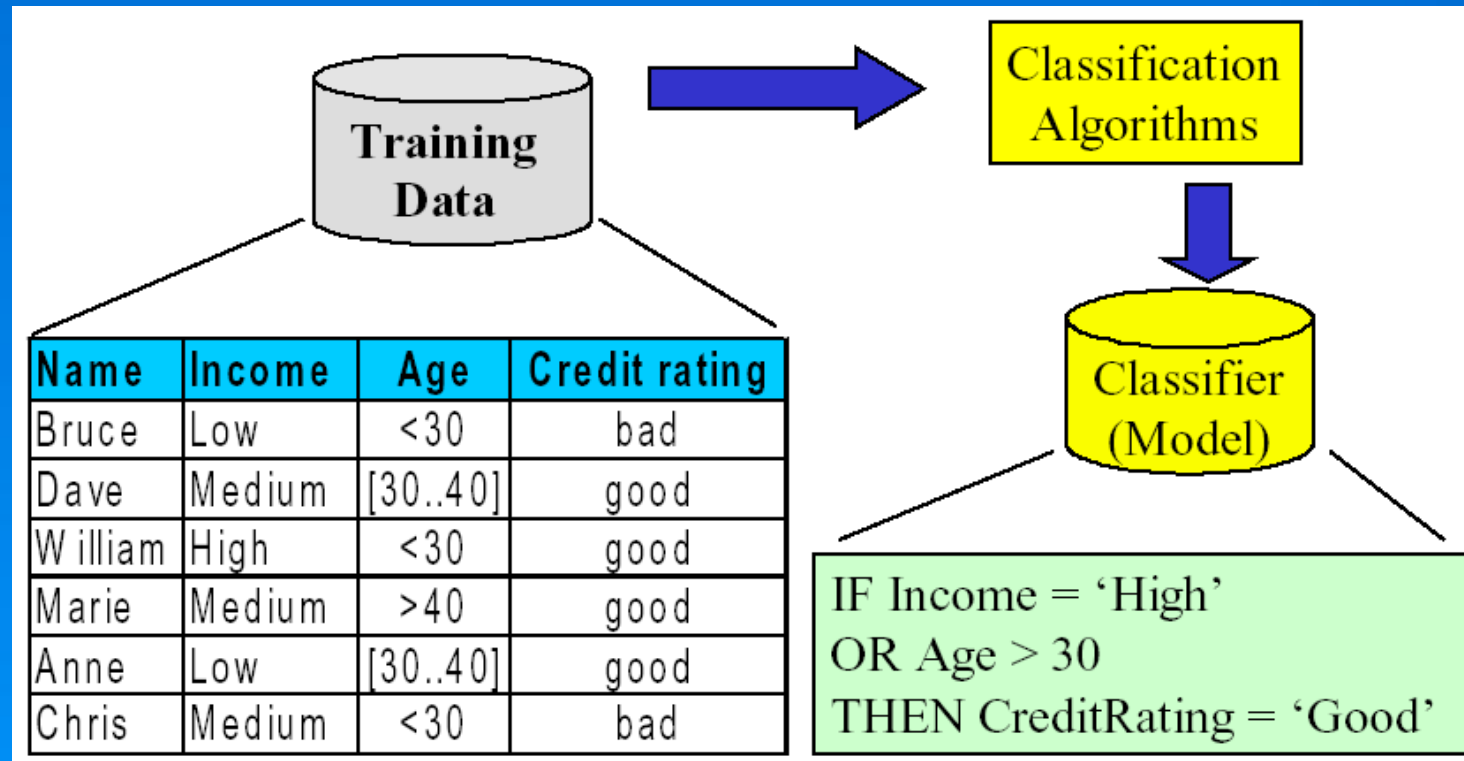
Cada tupla se supone que pertenece a una clase predefinida, dada por uno de los atributos, llamada **etiqueta de clase**

El conjunto de todas las tuplas usadas para la construcción del modelo se llama **conjunto de entrenamiento**

El modelo se representa mediante alguna de las siguientes formas:

- **Reglas de clasificación** (sentencias IF-THEN)
- **Árbol de decisión**
- **Fórmulas matemáticas**
- ...

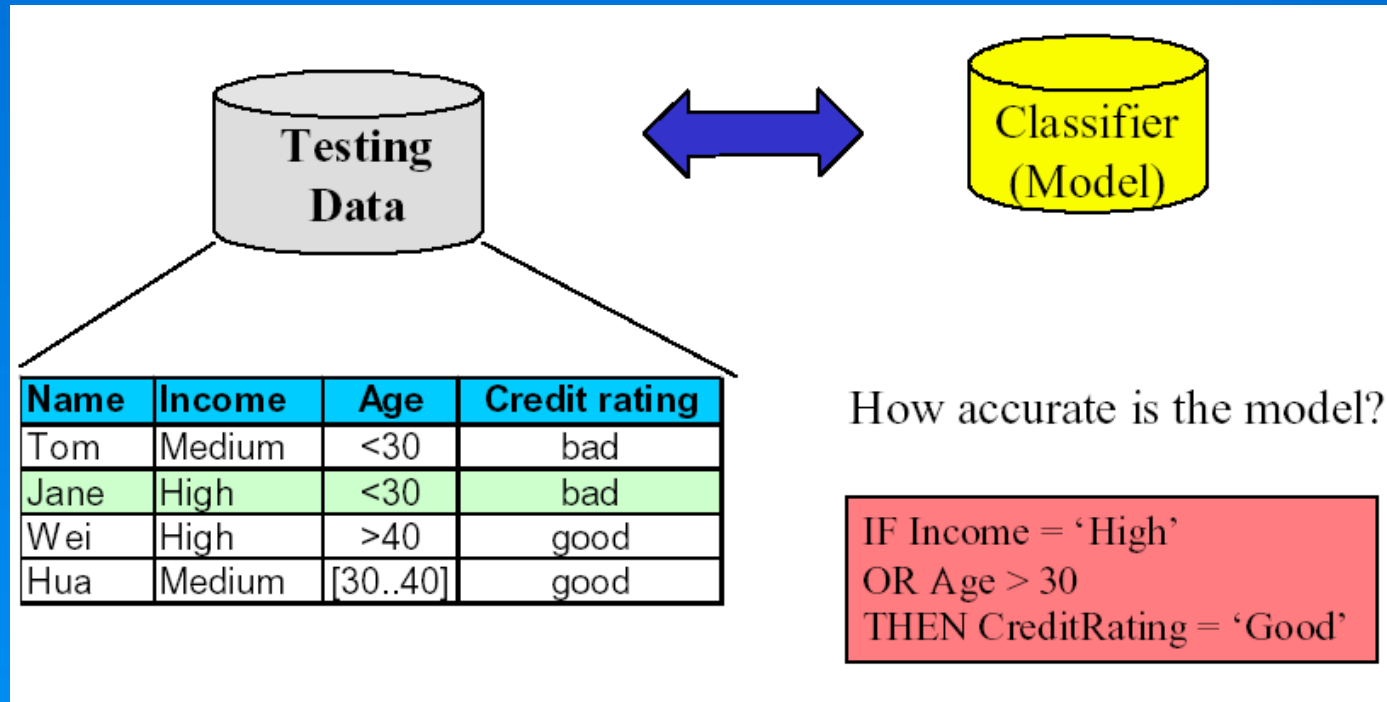
Aprendizaje



Etapas: Evaluación del modelo

- Se estima la exactitud del modelo basándose en un conjunto de test
 - Se compara la etiqueta conocida de una muestra de testeo con el resultado de aplicar el modelo de clasificación
 - *Accuracy rate* es el porcentaje de muestras del conjunto de test que son correctamente clasificadas por el modelo
 - El conjunto de test es independiente del conjunto de entrenamiento (método holdout)

Evaluación de Exactitud



Etapas: Evaluación del modelo

- **Holdout**
 - Los datos se particionan aleatoriamente en 2 conjuntos independientes: training set (2/3 de los datos) y test set (1/3 de los datos)
- **K-fold cross validation**
 - Datos iniciales particionados en k subconjuntos mutuamente excluyentes de aproximadamente igual tamaño. Se hace training y testing k veces, se calcula la exactitud promediando los resultados.
- **Stratified cross-validation**
 - Los subconjuntos son armados de tal manera que la distribución de clase de los ejemplos en cada uno es aproximadamente igual a la que tienen los datos iniciales

- Tasa/Coeficiente de Error

$$ce(h) = \frac{1}{n} \sum_{i=1}^n \|y_i \neq h(x_i)\|$$

- Precisión/Coeficiente de acierto

$$ca(h) = 1 - ce(h)$$

Evaluación del modelo: Matriz de Confusión

Etiqueta de clase	Predicciones C1	Predicciones C2	...	Predicciones Ck
Verdaderos C1	$M(C1,C1)$	$M(C1,C2)$	...	$M(C1,Ck)$
Verdaderos C2	$M(C2,C1)$	$M(C2,C2)$	...	$M(C2,Ck)$
...	...	...	...	...
Verdaderos Ck	$M(Ck,C1)$	$M(Ck,C2)$	...	$M(Ck,Ck)$



Clasificador ideal

- $M(C_i, C_i)$ Casos correctamente clasificados
- $M(C_i, C_j)$ $i \neq j$ Errores de clasificación

	C1	C2	...	Ck
C1	$M(C1, C1)$	0	...	0
C2	0	$M(C2, C2)$	...	0
...	...	...	...	0
Ck	0	0	...	$M(Ck, Ck)$

Evaluación del Modelo (Documentos)

- Precision

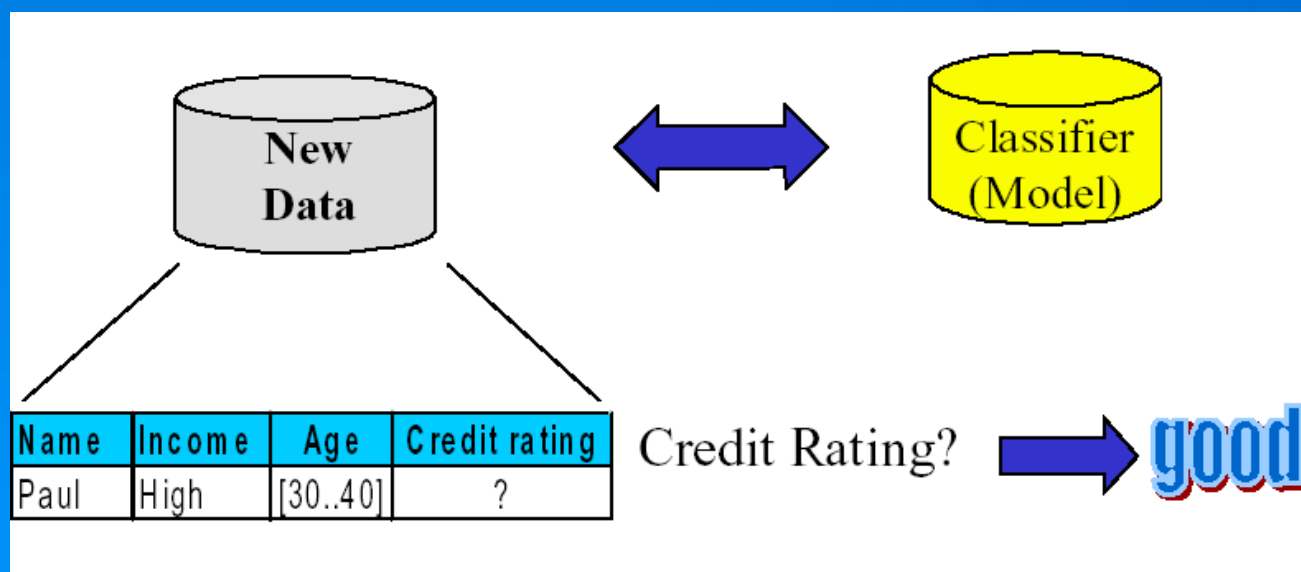
- De la cantidad de veces que se predijo una clase, cuántas fueron correctas?

- Recall

- Se encontraron todos los ejemplos que pertenecen a la clase?

Etapas: Uso del modelo - Clasificación

- El modelo se utiliza para clasificar nuevos objetos
 - Dar una etiqueta de clase a una nueva tupla
 - Predecir el valor de un atributo



Métodos de clasificación

- Inducción de árboles de decisión
 - Redes Neuronales
 - Clasificador Bayesiano
 - Clasificación basada en asociación
 - Vecinos más cercanos
 - Razonamiento Basado en Casos
 - Algoritmos Genéticos
 - Conjuntos difusos
 - Support Vector Machines
 - ...
-

Evaluación y comparación de métodos de clasificación

- **Exactitud de predicción**
 - Habilidad del modelo de predecir correctamente la etiqueta de clase de nuevos ejemplos
- **Velocidad**
 - Tiempo para construir el modelo
 - Tiempo para usar el modelo
- **Robustez**
 - Manejo de valores faltantes y ruido (predicciones correctas)
- **Escalabilidad**
 - Eficiencia en grandes bases de datos
- **Facilidad de interpretación**
 - Nivel de entendimiento provisto por el modelo
- **Forma de las reglas**
 - Tamaño del árbol de decisión
 - Qué tan compactas son las reglas de clasificación

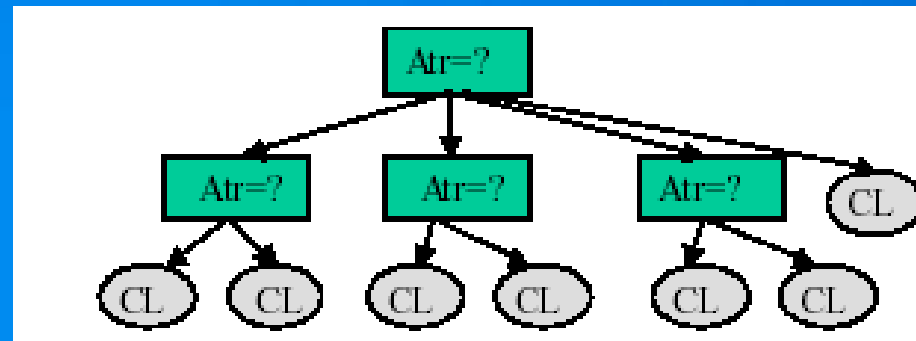
- Concepto
 - Clasificación
 - Predicción
 - Evaluación
 - Árboles de Decisión
 - Construcción
 - Uso
 - Poda
 - Clasificador Bayesiano
 - Ejemplos
-



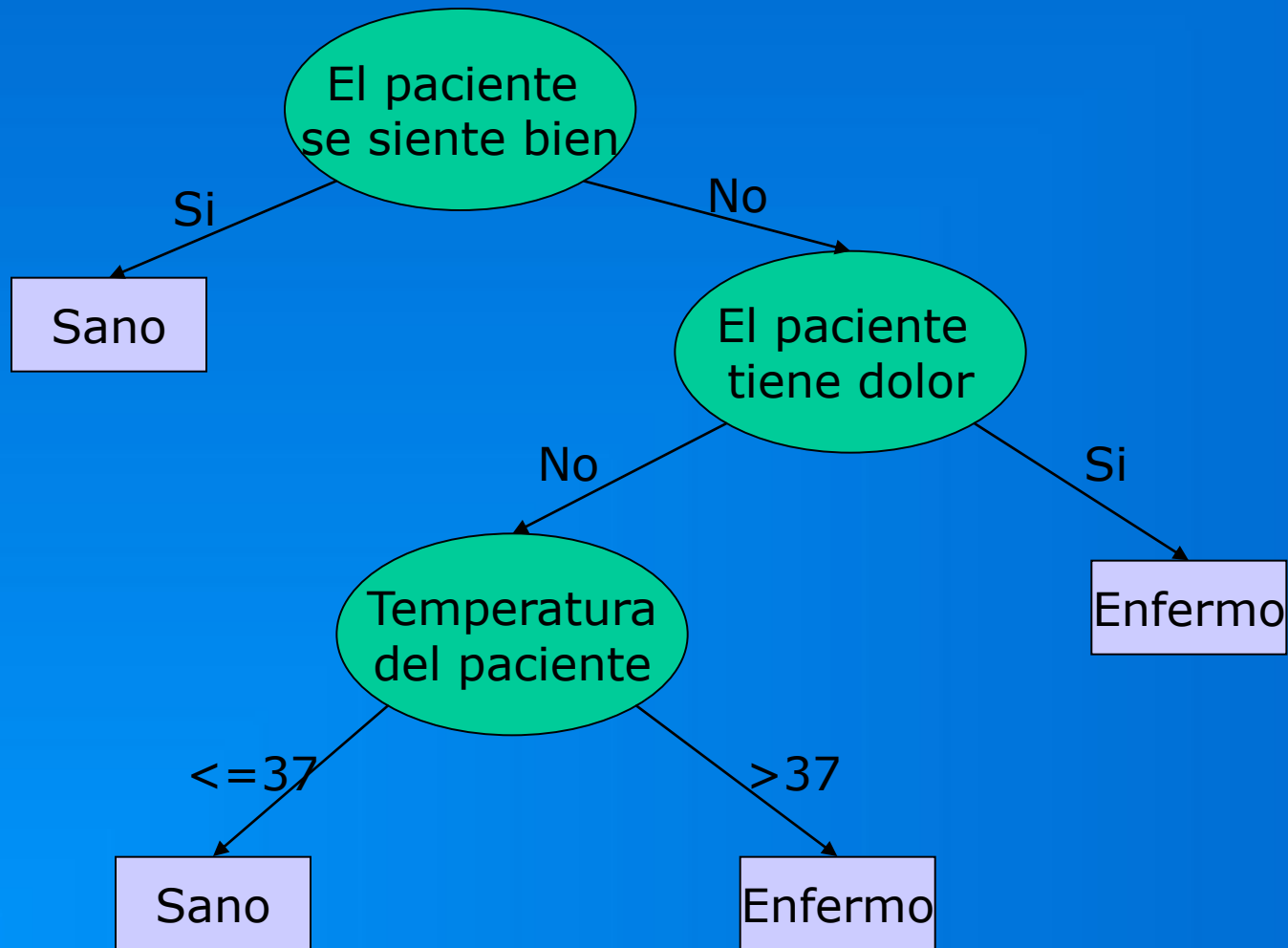
Árboles de Decisión

Un árbol de decisión consta de:

- nodos internos que denotan un test (comprobación) sobre un atributo
- Ramas que representan una salida del test
 - Todas las tuplas en una rama tienen el mismo valor para el atributo evaluado
- Nodos hojas que representan las etiquetas de clase



Ejemplo



Equivalente en reglas de clasificación

- **Si** El paciente se siente bien = **Si entonces**
 - Clase = Sano
- **Sino**
 - **Si** El paciente tiene dolor = **No entonces**
 - **Si** Temperatura del paciente ≤ 37 **entonces**
 - Clase = Sano
 - **Sino** (Temperatura del paciente > 37)
 - Clase = Enfermo
 - **Sino** (El paciente tiene dolor = Si)
 - **Clase = Enfermo**

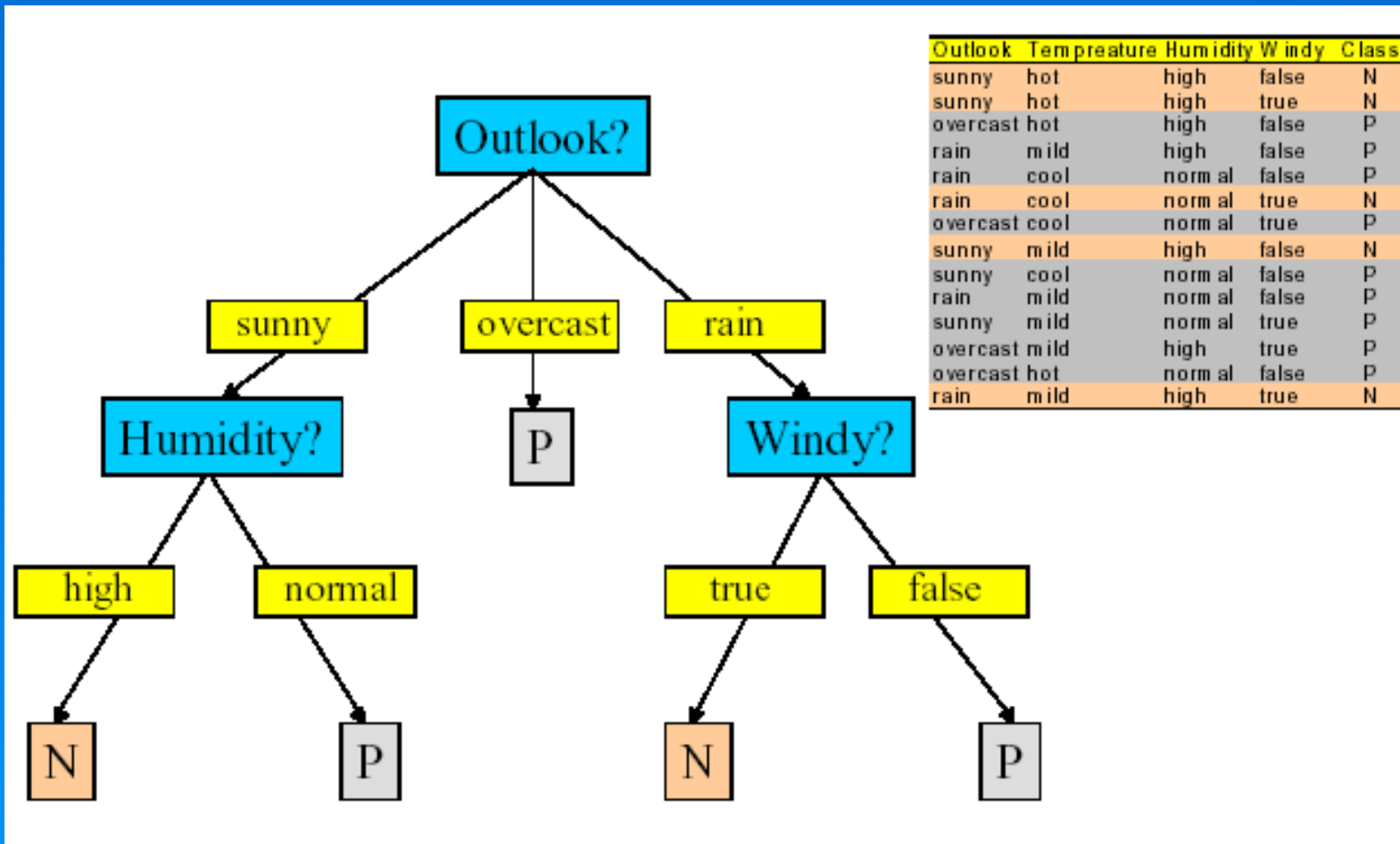
Equivalente en reglas de clasificación

- **Si** El paciente se siente bien = Si **entonces**
 - Clase = Sano
- **Si** El paciente se siente bien = No **and** El paciente tiene dolor = No **and** Temperatura del paciente ≤ 37 **entonces**
 - Clase = Sano
- **Si** El paciente se siente bien = No **and** El paciente tiene dolor = No **and** Temperatura del paciente > 37 **entonces**
 - Clase = Enfermo
- **Si** El paciente se siente bien = No **and** El paciente tiene dolor = Si **entonces**
 - Clase = Enfermo

Ejemplo: Datos de entrenamiento

Outlook	Tempreature	Humidity	Windy	Class
sunny	hot	high	false	N
sunny	hot	high	true	N
overcast	hot	high	false	P
rain	mild	high	false	P
rain	cool	normal	false	P
rain	cool	normal	true	N
overcast	cool	normal	true	P
sunny	mild	high	false	N
sunny	cool	normal	false	P
rain	mild	normal	false	P
sunny	mild	normal	true	P
overcast	mild	high	true	P
overcast	hot	normal	false	P
rain	mild	high	true	N

Ejemplo árbol de decisión



Métodos de clasificación con árboles de decisión

- La generación de árbol básica de arriba hacia abajo consiste de dos fases:
 1. **Construcción del árbol**
 - Al inicio todos los ejemplos de entrenamiento están en la raíz
 - La partición de los ejemplos se realiza recursivamente basándose en la selección de atributos
 2. **Podado del árbol**
 - Tiene por objetivo eliminar ramas del árbol que reflejen ruido en los datos de entrenamiento y lleven a errores cuando se clasifiquen los datos de test → mejorar la exactitud de clasificación

¿Cómo construir un árbol?

- Algoritmo

- Greedy

- Hacer una elección optimal en cada paso: seleccionar el mejor atributo para cada nodo del árbol

- Divide y conquista recursivo top-down

- De la raíz a las hojas
 - Partir un nodo en varias ramas
 - Para cada rama, correr recursivamente el algoritmo

Construcción del árbol de decisión

- El árbol comienza con un solo nodo que representa todos los datos de entrenamiento
- Si los ejemplos tienen todos la misma clase, entonces el nodo se convierte en una hoja con esa etiqueta de clase
- Si no, se **selecciona un atributo** que mejor separe la muestra en clases individuales. Este atributo se convierte en el atributo de test.
- Se crea una rama por cada valor conocido del atributo de test, y se particionan los ejemplos de acuerdo a estas ramas
- Se aplica el proceso recursivamente, para formar un árbol para los ejemplos en cada partición
- La recursión termina cuando:
 - Los ejemplos en el nodo corresponden a la misma clase
 - No quedan más atributos sobre los cuales separar (hoja con clase mayoritaria)
 - No hay ejemplos con el valor del atributo (para la rama)

Cuestiones principales en la construcción de árboles

- **Criterio de separación**
 - Usado para seleccionar el atributo para separar en un nodo del árbol durante la fase de generación del árbol
 - Diferentes algoritmos pueden usar distintas funciones: ganancia de información, gini, etc.
- **Esquema de ramificado**
 - Determinará la rama a la cual pertenece una muestra
 - Separación binaria (gini) versus separación múltiple (ganancia de información)
- **Decisiones de terminación**
 - Cuando detenerse y dejar de separar un nodo (medida de impureza) - Poda
- **Reglas de etiquetado**
 - Un nodo es etiquetado como la clase a la cual pertenecen la mayoría de los ejemplos que pertenecen a él

Elección del mejor atributo para clasificar

- Aleatoria
- Menos valores
- Más valores
- Ganancia de información (máxima)
- Índice Gini
- Razón de ganancia
- ...



Algoritmo ID3 (Induction Decision Trees)

- Algoritmo matemático para construir árboles de decisión
- Desarrollado por R. Quinlan en 1979
- Hace uso de Teoría de la Información desarrollada por Shannon
- Construye árbol de arriba hacia abajo
- Usa ganancia de la información para seleccionar atributo para particionar
- ID3, C4.5, J48...

Ganancia de Información

- El atributo que tiene **mayor ganancia de información** se utiliza para particionar
- Minimizar la información que se necesita para clasificar ejemplos en las particiones resultantes
- Mide qué tan bien separa los ejemplos de entrenamiento de un conjunto dado de acuerdo a su clasificación objetivo (impurezas)
- **Minimiza el número esperado de tests** que se necesitan para clasificar un objeto y garantiza la **construcción de un árbol simple**.

Entropía

-
- Medida del grado de incertidumbre asociado a una distribución de probabilidad.
 - En una distribución uniforme, todos los valores son igualmente probables $P_i = 1/N$ y por tanto la entropía es máxima, lo cual indica máxima incertidumbre.
 - Por el contrario, en una distribución pico en la que $P_i = 1$ y $P_j = 0$, para todo $j \neq i$ la entropía es mínima lo cual indica mínima incertidumbre o sea máxima información.

Entropía: mínimo teórico de bits promedio necesarios para transmitir un conjunto de mensajes sobre $\{x_1, x_2, \dots, x_n\}$ con distribución de probabilidades $P(X=x_i)=p_i$

$$\text{Entropy}(P) = -(p_1 \log(p_1) + p_2 \log(p_2) + \dots + p_n \log(p_n))$$

- Supongamos que hay 2 clases, P y N
 - Sea el conjunto de ejemplos S que contiene x elementos de la clase P e y elementos de la clase N
 - La cantidad de información que se necesita para decidir si un ejemplo arbitrario en S pertenece a P o N se define como

$$I(S_P, S_N) = -\frac{x}{x+y} \log_2 \frac{x}{x+y} - \frac{y}{x+y} \log_2 \frac{y}{x+y}$$

En gral.

$$I(s_1, s_2, \dots, s_n) = -\sum_{i=1}^n p_i \log_2(p_i)$$

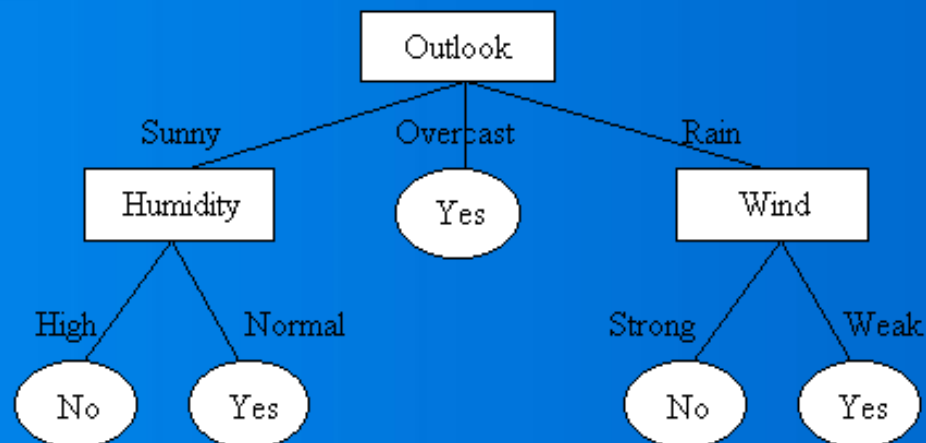
Ejemplo

$$E(S) = \sum_{i=1}^c -p_i \log_2 p_i$$

Play Golf	
Yes	No
9	5

$$\begin{aligned} \text{Entropy}(\text{PlayGolf}) &= \text{Entropy}(5,9) \\ &= \text{Entropy}(0.36, 0.64) \\ &= -(0.36 \log_2 0.36) - (0.64 \log_2 0.64) \\ &= 0.94 \end{aligned}$$

Atributos				Objetivo
Outlook	Temp	Humidity	Windy	Play Golf
Rainy	Hot	High	False	No
Rainy	Hot	High	True	No
Overcast	Hot	High	False	Yes
Sunny	Mild	High	False	Yes
Sunny	Cool	Normal	False	Yes
Sunny	Cool	Normal	True	No
Overcast	Cool	Normal	True	Yes
Rainy	Mild	High	False	No
Rainy	Cool	Normal	False	Yes
Sunny	Mild	Normal	False	Yes
Rainy	Mild	Normal	True	Yes
Overcast	Mild	High	True	Yes
Overcast	Hot	Normal	False	Yes
Sunny	Mild	High	True	No



Ganancia de información

- Criterio para elegir un split: seleccionar el split con la mayor ganancia de información (Gain)
- Equivalente a seleccionar el split con la **menor entropía ponderada**
- Medida de **cuanto ayuda el conocer el valor de una variable aleatoria X para conocer el verdadero valor de otra Y.**
- En nuestro caso, X es un atributo de un ejemplo dado mientras que Y es la clase a la que pertenece el ejemplo.
- Una **alta ganancia** implica que el **atributo X permite reducir la incertidumbre de la clasificación** del ejemplo de entrada.

Ganancia de Información: Ejemplo

- El atributo A es seleccionado tal que la ganancia de información

$$\text{Gain}(A) = I(S_P, S_N) - E(A)$$

sea maximal, es decir, $E(A)$ es minimal dado que $I(S_P, S_N)$ es la misma para todos los atributos del nodo (**reducción esperada en la entropía causada por conocer el valor del atributo A**)

- En el ejemplo dado, el atributo outlook es elegido para separar la raíz:

$$\text{Gain}(\text{outlook}) = 0.247$$

$$\text{Gain}(\text{temperature}) = 0.029$$

$$\text{Gain}(\text{humidity}) = 0.151$$

$$\text{Gain}(\text{windy}) = 0.048$$

Ejemplo

$E(\text{PlayGolf}, \text{Outlook})$

$$= p(\text{sunny}) * E(3,2) + p(\text{overcast}) * E(4,0) + p(\text{rainy}) * E(3,2) = \\ (5/14) * 0.971 + (9/14) * 0.971 = 0.693$$

$$G(\text{PlayGolf}, \text{Outlook}) = E(\text{PlayGolf}) - E(\text{PlayGolf}, \text{Outlook}) = \\ 0.940 - 0.693 = 0.247$$

Lo mismo se hace para los demás atributos: temperature, humidity y windy.

		Play Golf	
		Yes	No
Outlook	Sunny	3	2
	Overcast	4	0
	Rainy	2	3
		Gain = 0.247	

		Play Golf	
		Yes	No
Temp.	Hot	2	2
	Mid	4	2
	Cool	3	1
		Gain = 0.029	

		Play Golf	
		Yes	No
Humidity	High	3	4
	Normal	6	1
		Gain = 0.152	

		Play Golf	
		Yes	No
Windy	False	6	2
	True	3	3
		Gain = 0.048	

- Se calcula la entropía del nodo N

$$H(E_i) = p^+ \log_2(1/p^+) + p^- \log_2(1/p^-)$$

- Se calcula el valor medio de la entropía del nivel siguiente, generado por el atributo X al que se está aplicando el test

$$H(X, E_i) = p(X=v_1/E_i)H(E_{i1}) + \dots + \\ p(X=v_n/E_i)H(E_{in})$$

Utilización del árbol

- **Directamente**

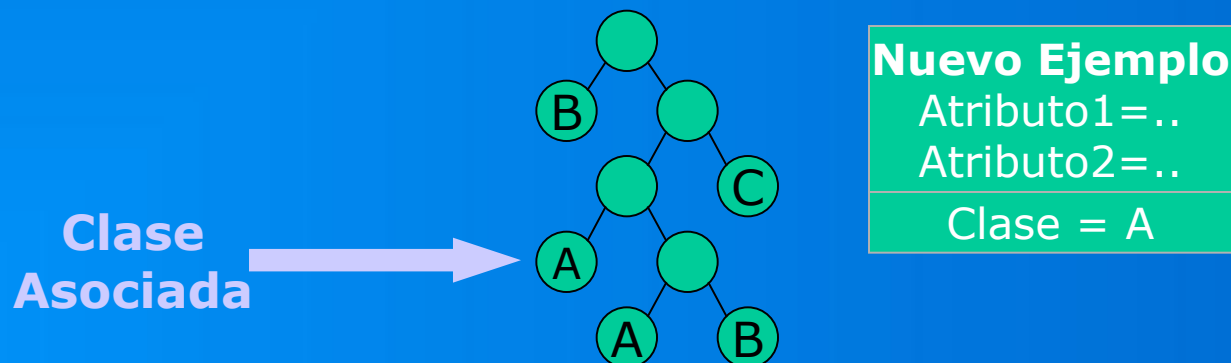
- Verificar el valor de un atributo de un ejemplo no conocido con el árbol
- Se sigue el camino desde la raíz a la hoja que posea la etiqueta

- **Indirectamente**

- El árbol de decisión se convierte en reglas de clasificación
- Se crea una regla por cada camino de la raíz a las hojas
- Las reglas IF-THEN son más fáciles de entender

Clasificación de nuevos ejemplos

- Partir desde la raíz
- Avanzar por los nodos de decisión hasta alcanzar una hoja
- La clase del nuevo ejemplo es la clase que representa la hoja.



Un árbol generado puede sobre-clasificar (super-ajustar) los ejemplos de entrenamiento debido a ruido o tamaño pequeño del conjunto de entrenamiento. Resulta en baja exactitud de clasificación de nuevos ejemplos.

- Existen 2 enfoques para evitar overfitting
 - **Parar antes (Pre-poda)**: detener el crecimiento del árbol antes que se produzca, decidiendo no particionar uno o más nodos
 - **Podado**: se construye todo el árbol, se permite overfitting y luego se poda el árbol. Se elige el árbol con menor tasa de error.

Podar el árbol: transformar un subárbol en una hoja

- Usar un conjunto de datos diferentes del conjunto de entrenamiento (conjunto de poda)
- En un nodo del árbol, si la exactitud sin particionar el nodo es más alta que la exactitud particionando el nodo, reemplazar el subárbol por una hoja, etiquetándolo con la clase mayoritaria

- Concepto
 - Clasificación
 - Predicción
 - Evaluación
 - Árboles de Decisión
 - Construcción
 - Uso
 - Poda
 - Clasificador Bayesiano
 - Ejemplos
-



Clasificador Bayesiano

- Es un clasificador estadístico basado en el **Teorema de Bayes**
- Usa aprendizaje probabilístico para calcular explícitamente las probabilidades de las hipótesis
- Un clasificador bayesiano simple o naive asume **independencia total entre atributos**
- Funciona bien con grandes conjuntos de datos y tiene alta precisión de clasificación
- El modelo es **incremental** es el sentido que cada ejemplo de entrenamiento puede aumentar o disminuir la probabilidad de que una hipótesis sea correcta. Conocimiento previo puede combinarse con datos observados

Clasificador Bayesiano

El objetivo de un clasificador es identificar a qué clase pertenece un objeto, basándose en ciertos atributos.

Consideremos un objeto X , cuyos atributos son $(a_1, a_2, \dots, a_n)$, y queremos clasificarlo dentro de un conjunto de clases S .

Sea A una clase, uno de los elementos de S . Entonces, trataremos de encontrar un A que maximice $P(A | a_1, a_2, \dots, a_n)$. Podemos decir entonces:

“Es muy probable que X pertenezca a la clase A porque para un objeto, dada la condición de tener los atributos $(a_1, a_2, \dots, a_n)$, la probabilidad de que la clase sea A es máxima.”

Teorema de Bayes

Dado un objeto X (descripción del clima) con etiqueta de clase desconocida, H es la hipótesis de que X pertenece a una clase C (apto para jugar golf)

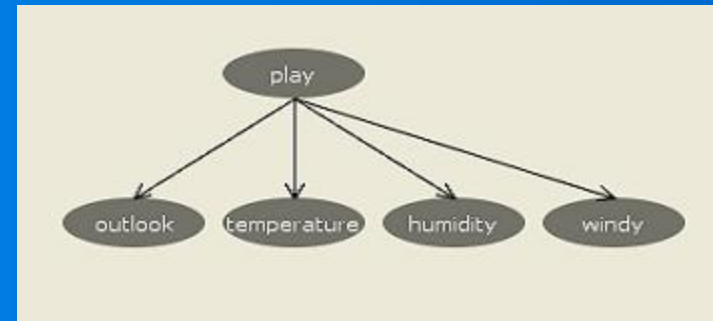
La probabilidad a posteriori de la hipótesis H , $P(H/X)$, probabilidad condicional de H dado X , sigue el teorema de Bayes

$$P(H/X) = \frac{P(X/H) P(H)}{P(X)}$$

$P(H/X)$ probabilidad a posteriori

$P(H)$, $P(X)$ probabilidad a priori (datos)

$P(X/H)$ probabilidad a posteriori (que el clima esté de cierta forma dado que juego golf)



Clasificador Naive Bayes

- Supongamos que tenemos n clases $C_1, C_2, \dots, C_n$. Dado un ejemplo desconocido A , el clasificador predecirá que $A=(a_1, \dots, a_m)$ pertenece a la clase con la mayor probabilidad a posteriori:

$A \in C_i$ si $P(C_i/A) > P(C_j/A)$ para $1 \leq j \leq n, j \neq i$

Maximizar $P(A/C_i) P(C_i) / P(A) \rightarrow$ Maximizar $P(A/C_i) P(C_i)$

$$P(C_i) = s_i/s$$

$$P(A/C_i) = P(a_1, \dots, a_m / C_i)$$

Suposición de Naive Bayes

$$P(a_1, a_2, \dots, a_n | A) = \prod_{i=1}^n P(a_i | A)$$

Siempre que $(a_1, a_2, \dots, a_n)$ sean **condicionalmente independientes** unas de otras.

Dado que $P(a_i | A)$ se obtiene fácilmente de los datos de entrenamiento, el clasificador tiene que encontrar una clase A que maximice la expresión:

$$P(A) \prod_{i=1}^n P(a_i | A)$$

Ejemplo

Tamaño	Precio	Peso	Color	Compro?
Grande	Barato	Pesado	azul	Si
Grande	Caro	Medio	Azul	No
Chico	Caro	Medio	Verde	Si
Chico	Barato	Liviano	Azul	No
Chico	Caro	pesado	Verde	No

(chico, barato, pesado, azul) compro?

Ejemplo

$P(C = \text{Si} \mid T = \text{Chico}, Pr = \text{Barato}, P = \text{Pesado}, C = \text{Azul})$ y

$P(C = \text{No} \mid T = \text{Chico}, Pr = \text{Barato}, P = \text{Pesado}, C = \text{Azul})$

Equivalente a:

(1): $P(C = S) P(T = \text{Chico} \mid C = S) P(Pr = \text{Barato} \mid C = S) P(P = \text{Pesado} \mid C = S) P(C = \text{Azul} \mid C = S)$

(2): $P(C = N) P(T = \text{Chico} \mid C = N) P(Pr = \text{Barato} \mid C = N) P(P = \text{Pesado} \mid C = N) P(C = \text{Azul} \mid C = N)$

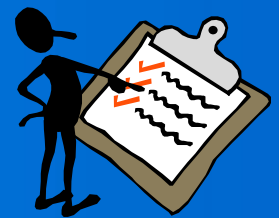
Resultados:

(1): $(2/5)(1/2)(1/2)(1/2)(1/2) = 0.025$

(2): $(3/5)(2/3)(1/2)(1/3)(2/3) = 0.044$

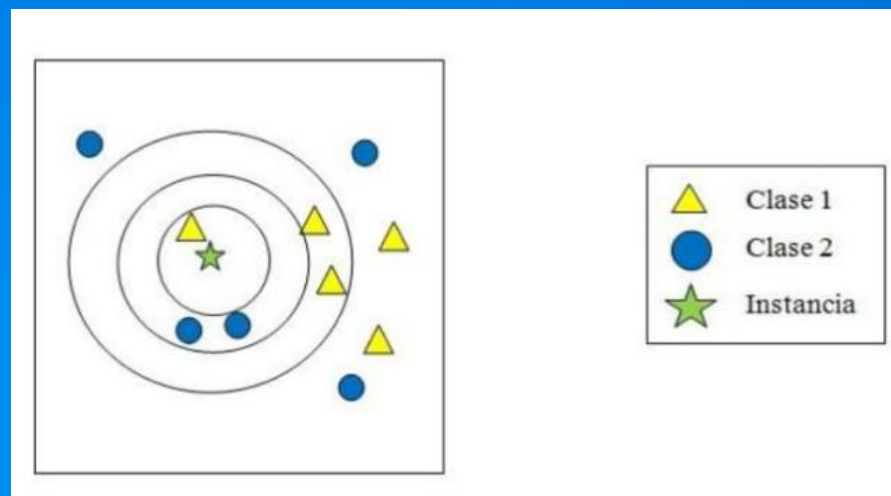
Clasificación

- Concepto
 - Clasificación
 - Predicción
 - Evaluación
- Árboles de Decisión
 - Construcción
 - Uso
 - Poda
- Clasificador Bayesiano
- Vecinos más cercanos



Vecinos más cercanos

Para $K=1$ (círculo más pequeño), la clase de la nueva instancia sería la Clase 1, ya que es la clase de su vecino más cercano, mientras que para $K=3$ la clase de la nueva instancia sería la Clase 2 pues habrían dos vecinos de la Clase 2 y solo 1 de la Clase 1



COMIENZO

Entrada: $D = \{(\mathbf{x}_1, c_1), \dots, (\mathbf{x}_N, c_N)\}$

$\mathbf{x} = (x_1, \dots, x_n)$ nuevo caso a clasificar

PARA todo objeto ya clasificado (x_i, c_i)

calcular $d_i = d(\mathbf{x}_i, \mathbf{x})$

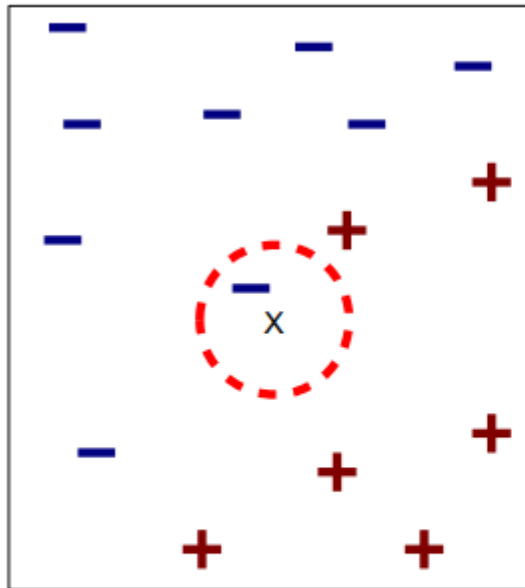
Ordenar $d_i (i = 1, \dots, N)$ en orden ascendente

Quedarnos con los K casos $D_{\mathbf{x}}^K$ ya clasificados más cercanos a $\mathbf{x}$

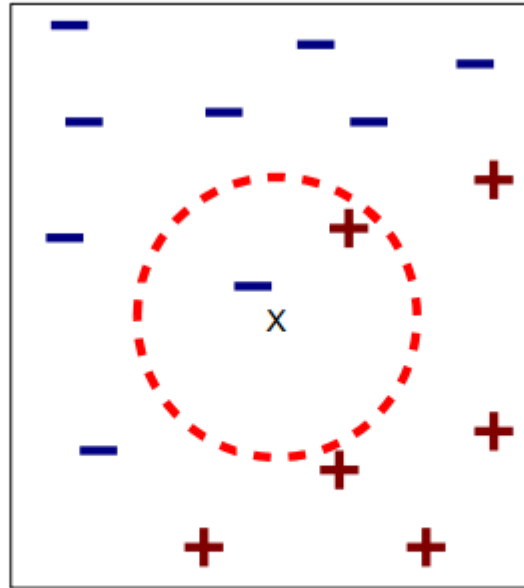
Asignar a $\mathbf{x}$ la clase más frecuente en $D_{\mathbf{x}}^K$

FIN

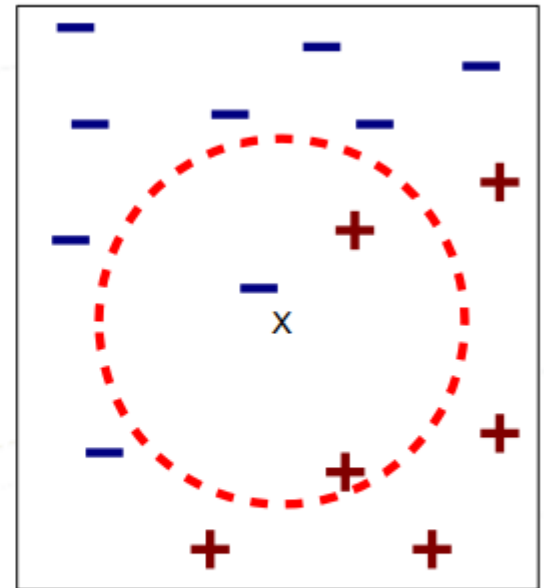
Cómo elegir K?



(a) 1-nearest neighbor



(b) 2-nearest neighbor



(c) 3-nearest neighbor

Escogiendo el valor de K:

- Si K es muy pequeño el modelo será muy sensitivo a puntos que son atípicos o que son ruido (datos corruptos, outliers)
- Si K es muy grande, el modelo tiende a asignar siempre a la clase más grande.

Cómo elegir K?

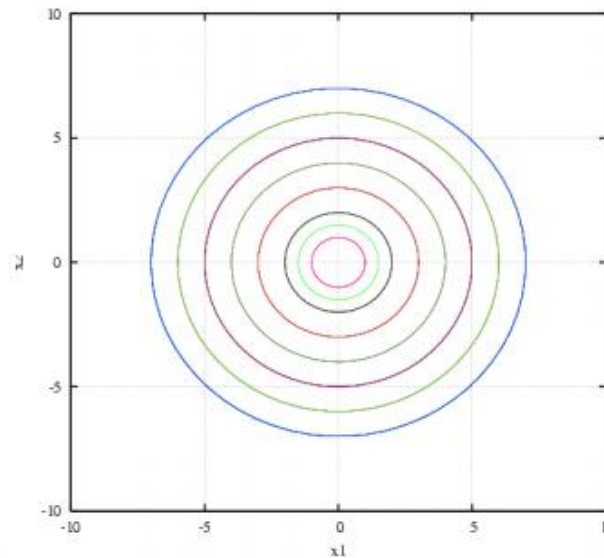
En algunos algoritmos, el modelo escogerá el valor de K que mejor clasificación logre en esta tabla, es decir, prueba con $K=1$, $K=2$,

Esto puede ser muy caro computacionalmente.

Cómo elegir la distancia?

$$d(A, B) \equiv \sqrt{\sum_{i=1}^n (A_i - B_i)^2} = \sqrt{(A - B)^T (A - B)}$$

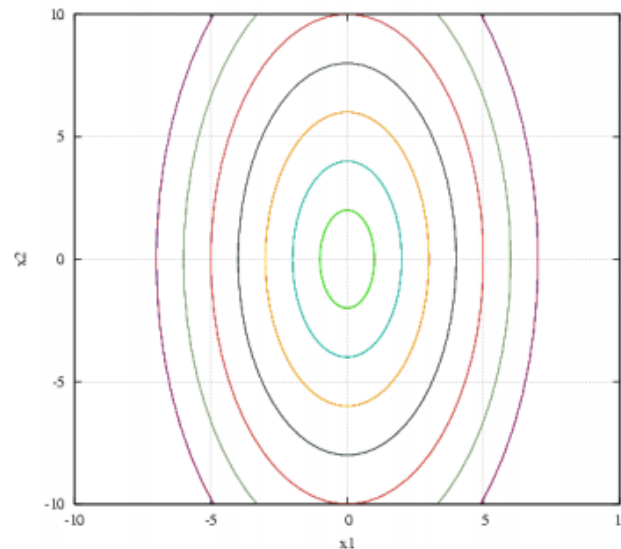
**Distancia
Euclídea**



Cómo elegir la distancia?

$$d(A, B) \equiv \sqrt{(A - B)^T M^T M (A - B)}$$

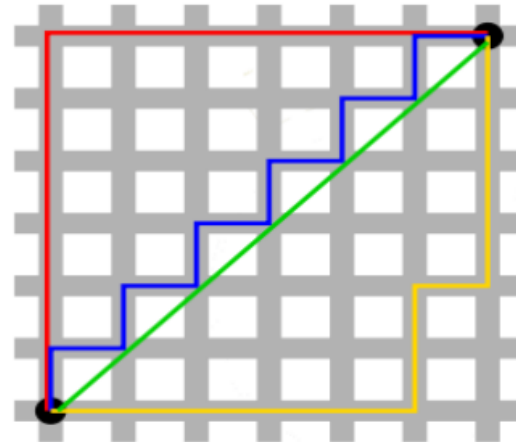
**Distancia
Euclídea
Ponderada**



Cómo elegir la distancia?

$$d(x, y) = \sum_{i=1}^n |x_i - y_i|$$

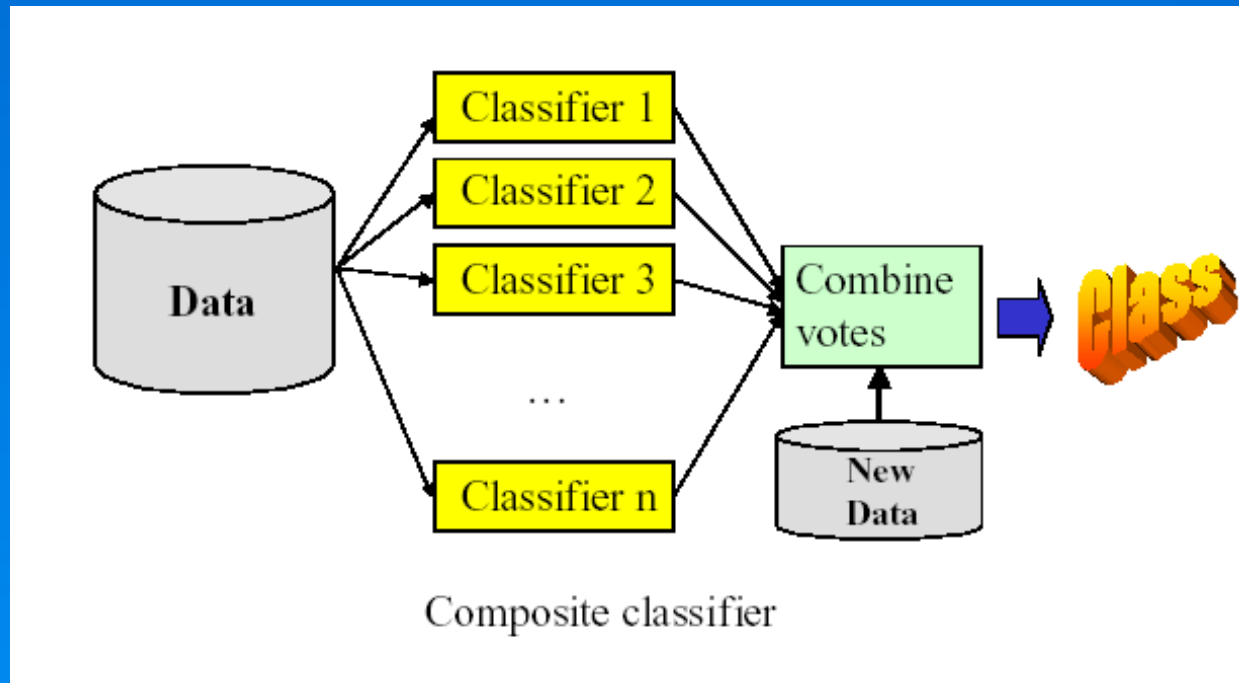
**Distancia
Manhattan
(city-block)**



Áreas de investigación en Árboles de Decisión

- Clases desbalanceadas
- Clasificadores multi-clase
- Aprendizaje semi-supervisado
- Combinación de clasificadores
- Selección de atributos
- ...

Mejorar la exactitud: Clasificadores compuestos



- **Bagging:** ej. consulto varios doctores y me quedo con la opinión mayoritaria (la que tenga más votos)
- **Boosting:** ej. pondero cada diagnóstico según la exactitud del médico (del clasificador)

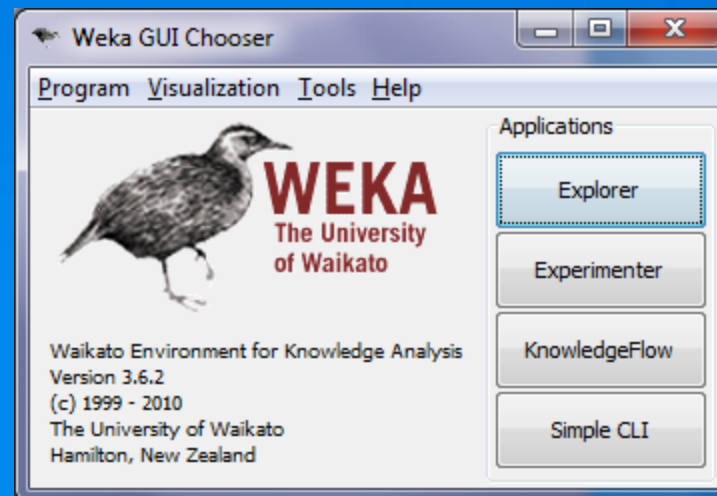
Clasificación

- Concepto
 - Clasificación
 - Predicción
 - Evaluación
- Árboles de Decisión
 - Construcción
 - Uso
 - Poda
- Clasificador Bayesiano
- Herramientas: Weka



Software: WEKA

- Machine Learning Software in Java
<http://www.cs.waikato.ac.nz/~ml/weka/>
- Data Mining: Practical Machine Learning Tools and Techniques - [Ian H. Witten](#), [Eibe Frank](#) (Third Edition) – Morgan Kaufmann - 2011



Weka 3.5.6 - Explorer

Program Applications Tools Visualization Windows Help

Explorer

Preprocess Classify Cluster Associate Select attributes Visualize

Open file... Open URL... Open DB... Generate... Undo Edit... Save...

Filter

Choose None Apply

Current relation

Relation: weather.symbolic
Instances: 14
Attributes: 5

Selected attribute

Name: outlook
Missing: 0 (0%)
Distinct: 3
Type: Nominal
Unique: 0 (0%)

Label	Count
sunny	5
overcast	4
rainy	5

Attributes

All None Invert Pattern

No.	Name
1	☒ outlook
2	☐ temperature
3	☐ humidity
4	☐ windy
5	☐ play

Remove

Class: play (Nom)

Visualize All



Status

OK

Log x 0

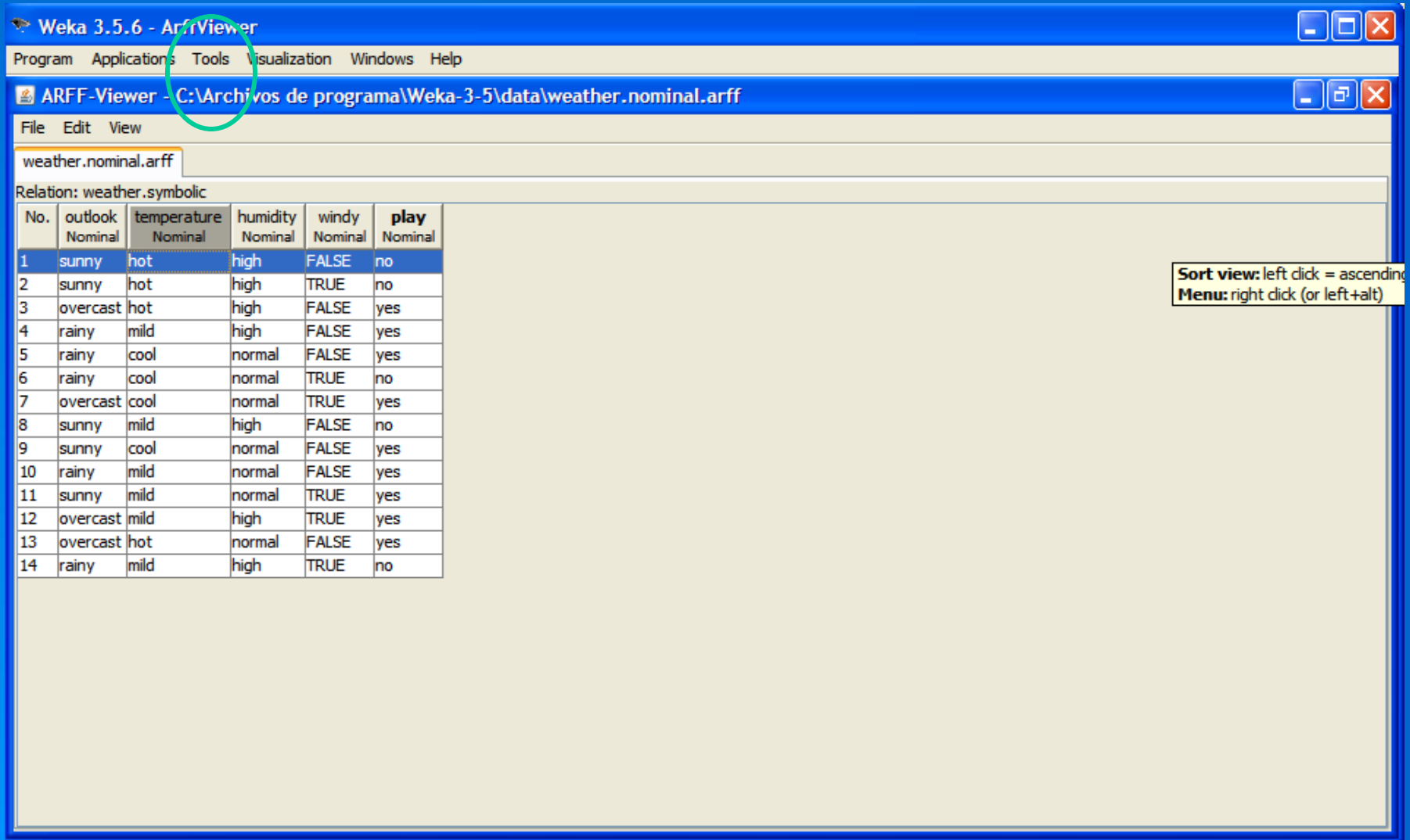
WEKA: Input file

```
@relation weather.symbolic
@attribute outlook {sunny, overcast, rainy}
@attribute temperature {hot, mild, cool}
@attribute humidity {high, normal}
@attribute windy {TRUE, FALSE}
@attribute play {yes, no}
```

```
@data
sunny,hot,high,FALSE,no
sunny,hot,high,TRUE,no
overcast,hot,high,FALSE,yes
rainy,mild,high,FALSE,yes
rainy,cool,normal,FALSE,yes
rainy,cool,normal,TRUE,no
overcast,cool,normal,TRUE,yes
sunny,mild,high,FALSE,no
sunny,cool,normal,FALSE,yes
rainy,mild,normal,FALSE,yes
sunny,mild,normal,TRUE,yes
overcast,mild,high,TRUE,yes
overcast,hot,normal,FALSE,yes
rainy,mild,high,TRUE,no
```



Weka: Arff viewer



The screenshot shows the Weka 3.5.6 - ARFF Viewer application. The title bar indicates the file being viewed is 'ARFF-Viewer - C:\Archivos de programa\Weka-3-5\data\weather.nominal.arff'. The 'Tools' menu is circled in green. The main window displays a table of weather data with the following columns: No., outlook, temperature, humidity, windy, and play. The data is sorted by the 'play' column in ascending order.

Relation: weather.symbolic

No.	outlook Nominal	temperature Nominal	humidity Nominal	windy Nominal	play Nominal
1	sunny	hot	high	FALSE	no
2	sunny	hot	high	TRUE	no
3	overcast	hot	high	FALSE	yes
4	rainy	mild	high	FALSE	yes
5	rainy	cool	normal	FALSE	yes
6	rainy	cool	normal	TRUE	no
7	overcast	cool	normal	TRUE	yes
8	sunny	mild	high	FALSE	no
9	sunny	cool	normal	FALSE	yes
10	rainy	mild	normal	FALSE	yes
11	sunny	mild	normal	TRUE	yes
12	overcast	mild	high	TRUE	yes
13	overcast	hot	normal	FALSE	yes
14	rainy	mild	high	TRUE	no

Sort view: left click = ascending
Menu: right click (or left+alt)

WEKA: Pre-processing

The screenshot displays the Weka 3.5.6 Explorer application window. The 'Preprocess' tab is active, showing a list of filters on the left and a detailed view of the selected attribute, 'outlook', on the right. The 'outlook' attribute is nominal with 3 distinct values: sunny (5), overcast (4), and rainy (5). The class is 'play (Nom)'. Below the attribute details, there are three stacked bar charts showing the distribution of the 'outlook' attribute across the 'play' class. The status bar at the bottom indicates 'OK' and shows the system clock as 11:41 a.m.

Weka 3.5.6 - Explorer

Program Applications Tools Visualization Windows Help

Explorer

Preprocess Classify Cluster Associate Select attributes Visualize

Open file... Open URL... Open DB... Generate... Undo Edit... Save...

Filter

- instance
 - Resample
 - SpreadSubsample
 - StratifiedRemoveFolds
- unsupervised
 - attribute
 - Add
 - AddCluster
 - AddExpression
 - AddID
 - AddNoise
 - AddValues
 - Center
 - ChangeDateFormat
 - ClassAssigner
 - ClusterMembership
 - Copy
 - Discretize
 - FirstOrder
 - InterquartileRange
 - KernelFilter
 - MakeIndicator
 - MathExpression
 - MergeTwoValues
 - MultiInstanceToPropositional

Attributes: 5

Invert Pattern

Selected attribute

Name: outlook
Missing: 0 (0%)
Distinct: 3
Type: Nominal
Unique: 0 (0%)

Label	Count
sunny	5
overcast	4
rainy	5

Class: play (Nom)

Visualize All

Status: OK

Log x 0

Inicio Gmail - Inbox (3360...) Your New Springer ... Cursada 2009 Microsoft PowerPoin... Weka 3.5.6 - Explorer Dibujo - Paint ES 11:41 a.m.

WEKA: Pre-processing

The screenshot displays the WEKA 3.5.6 Explorer application window. The 'Preprocess' tab is active, and the 'Discretize' filter is selected from the 'Filter' dropdown. The 'Current relation' is 'weather.symbolic' with 14 instances. The 'Attributes' list shows 'outlook', 'temperature', 'humidity', 'windy', and 'play'. The 'Discretize' dialog box is open, showing the following settings:

- Attribute Indices: first-last
- Bins: 10
- Desired Weight Of Instances Per Interval: -1.0
- Find Num Bins: False
- Ignore Class: False
- Invert Selection: False
- Make Binary: False
- Use Equal Frequency: False

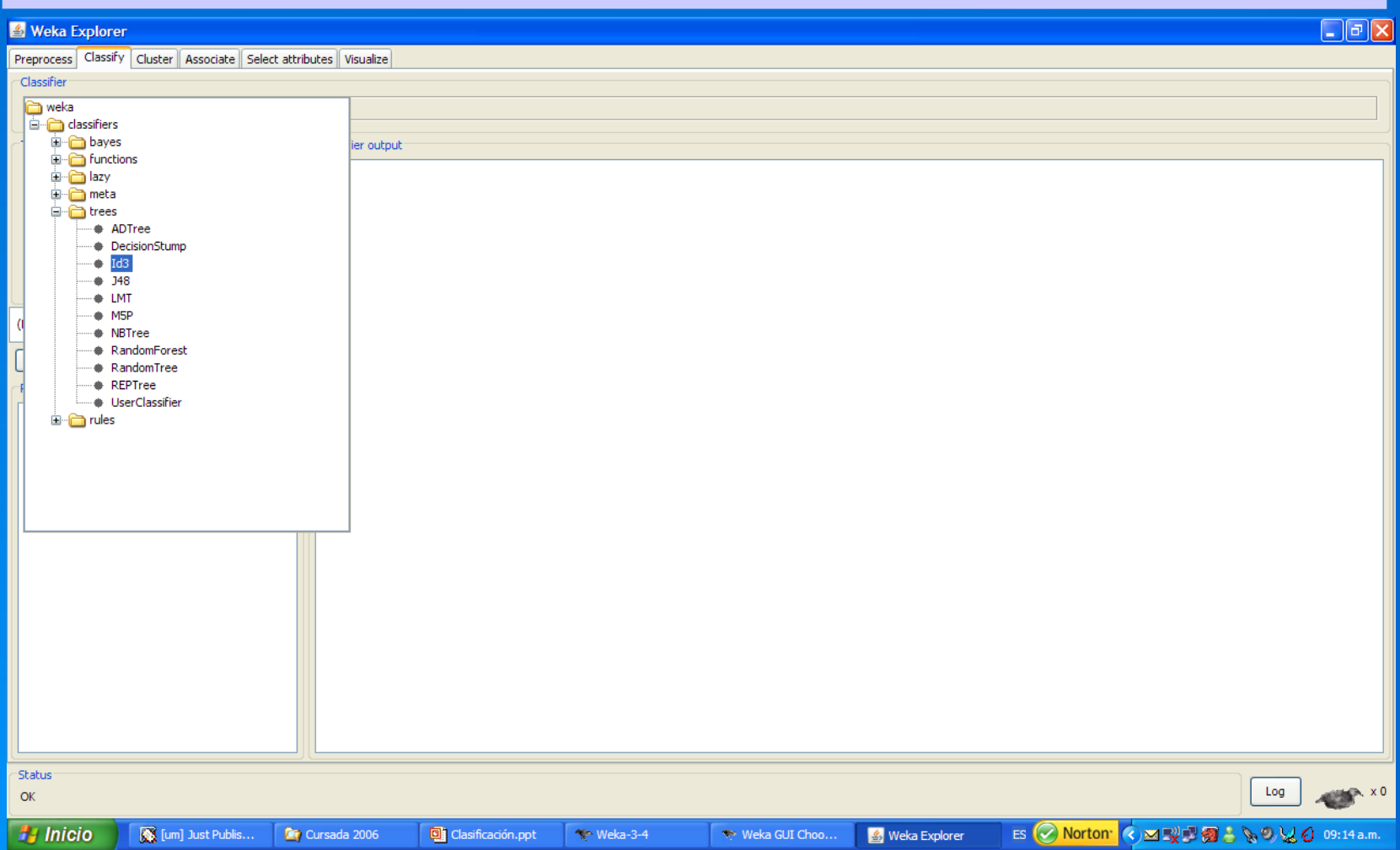
The dialog also includes 'More' and 'Capabilities' buttons. The background shows a bar chart with three bars, each divided into blue and red segments, representing the distribution of the discretized attribute.

Distinct: 3 **Type: Nominal**
Unique: 0 (0%)

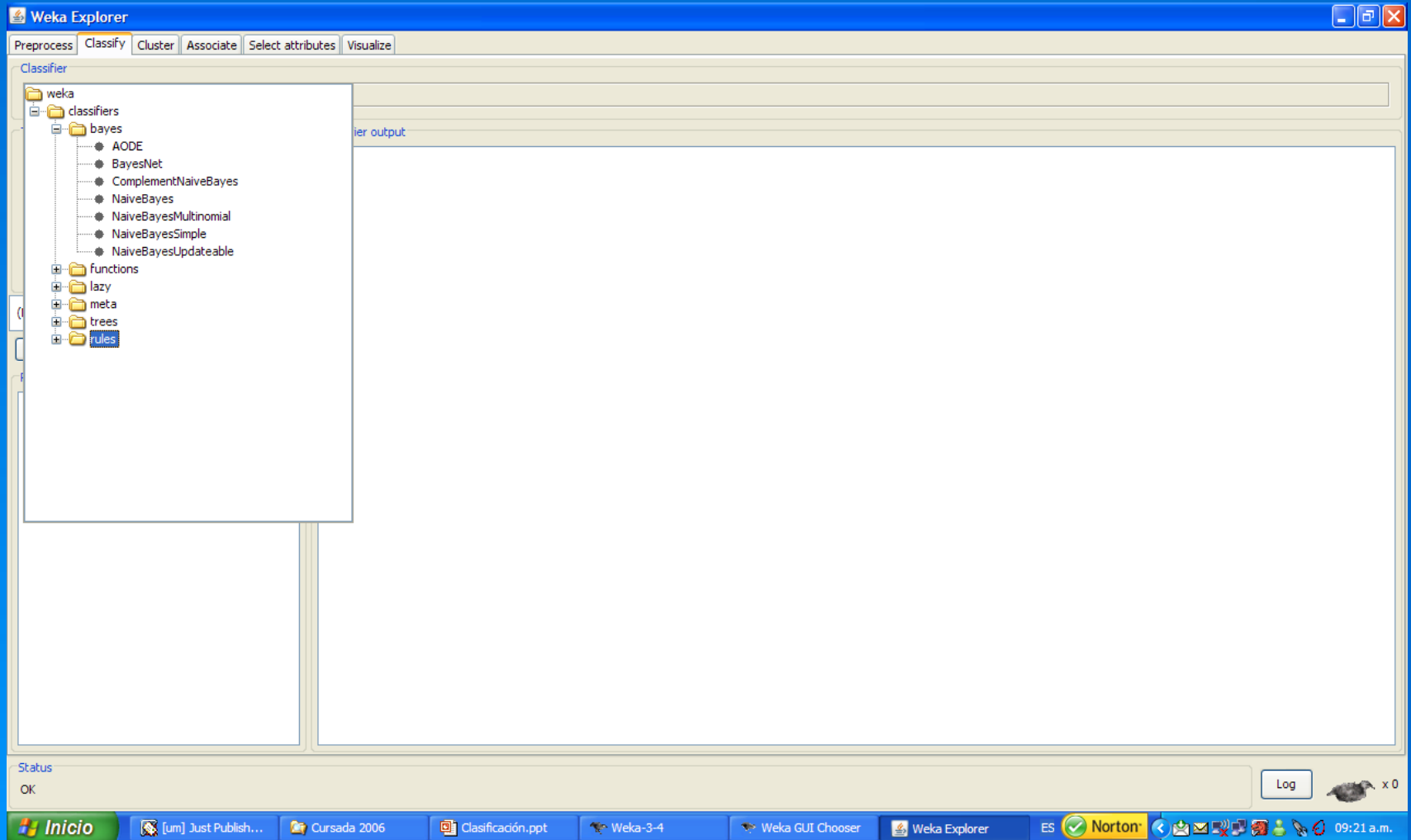
Distinct	Count
5	5
4	4
5	5

The status bar at the bottom shows 'OK' and a 'Log' button. The Windows taskbar at the very bottom includes icons for 'Inicio', 'Gmail - Inbox (...)', 'Your New Sprin...', 'Cursada 2009', 'Microsoft Powe...', 'Weka 3.5.6 - E...', 'weka.gui.Gene...', 'Dibujo - Paint', 'ES', and system icons for network, volume, and battery, along with the time '11:42 a.m.'.

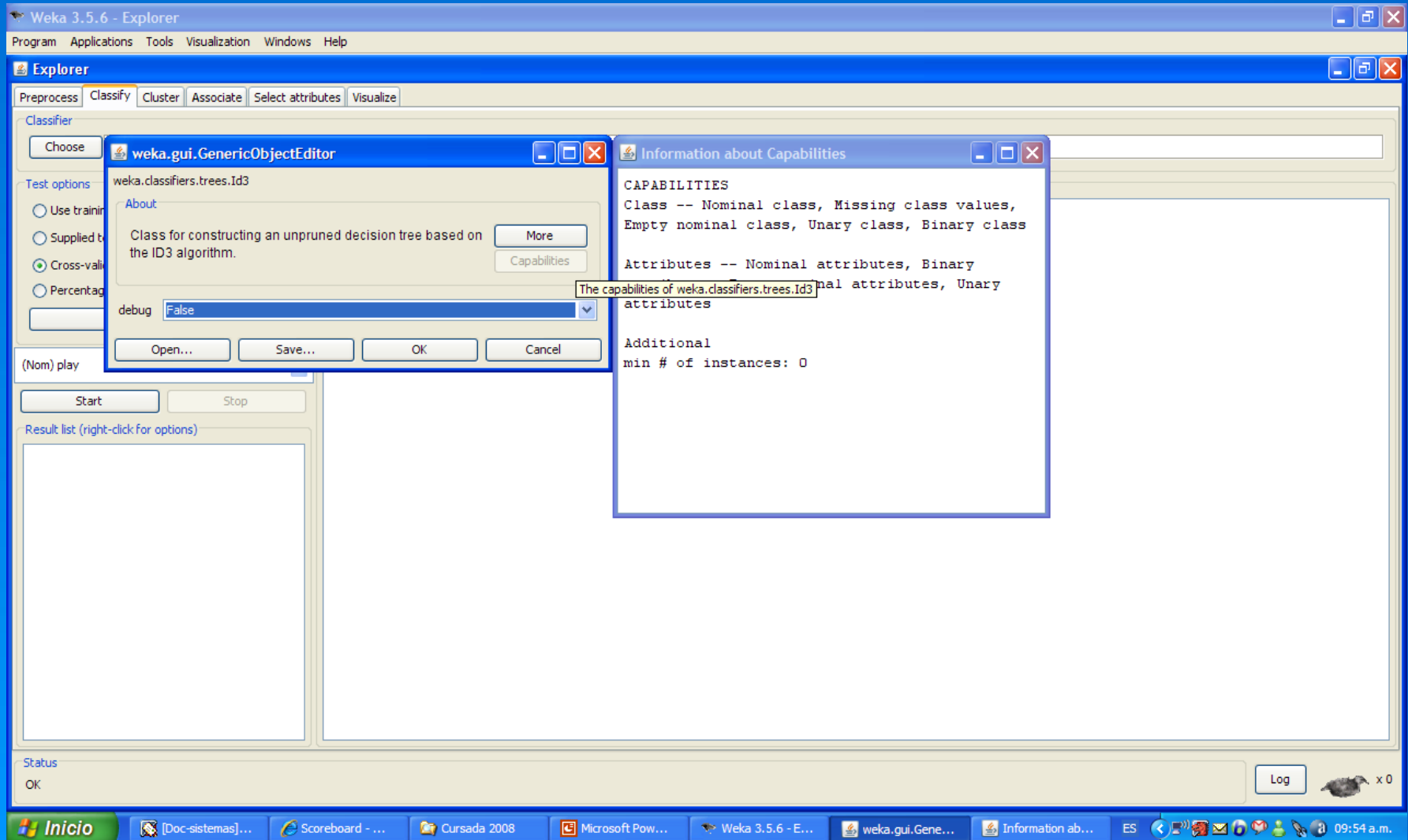
WEKA: Clasificadores



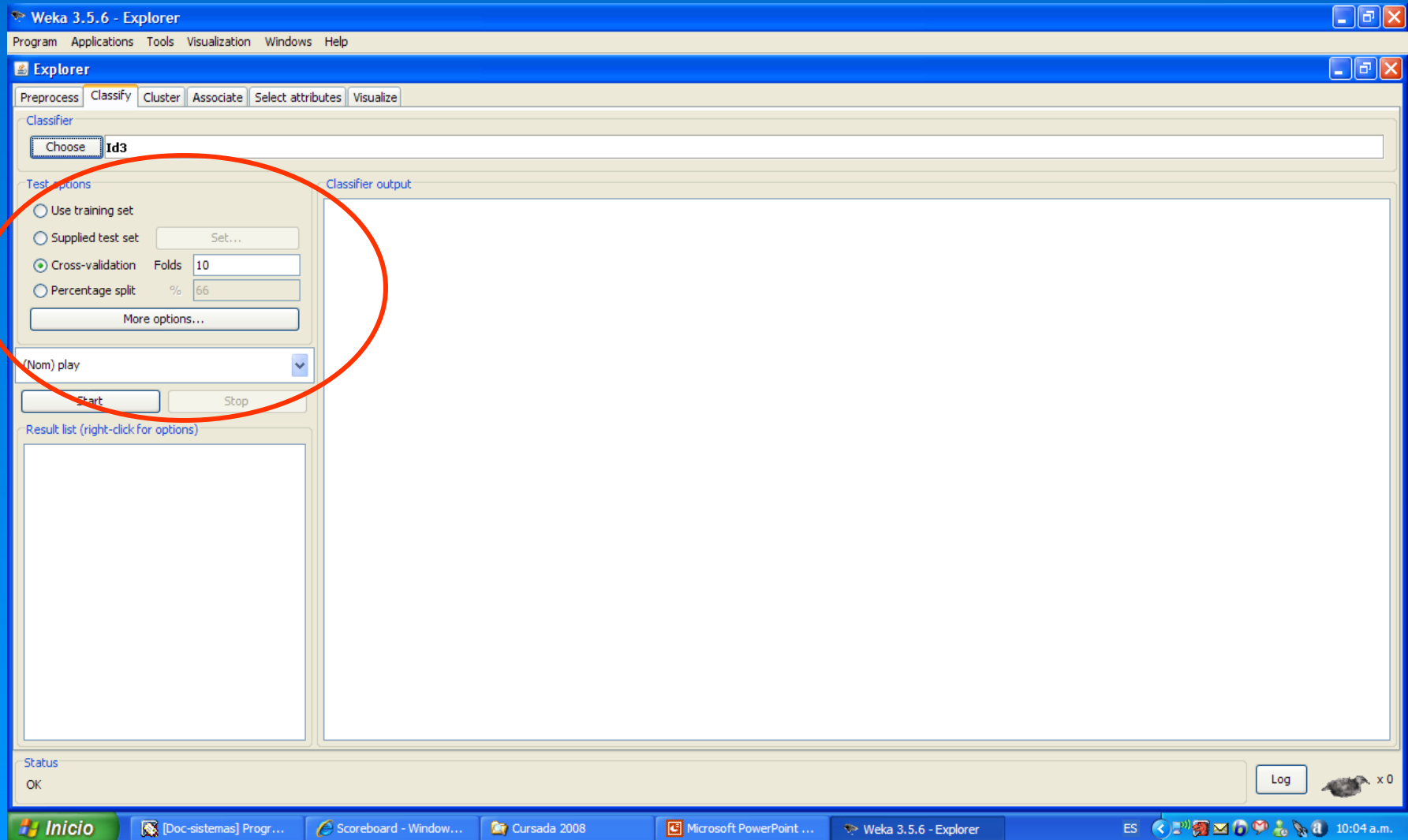
WEKA: Clasificadores



WEKA: Parámetros



WEKA: Training y Test



WEKA: Resultados

Weka Explorer

Preprocess **Classify** Cluster Associate Select attributes Visualize

Classifier

Choose **Id3**

Test options

☐ Use training set
☐ Supplied test set Set...
☒ Cross-validation Folds **10**
☐ Percentage split % **66**
More options...

(Nom) play

Start Stop

Result list (right-click for options)

09:15:06 - trees.Id3

Classifier output

```
=== Run information ===

Scheme:      weka.classifiers.trees.Id3
Relation:    weather.symbolic
Instances:   14
Attributes:  5
    outlook
    temperature
    humidity
    windy
    play
Test mode:   10-fold cross-validation

=== Classifier model (full training set) ===

Id3

outlook = sunny
| humidity = high: no
| humidity = normal: yes
outlook = overcast: yes
outlook = rainy
| windy = TRUE: no
| windy = FALSE: yes

Time taken to build model: 0 seconds

=== Stratified cross-validation ===
=== Summary ===

Correctly Classified Instances      12           85.7143 %
```

Status
OK

Log x 0

Inicio [um] Just Publi... Cursada 2006 Clasificación.ppt Weka-3-4 Weka GUI Choo... Weka Explorer ES Norton 09:15 a.m.

WEKA: Resultados

Weka Explorer

Preprocess **Classify** Cluster Associate Select attributes Visualize

Classifier: Choose **Prism**

Test options

- ☐ Use training set
- ☐ Supplied test set Set...
- ☒ Cross-validation Folds **10**
- ☐ Percentage split % **66**

More options...

(Nom) play

Start Stop

Result list (right-click for options)

- 10:06:10 - bayes.NaiveBayes
- 10:06:46 - trees.Id3
- 10:07:09 - trees.J48
- 10:08:35 - bayes.BayesNet
- 10:10:04 - bayes.NaiveBayesMultinomial
- 10:10:09 - bayes.NaiveBayesSimple
- 10:10:33 - rules.PART
- 10:10:45 - rules.MSRules
- 10:10:50 - rules.Prism**

Classifier output

Scheme: weka.classifiers.rules.Prism
Relation: weather.symbolic
Instances: 14
Attributes: 5
outlook
temperature
humidity
windy
play
Test mode: 10-fold cross-validation

=== Classifier Model (full training set) ===

Prism rules

If outlook = overcast then yes
If humidity = normal
and windy = FALSE then yes
If temperature = mild
and humidity = normal then yes
If outlook = rainy
and windy = FALSE then yes
If outlook = sunny
and humidity = high then no
If outlook = rainy
and windy = TRUE then no

Time taken to build model: 0 seconds

=== Stratified cross-validation ===
=== Summary ===

Status: OK

Log

WEKA: Visualización

Weka Explorer

Preprocess **Classify** Cluster Associate Select attributes Visualize

Classifier: Choose **J48 -C 0.25 -M 2**

Test options:
☐ Use training set
☐ Supplied test set Set...
☒ Cross-validation Folds **10**
☐ Percentage split % **66**
More options...

(Nom) play
Start Stop

Result list (right-click for options):
10:06:10 - bayes.NaiveBayes
10:06:46 - trees.Id3
10:07:09 - trees.J48

Classifier output:
Number of Leaves : 5
Size of the tree : 8

Weka Classifier Tree Visualizer: 10:07:09 - trees.J48 (weathe...)

Tree View

```
graph TD
    outlook((outlook)) -- sunny --> humidity((humidity))
    outlook -- overcast --> yes40[yes (4.0)]
    outlook -- rainy --> windy((windy))
    humidity -- high --> no30[no (3.0)]
    humidity -- normal --> yes20[yes (2.0)]
    windy -- TRUE --> no20[no (2.0)]
    windy -- FALSE --> yes30[yes (3.0)]
```

a b <-- classified as
5 4 | a = yes
3 2 | b = no

Status: OK

Log

WEKA: Visualización

The screenshot displays the Weka Explorer application window. The 'Classify' tab is active, showing the 'Classifier' section with 'BayesNet' selected. The 'Test options' section on the left includes radio buttons for 'Use training set', 'Supplied test set', 'Cross-validation' (selected), and 'Percentage split'. The 'Cross-validation' section shows 'Folds' set to 10 and 'Percentage split' set to 66. The 'Result list' on the left shows a list of models, with '10:08:35 - bayes.BayesNet' selected. The 'Classifier output' section on the right displays the following text:

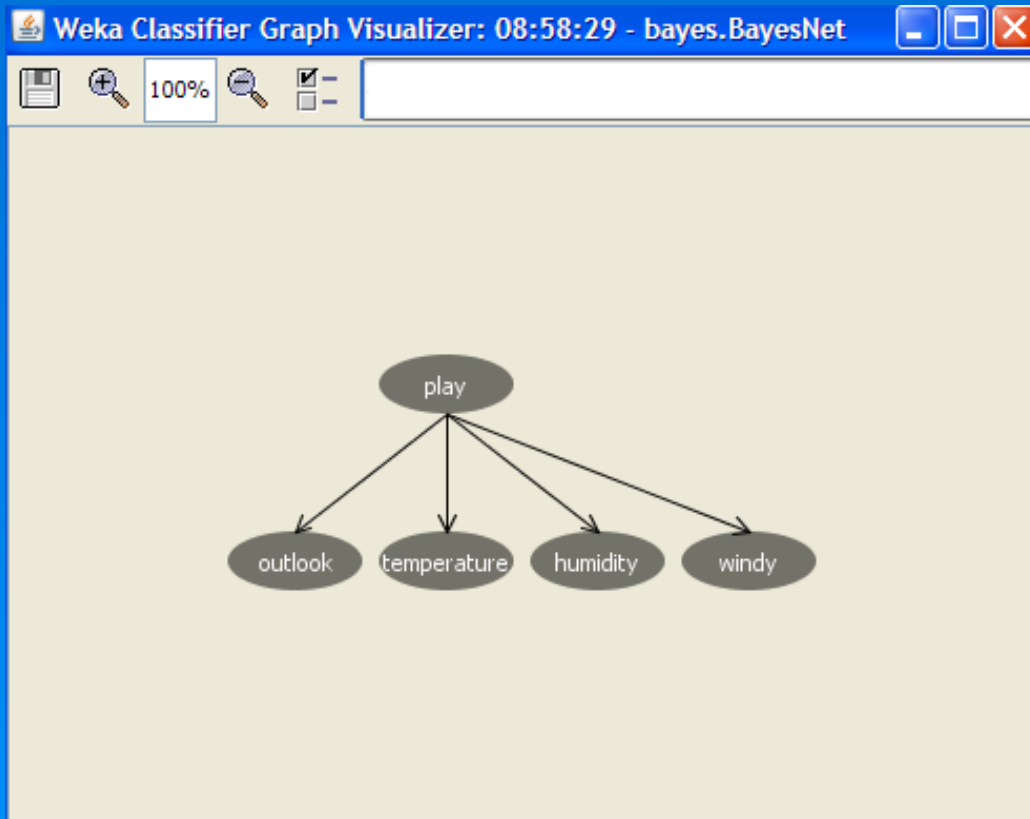
```
LogScore ENTROPY: -65.56181240647145
LogScore AIC: -78.56181240647145

Time taken to build model: 0.02 seconds
```

A 'Weka Classifier Graph Visualizer' window is overlaid on the main window, showing a directed acyclic graph (DAG) representing the BayesNet model. The graph has a root node 'play' which is connected to four child nodes: 'outlook', 'temperature', 'humidity', and 'windy'.

The status bar at the bottom of the Weka Explorer window shows 'Status OK' and a 'Log' button. The Windows taskbar at the bottom of the screen shows the 'Inicio' button and several open applications, including 'Weka-3-4', 'Weka GUT ...', 'Weka Expl...', 'Weka Clas...', and 'Weka Clas...'. The system clock shows '10:08 a.m.'.

Clasificador Bayesiano



Class yes: $P(C) = 0.625$

Attribute outlook		
sunny	overcast	rainy
0.25	0.41666667	0.33333333
Attribute temperature		
hot	mild	cool
0.25	0.41666667	0.33333333
Attribute humidity		
high	normal	
0.36363636	0.63636364	
Attribute windy		
TRUE	FALSE	
0.36363636	0.63636364	

WEKA: Evaluación

Weka Explorer

Preprocess **Classify** Cluster Associate Select attributes Visualize

Classifier: Choose **Id3**

Test options:
☐ Use training set
☐ Supplied test set (Set...)
☒ Cross-validation Folds: 10
☐ Percentage split %: 66
More options...

(Nom) play

Start Stop

Result list (right-click for options):
09:15:06 - trees.Id3

Classifier output:

```
outlook = overcast: yes
outlook = rainy
| windy = TRUE: no
| windy = FALSE: yes

Time taken to build model: 0 seconds

=== Stratified cross-validation ===
=== Summary ===

Correctly Classified Instances      12           85.7143 %
Incorrectly Classified Instances    2           14.2857 %
Kappa statistic                    0.6889
Mean absolute error                 0.1429
Root mean squared error             0.378
Relative absolute error             30 %
Root relative squared error         76.6097 %
Total Number of Instances          14

=== Detailed Accuracy By Class ===

TP Rate  FP Rate  Precision  Recall  F-Measure  Class
0.889    0.2     0.889     0.889   0.889     yes
0.8      0.111    0.8       0.8     0.8       no

=== Confusion Matrix ===

a b  <-- classified as
8 1 | a = yes
1 4 | b = no
```

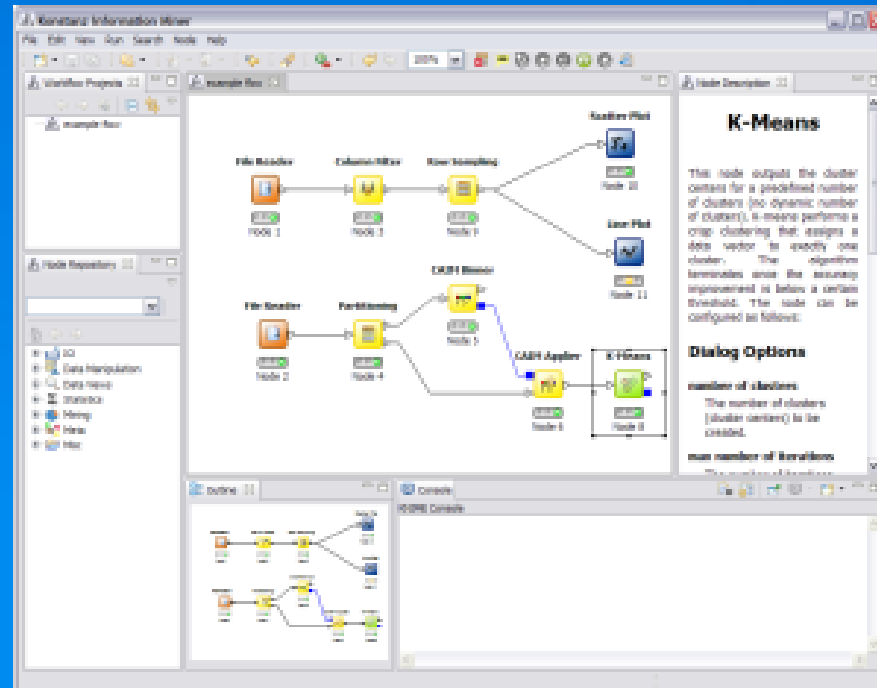
Status: OK

Log

Windows taskbar: Inicio, [um] Just Publi..., Cursada 2006, Clasificación.ppt, Weka-3-4, Weka GUI Choo..., Weka Explorer, ES, Norton, 09:18 a.m.

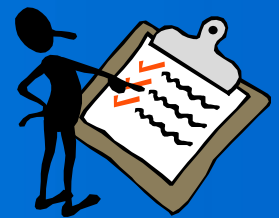
Otras herramientas

- Knime: <http://www.knime.org/>



Clasificación

- Concepto
 - Clasificación
 - Predicción
 - Evaluación
- Árboles de Decisión
 - Construcción
 - Uso
 - Poda
- Clasificador Bayesiano
- Ejemplo: Detección de líderes de equipo



Ejemplo: clasificación de mascotas

Ejemplos	Estatura	Pelaje	Ojos	Clase
e_1	Baja	Amarillo	Negros	A
e_2	Alta	Amarillo	Castaños	B
e_3	Alta	Cobrizo	Negros	A
e_4	Baja	Marrón	Negros	B
e_5	Alta	Marrón	Negros	B
e_6	Alta	Amarillo	Negros	A
e_7	Alta	Marrón	Castaños	B
e_8	Baja	Amarillo	Castaños	B

3 ejemplos en clase A; 5 ejemplos en clase B

- $H(E) = I(s_A, s_B) = \frac{3}{8} \log_2(8/3) + \frac{5}{8} \log_2(8/5) = 0.954$
 - $E_1 = \{\text{Ejemplos de } E \text{ que tienen pelaje=marrón}\}$
 - $E_2 = \{\text{Ejemplos de } E \text{ que tienen pelaje=cobrizo}\}$
 - $E_3 = \{\text{Ejemplos de } E \text{ que tienen pelaje=amarillo}\}$
-

ID3 Ejemplo

$$E = \{e_1, e_2, e_3, e_4, e_5, e_6, e_7, e_8\}$$

marrón

cobrizo

amarillo

$$E_2 = \{e_3: \text{alta, cobrizo, negros, A}\}$$

$$E_1 = \{e_4: \text{baja, marrón, negros, B}, e_5: \text{alta, marrón, negros, B}, e_7: \text{alta, marrón, castaños, B}\}$$

$$E_3 = \{e_1: \text{baja, amarillo, negros, A}, e_2: \text{alta, amarillo, castaños, B}, e_6: \text{alta, amarillo, negros, A}, e_8: \text{baja, amarillo, castaños, B}\}$$

- $H(\text{pelaje}, E) = p(\text{pelaje}=\text{marrón}/E) * H(E_1) +$
 $p(\text{pelaje}=\text{cobrizo}/E) * H(E_2) +$
 $p(\text{pelaje}=\text{amarillo}/E) * H(E_3)$

$$H(E_1) = p^A \log_2(1/p^A) + p^B \log_2(1/p^B) = 0 \log_2(1/0) + 1 \log_2(1/1) = 0$$

$$H(E_2) = p^A \log_2(1/p^A) + p^B \log_2(1/p^B) = 1 \log_2(1/1) + 0 \log_2(1/0) = 0$$

$$H(E_3) = p^A \log_2(1/p^A) + p^B \log_2(1/p^B) = 2/4 \log_2(4/2) + 2/4 \log_2(4/2) \\ = 0.5 + 0.5 \\ = 1$$

$$H(\text{pelaje}/E) = 3/8 \cdot 0 + 1/8 \cdot 0 + 4/8 \cdot 1 = 0.5$$

$$\text{Ganancia}(\text{pelaje}) = 0.954 - 0.5 = 0.454$$

ID3 Ejemplo

- $Ganancia(pelaje) = 0.454$
- $Ganancia(estatura) = 0.003$
- $Ganancia(ojos) = 0.347$

