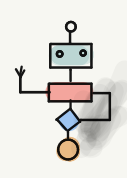


ALGORITHMEN



Sortieralgorithmen

- Bubblesort $O(n^2), O(n^2), O(n)$
Tauscht benachbarte Elemente
- Insertionsort $O(n^2), O(n^2), O(n)$
Verschiebt/Fügt Element in richtigem Platz ein
- Mergesort $O(n \log n), O(n \log n), O(n \log n)$
Teilt rekursiv in der Mitte und fügt aufsteigend zusammen
- Quicksort $O(n^2), O(n \log n), O(n \log n)$
Teilt um Pivot Element und schreibt alle kleinere davor

Optimierungsalgorithmen

- Anlagerungsverfahren (Bin-Packing):
First Fit: Objekt in ersten passenden Platz
Best Fit: Objekt in engsten passenden Platz
Next Fit: Objekt in zuletzt verwendeten Platz
- Random Interchange:
Tauscht zufällig Elemente und prüft auf verbesserung
- Kernighan-Lin:
Tauscht Knotenpaare und prüft auf verbesserung
- Greedy:
Nimmt größten wert unter Zielwert und zieht diesen davon ab.
wiederholt bis nichts übrig bleibt.
- Evaluationäre Algorithmen:
Biologisch inspiriert (Selektion, Rekombination, Mutation)

Suchalgorithmen

- Lineare Suche
Geht jedes Element durch, bis Lösung
- Binäre Suche
Sucht von der Mitte in die richtige Richtung
- Interpolationsuche
Interpoliert ungefähren Index mit
erstem und letztem Wert vorselektierter liste
- Breitensuche (Graphen)
(kürzester Pfad) Besucht immer die nächsten Knoten
- Tiefensuche (Graphen)
Folgt einem Pfad bis zum Ziel
- Dijkstra - Algorithmus
Wählt Nachbar mit geringster Distanz zum Start
und aktualisiert Distanzen aller Nachbarn. $O(n^2)$
- Kritischer Pfad
Längster Pfad im Ablauf

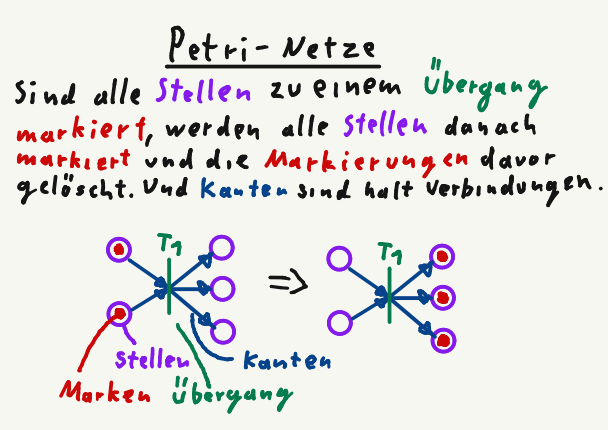
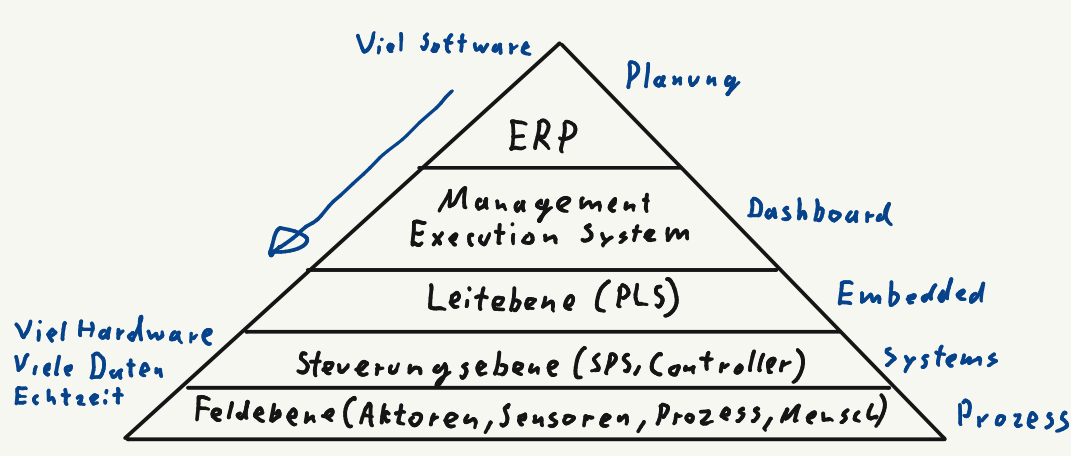
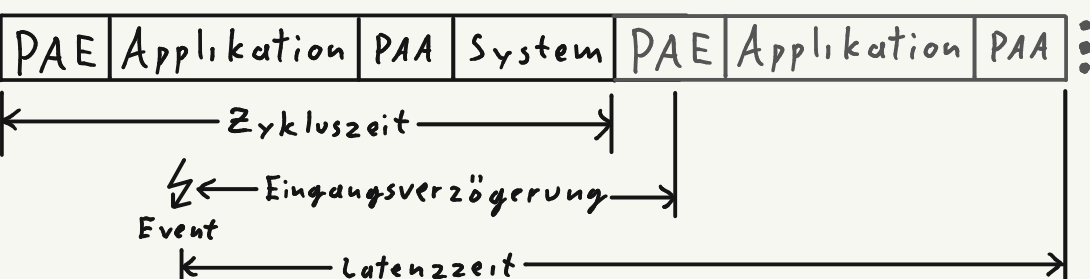
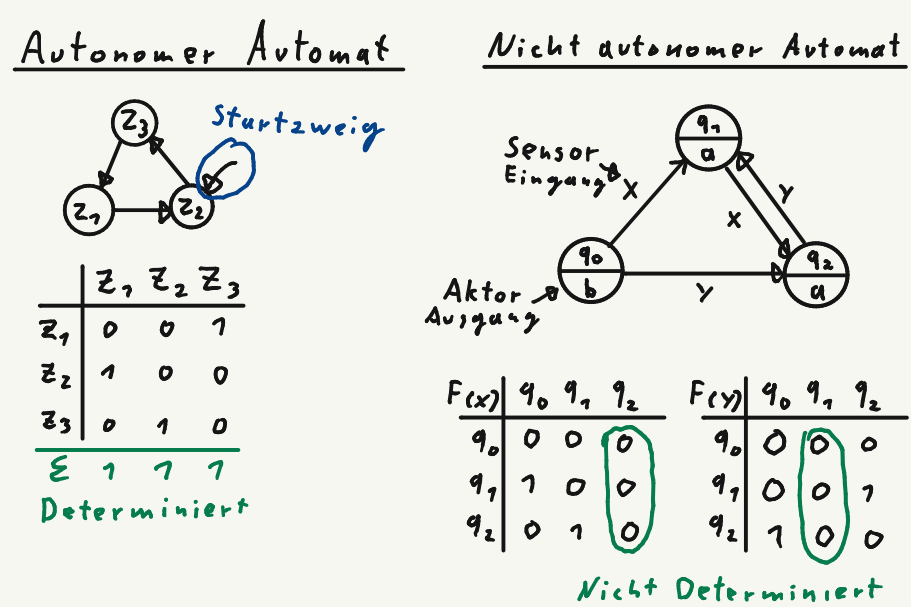
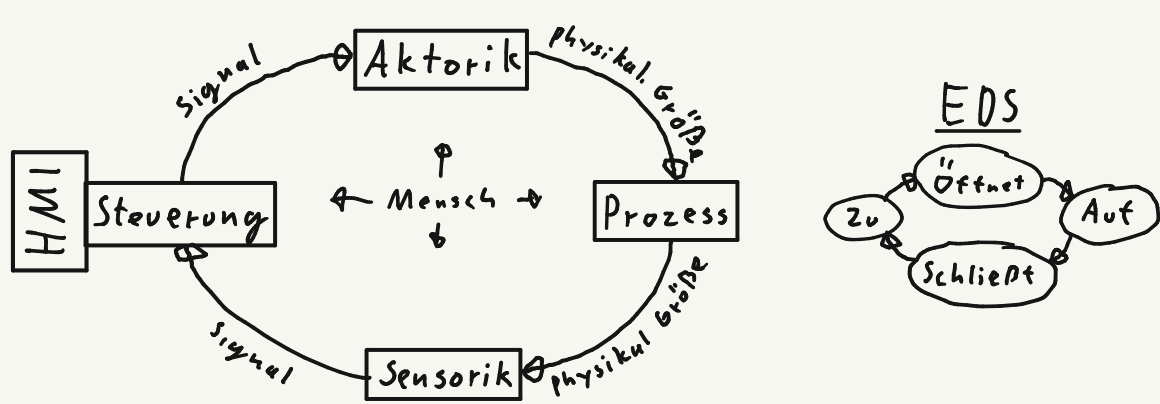
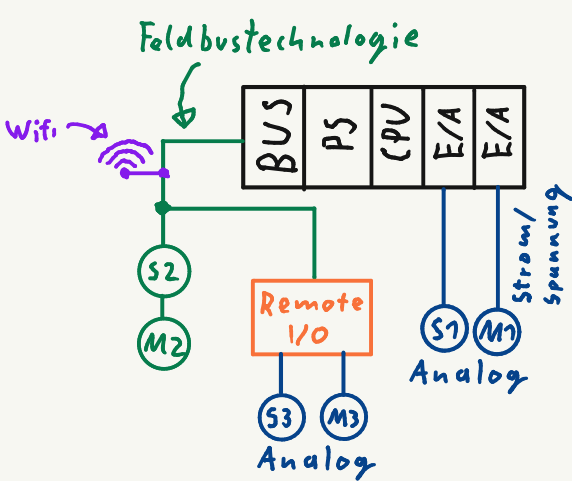
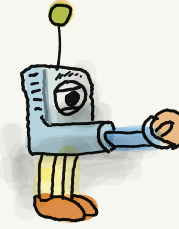
Laufzeit

- $O(n)$ Notation
- 7:50 Regel:
7% des Codes verursacht 50% der Laufzeit
- 90:10 Regel:
90% der Laufzeit ist durch 10% des Codes verursacht

Eigenschaften

- Stabil: Hält Reihenfolge von Gleichen
- Abbrechbar: kann frühzeitig beendet werden
- In-Place: Braucht kein extra Speicher
- Iterativ: Nutzt Schleife für Ergebnisberechnung
- Rekursiv: Ruft sich selbst mit Teilproblemen auf
- Heuristisch: Findet Näherungslösung
- Inkrementiell: Aktualisiert Ergebnis schrittweise
- Determiniert: Bei gleicher Eingabe immer gleiches Ergebnis
- Deterministisch: Schritte eindeutig (ohne Zufallsfaktor)
- Probabilistisch: Nutzt Zufallselemente
(Las-Vegas-Algorithmus: Ergebnis korrekt, Laufzeit zufällig)
(Monte-Carlo-Algorithmus: Ergebnis zufällig, Laufzeit vorhersehbar)

AUTOMATISIERUNGSTECHNIK



BIG DATA UND ML



Die 5 V's: Volumen, Variety, Velocity, Veracity, Value

- Menge der Daten
- Vielfalt der Daten
- Geschwindigkeit der Datengenerierung
- Zuverlässigkeit / Sinnhaftigkeit
- Wirtschaftlicher Wert

Branchen:

- Personalservice
- Produktionsentwicklung (Marktlücken & Trends)
- Distribution & Logistik
- Gesundheitswesen
- Produktion (Wartung)
- Marketing & Vertrieb (Werbung)
- Versicherung (Betriebskenntnis)
- Finanzen (Vorkursen)
- Wirtschaft & Forschung

Skalentypen:

- Nominal: Kategorien
- Ordinal: Reihenfolge mit Rangordnung
- Intervall: Reihenfolge + Gleichabständigkeit
- Ratioskala: Intervall + Echte Null
- Numerisch (Diskret): Ganze Zahlen
- Kardinal: Quantitative Vergleiche

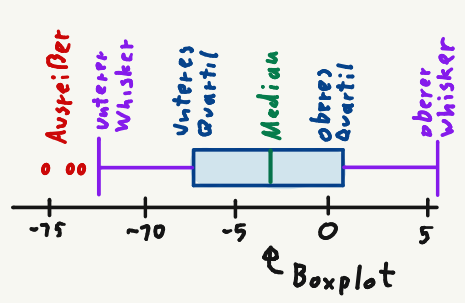
Arten der Datenanalyse

- Descriptive Analytics: Beschreibt Vergangenheit
- Diagnostic Analytics: Diagnostiziert Ursachen
- Predictive Analytics: Prognostiziert Zukunft
- Prescriptive Analytics: Empfiehlt Handlung

Statistikwerte

- Mean: $\bar{x} = \frac{1}{n} \sum x_i$
- Median: $\begin{cases} \text{even: } \bar{x} = \frac{x_{(n/2)+1} + x_{(n/2)}}{2} \\ \text{uneven: } \bar{x} = x_{(n+1)/2} \end{cases}$
- Modus: $\hat{x}$ = Häufigster Wert von x
- Std-Abweichung: $std = \sqrt{\frac{2}{n-1} \sum (x_i - \bar{x})^2}$
- Varianz: std^2

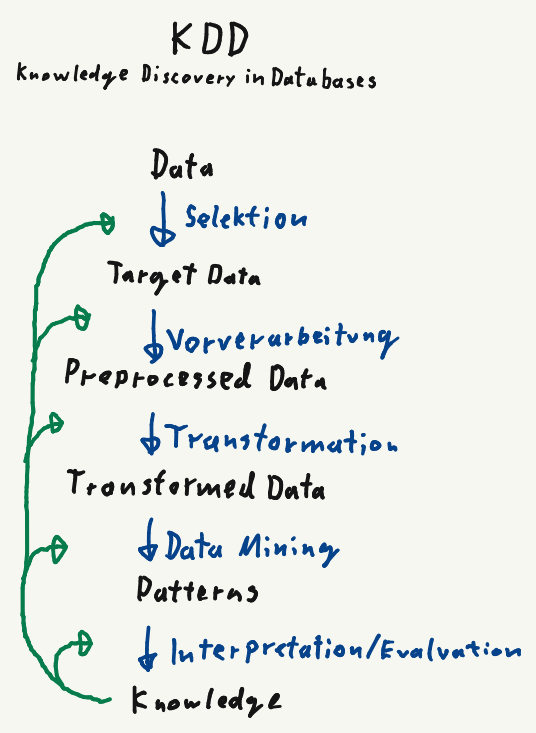
- 25%-Percentil = 1. Quartil (= Q1)
- 50%-Percentil = 2. Quartil (= Median)
- 75%-Percentil = 3. Quartil (= Q3)
- Unterer Whisker = $Q1 - 1,5 \cdot (Q3 - Q1)$
- Oberer Whisker = $Q3 + 1,5 \cdot (Q3 - Q1)$



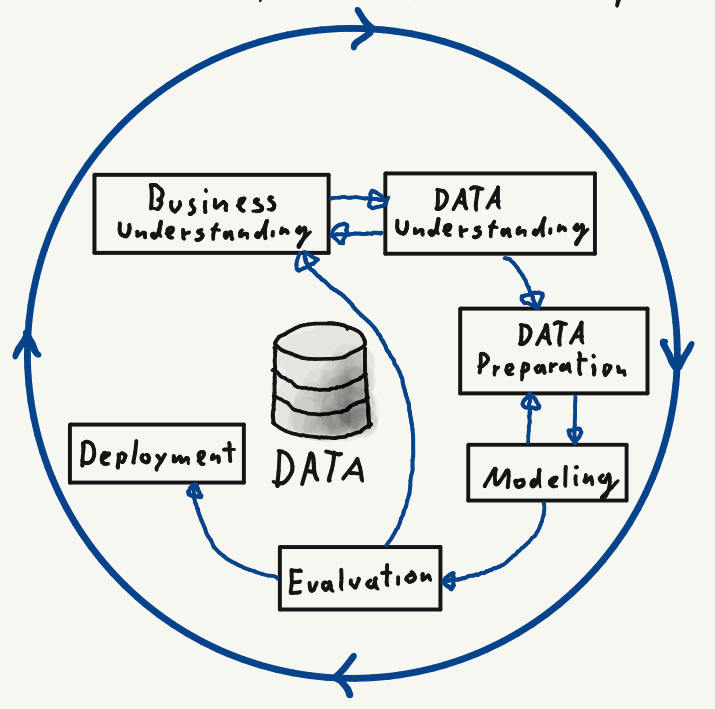
Data Preparation

- Bereinigung fehlender Werte: Interpolation
- Ausreißererkennung: $\begin{cases} \text{Obere Grenze} = \text{mean} + k \cdot \text{std} \\ \text{Untere Grenze} = \text{mean} - k \cdot \text{std} \end{cases}$
- Lineare Interpolation: $y = y_1 + ((x - x_1) \cdot (y_2 - y_1) / (x_2 - x_1))$
- Normierung: $\begin{cases} \text{Z-Score: } z = \frac{(x - \text{mean})}{\text{std}} \\ \text{Min-Max: } x_{\text{neu}} = \frac{(x - x_{\text{min}})}{(x_{\text{max}} - x_{\text{min}})} \end{cases}$

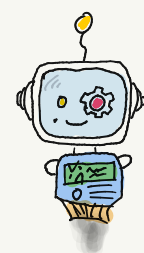
Prozesse



CRISP-DM

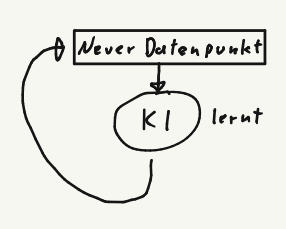


MASCHINELLES LERNEN



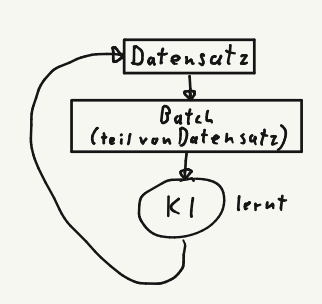
Incremental (Online) Learning

Kontinuierliches, echtzeit-lernen, flexibel, erfordert stetige Daten und schnelle Anpassung.

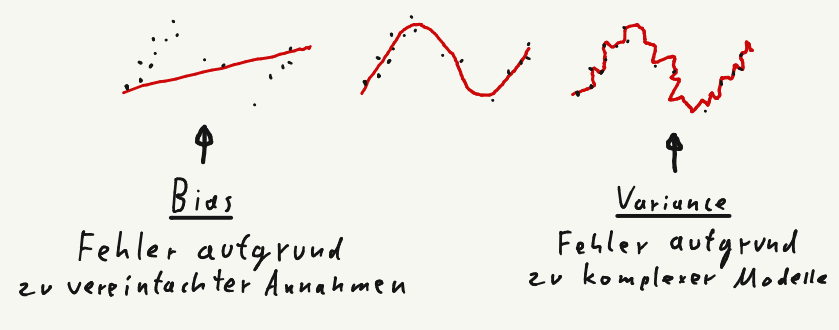


Batch (Offline) Learning

Lernen mit testeten Daten / tiefergehend, erfordert umfassende Datenverarbeitung, kann nicht sofort auf neue Daten reagieren



Underfitting GOOD Overfitting



Regularisierung: Vermeidung von Overfitting durch hinzufügen eines Komplexitätsstrahs zur Verlustfunktion

Validierungs- und Evaluierungsmethoden

CONFUSION MATRIX

$$TPR = \frac{\sum \text{True Positives}}{\sum \text{Condition Positives}}$$

$$FPR = \frac{\sum \text{False Positives}}{\sum \text{Condition Negatives}}$$

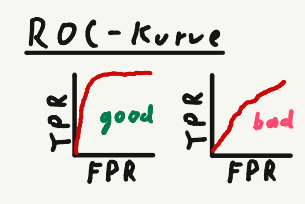
$$FNR = \frac{\sum \text{False Negatives}}{\sum \text{Condition Positives}}$$

$$TNR = \frac{\sum \text{True Negatives}}{\sum \text{Condition Negatives}}$$

Hold-out cross validation
Teil von Dataset wird nicht zum training sondern nur zum testen verwendet

K-fold cross validation
Teilt in k-teile und verwendet jedes mal anderen Teil zum Testen

Leave-one-out cross validation



Unsupervised Learning

Clustering

- Hierarchisch: Dendrogramm
- Partitionierend: k-Means
- Dichtebasiert: DBSCAN
- Gitterbasiert: Gitter-IOL

Dimension Reduction

Bsp. Temp + Luftfeuchtigkeit → Wohlfühlindex

Vorteile:

- Effizienz
- weniger Overfitting
- Bessere Visualisierung
- Rauschreduktion

DBSCAN

Parameter: Max Distanz, min. Samples per Cluster

Pro: Komplexe Clusterformen, Rauscherkennung, Dynamische Anzahl Cluster

Con: Sensible Parameter

K-Means

Parameter: Anzahl Cluster

Pro: Schnell, Datenunabhängig

Con: Anzahl Cluster muss gewählt werden

Dendrogramm

Parameter: (Anzahl Cluster/Abstand)

Pro: Cluster Hierarchie, Keine Parameterwahl

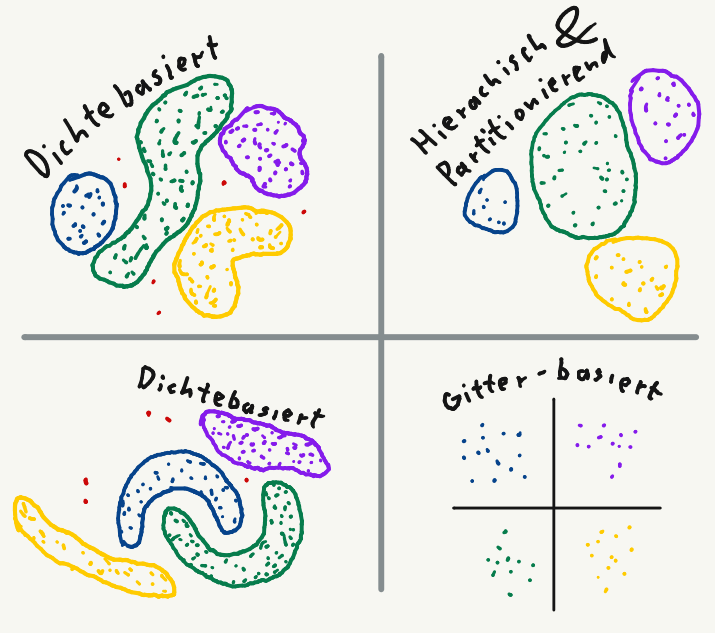
Con: Keine Iteration, Keine Optimierung

Gitterbasierend

Parameter: Gittergröße

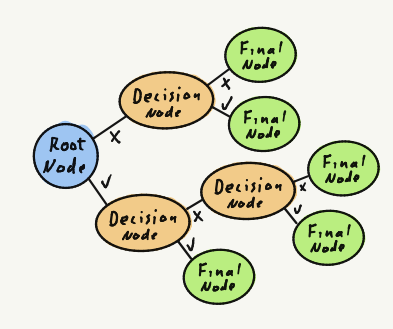
Pro: Schnell, effizient, einfach

Con: Komplexe Cluster, Parameterwahl

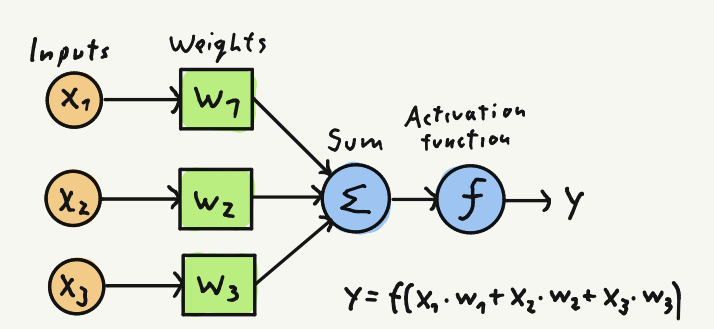


Supervised Learning

Decision Tree Algorithm



Perceptron Lernalgorithmus



Pattern Recognition

Image → Feature extraction → Classifier

(Deep Learning: Low, mid & high level Features)

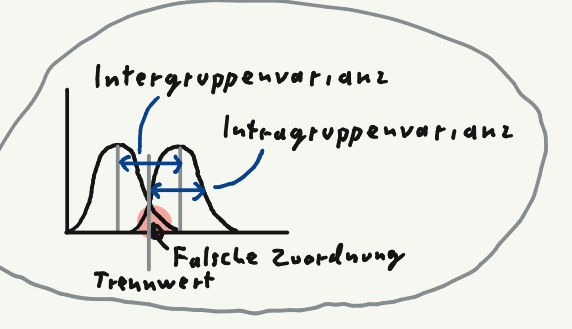
Regression

K-nearest-neighbor

Mehrheit aus K nächsten Nachbarn

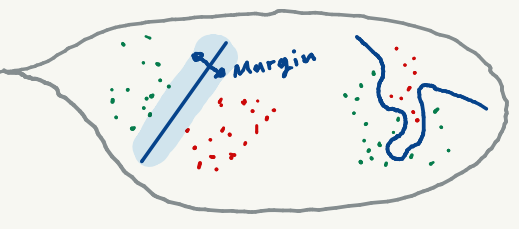
Klassifizierung

• Diskriminanzanalyse
Intergruppenvarianz möglichst groß
Intragruppenvarianz möglichst klein



Support Vector Machine

Margin möglichst groß



Regression

- Gerade
- Kurve

Ziel: Minimierung des mittleren (quadratischen) Fehlers

Semi-Supervised Learning

Teilweise gelabelte Daten.
Modell wird mit gelabelten Daten trainiert,
Ungelabelte Daten werden von ML-Modell
gelabelt und es wird erneut auf alle Daten trainiert.

Pro:

- Genauer als Unsupervised Learning
- Weniger Aufwand wegen Labels

Con:

- Ungenauer als supervised Learning

Reinforcement Learning

Agent führt Aktionen in Umgebung aus um Belohnung zu maximieren

Anwendung: Q-Learning & Monte-Carlo-Methode

