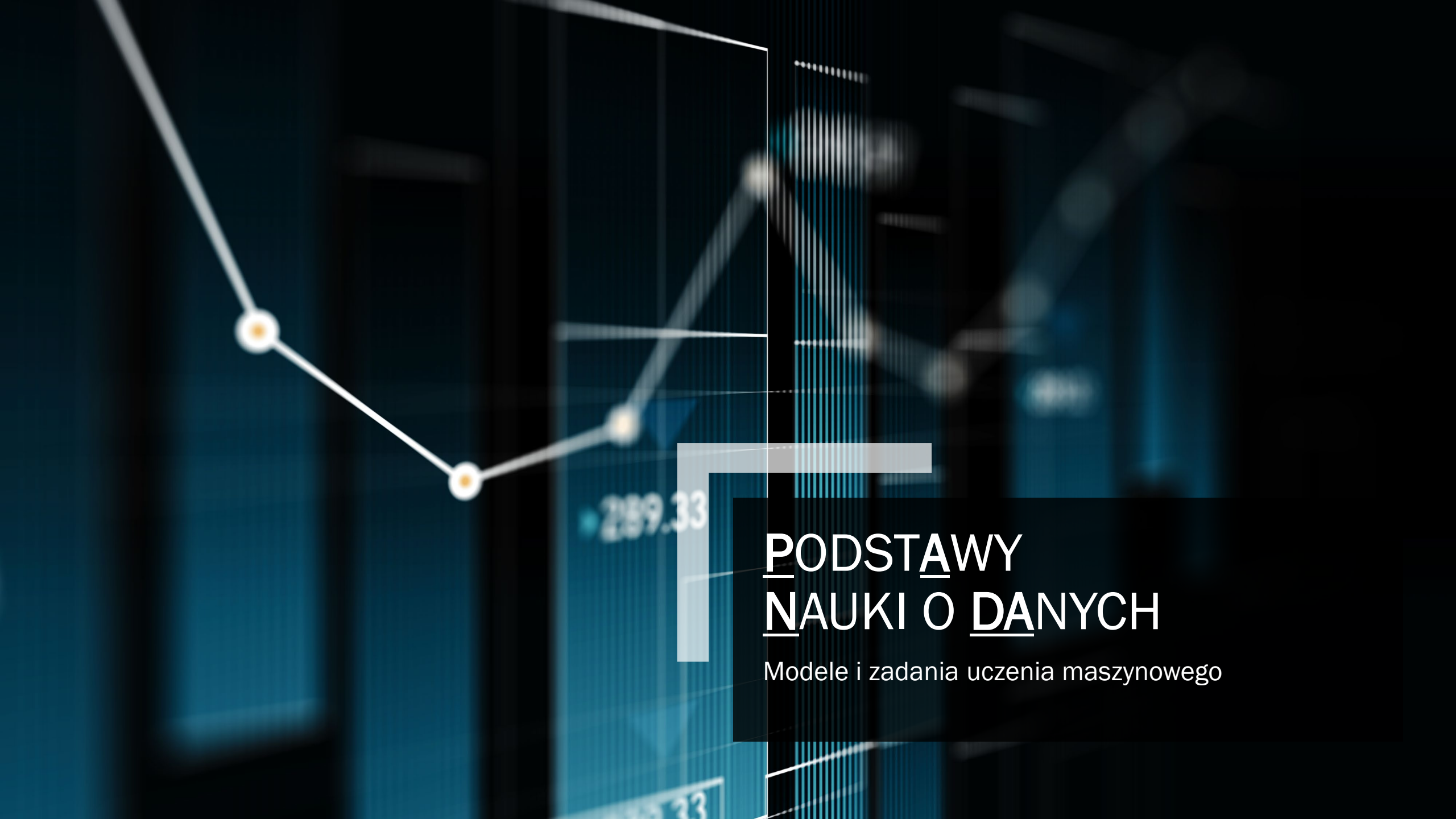
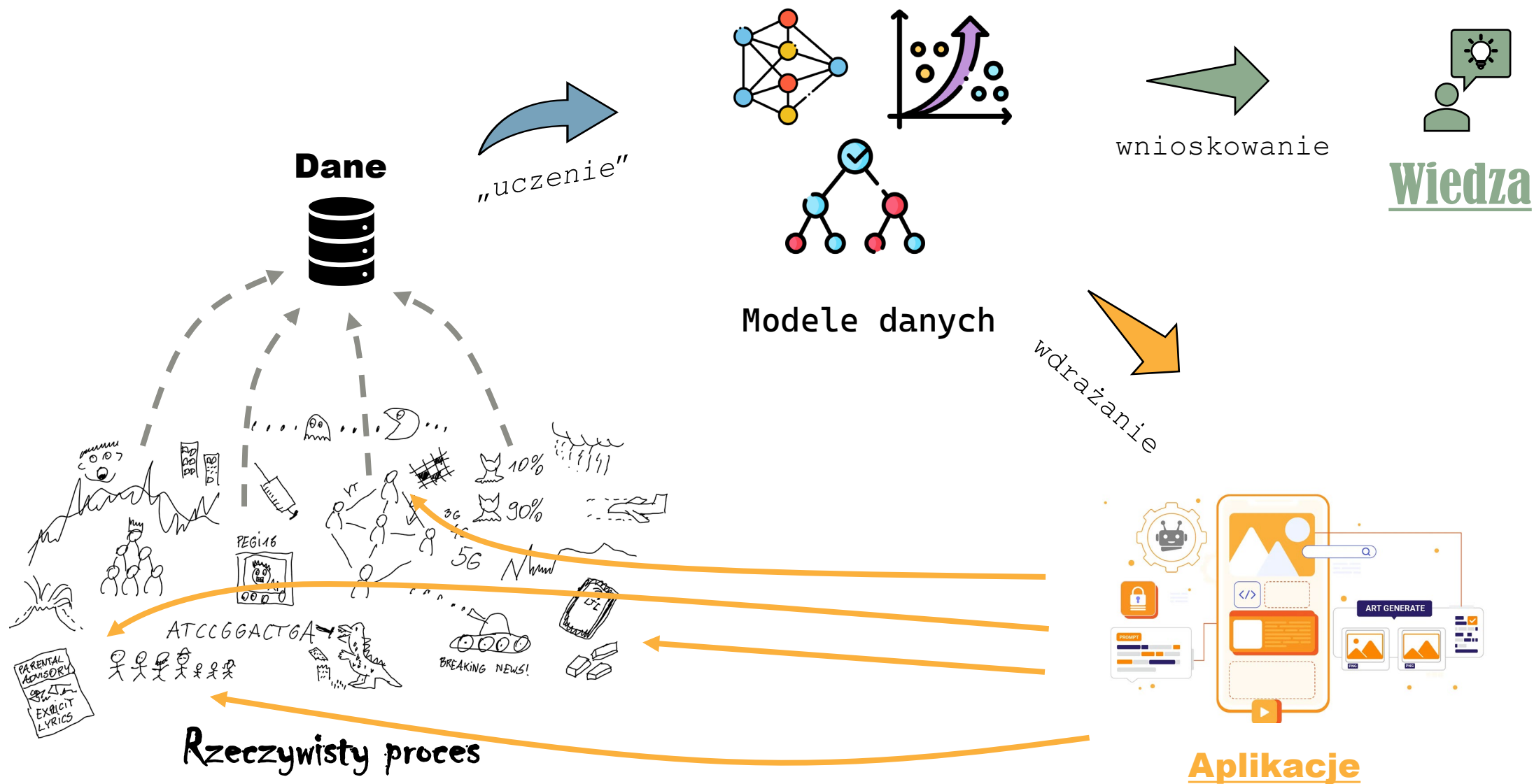




PODSTAWY NAUKI O DANYCH

Modele i zadania uczenia maszynowego



Czym jest uczenie maszynowe

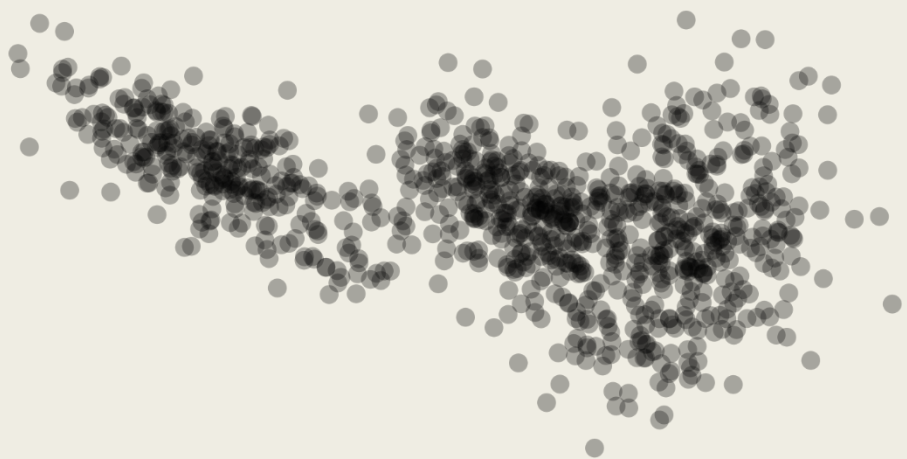
- ang. *machine learning* *learning* ≠ uczenie

I've learned that you are a confident person

- modelowanie statystycznych regularności w obserwacjach / pomiarach
- modelowanie danych ≠ modelowanie procesu
- dopasowanie modelu do danych (zadanie estymacji)
- **Def.** Proces automatycznego poszukiwania lepszej **reprezentacji danych**
w ramach zdefiniowanej przestrzeni możliwości
na podstawie sygnału **informacji zwrotnej**

PRZYKŁADY ZADAŃ UCZENIA MASZYNOWEGO

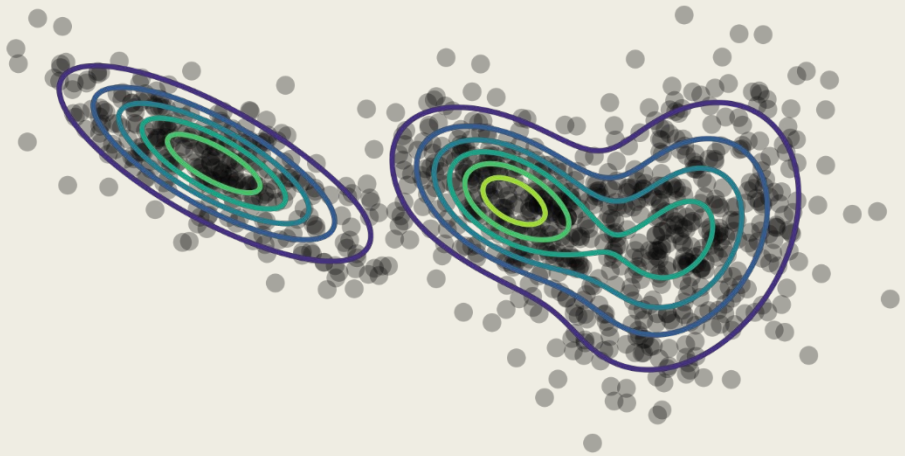




Mieszanina rozkładów Gaussa

$$p(\mathbf{x}|\theta) = \sum_{m=1}^M \pi_m \mathcal{N}(\mathbf{x}|\boldsymbol{\mu}_m, \boldsymbol{\Sigma}_m)$$

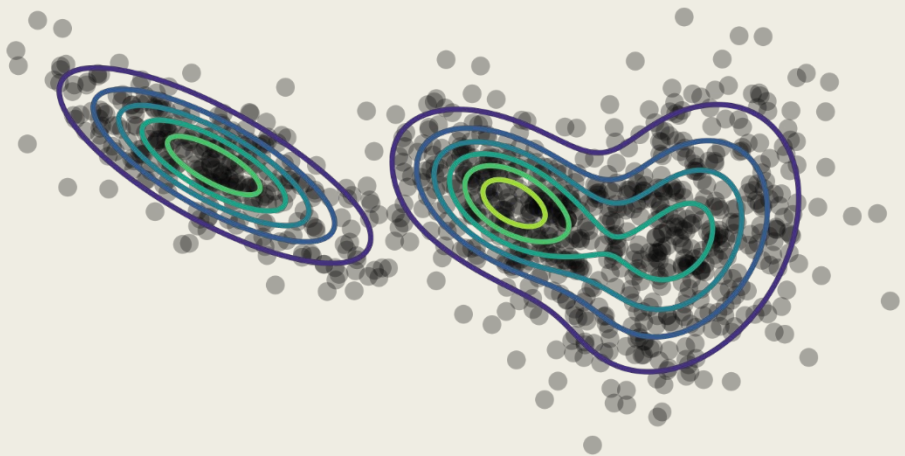
$$0 \leq \pi_m \leq 1, \quad \sum_{m=1}^M \pi_m = 1$$

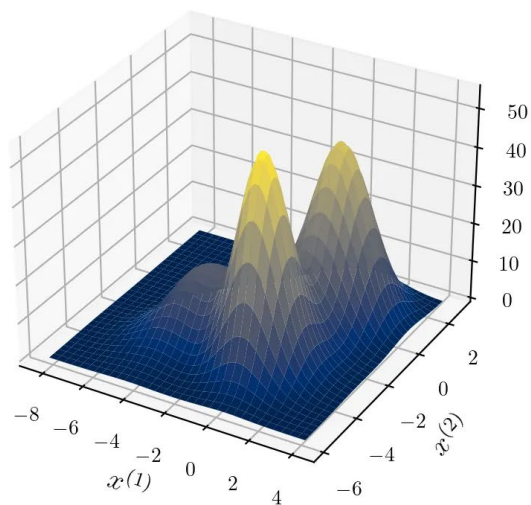
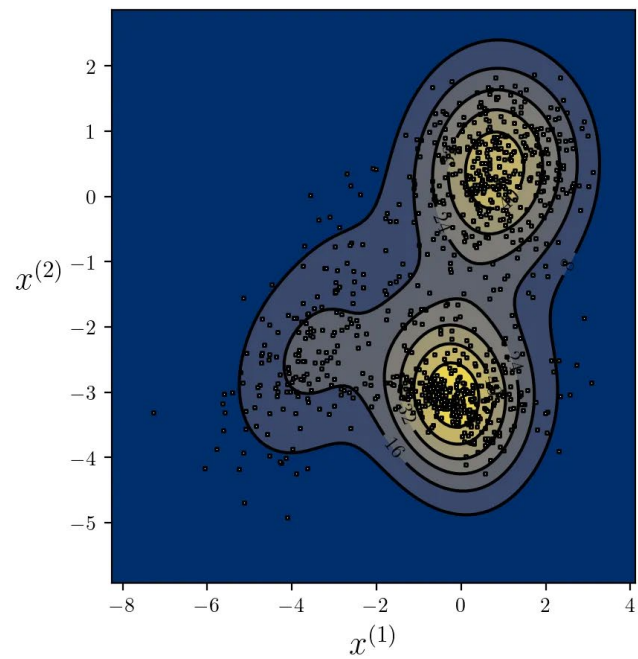
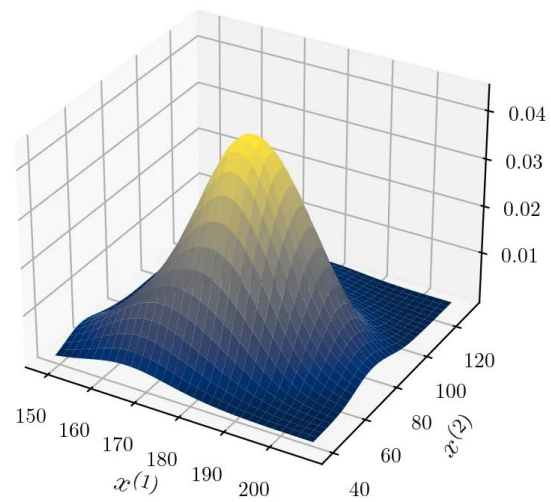
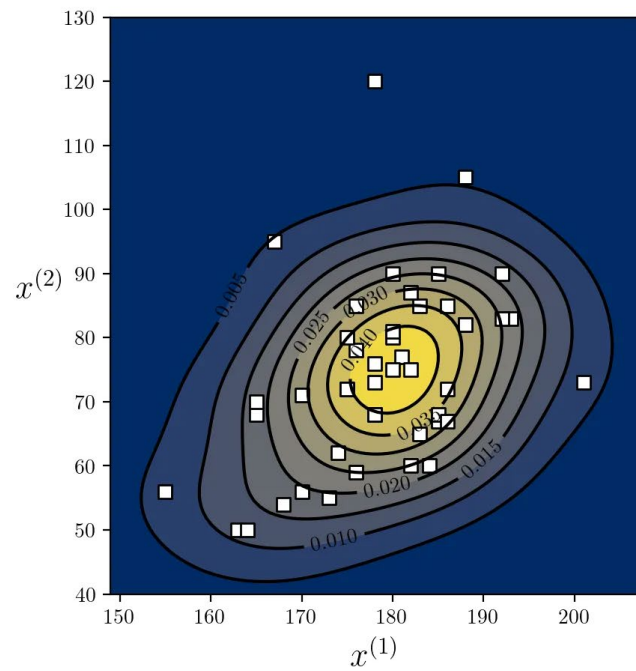


Mieszanina rozkładów Gaussa

$$p(\mathbf{x}|\theta) = \sum_{m=1}^M \pi_m \mathcal{N}(\mathbf{x}|\boldsymbol{\mu}_m, \boldsymbol{\Sigma}_m)$$

$$0 \leq \pi_m \leq 1, \quad \sum_{m=1}^M \pi_m = 1$$





Przypominajka statystyczna

$p(X; \theta)$ – rozkład zmiennej losowej X , którego parametrem jest θ

zwróć uwagę na tę różnicę

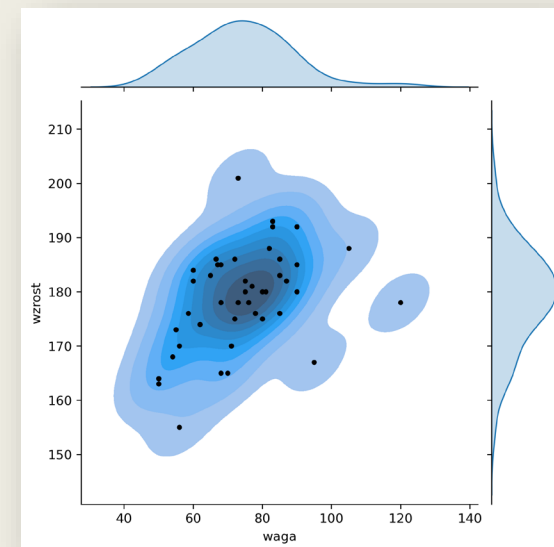
$p(X, \theta)$ – **rozkład łączny** dwóch zmiennych losowych X oraz θ

ang. *joint probability distribution*

$p(X|\theta)$ – **rozkład warunkowy** zmiennej losowej X pod warunkiem θ

ang. *conditional PDF*

$p(\theta|X)$ – **rozkład warunkowy** zmiennej losowej θ pod warunkiem X



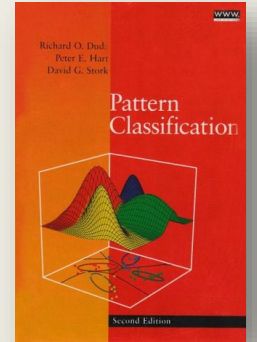
$$p(X, \theta) = p(\theta, X)$$

Probability Distribution Function

Probability Density Function

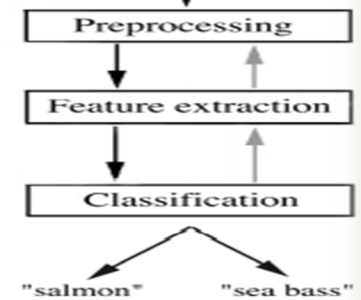
➤ PDF

Klasyfikacja

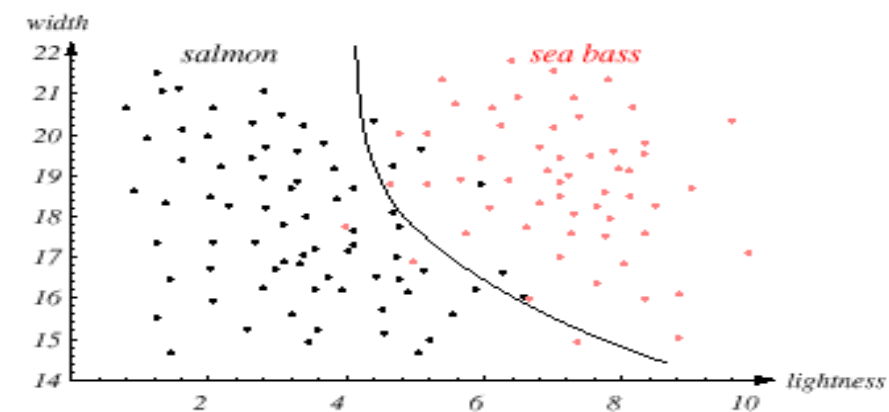
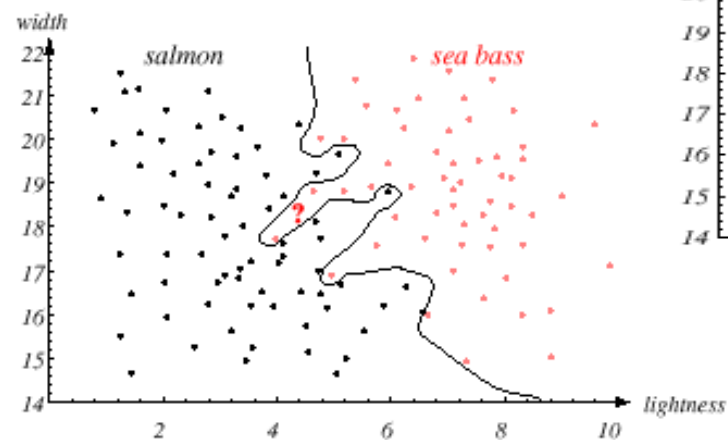
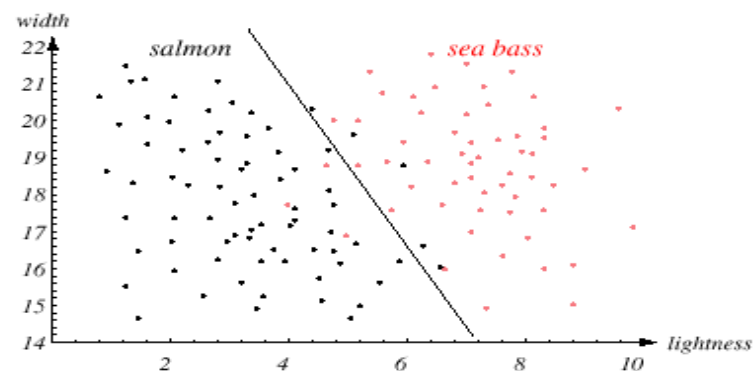
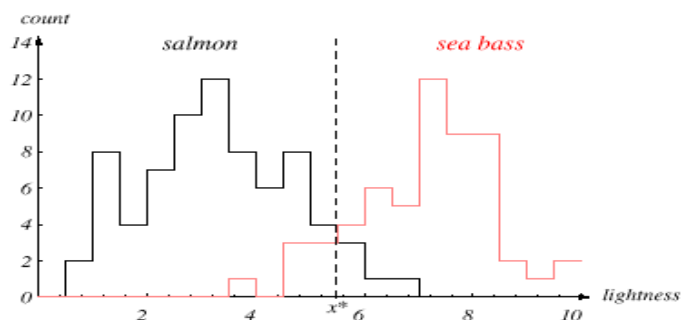
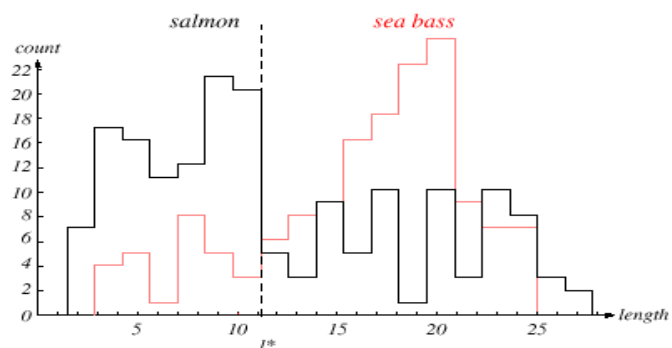
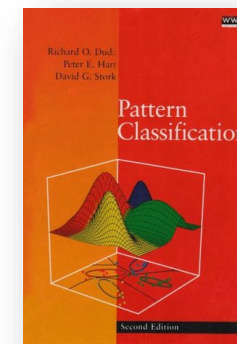


Cechy:

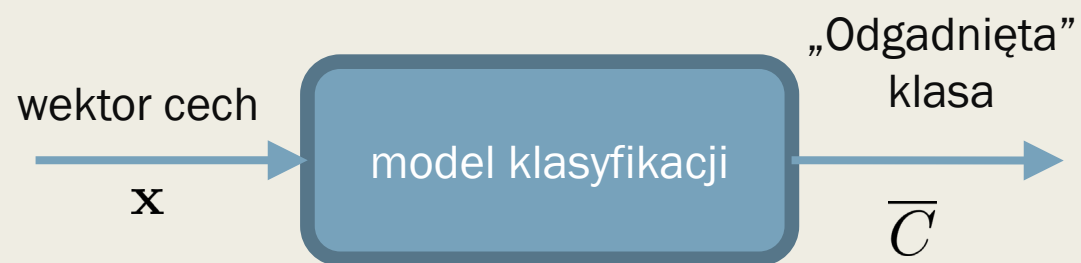
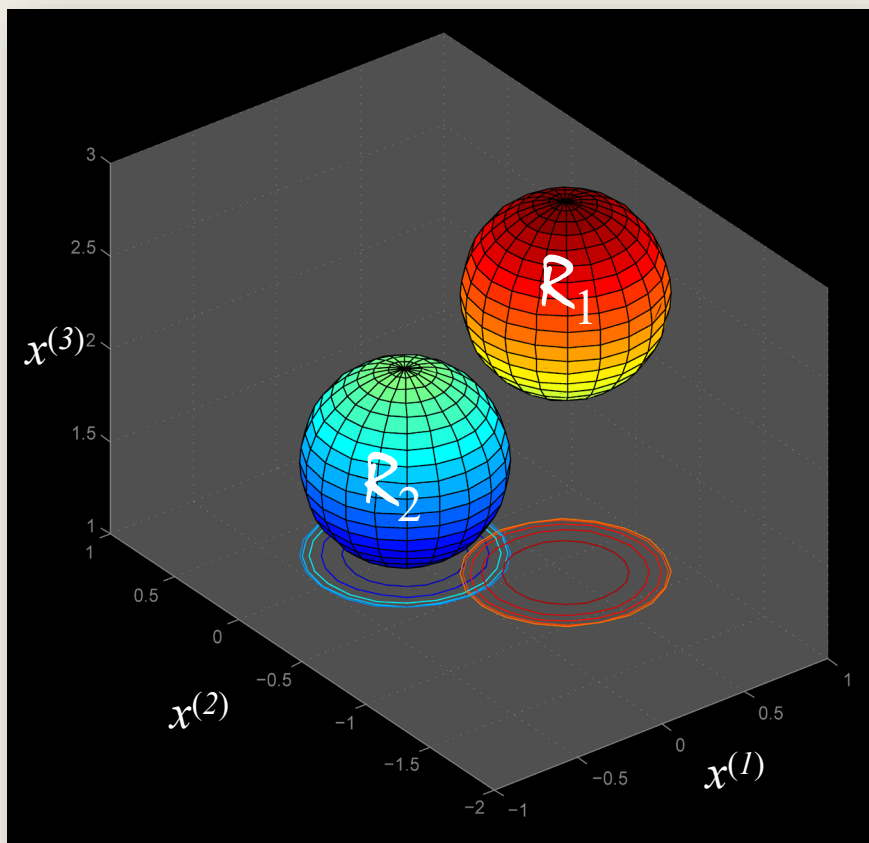
- ✓ *długość*
- ✓ *jasność*
- ✓ *szerokość*
- ✓ *liczba i kształt płetw*
- ✓ *położenie otworu gębowego*



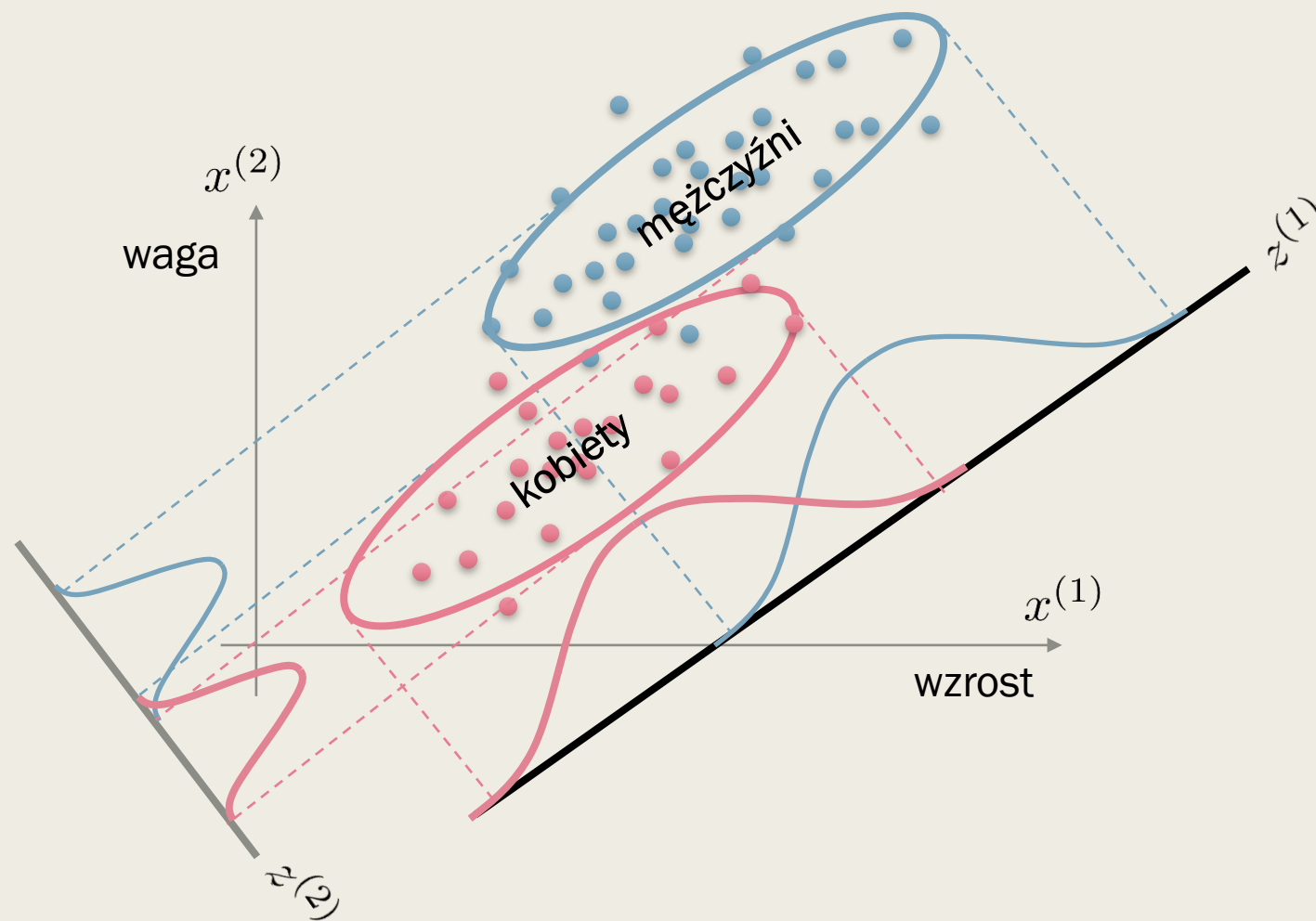
Klasyfikacja



Klasyfikacja



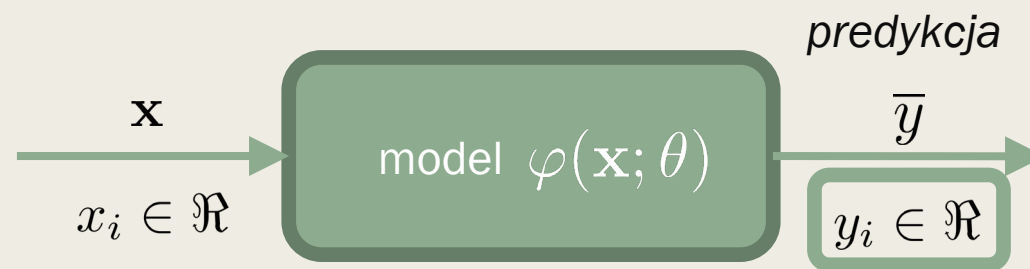
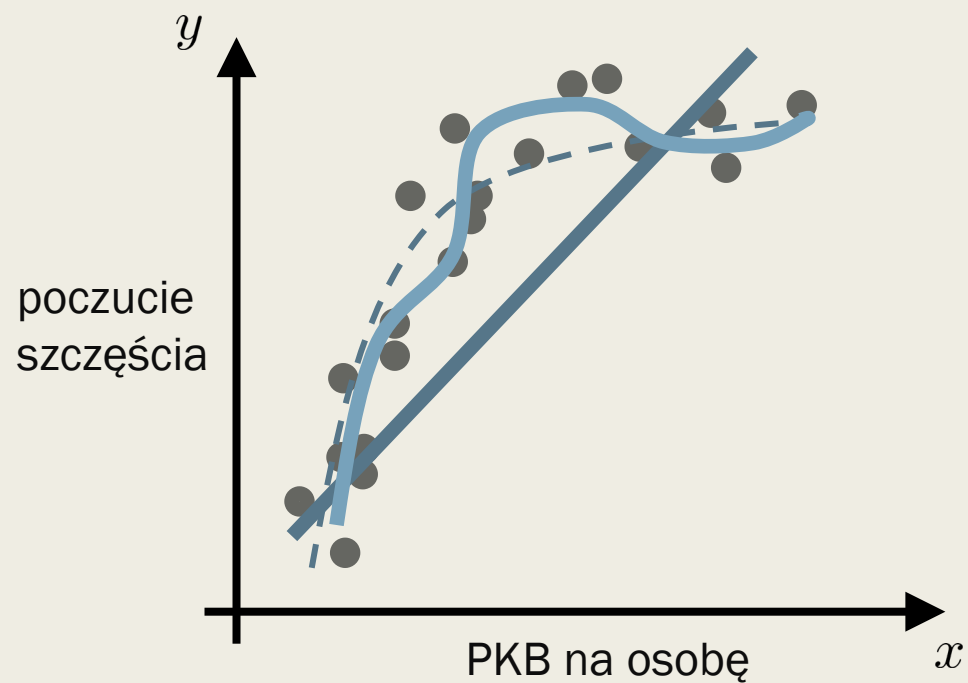
Transformacje



Przykłady modeli

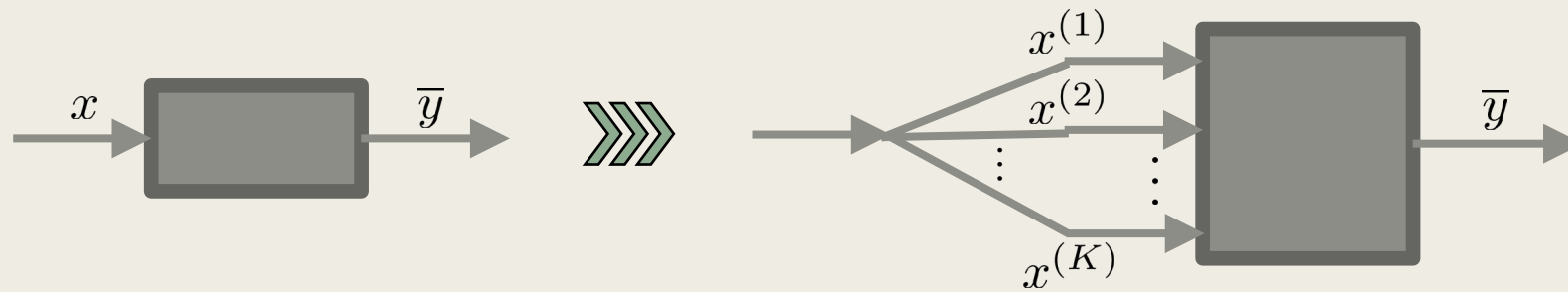
- k -najbliższych sąsiadów
- klasyfikator liniowy
- regresja logistyczna
- maszyny wektorów wspierających (ang. *Support Vector Machines*)
- drzewa decyzyjne (lasy)
- sieci neuronowe
- ...

Regresja



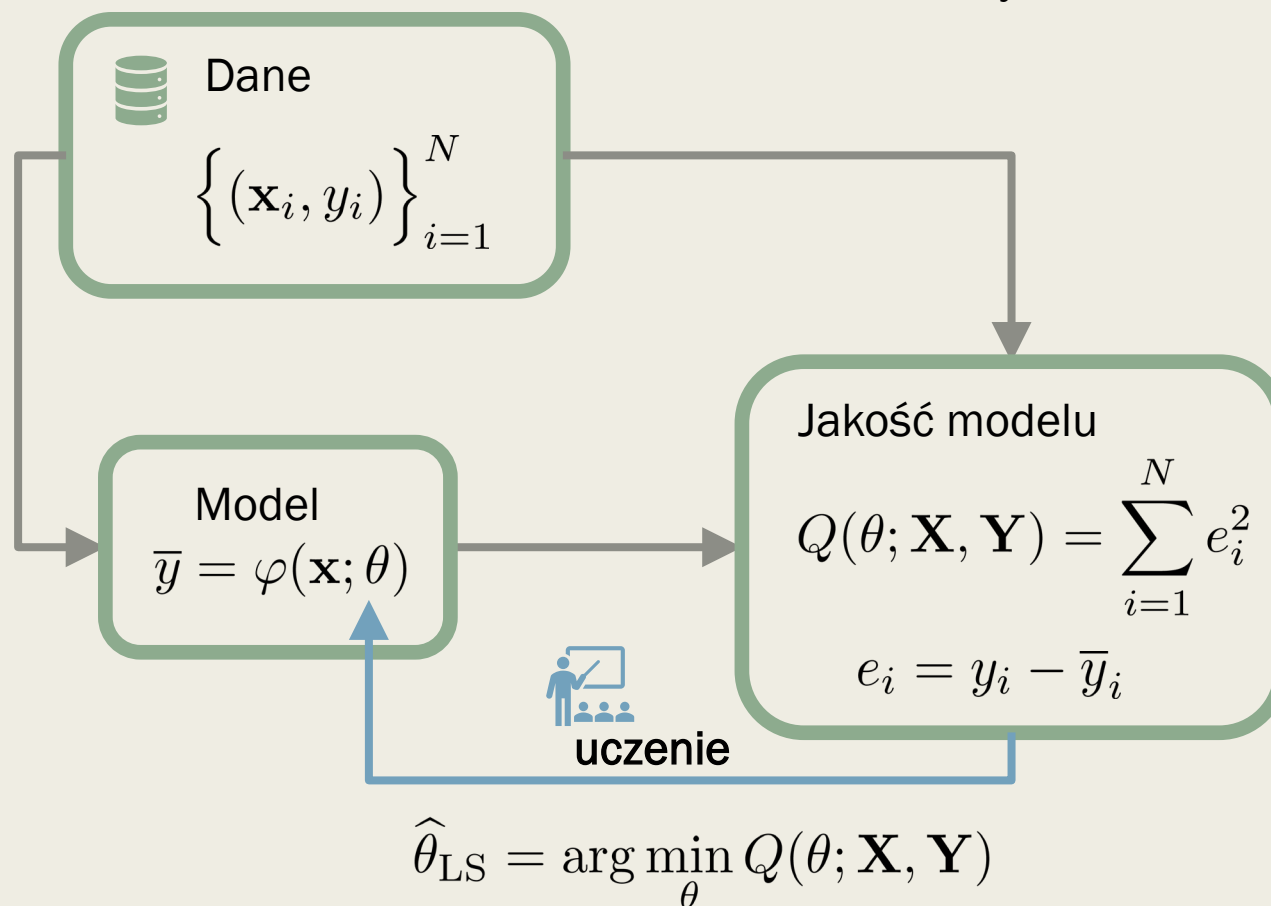
→ Model odniesienia / naiwny

Uogólnione modele liniowe



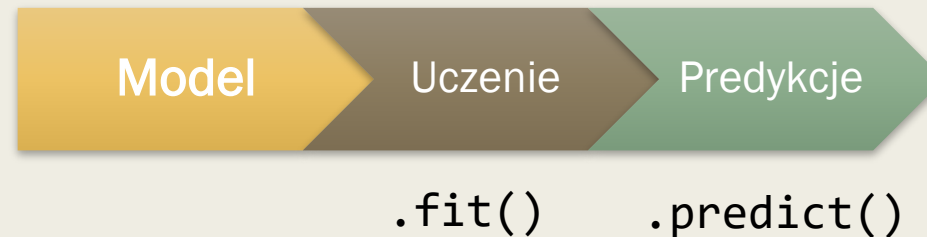
Regresja

Dopasowanie
modelu do
danych = uczenie

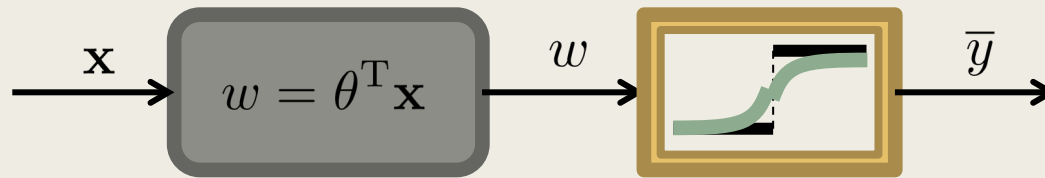


Przykłady modeli

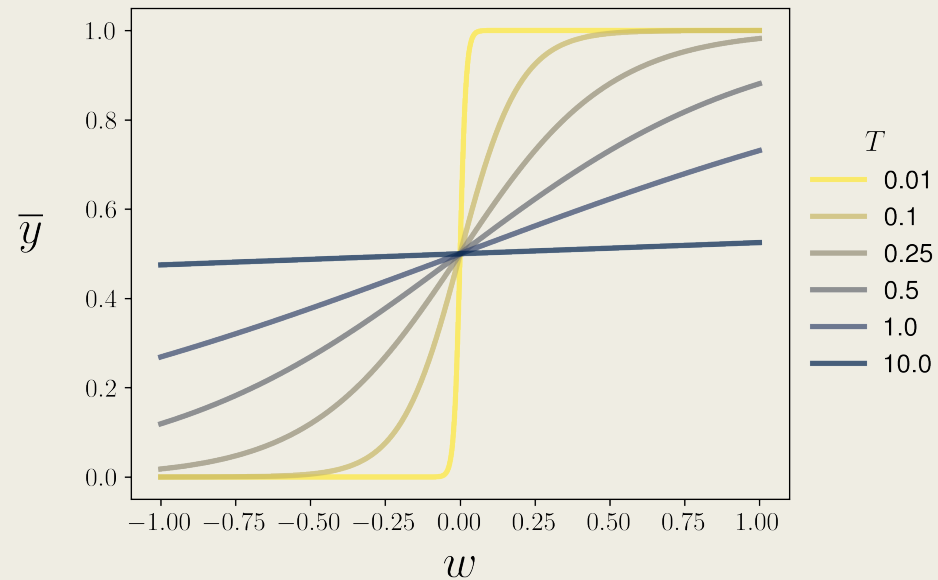
- regresja liniowa
- uogólnione modele liniowe (ang. *Generalized linear models*)
- drzewa decyzyjne (lasy)
- regresja bayesowska
- maszyny wektorów wspierających (ang. *Support Vector Machines*)
- sieci neuronowe
- ...



Model regresji jako klasyfikator

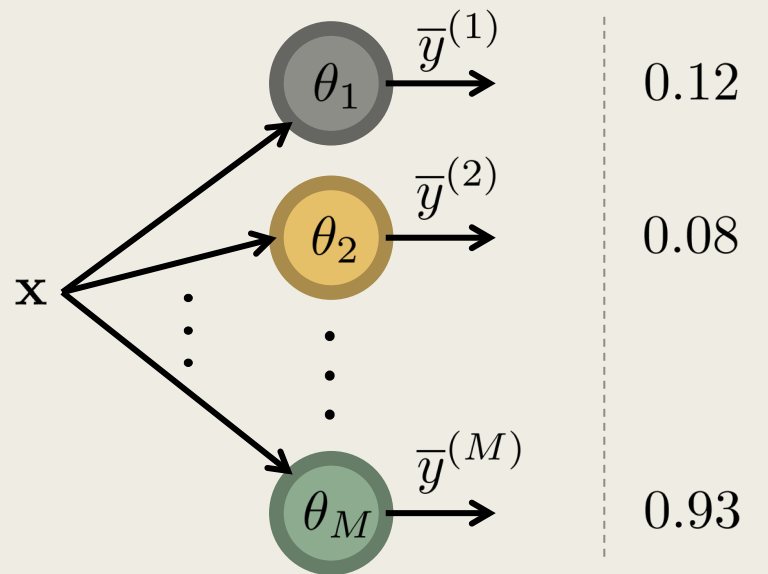


$$\bar{y} = \frac{1}{1 + e^{-w}}$$



Model regresji jako klasyfikator

Wiele klas



$$\bar{y}^{(m)} = \frac{1}{1 + e^{-\theta_m^T \mathbf{x}}} \quad (m = 1, 2, \dots, M)$$

vs

Softmax

$$\bar{y}^{(m)} = \frac{e^{\theta_m^T \mathbf{x}}}{\sum_{l=1}^M e^{\theta_l^T \mathbf{x}}} \quad (m = 1, 2, \dots, M)$$



Ujęcie probabilistyczne klasyfikacji

wektor cech

$$\mathbf{x} = \begin{bmatrix} x^{(1)} \\ \vdots \\ x^{(K)} \end{bmatrix}$$

etykiety klas

$$\{C_1, C_2, \dots, C_M\}$$

$$p(\mathbf{x}, C)$$

$$p(x^{(1)}|C)$$

$\vdots$

$$p(x^{(K)}|C)$$

$p(\mathbf{x})$ – rozkład cech

$p(C)$ – rozkład klas

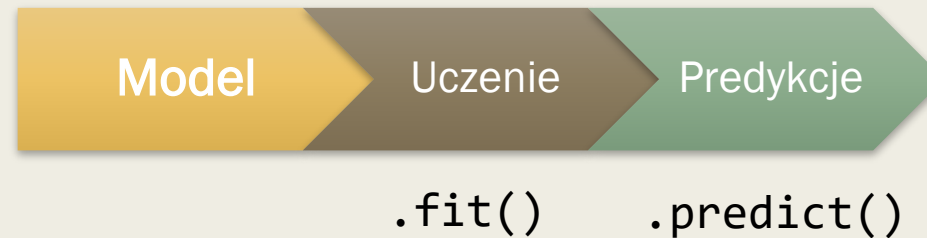
$p(\mathbf{x}|C)$ – rozkład cech w klasach

→ Modele generatywne

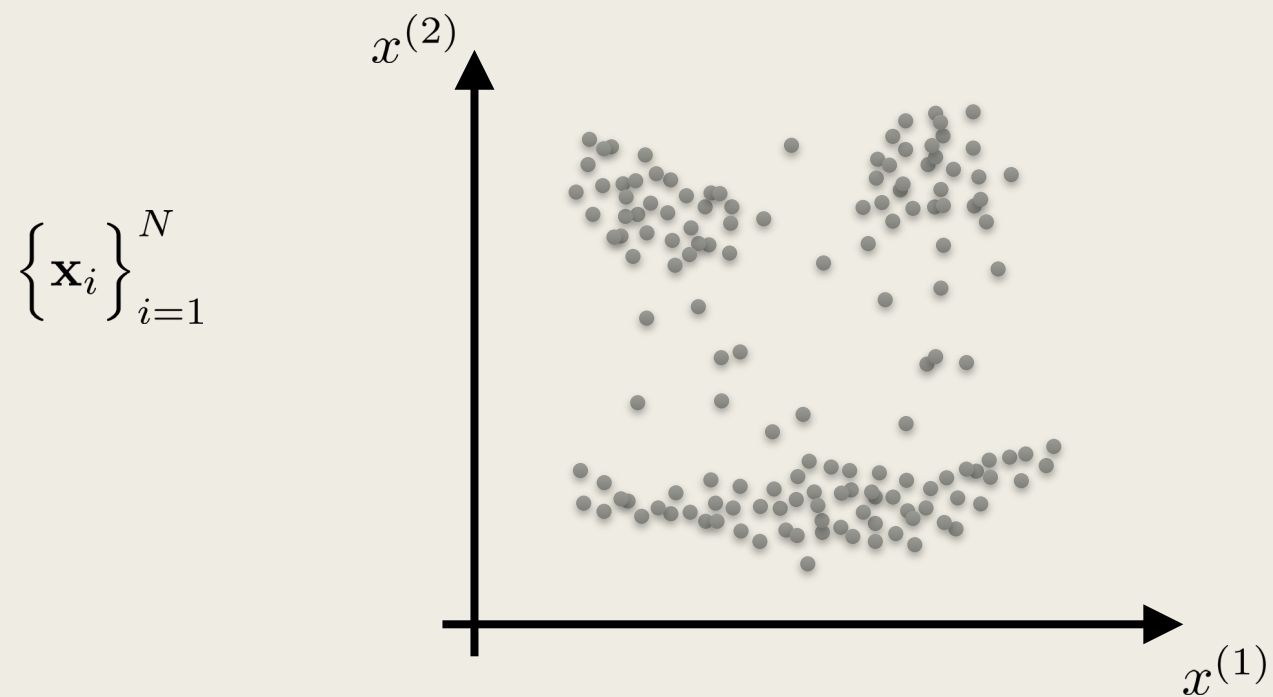
$$p(C|\mathbf{x}) = \frac{p(\mathbf{x}|C) \cdot p(C)}{p(\mathbf{x})}$$

Przykłady modeli

- optymalny klasyfikator Bayesa
- naiwny klasyfikator Bayesa (ang. *Naive Bayes Classifier*)
- ukryte modele Markowa (ang. *Hidden Markov Models*)
- ...



Grupowanie



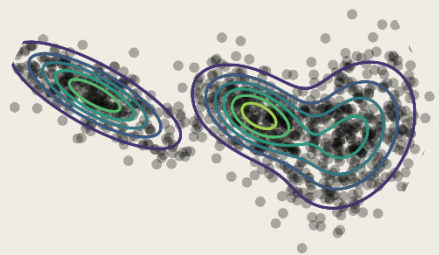
Przykłady modeli / algorytmów

- *k*-średnich (ang. *k-means*) / centroidów
- mieszanina rozkładów Gaussa
 - maksymalizacja oczekiwania (ang. *Expectation–Maximization*)
- DBSCAN
- widmowa analiza skupień
- grupowanie twarde i miękkie
 - rozmyte *c*-średnich (ang. *fuzzy c-means*)
- metody aglomeracyjne

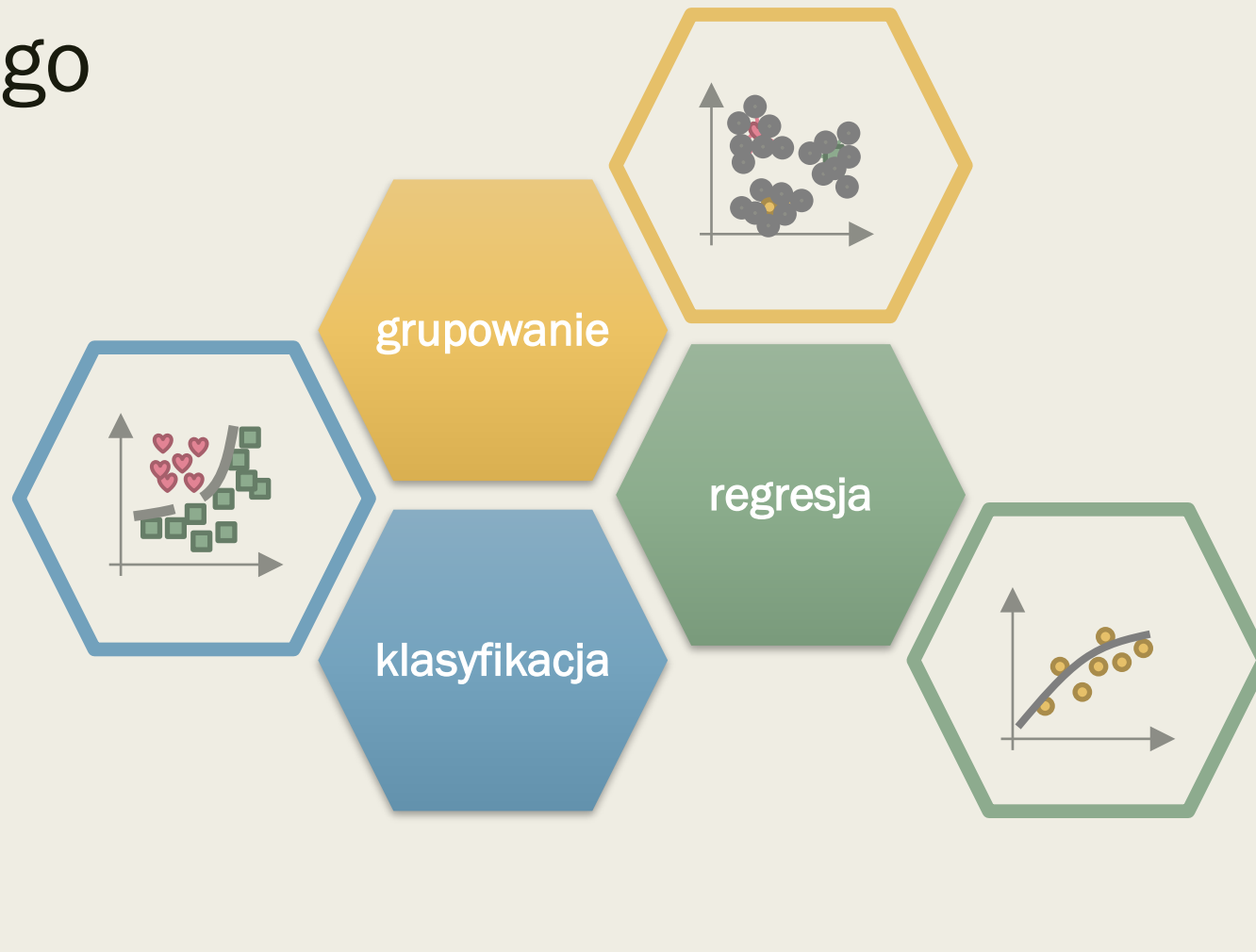
```
sklearn.mixture.GaussianMixture  
                .BayesianGaussianMixture
```

```
sklearn.cluster.Kmeans  
                .DBSCAN  
                .AgglomerativeClustering  
                .SpectralClustering  
scipy.cluster.hierarchy.linkage  
                      .dendrogram
```

Podstawowe zadania uczenia maszynowego



Estymacja
rozkładu



$$\left\{ (\mathbf{x}_i, y_i) \right\}_{i=1}^N$$

uczenie nadzorowane
(ang. *supervised learning*)

etykieta
„prawdziwa
odpowiedź”

regresja
wyjście ciągłe

klasyfikacja
wyjście
kategorialne

$$\left\{ \mathbf{x}_i \right\}_{i=1}^N$$

uczenie nienadzorowane
(ang. *unsupervised learning*)

brak etykiet

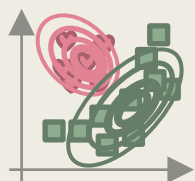
grupowanie

uczenie
półnadzorowane
(ang. *semisupervised
learning*)

uczenie ze
wzmocnieniem
(ang. *reinforcement
learning*)

Uczenie maszynowe

Estymacja rozkładu



- rozkład dedykowany
- estymator jądrowy
- mieszanina rozkładów Gaussa

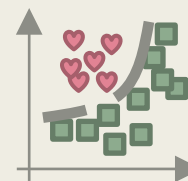
Uczenie nadzorowane

Regresja



- regresja liniowa
- uogólnione modele liniowe
- SVR
- drzewa decyzyjne
- sieci neuronowe

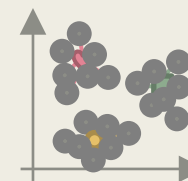
Klasyfikacja



- k -najbliższych sąsiadów
- naiwny Bayes
- regresja logistyczna
- SVM
- drzewa decyzyjne
- lasy losowe
- sieci neuronowe

Uczenie nienadzorowane

Grupowanie



- k -średnich
- rozmyte c-średnich
- mieszanina rozkładów Gaussa

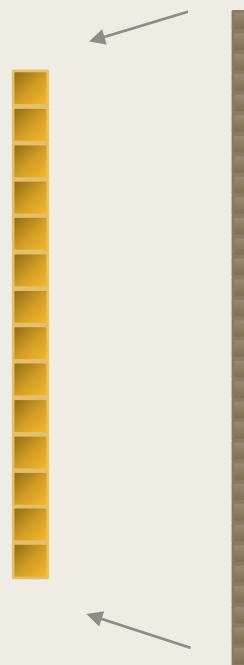
Uczenie głębokie

x



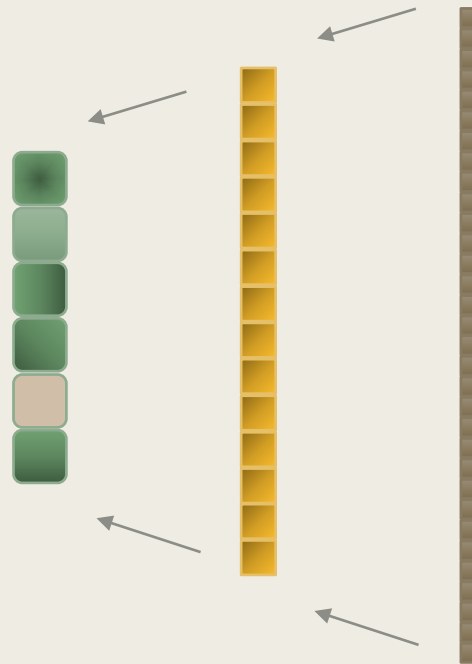
Uczenie głębokie

$$\varphi_0(\mathbf{W}_0 \cdot \mathbf{x})$$



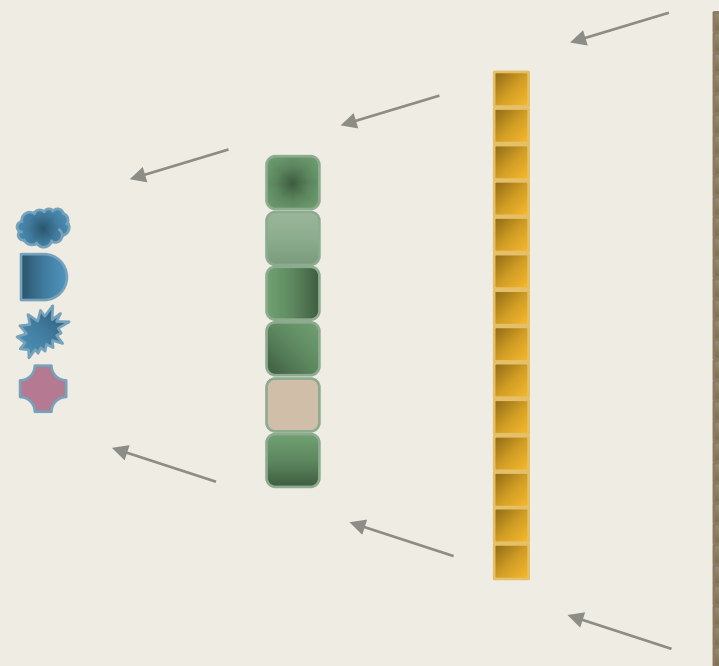
Uczenie głębokie

$$\varphi_1(\mathbf{W}_1 \cdot \varphi_0(\mathbf{W}_0 \cdot \mathbf{x}))$$



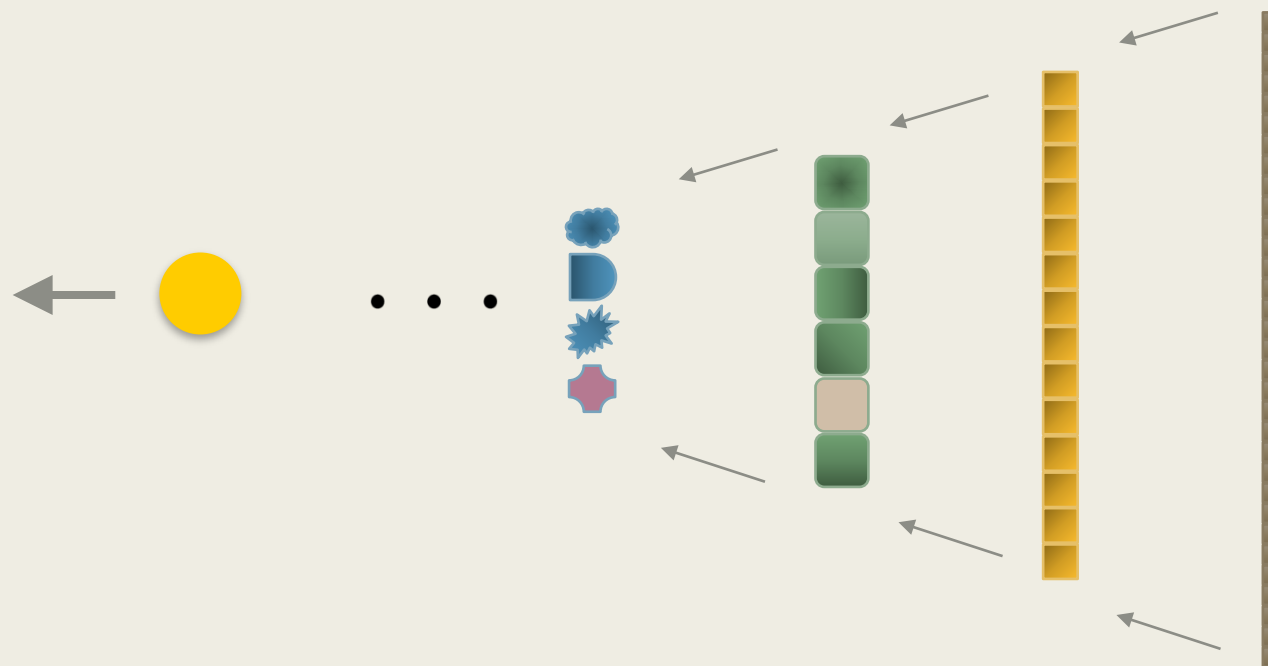
Uczenie głębokie

$$\varphi_2\left(\mathbf{W}_2 \cdot \varphi_1\left(\mathbf{W}_1 \cdot \varphi_0\left(\mathbf{W}_0 \cdot \mathbf{x}\right)\right)\right)$$



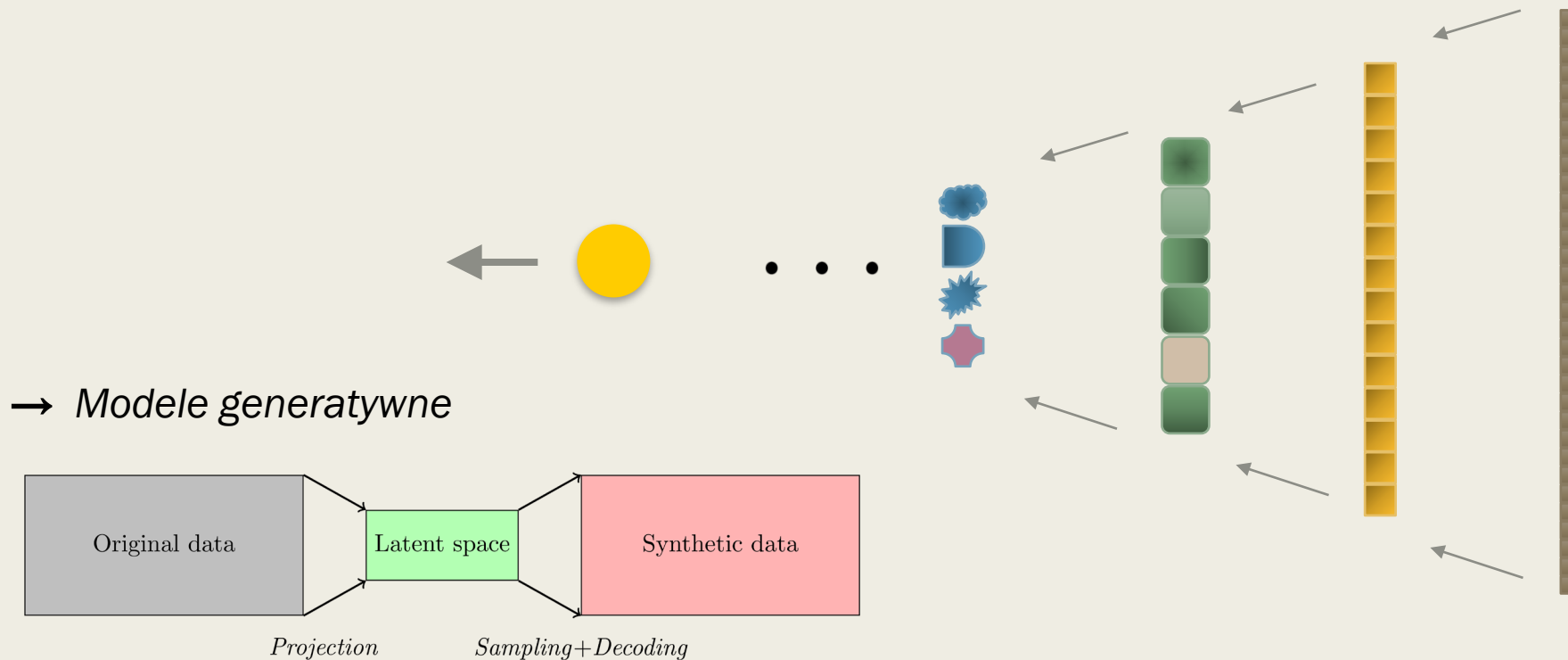
Uczenie głębokie

$$y = \varphi_K \left(\mathbf{W}_K \cdot \dots \cdot \varphi_2 \left(\mathbf{W}_2 \cdot \varphi_1 \left(\mathbf{W}_1 \cdot \varphi_0 (\mathbf{W}_0 \cdot \mathbf{x}) \right) \right) \dots \right)$$

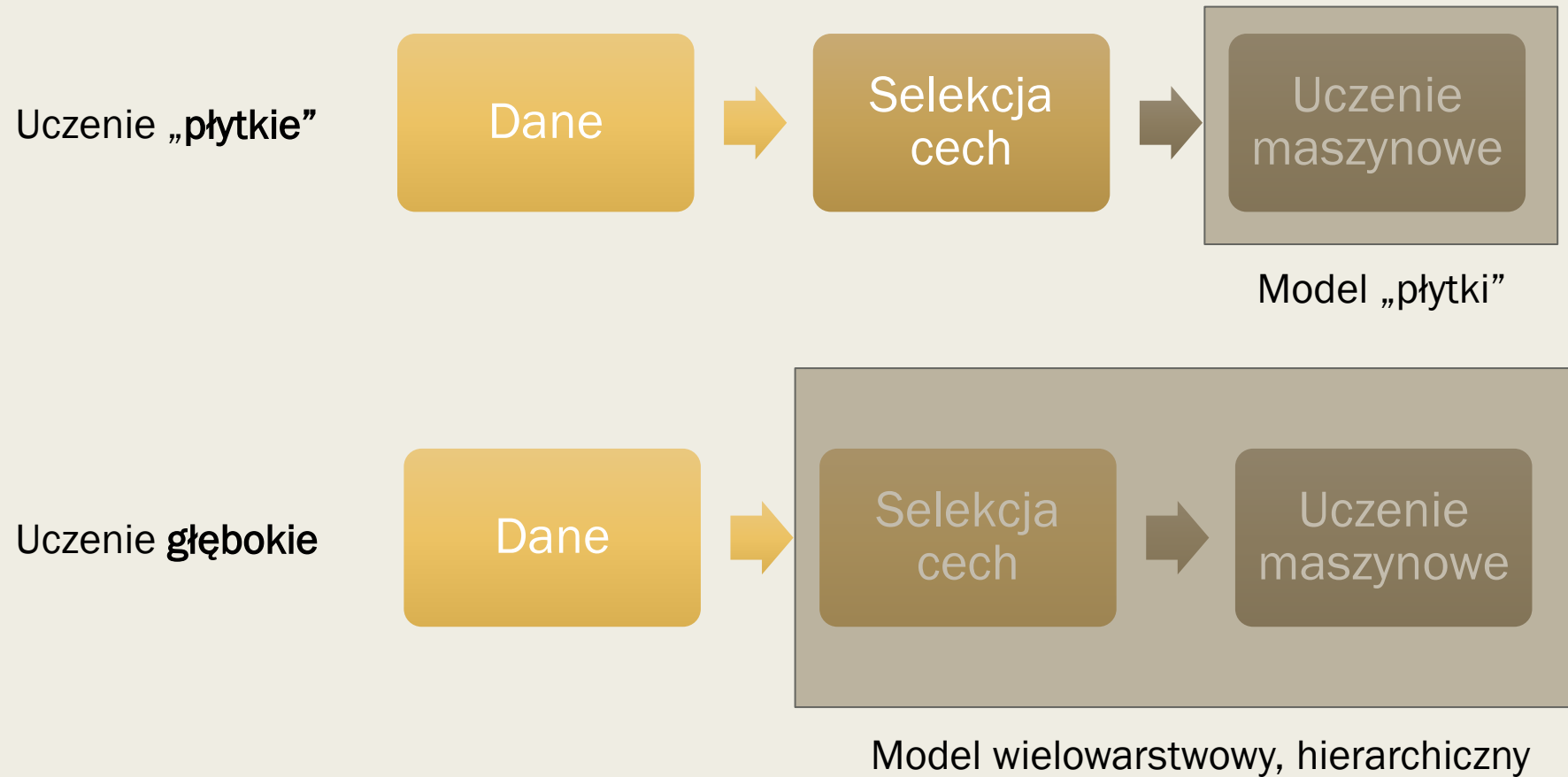


Uczenie głębokie

$$y = \varphi_K \left(\mathbf{W}_K \cdot \dots \cdot \varphi_2 \left(\mathbf{W}_2 \cdot \varphi_1 \left(\mathbf{W}_1 \cdot \varphi_0 \left(\mathbf{W}_0 \cdot \mathbf{x} \right) \right) \right) \dots \right)$$

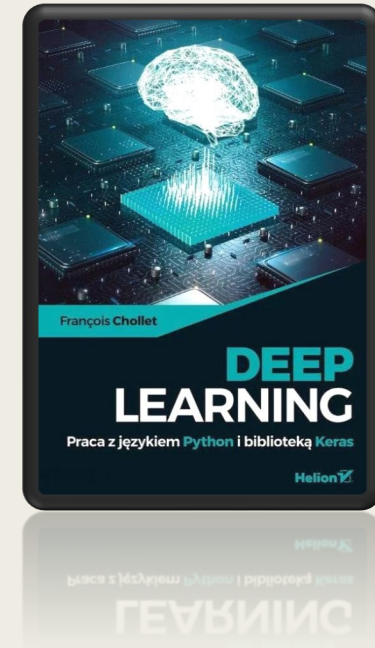


Uczenie głębokie



Uczenie głębokie

- Uczenie reprezentacji warstwowych i hierarchicznych
- Model samodzielnie uczy się cech
- Uczenie transferowe
- Przestrzeń ukryta
- Zdominowane przez sieci neuronowe
- Mistyczna otoczka tego, że technologia ta przypomina działanie mózgu



Parametryczne

Ustalona liczba parametrów



- ☐ klasyfikator liniowy
- ☐ regresja liniowa i logistyczna
- ☐ mieszanina rozkładów Gaussa
- ☐ liniowy SVM
- ☐ naiwny klasyfikator Bayesa
- ☐ ...

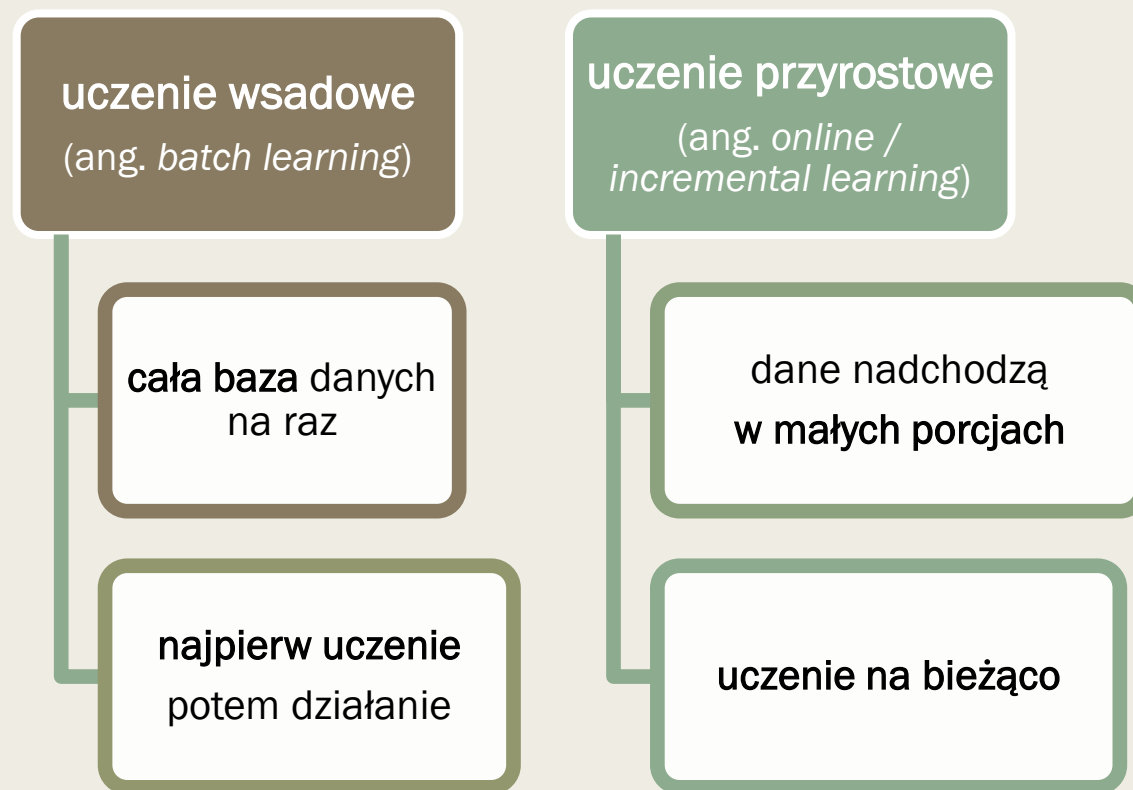
Nieparametryczne

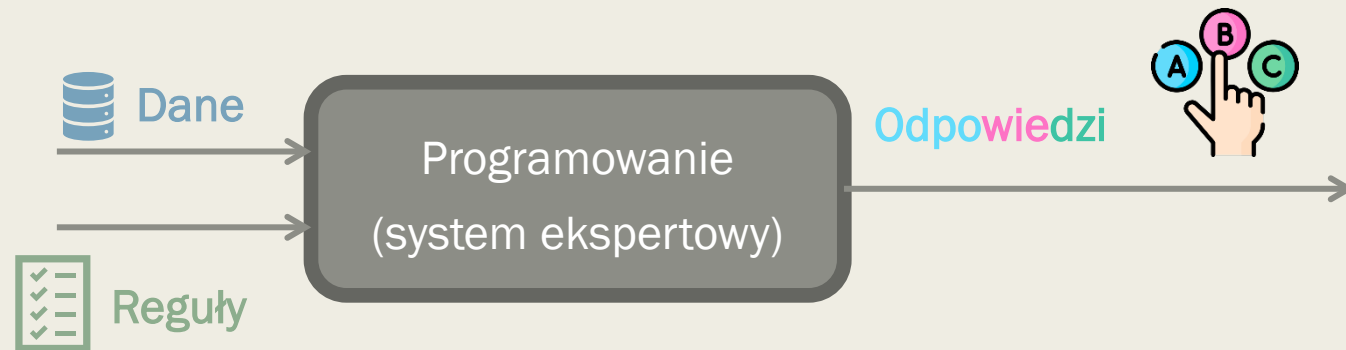
Postać modelu i liczba parametrów zależą od danych



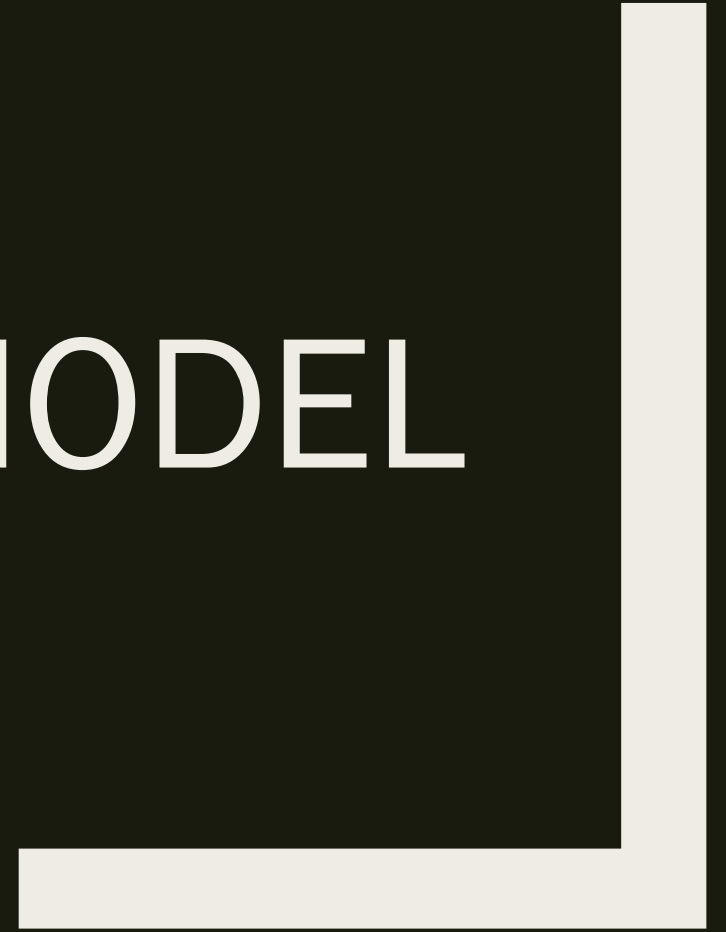
- ☐ k -NN
- ☐ drzewa decyzyjne
- ☐ estymatory jądrowe (KDE)
- ☐ sieci neuronowe
- ☐ ...

Sposoby korzystania z bazy danych

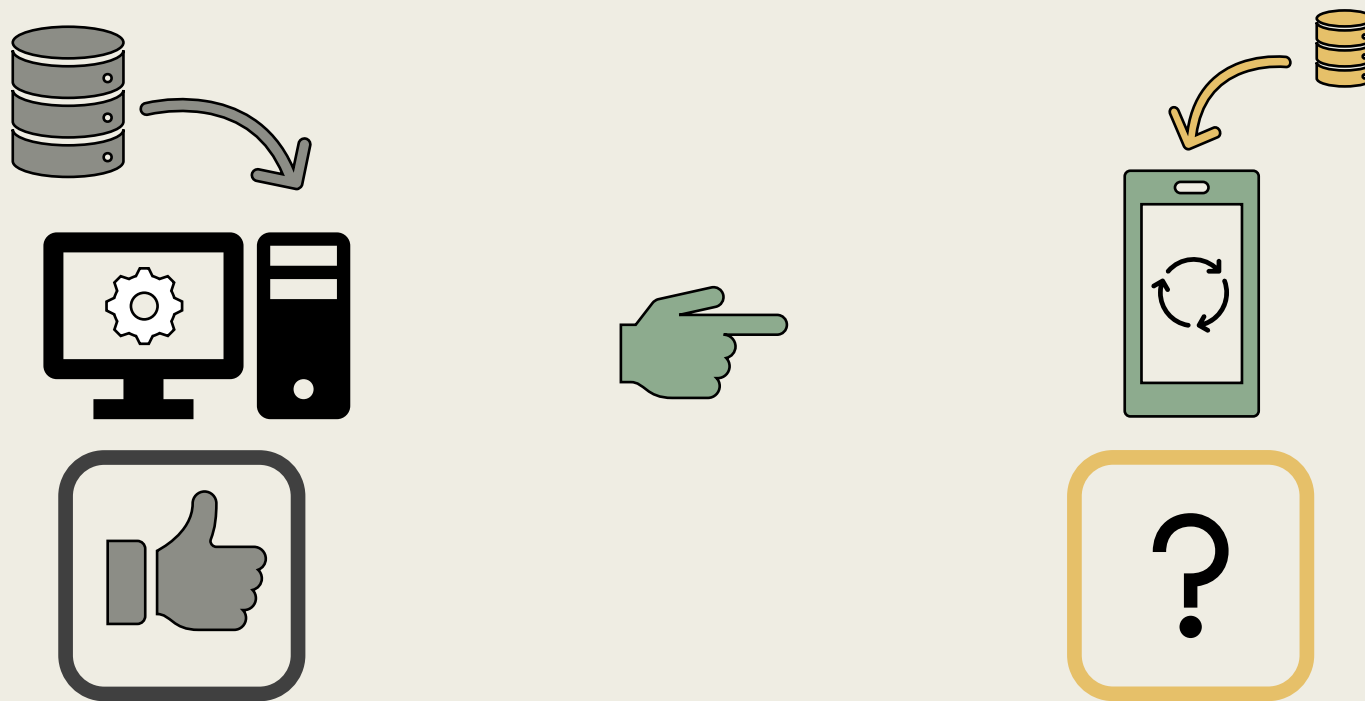




DOBRY MODEL



Jaki model jest dobry?



Jakość na danych treningowych

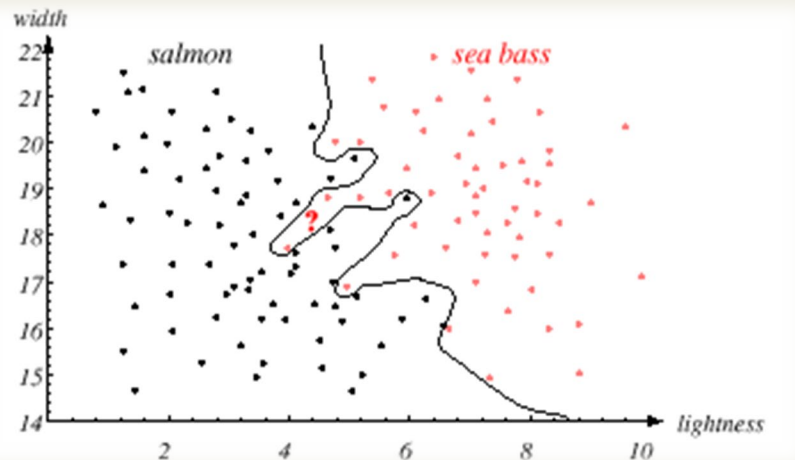


Jakość na **nowych danych**

Złożoność modelu

kompromis obciążenie – wariancja
(ang. *bias-variance dilemma*)

duża wariancja

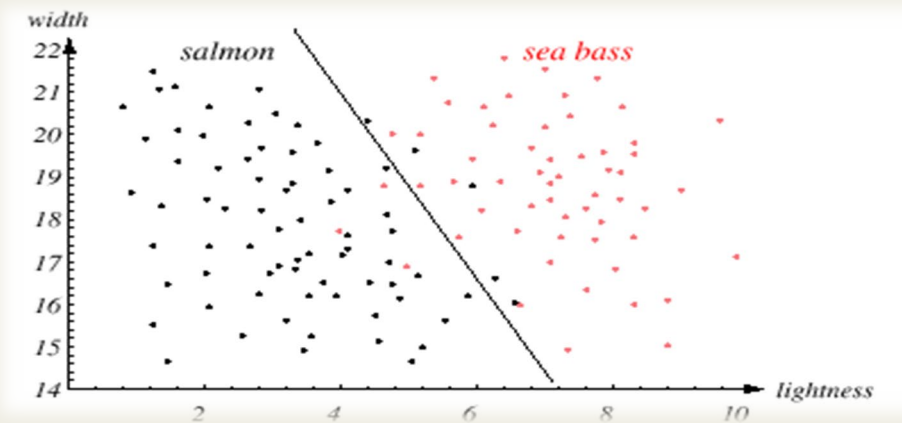


nadmiernie dopasowany do tych danych

dane

wiedza

duże obciążenie

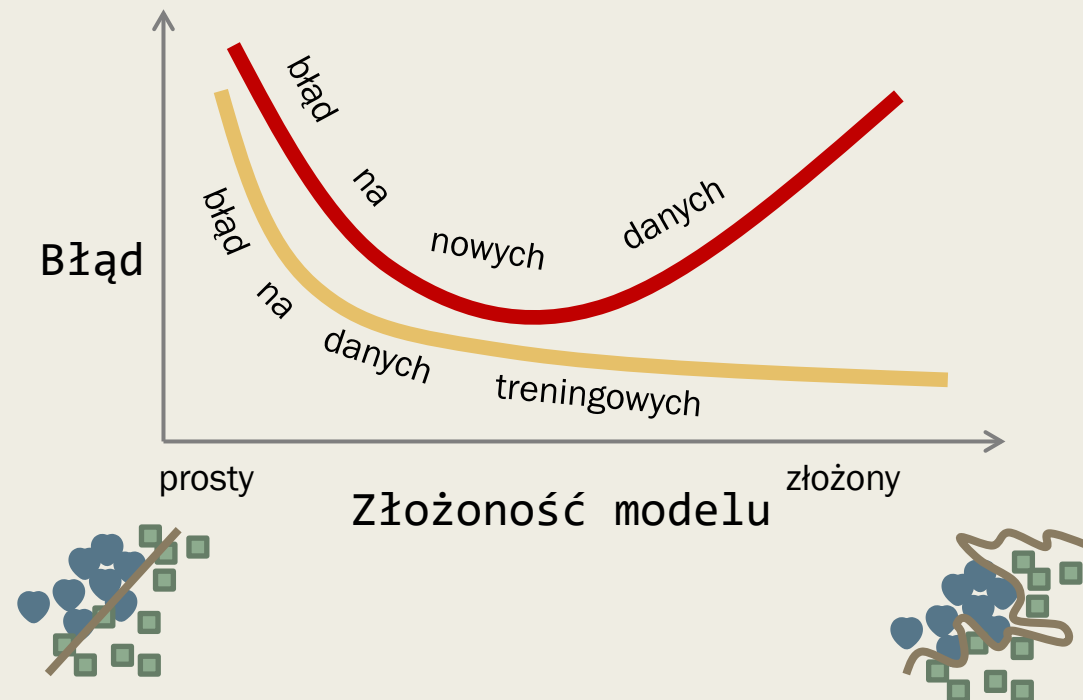


mocno trzymający się założeń

Im bardziej model elastyczny, tym lepiej dopasuje się do aktualnych danych

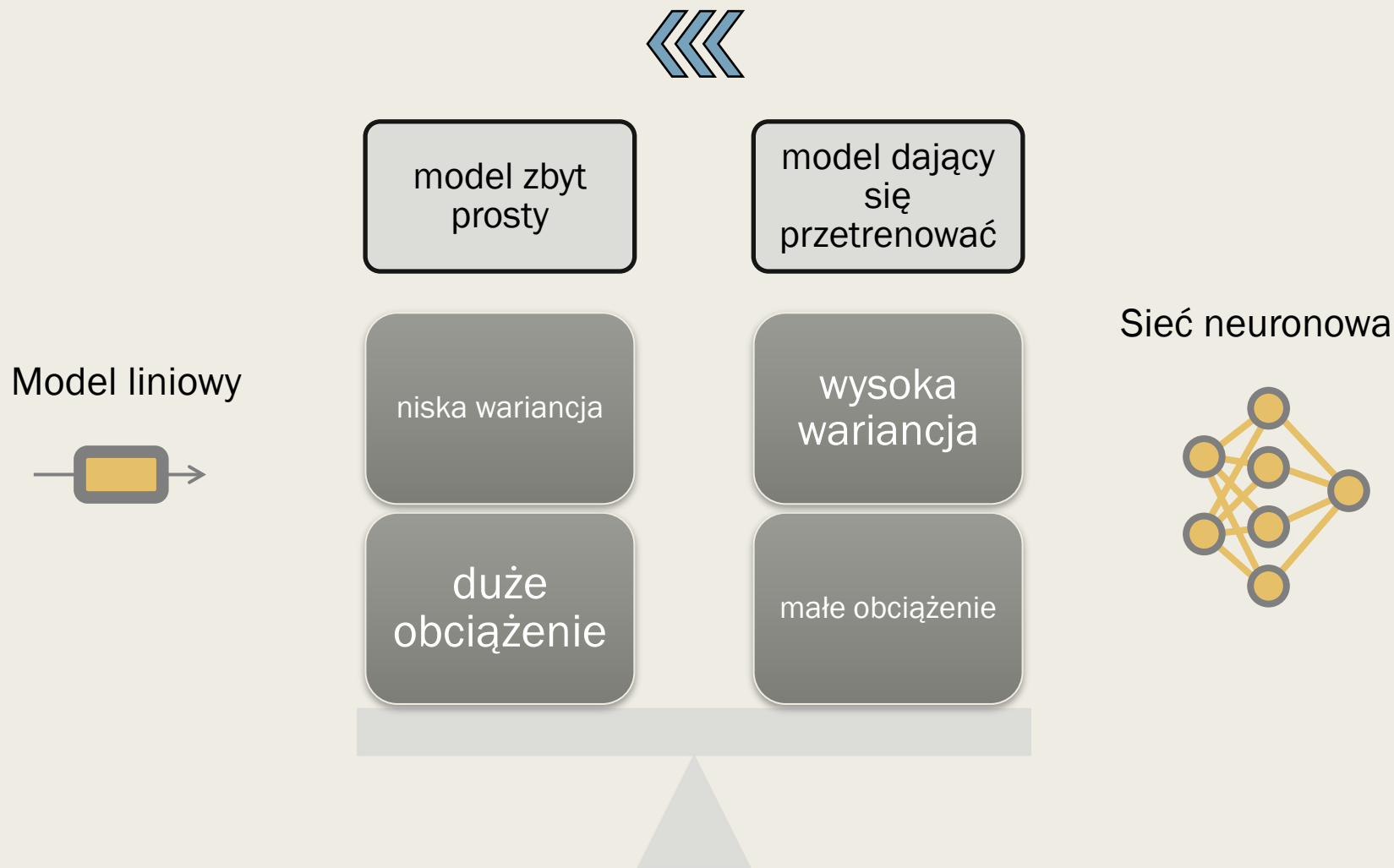
Złożoność modelu

- mały błąd na danych uczących nie gwarantuje małego błędu na nowych obserwacjach
- przeuczenie / nadmierne dopasowanie / przetrenowanie (ang. *overfitting*)
 - *model mapuje jak słownik*
- niedouczenie / niedotrenowanie (ang. *underfitting*)
- błąd uogólniania (ang. *generalization error*)



Regularyzacja

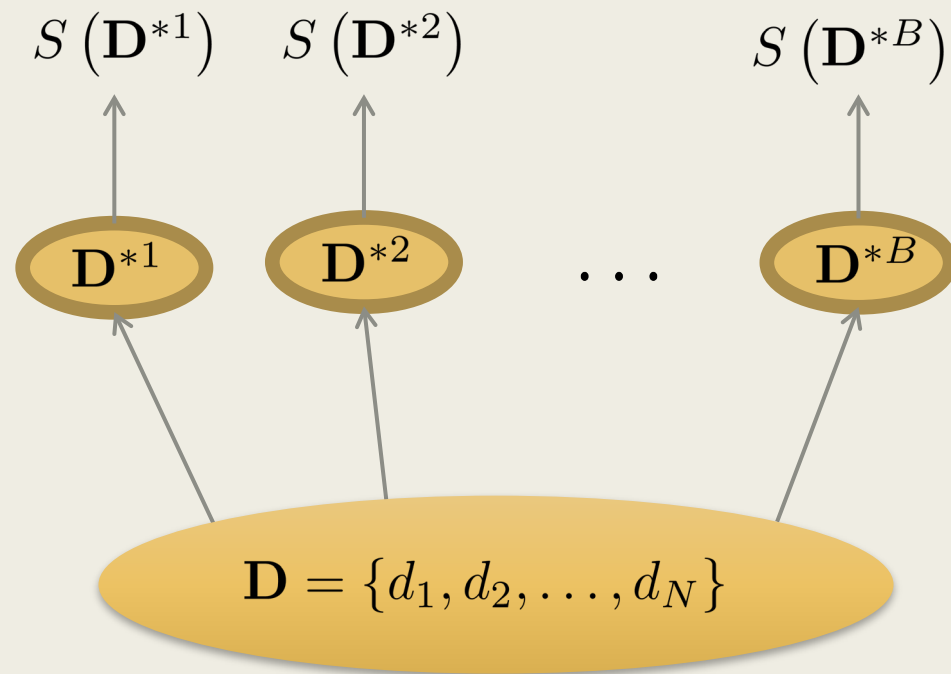
$$\text{Ocena modelu} = \text{błąd dopasowania} + \lambda \cdot \text{kara za złożoność}$$



OCENA JAKOŚCI MODELU



Metoda bootstrap



replikacje *bootstrap*owe

N – elementowe próby
*bootstrap*owe losowane
ze zwracaniem

Estymacja Monte-Carlo dokładności wyznaczenia

$$\bar{S}^* = \frac{1}{B} \sum_{b=1}^B S(D^{*b})$$

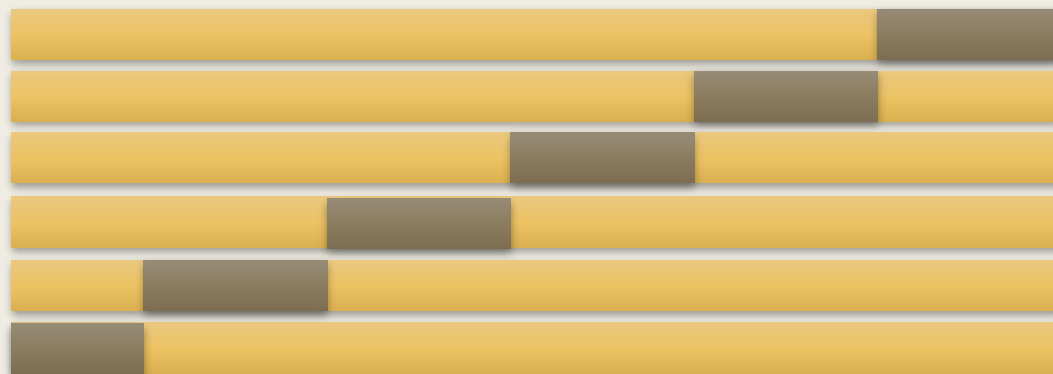
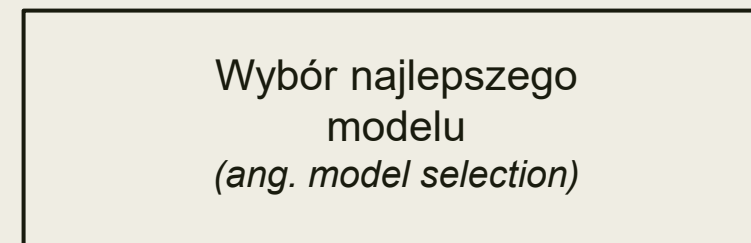
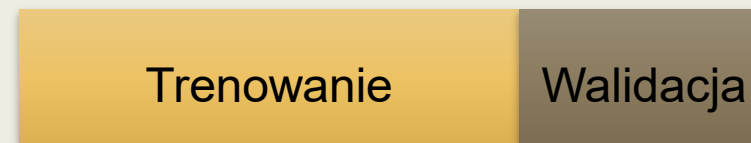
$$\widehat{\text{Var}}[S(D)] = \frac{1}{B-1} \sum_{b=1}^B \left(S(D^{*b}) - \bar{S}^* \right)^2$$



Walidacja krzyżowa

(ang. *Cross-Validation* – CV)

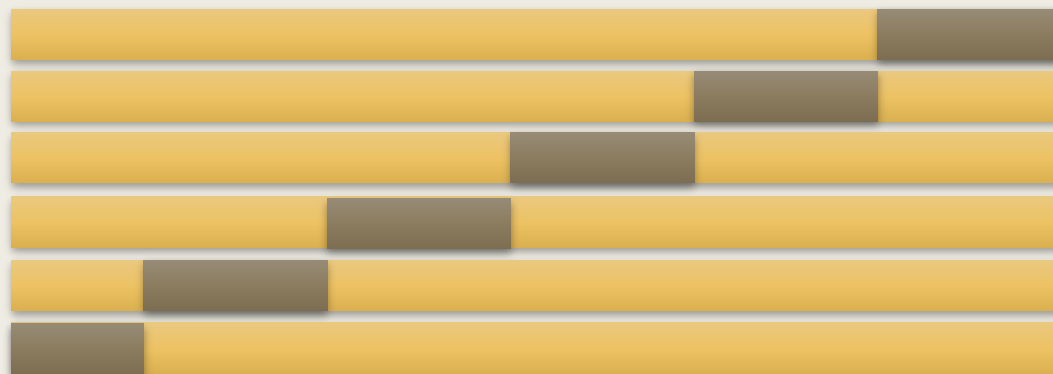
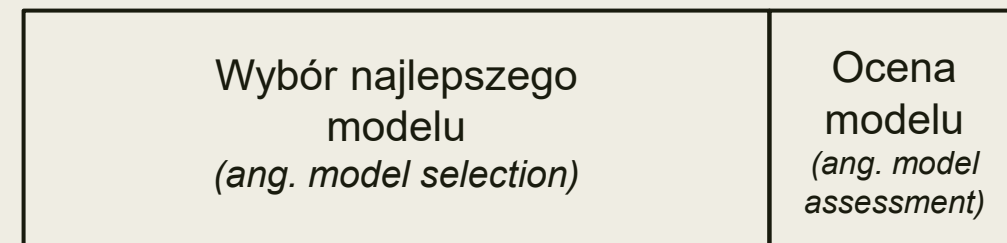
- Losowy podział
- Wielokrotny podział



Walidacja krzyżowa

(ang. *Cross-Validation* – CV)

- Losowy podział
- Wielokrotny podział
- Oszacowanie błędu na zbiorze testującym

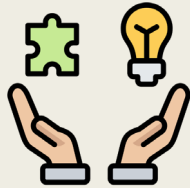


Potok wydobywania wiedzy z danych

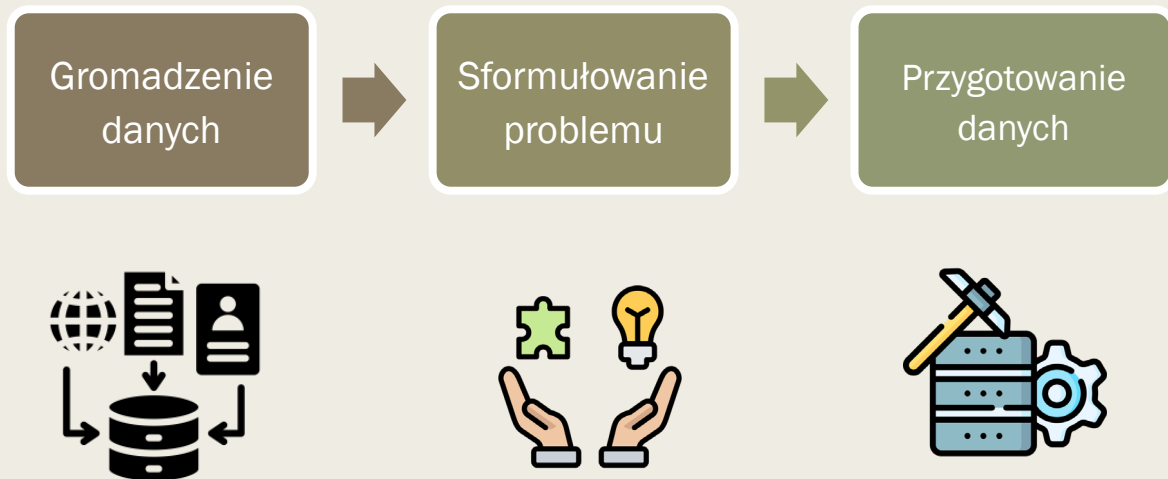
Gromadzenie
danych



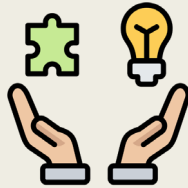
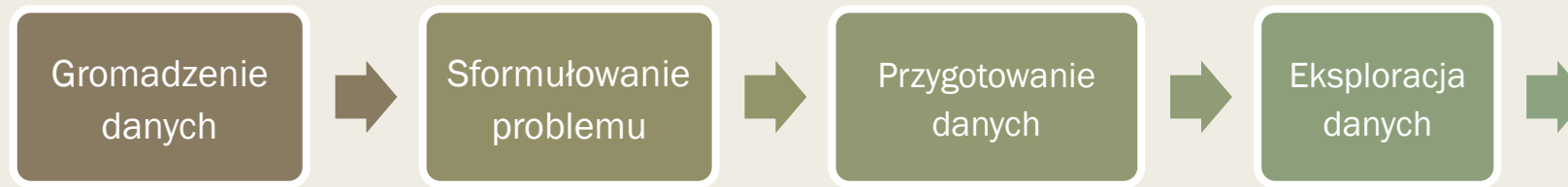
Potok wydobywania wiedzy z danych



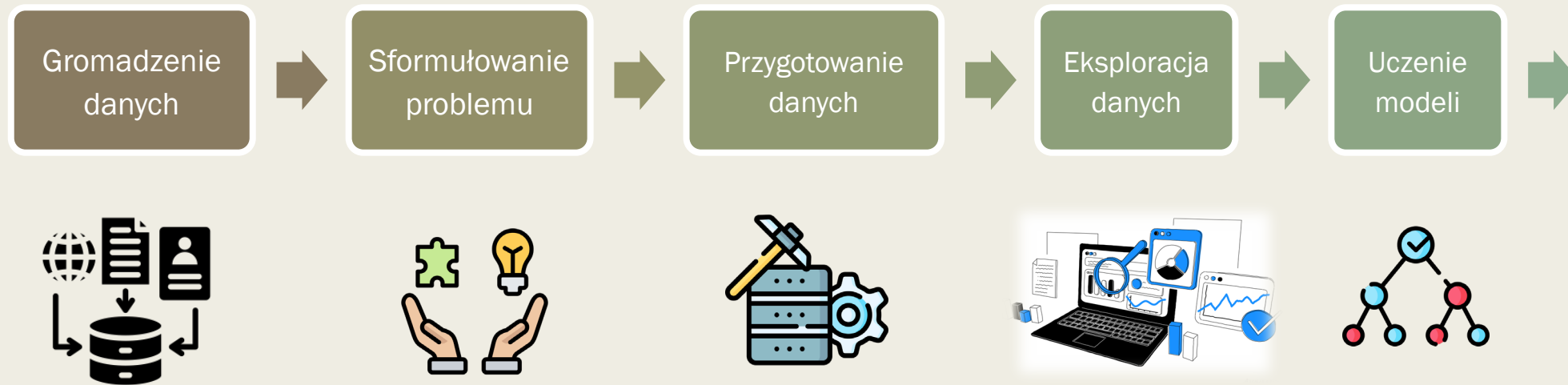
Potok wydobywania wiedzy z danych



Potok wydobywania wiedzy z danych



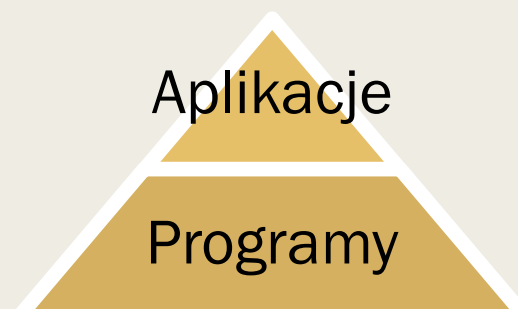
Potok wydobywania wiedzy z danych



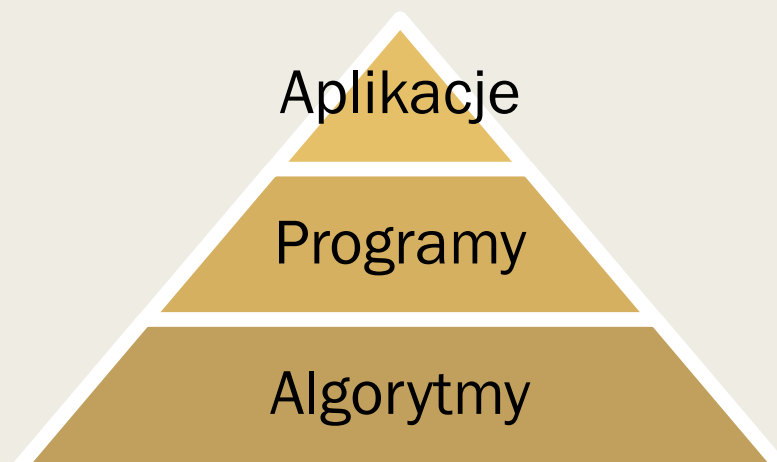
Potok wydobywania wiedzy z danych



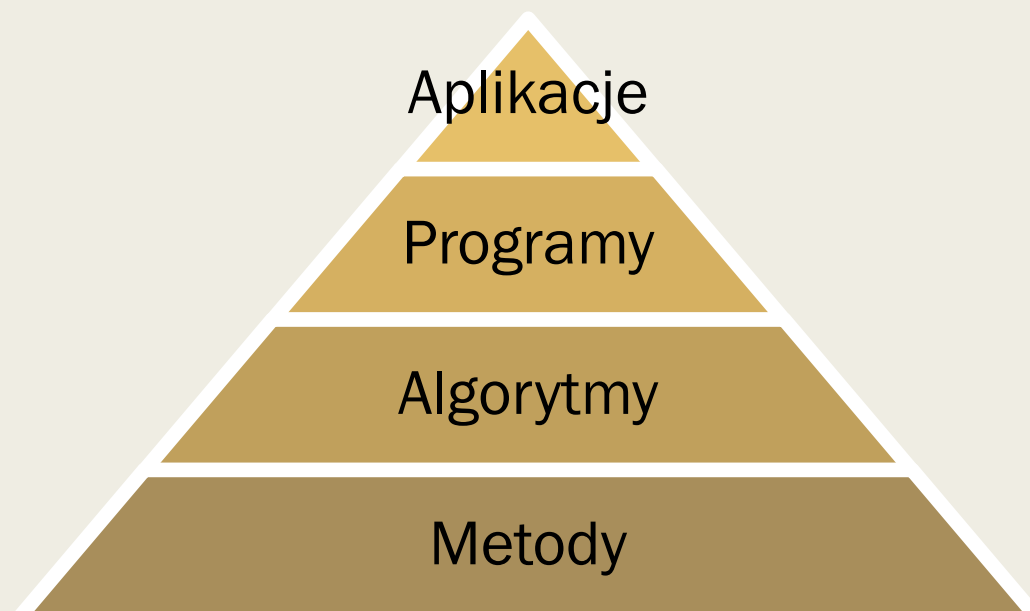
Skąd się biorą aplikacje ze sztuczną inteligencją na pokładzie?



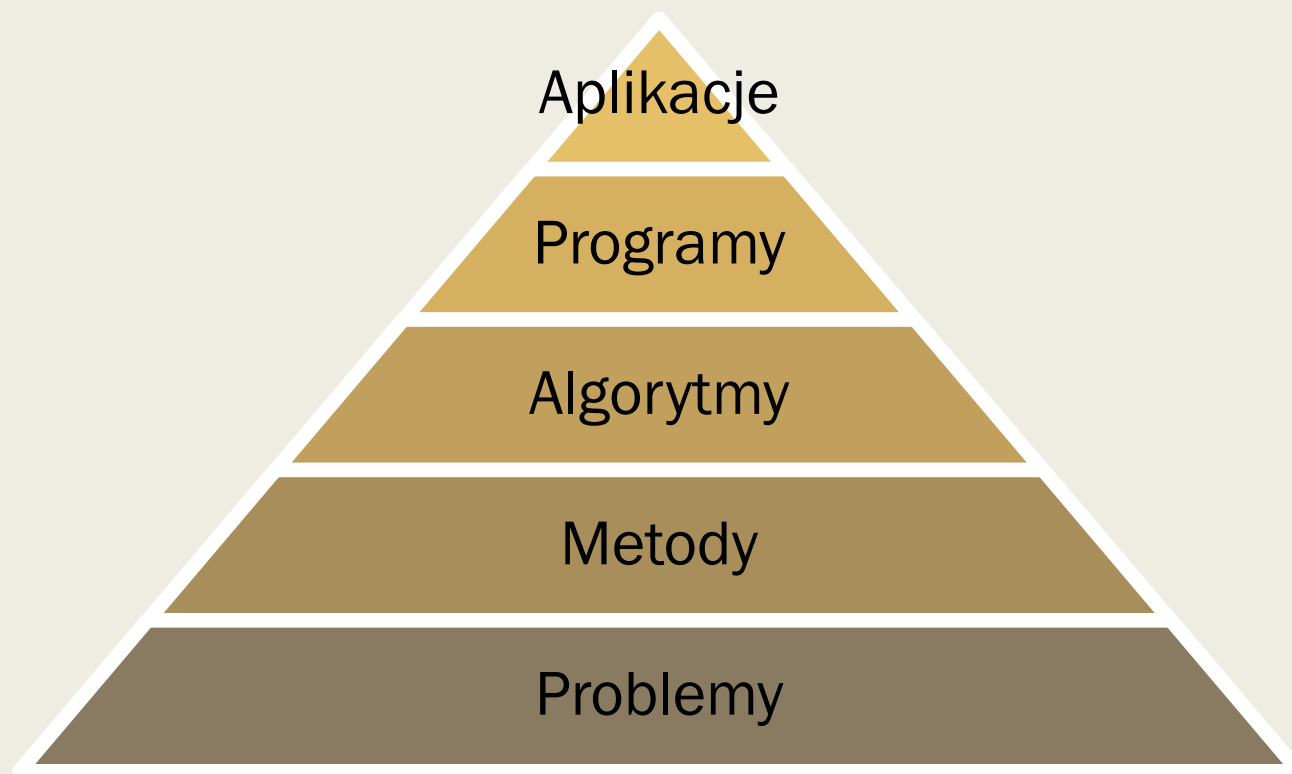
Skąd się biorą aplikacje ze sztuczną inteligencją na pokładzie?



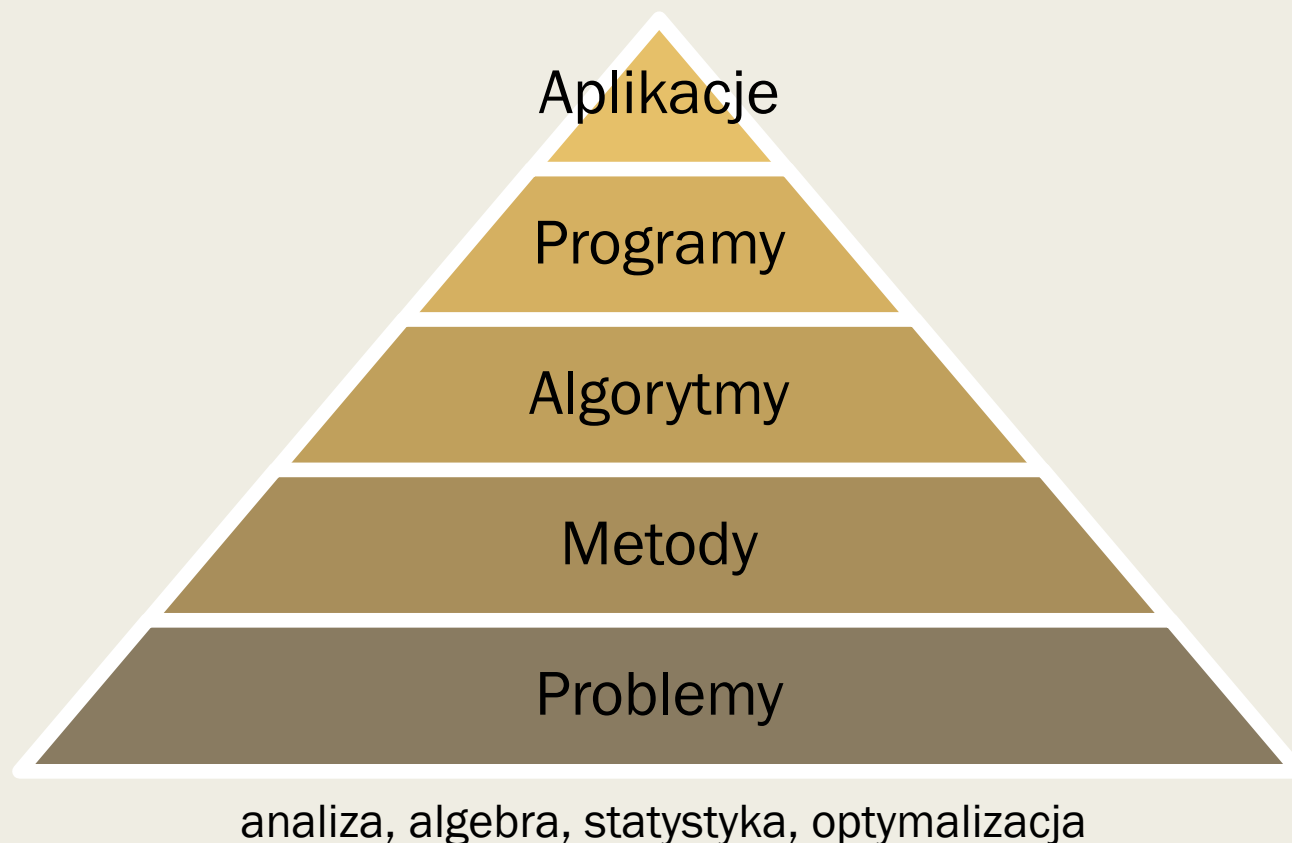
Skąd się biorą aplikacje ze sztuczną inteligencją na pokładzie?



Skąd się biorą aplikacje ze sztuczną inteligencją na pokładzie?



Skąd się biorą aplikacje ze sztuczną inteligencją na pokładzie?



koniec