

dr inż. Artur Klepaczko

Eksploracja danych Drzewa i reguły decyzyjne

Zadanie nr 13 – Studia podyplomowe „Przetwarzanie i analiza obrazów biomedycznych”



KAPITAŁ LUDZKI
NARODOWA STRATEGIA SPÓJNOŚCI

UNIA EUROPEJSKA
EUROPEJSKI
FUNDUSZ SPOŁECZNY



Prezentacja multimedialna
współfinansowana przez Unię Europejską
w ramach Europejskiego Funduszu Społecznego
w projekcie

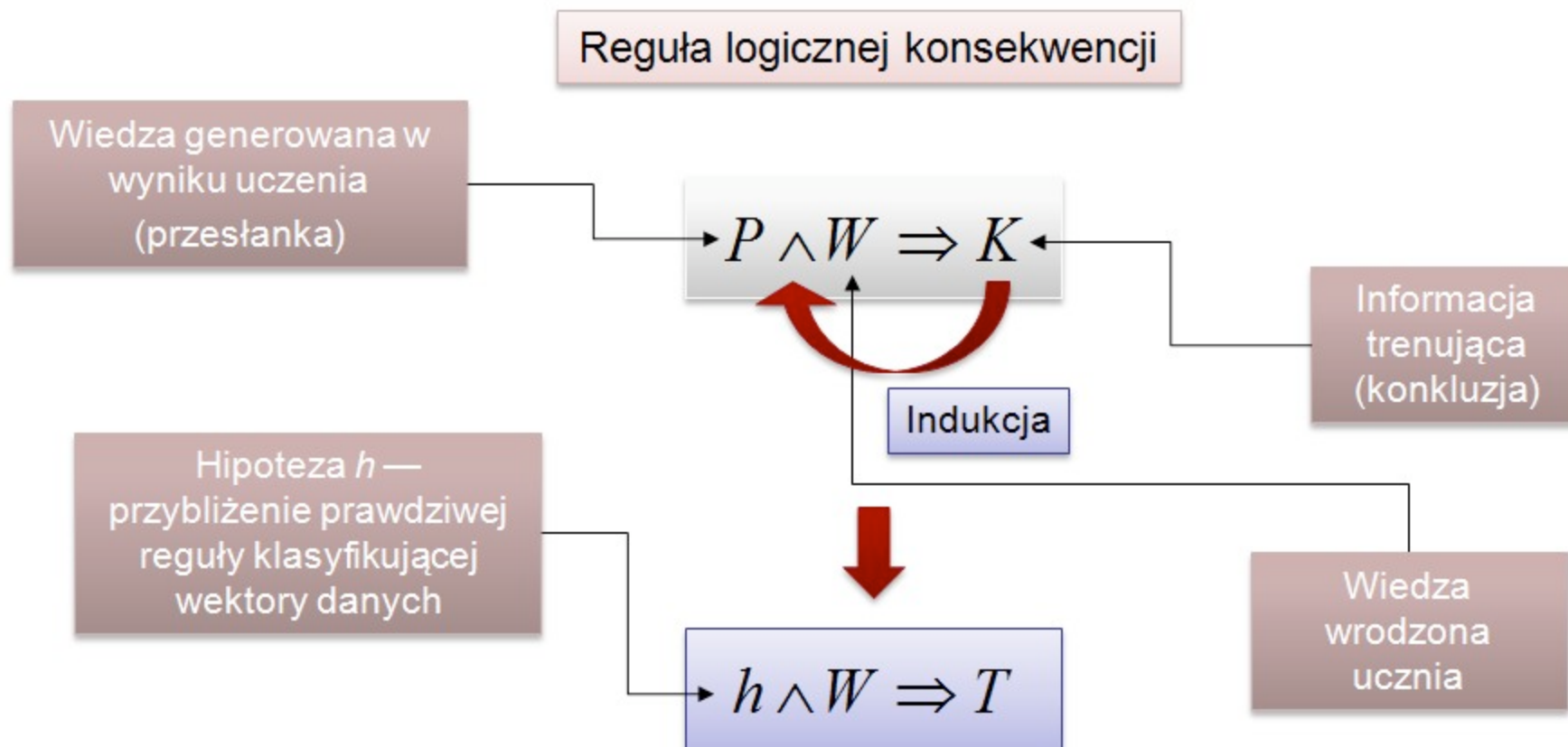
*„Innowacyjna dydaktyka bez ograniczeń
– zintegrowany rozwój Politechniki Łódzkiej –
zarządzanie Uczelnią,
nowoczesna oferta edukacyjna
i wzmacniania zdolności do zatrudniania
osób niepełnosprawnych”*



Politechnika Łódzka
Instytut Elektroniki

90-924 Łódź, ul. Żeromskiego 116,
tel. 042 631 28 83
www.kapitalludzki.p.lodz.pl

Uczenie przez indukcję

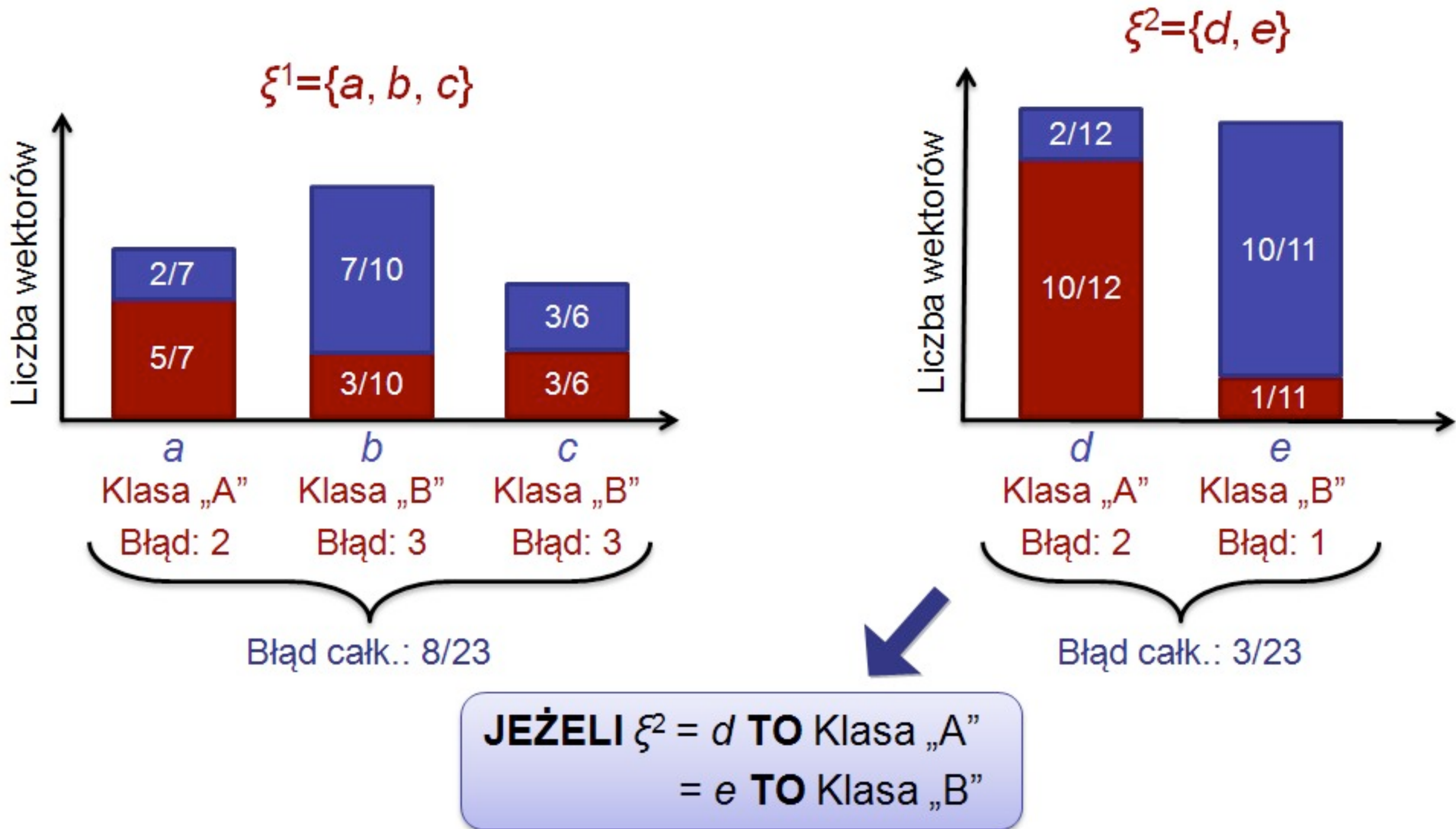


We wnioskowaniu indukcyjnym zakładamy, że informacja trenująca dostępna w procesie uczenia jest logiczną konsekwencją hipotezy h oraz być może niepustej wiedzy wrodzonej ucznia W .

Indukcja reguł decyzyjnych



Reguły proste – algorytm 1-R

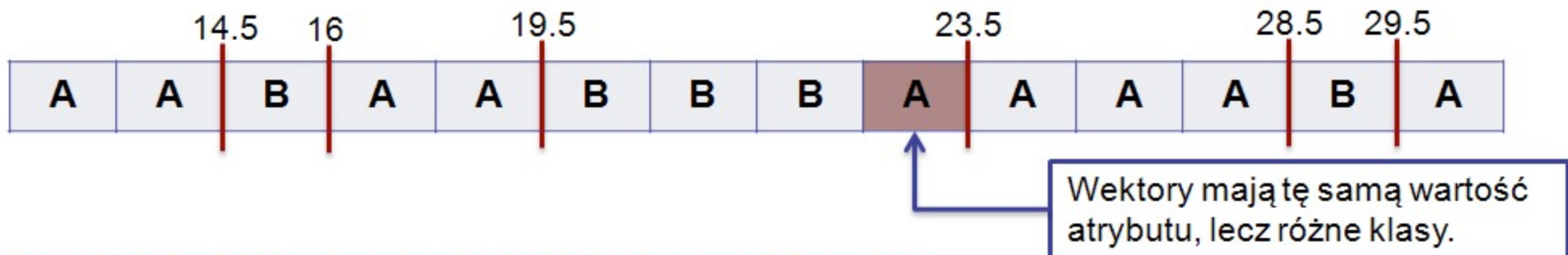


Dyskretyzacja atrybutów ilościowych

→ Porządkowanie według atrybutu numerycznego

12	14	15	17	19	20	21	22	22	25	27	28	29	30
A	A	B	A	A	B	B	B	A	A	A	A	B	A

→ Wydzielenie przedziałów wartości atrybutu, dla których klasa nie zmienia się



→ Przeciwdziałanie nadmiernemu dopasowaniu

A	A	B	A	A	B	B	B	A	A	A	A	B	A
---	---	---	---	---	---	---	---	---	---	---	---	---	---

Zbiór danych — *Iris*

Iris – zbiór pochodzący z repozytorium wzorcowych zbiorów danych UCI*



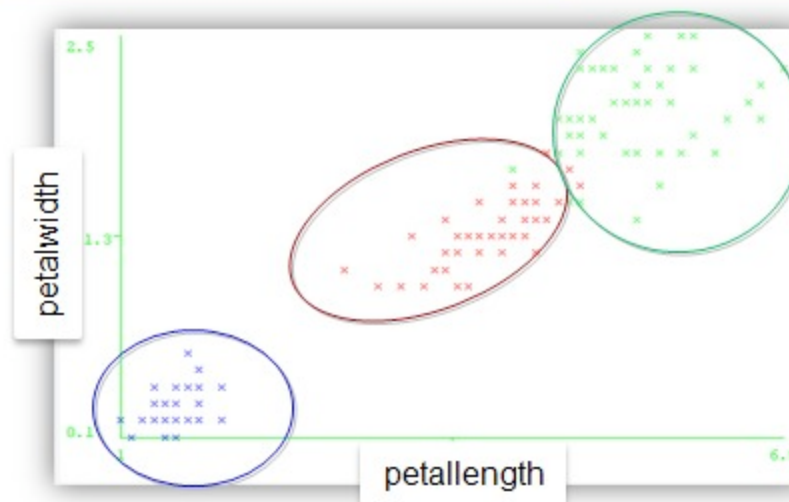
Źródło: pl.wikipedia.org.

Grafika dostępna na zasadach GFDL

Liczba wektorów danych: 150

Liczba klas: 3

Liczba cech 4 (długość i szerokość działki kielicha (ang. *sepal*) oraz płatek kwiatowego (ang. *petal*))



*UCI – University of California, Irvine. Machine Learning Repository. <http://archive.ics.uci.edu/ml/index.html>

Algorytm 1-R w zastosowaniu do zbioru danych *Iris*

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

petallength:

< 2.45 -> Iris-setosa
< 4.85 -> Iris-versicolor
≥ 4.85 -> Iris-virginica

(143/150 instances correct)

Time taken to build model: 0 seconds

=== Evaluation on training set ===

=== Summary ===

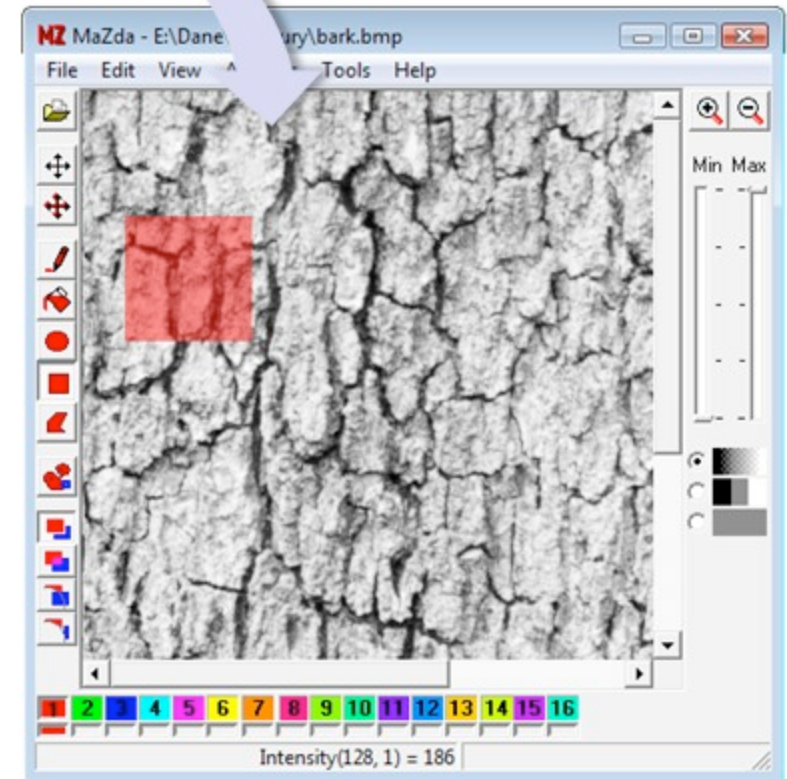
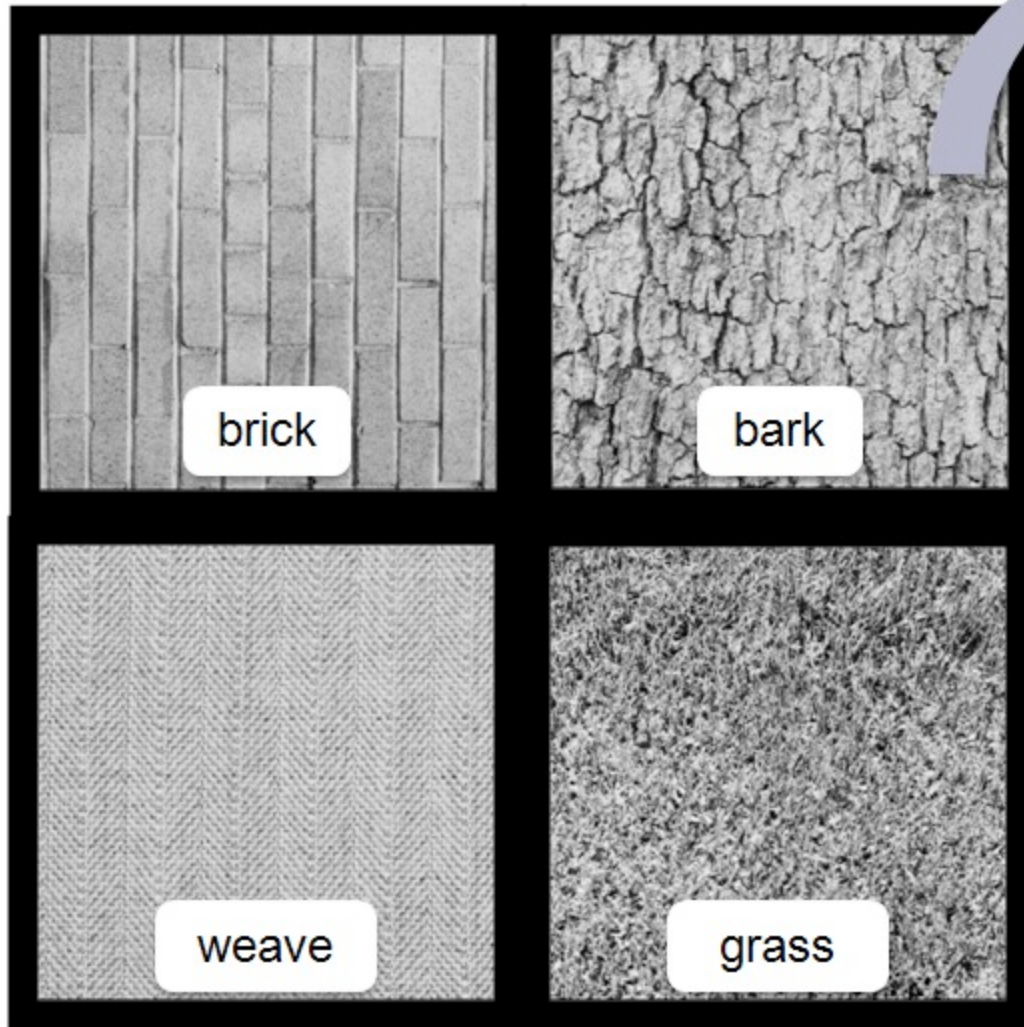
Correctly Classified Instances	143	95.3333 %
Incorrectly Classified Instances	7	4.6667 %

Wynik uzyskany przy użyciu programu Weka Explorer

Algorytm 1-R

- o skuteczny w przypadku wielu rzeczywistych zbiorów danych
- o wydajny (ze względu na krótki czas obliczeń)
- o prowadzi do prostej, łatwej w interpretacji hipotezy prawdziwej reguły decyzyjnej

Zbiór danych — *Tekstury Brodatza**



MaZda
wyznaczenie i selekcja
cech tekstury

*Brodatz P., *A photographic album for artists and designers*, Dover, New York 1966.

Algorytm 1-R – klasyfikacja tekstur

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

S(0,4)DifEntrp:

< 1.3324805 -> brick
< 1.4217385 -> calf_leather
< 1.4260485 -> grass
< 1.4309895 -> calf_leather
< 1.441526 -> grass
< 1.4443815 -> calf_leather
< 1.4808575 -> grass
≥ 1.4808575 -> herringbone_weave

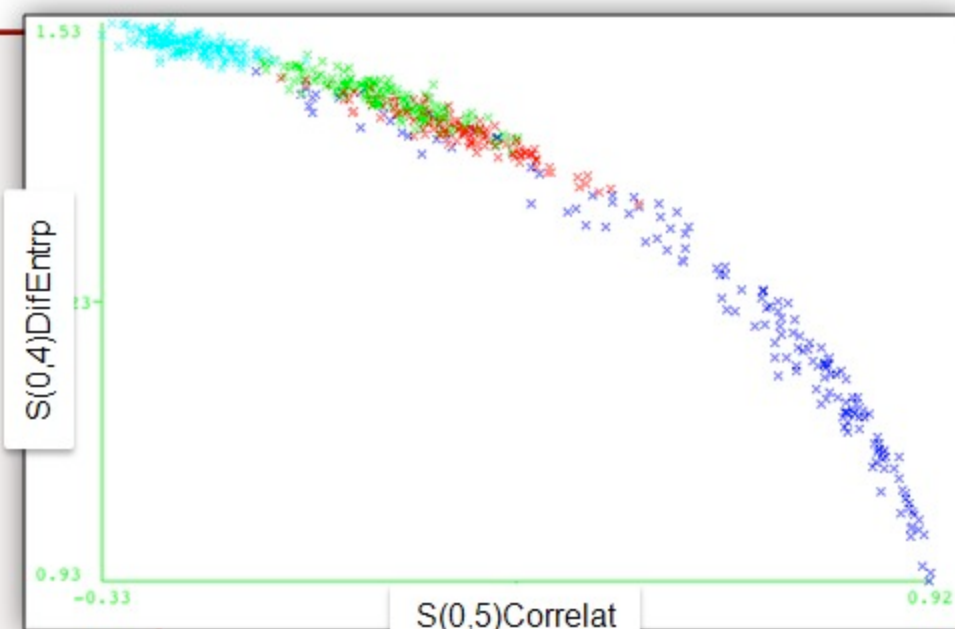
(533/640 instances correct)

Time taken to build model: 0 seconds

=== Evaluation on training set ===

=== Summary ===

Correctly Classified Instances	533	83.2813 %
Incorrectly Classified Instances	107	16.7188 %



- Wektory należące do różnych klas tekstur nie tworzą wyraźnie odrębnych skupień.
- Klasyfikator 1-R mimo to osiąga dokładność na poziomie 80%.

Walidacja krzyżowa – ocena skuteczności uczenia

Zbiór wektorów danych o liczności Q



Podział na trzy podzbiory

$S_1 (1/3 \cdot Q)$

$S_2 (1/3 \cdot Q)$

$S_3 (1/3 \cdot Q)$

Trzy cykle uczenia:

Nr	Podzbiór uczący	Podzbiór testowy	Błąd
1	S_1, S_2	S_3	e_1
2	S_1, S_3	S_2	e_2
3	S_2, S_3	S_1	e_3

Błąd na zbiorze uczącym:

$$e = (e_1 + e_2 + e_3) / 3$$

Uwiarygodniona ocena klasyfikatora – przykład

Ocena ze zbiorym uczącym

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

S(0,4)DifEntrp:

< 1.3324805 -> brick
< 1.4217385 -> calf_leather
< 1.4260485 -> grass
< 1.4309895 -> calf_leather
< 1.441526 -> grass
< 1.4443815 -> calf_leather
< 1.4808575 -> grass
>= 1.4808575 -> herringbone_weave

(533/640 instances correct)

=== Evaluation on training set ===

=== Summary ===

Correctly Classified Instances	533	83.2813 %
Incorrectly Classified Instances	107	16.7188 %

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

S(0,4)DifEntrp:

< 1.3324805 -> brick
< 1.4217385 -> calf_leather
< 1.4260485 -> grass
< 1.4309895 -> calf_leather
< 1.441526 -> grass
< 1.4443815 -> calf_leather
< 1.4808575 -> grass
>= 1.4808575 -> herringbone_weave

(533/640 instances correct)

=== Stratified cross-validation ===

=== Summary ===

Correctly Classified Instances	510	79.6875 %
Incorrectly Classified Instances	130	20.3125 %

Walidacja krzyżowa

Reguła pokrycia

Przykład

Poprawność reguły

Zbiór danych

Reguła dla klasy „A”

Bardziej poprawna
reguła dla „A”Pokrywa część zbioru
wektorów danychWybór opcji 4. –
największa poprawność

Etap 1.

1. JEŚLI $\xi^1=a$ TO Klasa „A”	5/7
2. JEŚLI $\xi^1=b$ TO Klasa „A”	3/10
3. JEŚLI $\xi^1=c$ TO Klasa „A”	3/6
4. JEŚLI $\xi^2=d$ TO Klasa „A”	10/12
5. JEŚLI $\xi^2=e$ TO Klasa „A”	1/11

Etap 2.

1. Jeśli $\xi^2=d$ & $\xi^1=a$ TO Klasa „A”	5/5
2. Jeśli $\xi^2=d$ & $\xi^1=b$ TO Klasa „A”	3/5
3. Jeśli $\xi^2=d$ & $\xi^1=c$ TO Klasa „A”	2/2

Wybór opcji 1. – największa
poprawność, największe pokrycie

Reguła pokrycia w zastosowaniu do zbioru *Iris*

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

Prism rules

If petal

If petal

If petal

If petal

If petal

Time taken to build model: 0.01 seconds

=== Evaluation on training set ===

=== Summary ===

Correctly Classified Instances	147	98%
Incorrectly Classified Instances	3	2%

Algorytm pokrycia

- prowadzi do wielu, skomplikowanych reguł
- niekoniecznie zapewnia mniejszy błąd klasyfikacji niż reguła 1R
- poszczególne reguły mogą pokrywać ten sam podzbiór wektorów danych

Wynik uzyskany przy użyciu programu Weka Explorer

Przycinanie reguł

Wiele rozbudowanych reguł – czy warto?

TAK

Zwiększona dokładność klasyfikacji

NIE

Niewielka poprawa, *Overfitting*

IREP* – inkrementacyjne przycinanie
ze zmniejszeniem błędu

Zbiór rozrostu → konstruowanie reguły – np. algorytm pokrycia

Zbiór przycinający → przycinanie – usuwanie poszczególnych warunków i test reguły

$$\frac{[p + (N - n)]}{T}$$

T – liczność zbioru przycinającego

p – liczba prawidłowo pokrytych wektorów

n – liczba nieprawidłowo pokrytych wektorów

N – liczba wszystkich negatywnych wektorów

*IREP – *incremental reduced-error pruning*

Reguły z wyjątkami



Przycięte reguły z wyjątkami w zastosowaniu do zbioru *Iris*

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

Ripple Down Rule Learner(Ridor) rules

class = Iris-setosa (150.0/100.0)

Except (petallength > 2.45) => class = Iris-versicolor (67.0/0.0) [33.0/0.0]

Except (petalwidth > 1.75) and (petallength > 4.85) => class = Iris-virginica (29.0/0.0) [14.0/0.0]

Total number of rules (incl. the default rule): 3

Time taken to build model: 0 seconds

=== Stratified cross-validation ===

=== Summary ===

Correctly Classified Instances	141	94 %
Incorrectly Classified Instances	9	6 %

Algorytm Ridor

- o reguły na najniższym poziomie są ignorowane
- o reguły na wyższych poziomach są przycinane (za pomocą metody IREP)
- o można ustalić minimalny stosunek liczności podzbiorów danych pokrytego i wykluczonego przez regułę

Decyzje na podstawie prawdopodobieństwa

	ξ^1			ξ^2		Klasa	
	A	B		A	B	A	B
a	5/11	2/12	d	10/11	2/12	11/23	12/23
b	3/11	7/12	e	1/11	10/12		
c	3/11	3/12					

Prawdopodobieństwo tego, że dla wektora danych należącego do klasy „A” atrybut $\xi^1=a$.

Prawdopodobieństwo tego, że wektor danych należy do klasy „B”

Jaka będzie klasa nowego wektora $x_{\text{test}} = \{b, d\}$?

Naiwny klasyfikator Bayes'a

Wektor testowy: $\mathbf{x}_{\text{test}} = \{b, d\}$

Prawdopodobieństwo *a priori* klasy „A”

$$\Pr(A) = 11/23$$

Prawdopodobieństwo *a priori* klasy „B”

$$\Pr(B) = 12/23$$

Prawdopodobieństwo *a posteriori* klasy „A”

$$\Pr(\text{Klasa „A”} \mid \mathbf{x}_{\text{test}})$$

Prawdopodobieństwo *a posteriori* klasy „B”

$$\Pr(\text{Klasa „B”} \mid \mathbf{x}_{\text{test}})$$

Iloczyny prawdopodobieństw atrybutów

$$3/11 \times 10/11 \times 11/23 \approx 0,12$$

$$7/12 \times 2/12 \times 12/23 \approx 0,05$$

Obliczenia są poprawne przy (naiwnym) założeniu niezależności atrybutów.

Klasyfikator Bayesa dla atrybutów ciągłych

- Zamiast prawdopodobieństw atrybutów oblicza się gęstość prawdopodobieństwa
- Należy założyć określony typ rozkładu i funkcji gęstości prawdopodobieństwa, np. dla rozkładu normalnego →

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

- Dla każdej klasy C
 - Dla każdego atrybutu ξ^j
 - Wartość średnia $\mu^j \rightarrow$

$$\frac{\text{suma wartości atrybutu } \xi^j \text{ wektorów z klasy } C}{\text{liczność klasy } C}$$

- Odchylenie standardowe $\sigma^j \rightarrow$

$$\sqrt{\frac{\text{suma } (\xi^j - \mu^j)^2 \text{ dla wektorów z klasy } C}{\text{liczność klasy } C - 1}}$$

- Klasę dla wektora testowego wyznacza się obliczając funkcje gęstości prawdopodobieństw poszczególnych atrybutów

Naiwny klasyfikator Bayesa w zastosowaniu do zbioru tekstur

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

Naive Bayes Classifier

Attribute	Class		
	brick	calf_leather	grass
	(0.25)	(0.25)	(0.25)

S(0,5)DifEntrp

mean	1.2085	1.4297	1.4
std. dev.	0.1393	0.0253	0.0

...

Time taken to build model: 0.01 seconds

=== Stratified cross-validation ===

=== Summary ===

Correctly Classified Instances	512	80	%
Incorrectly Classified Instances	128	20	%

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

S(0,4)DifEntrp:

- < 1.3324805 -> brick
- < 1.4217385 -> calf_leather
- < 1.4260485 -> grass
- < 1.4309895 -> calf_leather
- < 1.441526 -> grass
- < 1.4443815 -> calf_leather
- < 1.4808575 -> grass
- >= 1.4808575 -> herringbone_weave

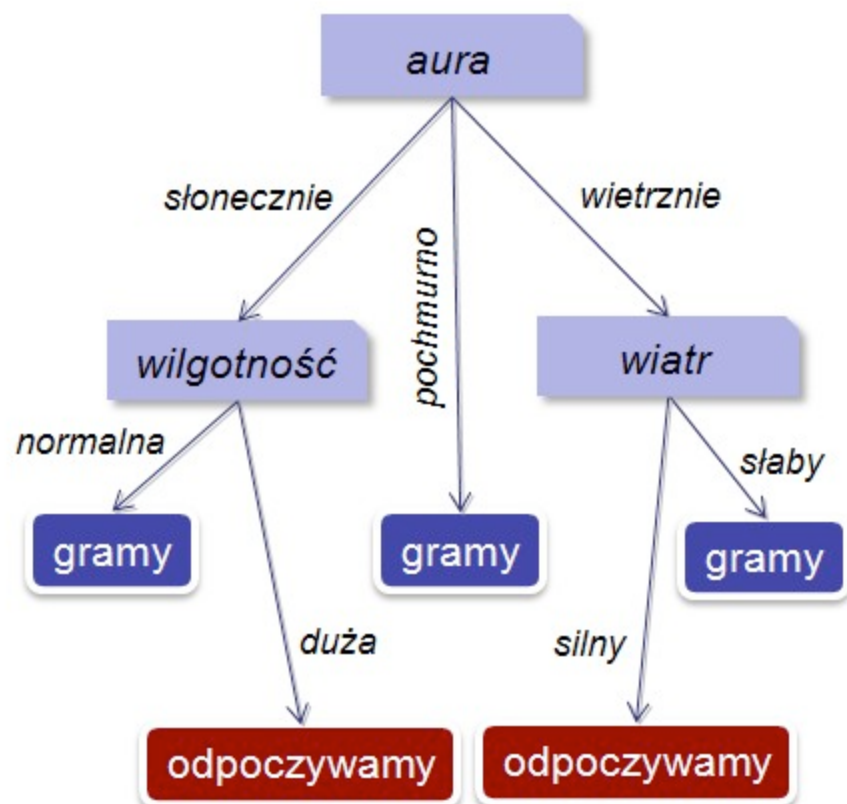
(533/640 instances correct)

=== Stratified cross-validation ===

=== Summary ===

Correctly Classified Instances	510	79.6875 %
Incorrectly Classified Instances	130	20.3125 %

Drzewa decyzyjne – metoda „dziel-i-rządź”

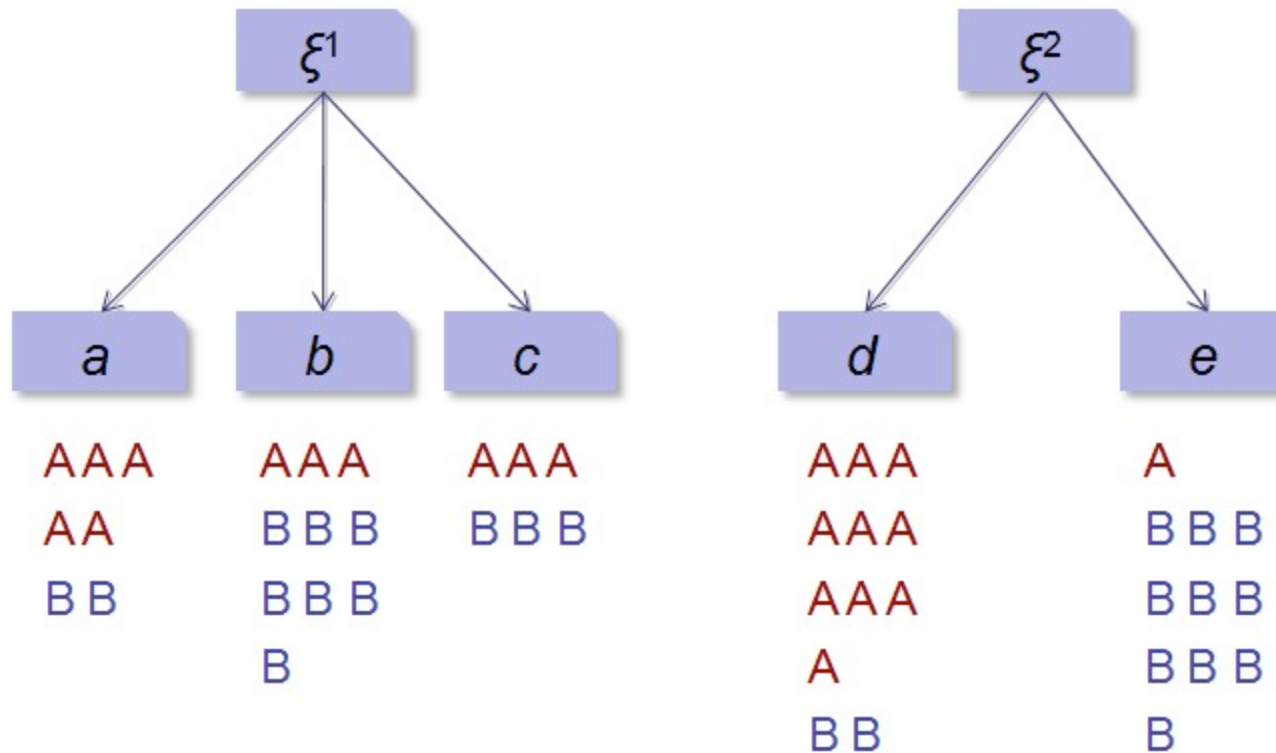


Jak to zostało skonstruowane?

- Wybierz atrybut, który zostanie umieszczony w węźle–**korzeniu**.
- Podziel** zbiór wektorów danych na podzbiory odpowiadające poszczególnym wartościom wybranego atrybutu.
- Jeśli w utworzonych podzbiorach wszystkie wektory należą do tej samej klasy, to reprezentujący je węzeł staje się **liściem**.
- W przeciwnym przypadku, węzeł staje się **rodzicem**, w którym należy wybrać kolejny atrybut rozdzielający wektory danych na podzbiory.

Konstruowanie drzew decyzyjnych – wybór atrybutu

Dwie możliwości – który atrybut lepiej nadaje się na *korzeń*?



Potrzebna jest obiektywna miara do oceny, jak poszczególne atrybuty nadają się do podziału wektorów w danym węźle drzewa.

Entropia – chaos i porządek

Definicja termodynamiczna

- Miara stopnia nieuporządkowania układu makroskopowego*

Definicja teorii informacyjnej

- Ilość informacji zawartej w wiadomości wysłanej przez źródło*

Chaos

AAAAA
BB

AAA
BBB

AAA
BBBBBBBB

Informacja niejednoznaczna — obie klasy są równie prawdopodobne.

Porządek

A

BBBBBBBBBB

AAAAAAAAAA
BB

Informacja jednoznaczna — klasa „A” jest bardziej prawdopodobna.

Zysk informacyjny

Entropia informacyjna

$$\bullet \text{Entrp}(p_1, p_2, \dots, p_n) = -p_1 \log p_1 - p_2 \log p_2 \dots - p_n \log p_n$$

Entropia zbioru w węźle-korzeniu (11 wektorów klasy „A”, 12 wektorów klasy „B”)

$$\bullet \text{Entrp}(11/23, 12/23) = -11/23 \times \log(11/23) - 12/23 \times \log(12/23) = \mathbf{0,99 \text{ bit}}$$

Średnia entropia w węźle ξ^1

$$\bullet 7/23 \times \text{Entrp}(2/7, 5/7) + 6/23 \times \text{Entrp}(3/6, 3/6) + 10/23 \times \text{Entrp}(3/10, 7/10) = \mathbf{0,26 + 0,26 + 0,38 = 0,90 \text{ bit}}$$

Średnia entropia w węźle ξ^2

$$\bullet 11/23 \times \text{Entrp}(1/11, 10/11) + 12/23 \times \text{Entrp}(10/12, 2/12) = \mathbf{0,21 + 0,34 = 0,55 \text{ bit}}$$

Zysk informacyjny w przypadku wyboru atrybutu ξ^1

$$\bullet 0,99 - 0,90 = \mathbf{0,09 \text{ bit}}$$

Zysk informacyjny w przypadku wyboru atrybutu ξ^2

$$\bullet 0,99 - 0,55 = \mathbf{0,44 \text{ bit}}$$

Drzewa decyzyjne dla atrybutów rzeczywistych

12	14	15	17	19	20	21	22	22	25	27	28	29	30
B	A	B	A	B	B	B	A	A	A	A	A	B	A

- Porządkowanie wektorów danych według atrybutu numerycznego
- Podział wektorów na dwa podzbiory
 - Drzewa binarne
 - Kryterium podziału – zysk informacyjny
- Możliwe są kolejne testy tego samego atrybutu
 - Wada: konstruowane drzewa są zbyt rozbudowane i trudne w interpretacji
- Alternatywnie: drzewa niebinarne (więcej niż dwa rozgałęzienia w węźle)
 - Wada: trudniejsze w implementacji
- Alternatywnie: dyskretyzacja atrybutu
 - Wada: utrata części informacji

Indukcja drzewa decyzyjnego – przykład

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

J48 unpruned tree

S(0,5)Correlat <= -0.063539
| S(0,5)Correlat <= -0.0943: herringbone_weave (153.0/1.0)
| S(0,5)Correlat > -0.0943
| | S(0,4)DifEntrp <= 1.485315
| | | S(0,5)DifEntrp <= 1.492711: herringbone_weave (3.0/1.0)
| | | S(0,5)DifEntrp > 1.492711: grass (5.0)
| | S(0,4)DifEntrp > 1.485315: herringbone_weave (5.0)

Number of Leaves : 15

Size of the tree : 29

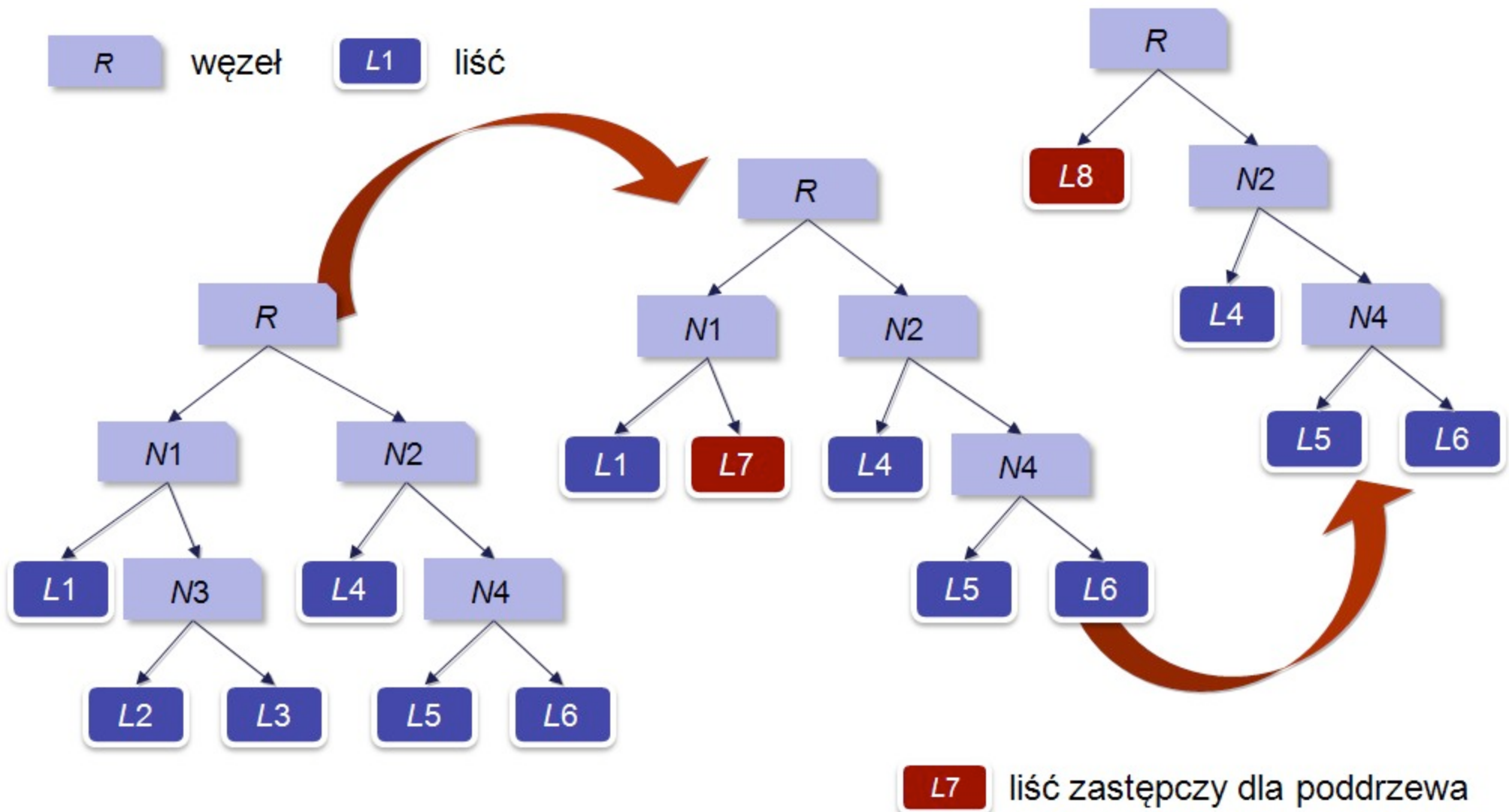
=== Summary ===

Correctly Classified Instances	522	81.5625 %
Incorrectly Classified Instances	118	18.4375 %

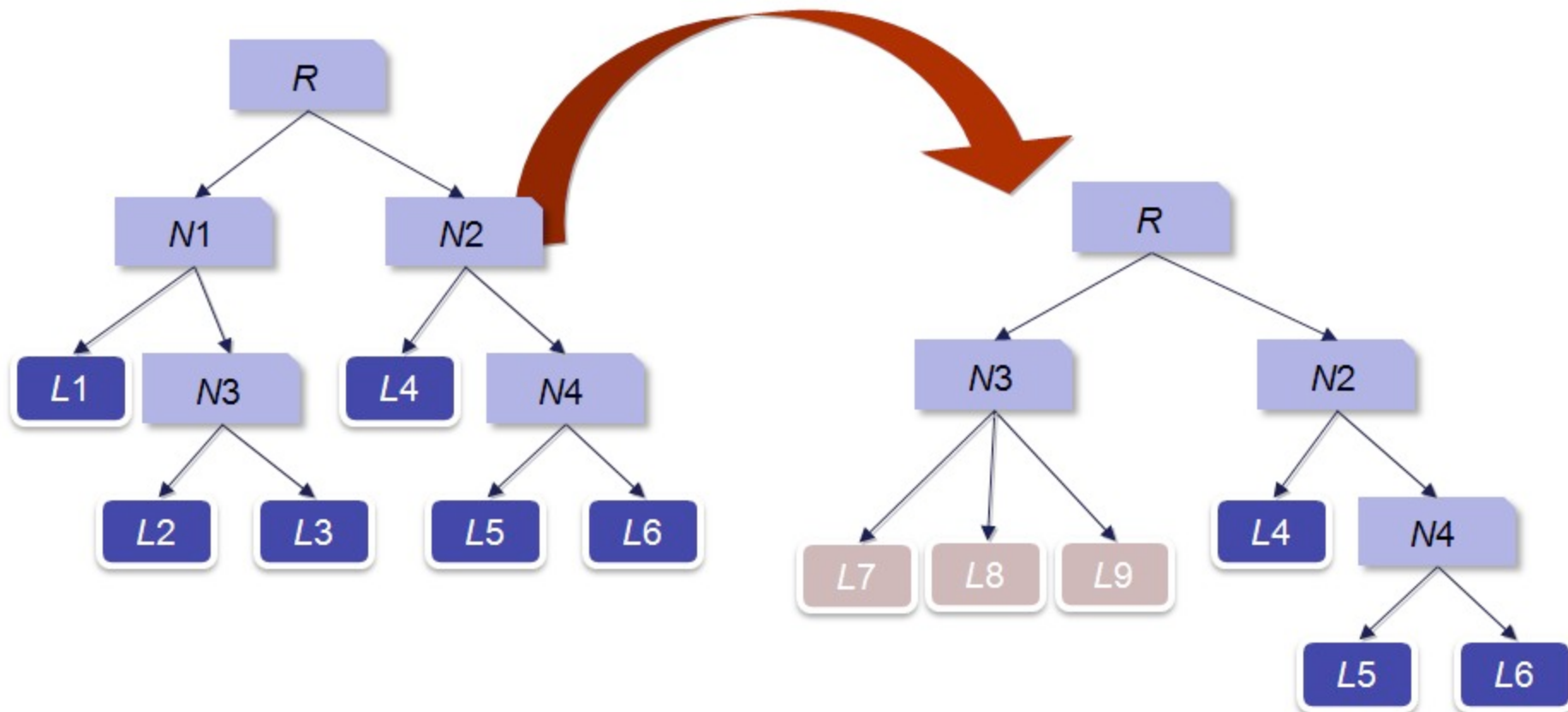
Proste drzewa decyzyjne

- o w przypadku rzeczywistych zbiorów danych powstaje wiele rozgałęzień
- o istnieje duże ryzyko nadmiernego dopasowania drzewa do zbioru treningowego
- o potrzebny jest mechanizm uproszczenia drzewa — przycinanie

Przycinanie drzewa — *subtree replacement*



Przycinanie drzewa — *subtree raising*



Kiedy przycina się drzewa?

L7 nowy liść poddrzewa

Kryterium przycinania

Przycinanie redukujące błąd

- Oddzielny zbiór testowy
- Duży zbiór trenujący, z którego można wydzielić podzbiór testowy
- Przycinanie następuje, jeśli błąd na zbiorze testowym nie zwiększy się.

Przycinanie pesymistyczne

- Brak oddzielnego zbioru testującego
- Zbiór treningowy jest zbyt mały, aby rezygnować z części wektorów danych w trakcie uczenia.
- Przycinanie następuje, jeśli nie zwiększy się pesymistyczne oszacowanie błędu rzeczywistego na podstawie zbioru treningowego.

$$e_T(l) \leq e_T(n) + \sqrt{\frac{e_T(n)(1 - e_T(n))}{|T|}}$$

Przycięte drzewa w zastosowaniu do zbioru tekstur

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

J48 pruned tree

```
S(0,5)Correlat <= -0.063539
| S(0,5)Correlat <= -0.0943: herringbone_weave (153.0/1.0)
| S(0,5)Correlat > -0.0943
| | S(0,4)DifEntrp <= 1.485315
| | | S(0,5)DifEntrp <= 1.492711: herringbone_weave (3.0/1.0)
| | | S(0,5)DifEntrp > 1.492711: grass (5.0)
| | S(0,4)DifEntrp > 1.485315: herringbone_weave (5.0)
```

Number of Leaves : 12

Size of the tree : 23

=== Summary ===

Correctly Classified Instances	527	82.3438 %
Incorrectly Classified Instances	113	17.6563 %

Algorytm C4.5

- o Redukcja liczby rozgałęzień
- o Poprawa (niekiedy znaczna) jakości klasyfikacji
- o Implementacja zarówno metody *subtree replacement* jak i *subtree raising*
- o Skalę przycięcia można kontrolować poprzez zmianę jednego z parametrów