

L'intelligenza artificiale e le imprese: conoscenza, consapevolezza, ma soprattutto come usarla

Paolo Nesi, paolo.nesi@unifi.it

Gianni Pantaleo, gianni.pantaleo@unifi.it

Dipartimento di Ingegneria dell'Informazione, DISIT Lab

<https://www.disit.dinfo.unifi.it>

Giovedì 9 ottobre 2025 | ore 15.00



UNIVERSITÀ
DEGLI STUDI
FIRENZE

DINFO
DIPARTIMENTO DI
INGEGNERIA
DELL'INFORMAZIONE

DISIT
DISTRIBUTED SYSTEMS AND
INTERNET TECHNOLOGIES LAB
DISTRIBUTED DATA INTELLIGENCE
AND TECHNOLOGIES LAB



Agenda

- **Ragionamento semantico e simbolico**
- Dalla statistica al machine learning e all'intelligenza artificiale
- Apprendimento supervisionato e non supervisionato
- Dall'elaborazione del linguaggio naturale ai modelli linguistici di grandi dimensioni (LLM)
- Intelligenza artificiale per previsioni, rilevamento di anomalie, classificazione
- Applicazioni industriali: raccomandazioni
- Esercitazione su LLM
- Test

Ragionamento Simbolico

Ragionamento: attività per produrre deduzioni

- Simbolico → basato su regole e rappresentazioni esplicite
 - Logica formale e matematica
 - Regole di inferenze per produrre deduzioni tramite la soluzione di «teoremi»
 - Esempi
 - $A > B$ and $B > C \Rightarrow A > C$
 - A vero in $[3, 4]t$ and B vero dopo $A \Rightarrow B$ può essere in > 3
- **Sono Simbolici:** i modelli matematici, le risoluzioni matematiche, i modelli a grafo, etc.

Ragionamento Semantico

- **Semantico** → basato su significati, relazioni e contesto
 - Il gatto è un felino, Miu è un gatto → Miu è un felino
- Modelli per triple: <Subject><Predicate><Object> o quadruple,
 - generic definition for the ontology
 - «cat» «is a specialization of» «animal»
 - Istanze di definizioni dell'ontologia
 - «micu» «is an instance of» «cat»
 - «leg1» «is part of the» «Mytable»
 - «carl» «is belong to class» «5c»
 - «carl» «is child of» «frank»
 - «frank» «is father of» «carl»
 - Ontologia + instances → Knowledge Base

Km4City Ontology elements 1.6.8

- **DCTERMS**: for metadata Dublin Core Metadata Initiative
- **FOAF**: friends of a friends
- **Good Relation**: entities relationships
- **iot-lite**: IOT Vocabulary
- **OTN**: Ontology of Transportation Networks
- **OWL-Time**: time reasoning
- **SAREF** Smart Appliances REference extension for building devices available at <https://saref.etsi.org/saref4bldg/>
- **Schema.org** for people and organizations
- **SSN**: Semantic Sensor Network Ontology (see <https://www.w3.org/TR/vocab-ssn/>)
- **WGS84** Datum of Geo-Objects
- **GTFS**, General Transit Feed Specification, and **Transmodel**, for public transport infrastructures: lines/rides time schedules, real-time records, paths, etc.;
- **BOT**: Building Topology Ontology. <https://w3c-lbd-cg.github.io/bot/>
- **S4CITY**: SAREF extension for Smart City. <https://saref.etsi.org/saref4city/v1.1.2/>



Km4City: Knowledge Base



- Multiple DOMAINS
- Geospatial reasoning
- Temporal reasoning
- Metadata
- Statistics
- Risk and Resilience
- Licensing
- Open and Private Data
- Static and Real time
- IOT/IOE

- Street-Guide
- Mobility and transport
- Points of interest
- Sensors, IOT, ...
- Energy
- Administration
- Citations from strings

Big Data Tools

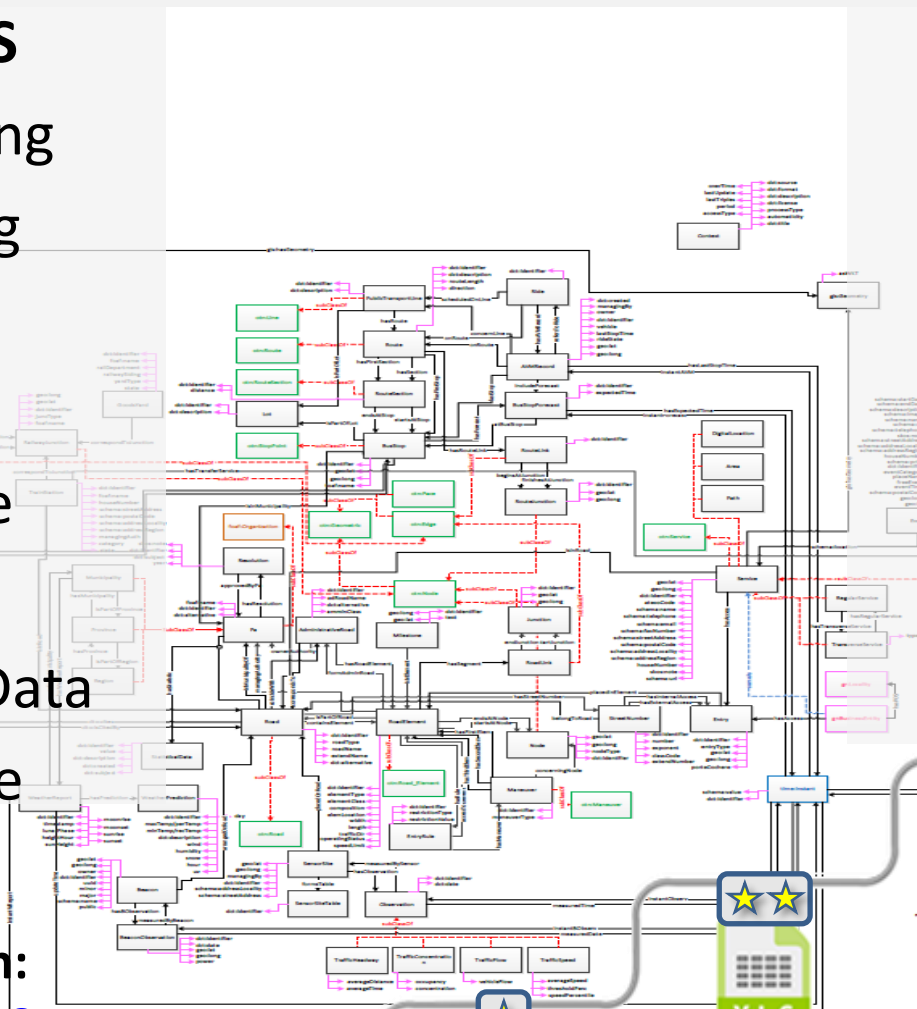


LOD and reasoners



Schema: <http://www.disit.org/km4city/schema>

RDF version: <http://www.disit.org/km4city.rdf>



Agenda

- Ragionamento semantico e simbolico
- **Dalla statistica al machine learning e all'intelligenza artificiale**
- **Apprendimento supervisionato e non supervisionato**
- Dall'elaborazione del linguaggio naturale ai modelli linguistici di grandi dimensioni (LLM)
- Intelligenza artificiale per previsioni, rilevamento di anomalie, classificazione
- Applicazioni industriali: raccomandazioni
- Esercitazione su LLM
- Test

Dalla Statistica all'Intelligenza Artificiale Generativa

- ***Statistica***: studio di fenomeni osservabili attraverso la raccolta, descrizione e analisi di dati qualitativi e quantitativi misurabili (attraverso medie, variazioni, correlazioni...);
- ***Machine Learning***: si apprende dai dati per produrre modelli predittivi, di classificazione, processi decisionali;
- ***Deep Learning***: si apprendono rappresentazioni complesse tramite reti neurali profonde.
- ***AI generativa***: modelli addestrati su grandi quantità di dati che producono contenuti (testo, immagini, audio, video...).

Machine Learning: Definizioni

- « Machine Learning is a field of study that gives computers the **ability to learn without being explicitly programmed**. » (A. L. Samuel)
 - ✓ « A computer **can be programmed so that it will learn to play a better game** of checkers **than can be played by the person** who wrote the program ».
...e...
 - ✓ « **Programming computers to learn from experience** should eventually eliminate the need for much of this detailed programming effort. »
(A. L. Samuel, "Some Studies in Machine Learning Using the Game of Checkers", 1959)
- « A computer program is said to learn from **experience E** with respect to some class of **tasks T** and **performance measure P** , if its performance at tasks in **T** , as measured by **P** , improves with experience **E** . »
(T. M. Mitchell, "Machine Learning", McGraw Hill, 1997)

Machine Learning: Apprendere dai Dati

- Invece di programmare sistemi a regole, gli algoritmi di Machine Learning imparano a inferire tali regole (pattern, relazioni nascoste tra i dati ecc...) dai dati stessi.
- Program complexity VS Data complexity

Sistemi a Regole



Machine Learning



Machine Learning: Il Problema dell'Apprendimento

- Dati due spazi di variabili: *input* $\mathbf{x} = x_1, x_2, \dots, x_N$ e *output* $\mathbf{y} = y_1, y_2, \dots, y_N$
- Supponendo che esista una relazione f non nota (*target function*) tra $\mathbf{x}$ e $\mathbf{y}$:

$$\mathbf{y} = f(\mathbf{x})$$

- Attraverso l'algoritmo di apprendimento (Machine Learning) si vuole determinare una funzione h che approssimi al meglio la funzione target f :

$$\hat{\mathbf{y}} = h(\mathbf{x}) + \varepsilon$$

dove ε è una misura dell'errore tra $\hat{\mathbf{y}}$ (valore predetto dal modello) e $\mathbf{y}$.

- Questo vuol dire **imparare dai dati**, senza utilizzare regole o funzioni predefinite da applicare alle variabili di ingresso (input) per ottenere i risultati in uscita (output).
- L'obiettivo, dunque, è quello di minimizzare l'errore ε , in modo che la funzione h possa essere utilizzata per predire nuovi valori di output a partire da un nuovo insieme di variabili di input non osservate in precedenza.

Classificazione degli Algoritmi di Apprendimento

Classificazione dei processi di apprendimento attraverso cui stimare una funzione (non nota) f che descriva pattern e/o relazioni nascoste tra data set di input e output attesi:

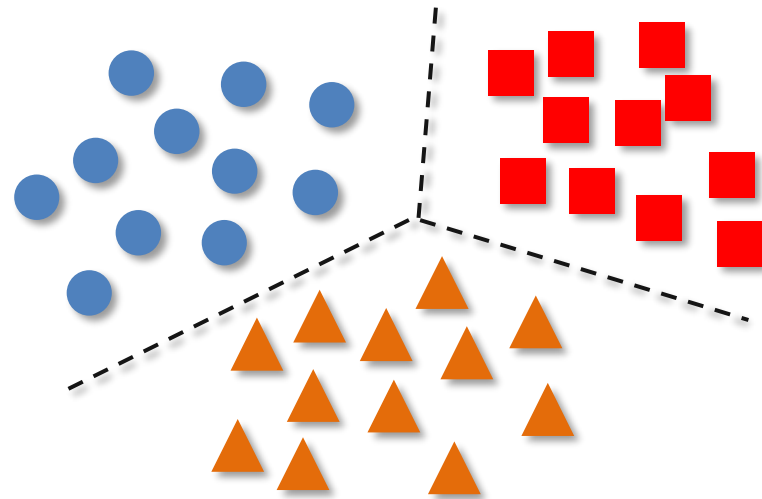
- **Supervised learning:** l'algoritmo produce un modello h (che approssima f secondo opportuni criteri di minimizzazione dell'errore) addestrato con un data set (*training* data set) annotato $\mathbb{T} = \{(x_i, y_i)\}_{i=1}^N$:

$$\mathbf{y} = h(\mathbf{x}) + \varepsilon$$

$$\mathbf{x} = x_1, x_2, \dots, x_N; \quad \mathbf{y} = y_1, y_2, \dots, y_N; \quad \hat{\mathbf{y}} = \hat{y}_1, \hat{y}_2, \dots, \hat{y}_N$$

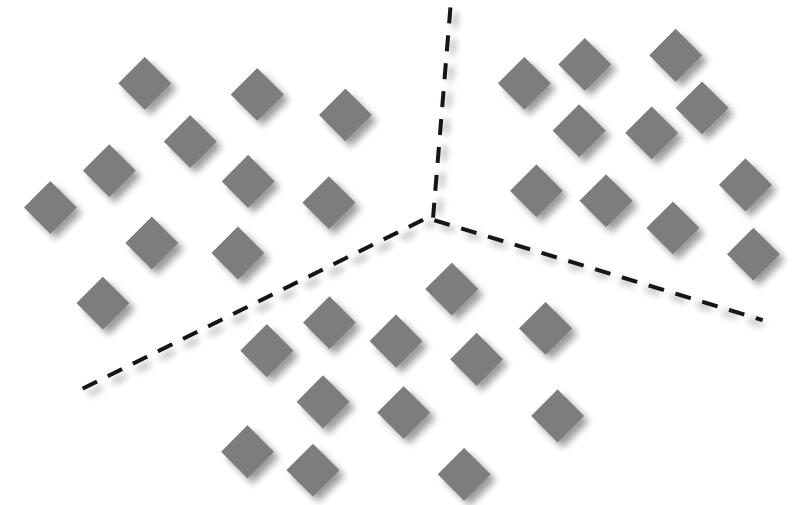
- **Unsupervised learning:** l'algoritmo riceve in input un set di dati non annotati $\{x_i\}_{i=1}^N$ con l'intento di trovare al suo interno pattern, strutture e relazioni latenti per eseguire task di classificazione, raggruppamento ecc.
- **Semi-Supervised Learning:** l'algoritmo produce un modello sulla base di un data set in cui sono presenti una minor parte di dati annotati e una maggior parte di dati non annotati.

Classificazione degli Algoritmi di Apprendimento



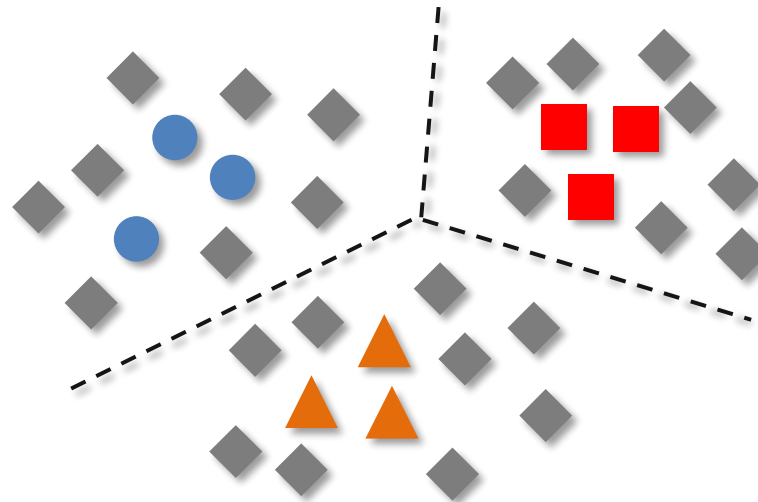
Supervised Learning

Apprendere da esempi di cui conosciamo già la risposta. Ad esempio, addestrare un algoritmo di classificazione usando un dataset annotato, ovvero con dati di cui già abbiamo una classificazione)



Unsupervised Learning

scoprire strutture e relazioni nei dati senza risposte predefinite (cioè avendo a disposizione dati non annotati precedentemente)



Semi-Supervised Learning

Unsupervised Learning

- Esempio di Unsupervised Learning → **Clustering K-Means**: partiziona un insieme di dati in K gruppi ("cluster")
- **Task**: Raggruppare i clienti in segmenti omogenei partendo solo dai loro comportamenti d'acquisto (nessuna etichetta nota).
- **Feature usate**:
 - `recency_days`: giorni dall'ultimo acquisto
 - `orders_yr`: numero ordini negli ultimi 12 mesi
 - `avg_basket`: importo medio di acquisto (€)

Unsupervised Learning

- Metodo: **K-Means**: i dati vengono raggruppati in un numero K di cluster, minimizzando la distanza tra i «punti» (ovvero, la rappresentazione dei dati nello spazio delle features) e il centro del loro gruppo.
- Apprendimento iterativo: ad ogni iterazione viene assegnato ogni punto al centro (*centroide*) più vicino e vengono ricalcolati i centroidi, di solito usando la distanza euclidea.
- Risultato atteso: 3 cluster interpretabili, ad esempio:
 - Cluster A (client "Gold"): recency bassa, molti ordini, importi di acquisto alti;
 - Cluster B (client "Silver"): valori intermedi;
 - Cluster C (client "Bronze"): recency alta, pochi ordini, importi di acquisto bassi.
- Esempio tramite uso dell'applicativo **Weka** (app open-source per Machine Learning e Data Mining): <https://ml.cms.waikato.ac.nz/weka/>

Unsupervised Learning

- Input per Weka: file **customers_k-means.arff**

```
@relation customers_k-means

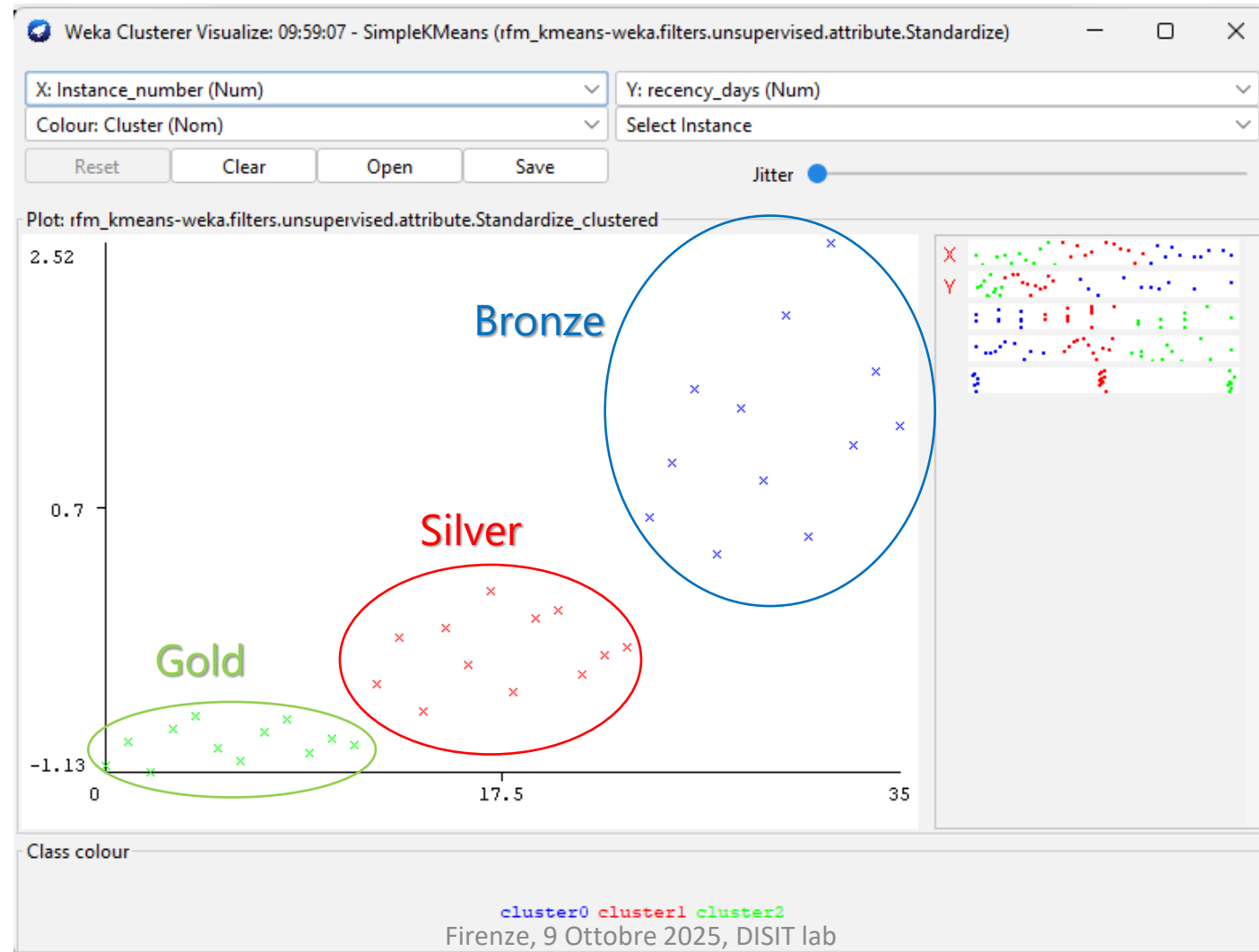
@attribute recency_days numeric
@attribute orders_yr      numeric
@attribute avg_basket     numeric

@data
15,9,105
28,8,110
12,10,130
...
```

Unsupervised Learning

Visualizzazione
del Clustering
basata sulla
feature:

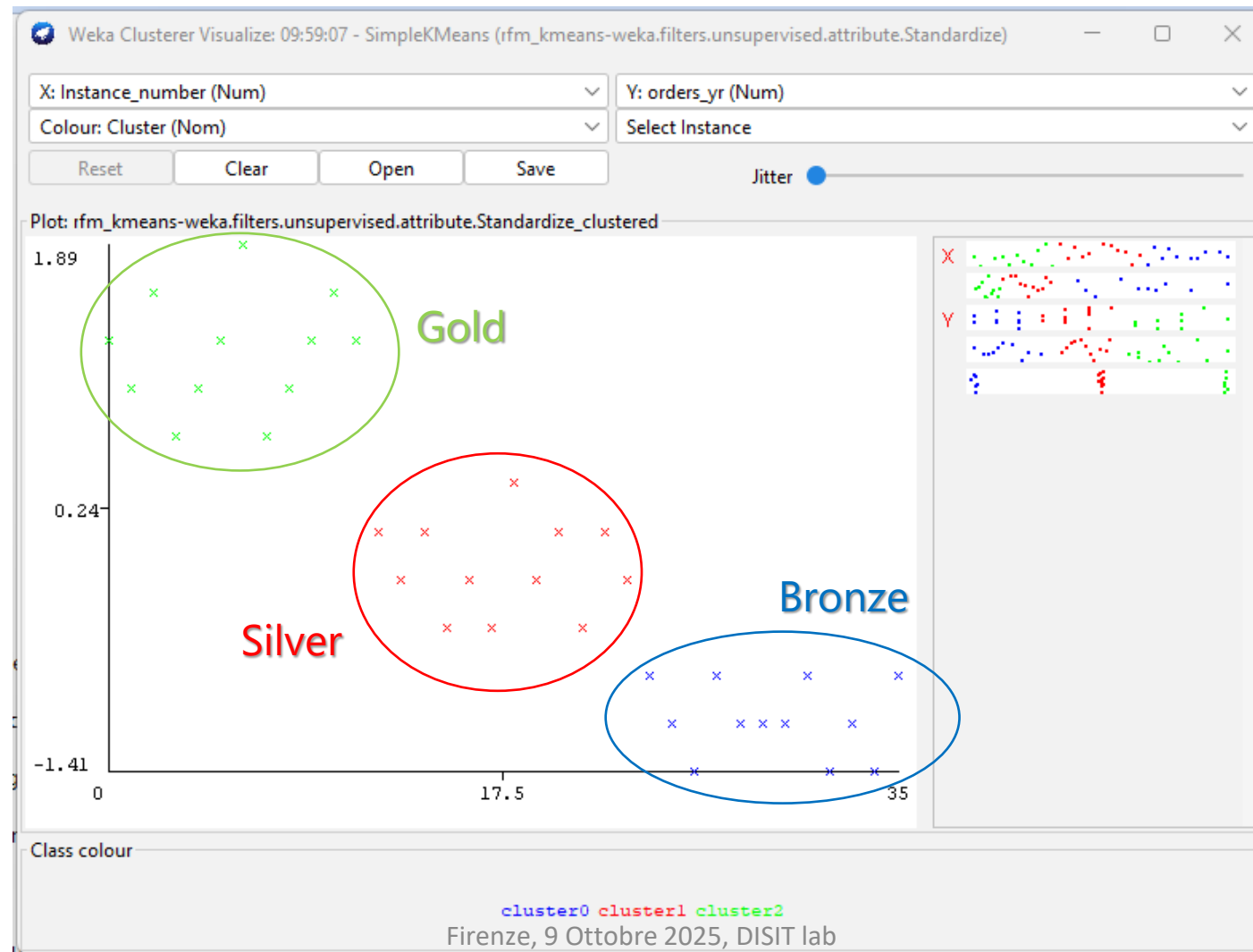
recency_days



Unsupervised Learning

Visualizzazione
del Clustering
basata sulla
feature:

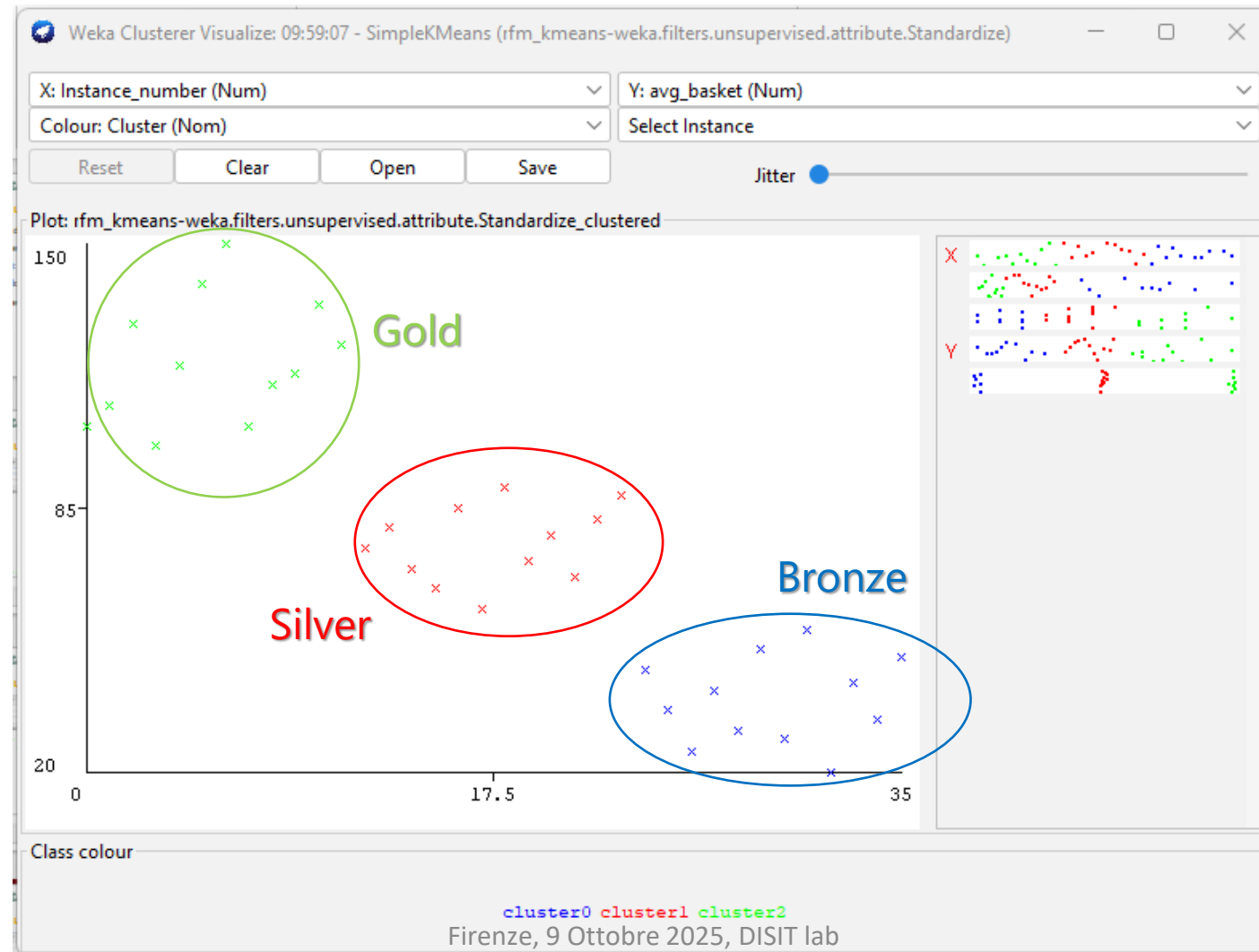
orders_yr



Unsupervised Learning

Visualizzazione
del Clustering
basata sulla
feature:

average_basket

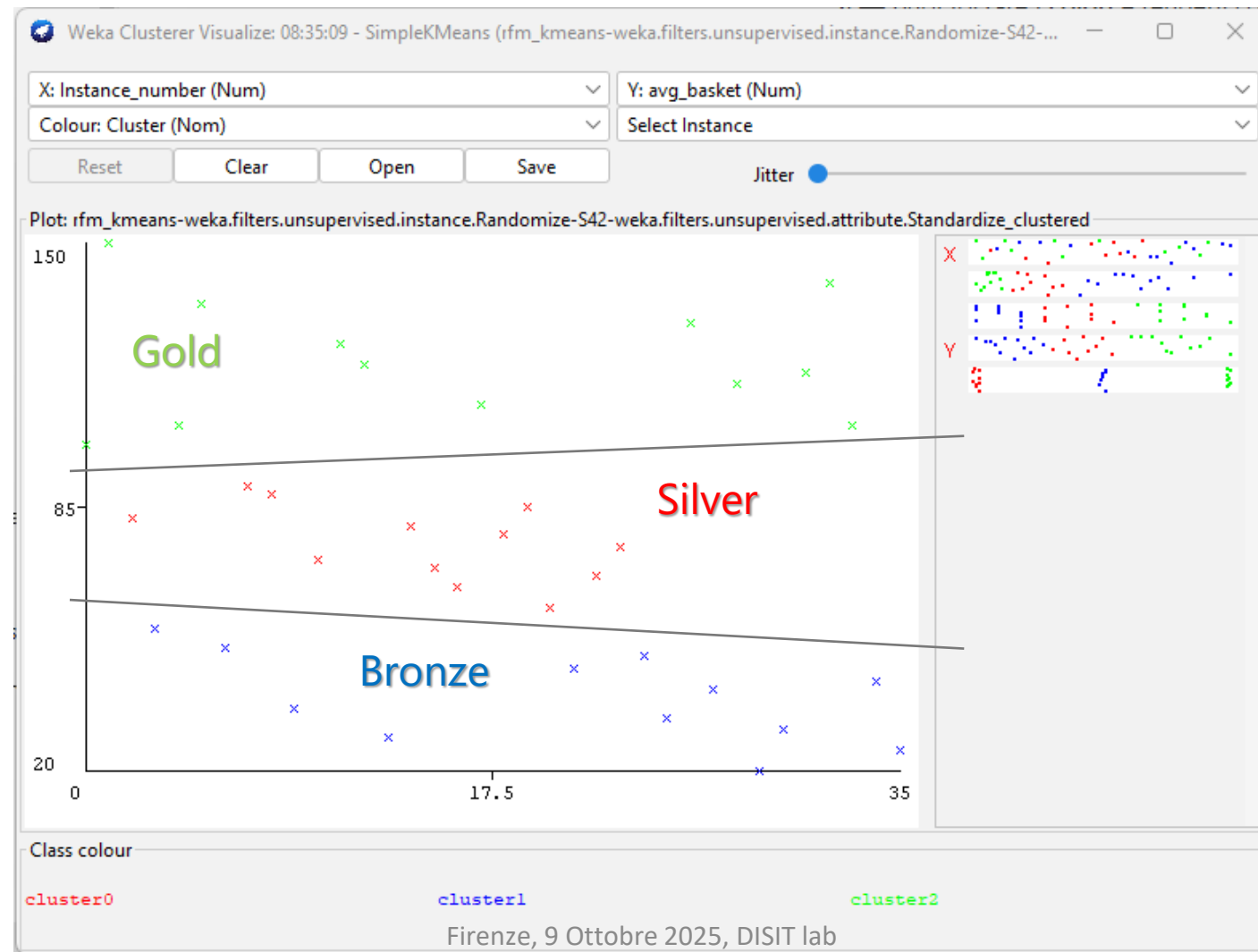


Unsupervised Learning

Visualizzazione
del Clustering
basata sulla
feature:

average_basket

...ma con
differente
ordinamento dei
dati



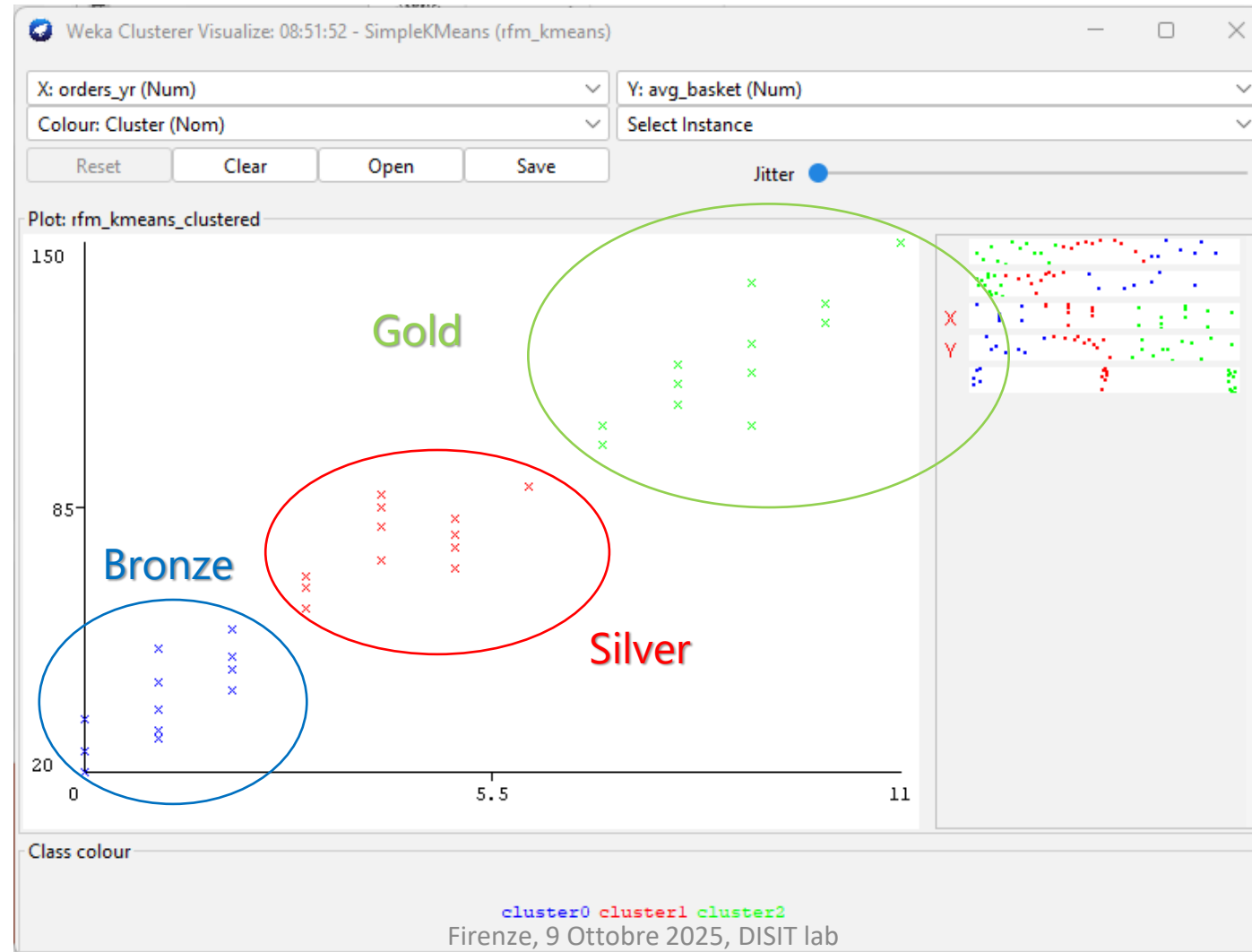
Unsupervised Learning

Visualizzazione
del Clustering
basata sulle
features:

average_basket

e

orders_yr



Informazione
combinata che si può
dedurre tra le due
feature
average_basket e
orders_yr):

forte correlazione
positiva:
più ordini ⇒
scontrino medio più
alto.

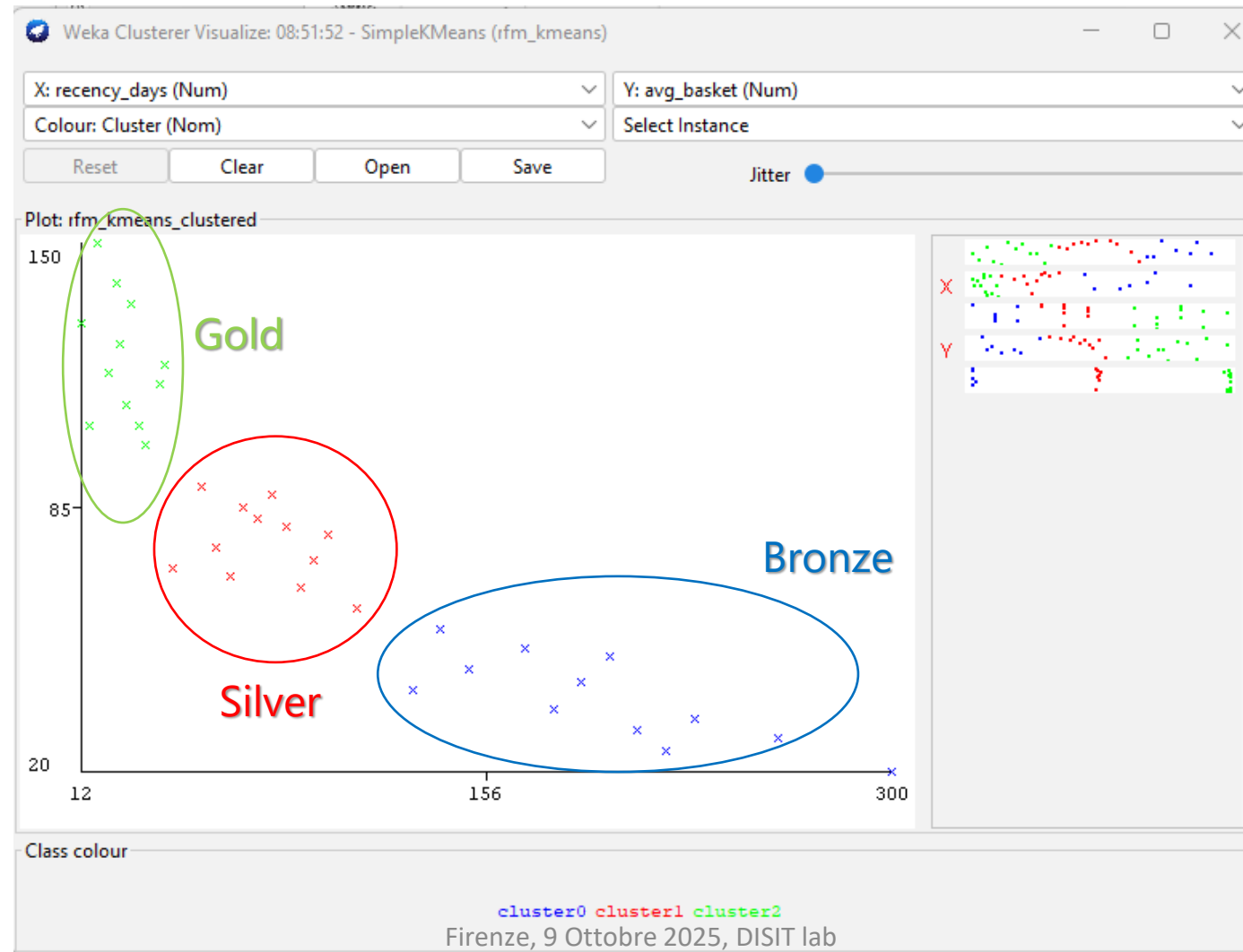
Unsupervised Learning

Visualizzazione
del Clustering
basata sulla
feature:

average_basket

e

recency_days



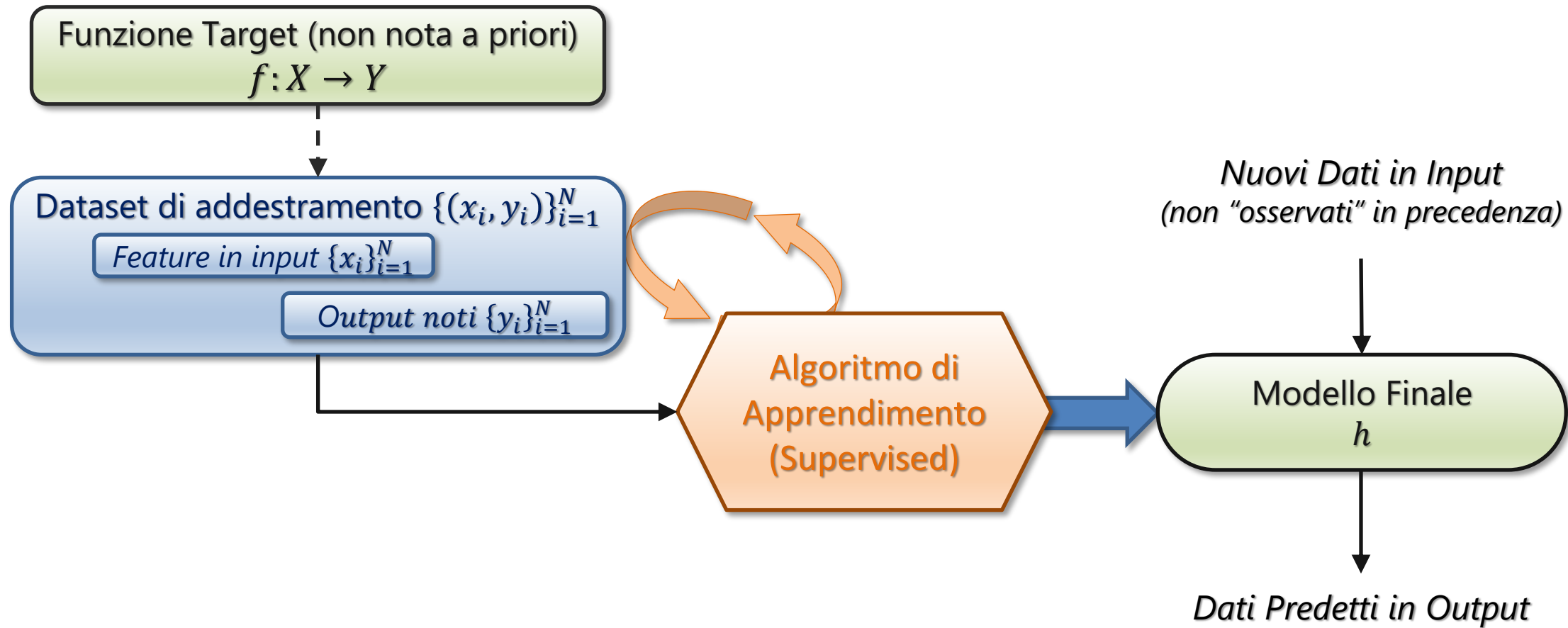
Informazione
combinata che si può
dedurre tra le due
feature
average_basket e
recency_days):

correlazione
negativa:
più giorni dall'ultimo
acquisto $\Rightarrow$ scontrino
medio più basso.

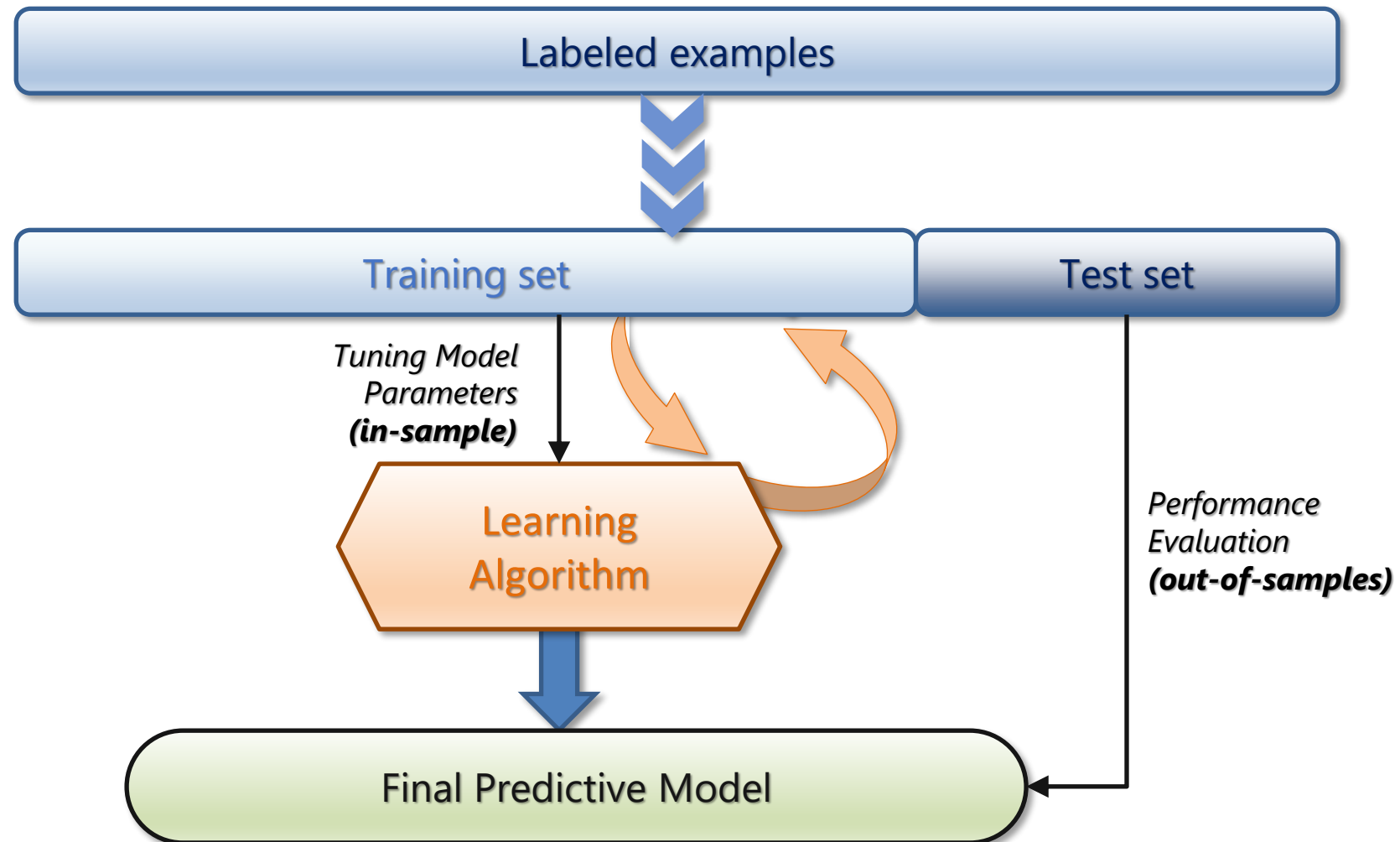
Supervised Learning – Tipologie di Problemi e Modelli

- **Regressione:** La variabile da predire è di tipo continuo. Il numero di feature in input può essere uno (regressione uni-variata) o maggiore di uno (regressione multi-variata)
 - Prezzo di acquisto di un appartamento data l'estensione in metri quadrati, numero di stanze, numero di piani (regressione multivariata)
 - Voti agli esami date le ore di studio (regressione univariata)
 - ...
- **Classificazione:** La variabile da predire è di tipo discreto (numero finito di classi: binary-class / multi-class...)
 - Classificazione delle email: spam / non spam (classificazione binaria)
 - Alberi Decisionali (Decision Trees)
 - ...

Supervised Learning



Supervised Learning: Training, Validation, Test



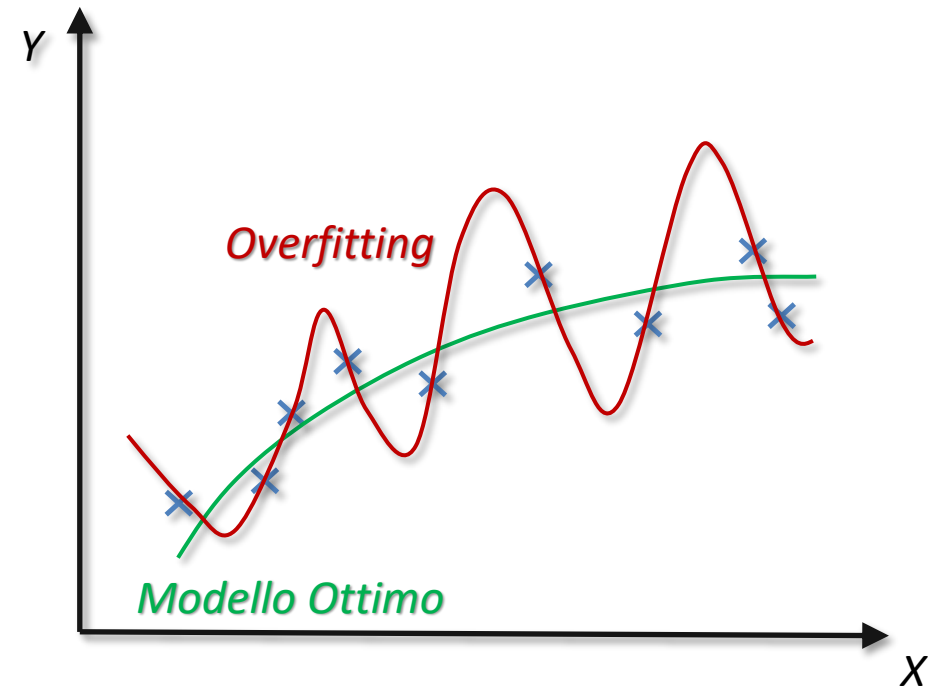
Supervised Learning - Overfitting

Il problema dell'Overfitting

Aumentare la complessità del modello in modo da ottenere ottime performance (in termini di capacità predittiva) in fase di training spesso porta al fenomeno dell'**Overfitting**.

Questo si traduce in una diminuzione dell'errore del modello sul dataset di training ma in un aumento dell'errore sul dataset di test (il modello si comporta molto bene in fase di training ma non riesce a generalizzare in maniera soddisfacente su dati out-of-sample).

Ciò accade perché il modello, in questi casi, viene addestrato ad adattarsi anche il rumore, oltre che ai dati.

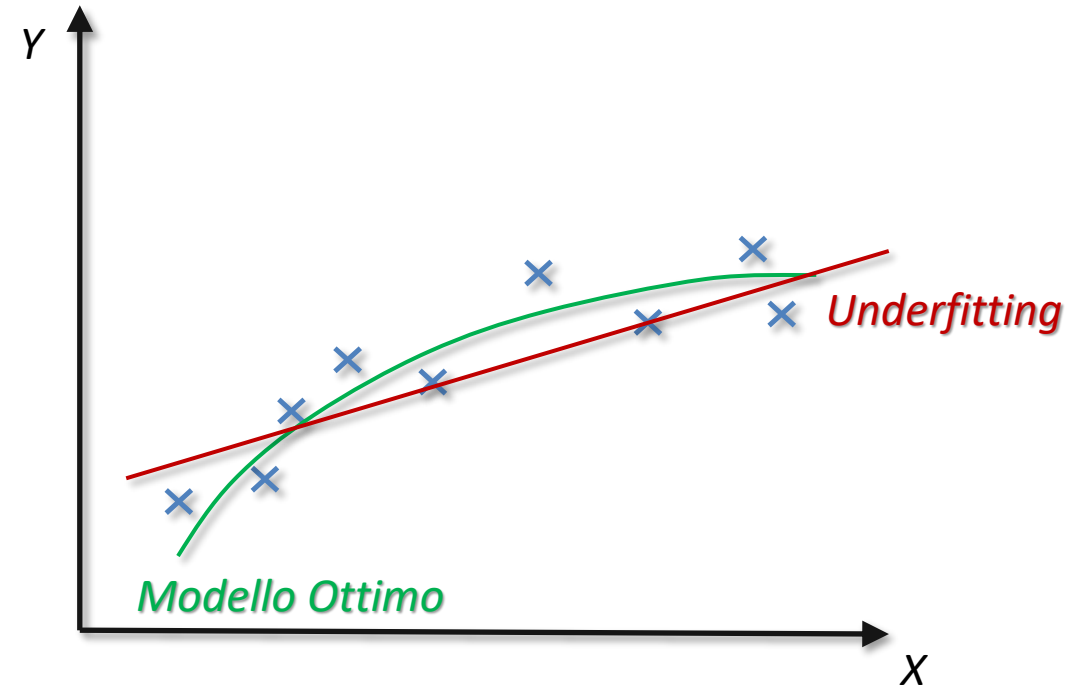


Supervised Learning - Underfitting

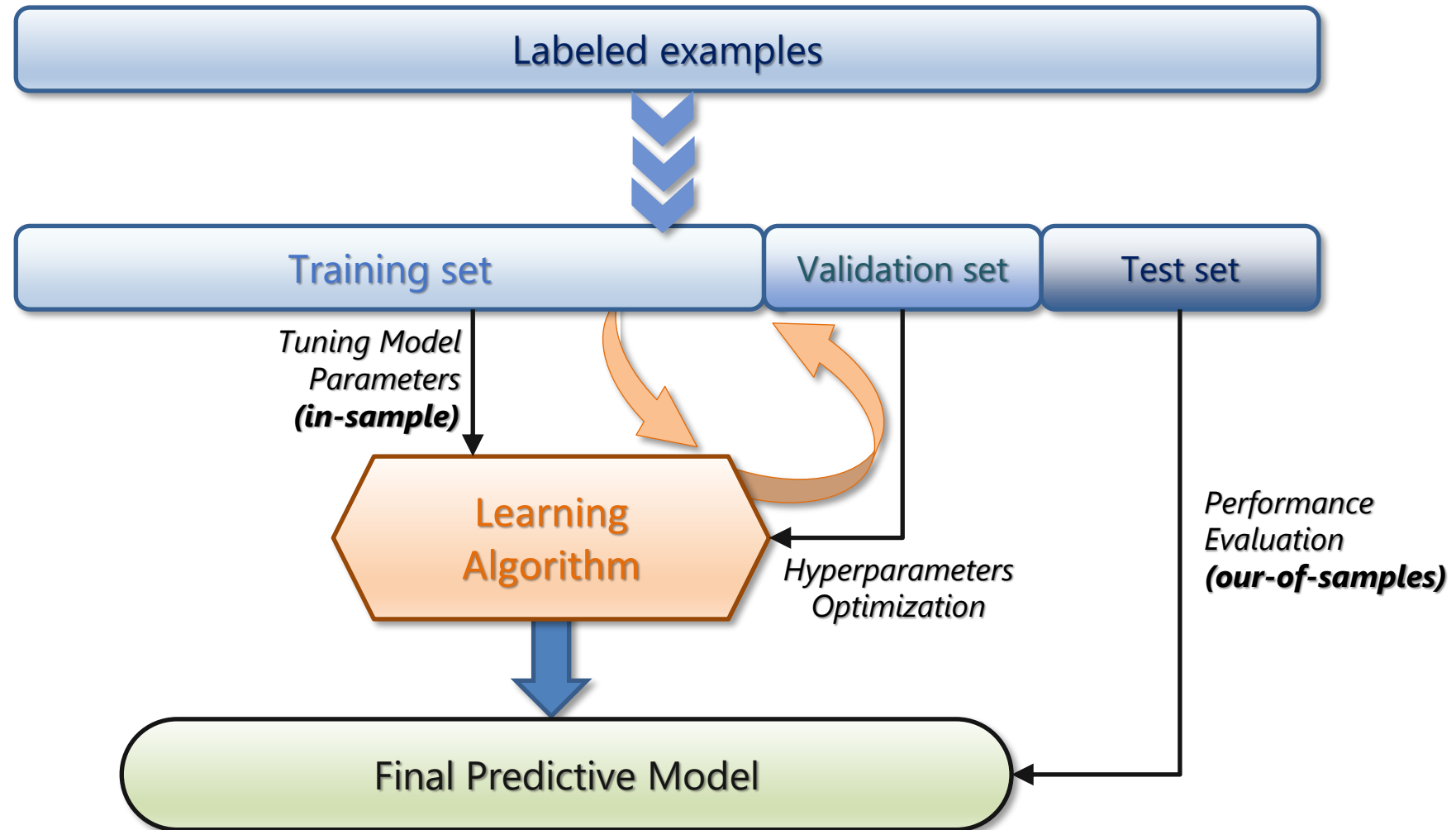
Il problema dell'Underfitting

Il problema opposto all'Overfitting, ovvero l'**Underfitting**, si incontra quando il modello è troppo semplice (basso numero di parametri/iterazioni, pochi dati nel data set di training...)

In questo caso il modello non riesce ad apprendere pattern, trend dai dati e tipicamente non raggiunge prestazioni soddisfacenti (in termini di capacità predittive o di classificazione).

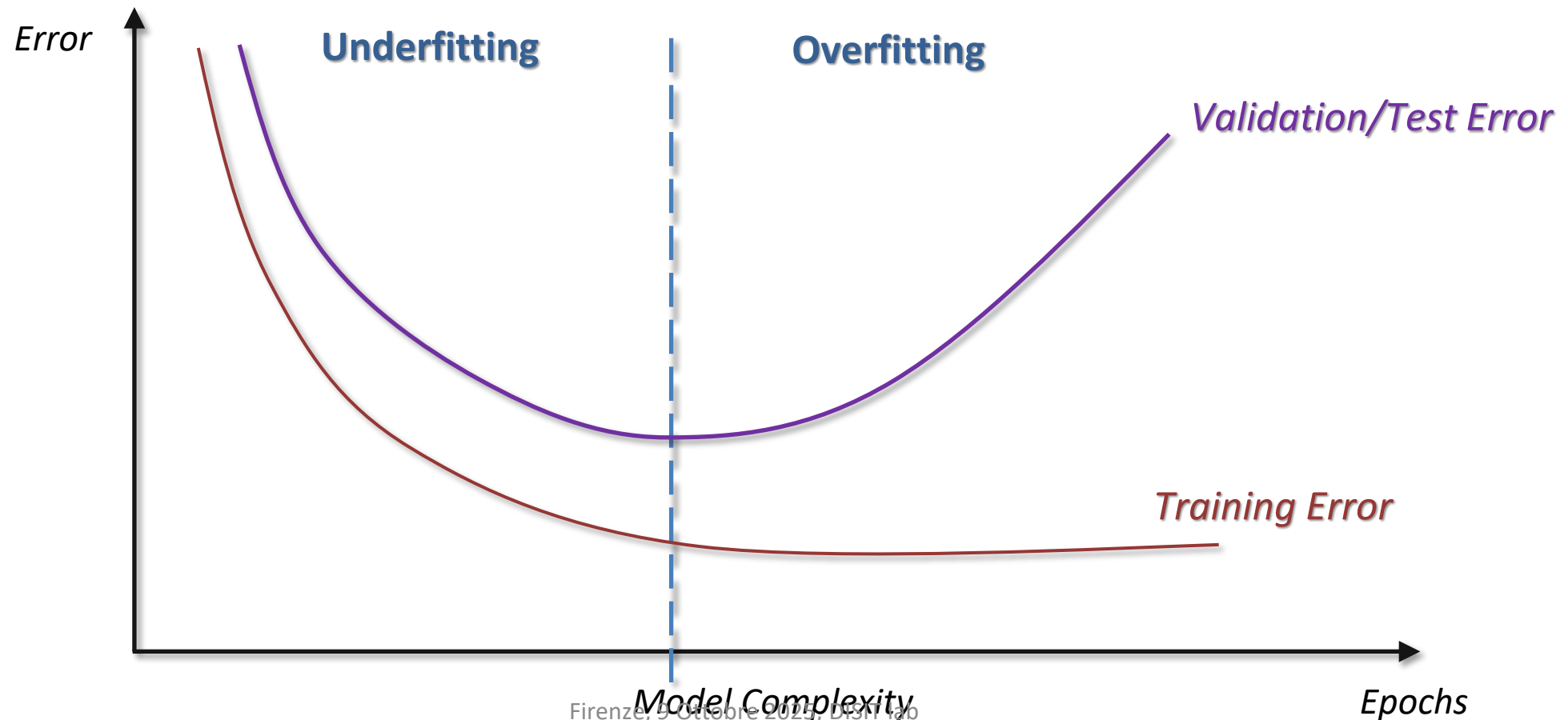


Supervised Learning: Training, Validation, Test



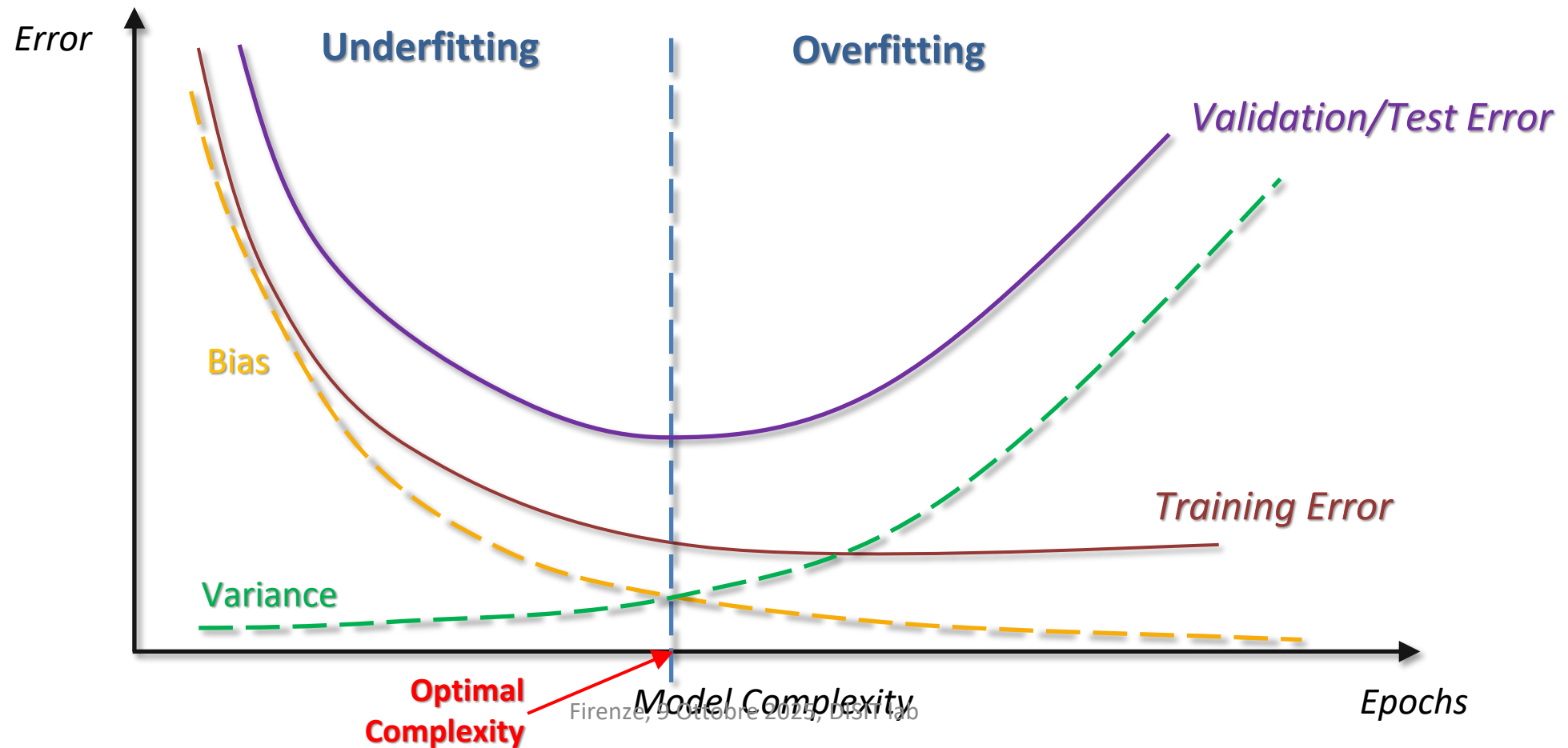
Supervised Learning – Overfitting VS Underfitting

Underfitting, Overfitting – Errore di Training ed Errore di Test



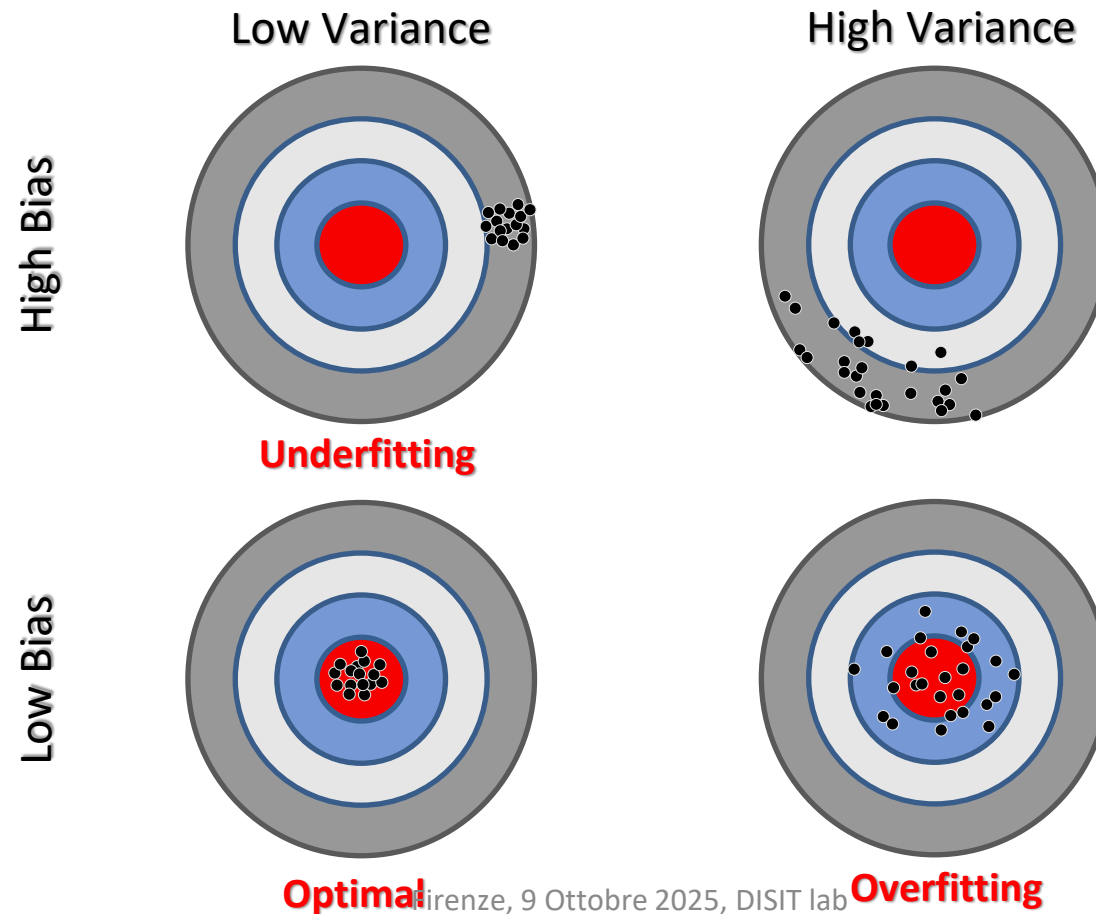
Supervised Learning – Overfitting VS Underfitting

Underfitting, Overfitting – Trade-off tra Bias e Varianza



Supervised Learning – Overfitting VS Underfitting

Underfitting, Overfitting – Trade-off tra Bias e Varianza



Supervised Learning – Decision Trees

- Problema: Segmentazione clienti.
- Task: assegnare ogni cliente a una delle seguenti tre categorie: Bronze, Silver, Gold.
- Features:
 - **recency_days:** giorni dall'ultimo acquisto
 - **orders_yr:** numero ordini negli ultimi 12 mesi
 - **avg_basket:** importo medio di acquisto (€)
 - **coupon_user:** ha usato coupon nell'ultimo anno? (0/1)

Supervised Learning – Decision Trees

recency_days	orders_yr	avg_basket (€)	coupon_user	segment
15	9	95	1	Gold
28	8	110	1	Gold
12	10	130	1	Gold
35	7	100	0	Gold
42	8	120	1	Gold
25	9	140	1	Gold
18	11	150	1	Gold
33	7	105	0	Gold
40	8	115	1	Gold
22	9	98	1	Gold
60	5	75	0	Silver
85	4	80	1	Silver
45	5	70	0	Silver
90	3	65	0	Silver
70	4	85	1	Silver

recency_days	orders_yr	avg_basket (€)	coupon_user	segment
110	3	60	0	Silver
55	6	90	1	Silver
95	4	72	0	Silver
100	5	78	0	Silver
65	3	68	1	Silver
150	2	45	0	Bronze
180	1	35	0	Bronze
220	0	25	0	Bronze
130	2	40	0	Bronze
210	1	30	0	Bronze
170	3	50	1	Bronze
260	1	28	0	Bronze
140	2	55	0	Bronze
300	0	20	0	Bronze
190	1	42	0	Bronze

Supervised Learning – Decision Trees

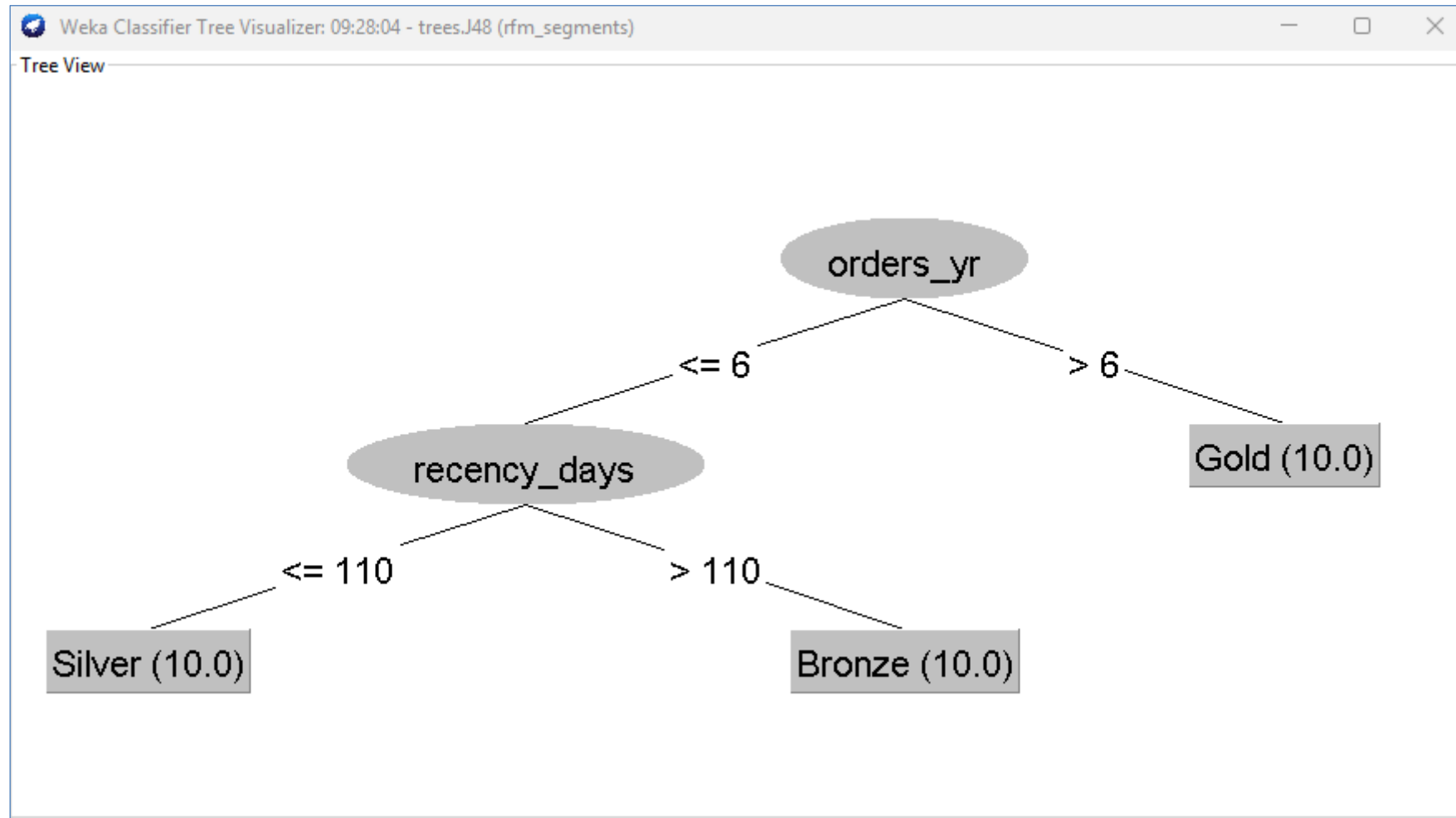
- Input per Weka: file **customers_decision-trees.arff**

```
@relation customers_decision-trees

@attribute recency_days numeric
@attribute orders_yr      numeric
@attribute avg_basket     numeric
@attribute coupon_user    {0,1}
@attribute segment        {Bronze,Silver,Gold}

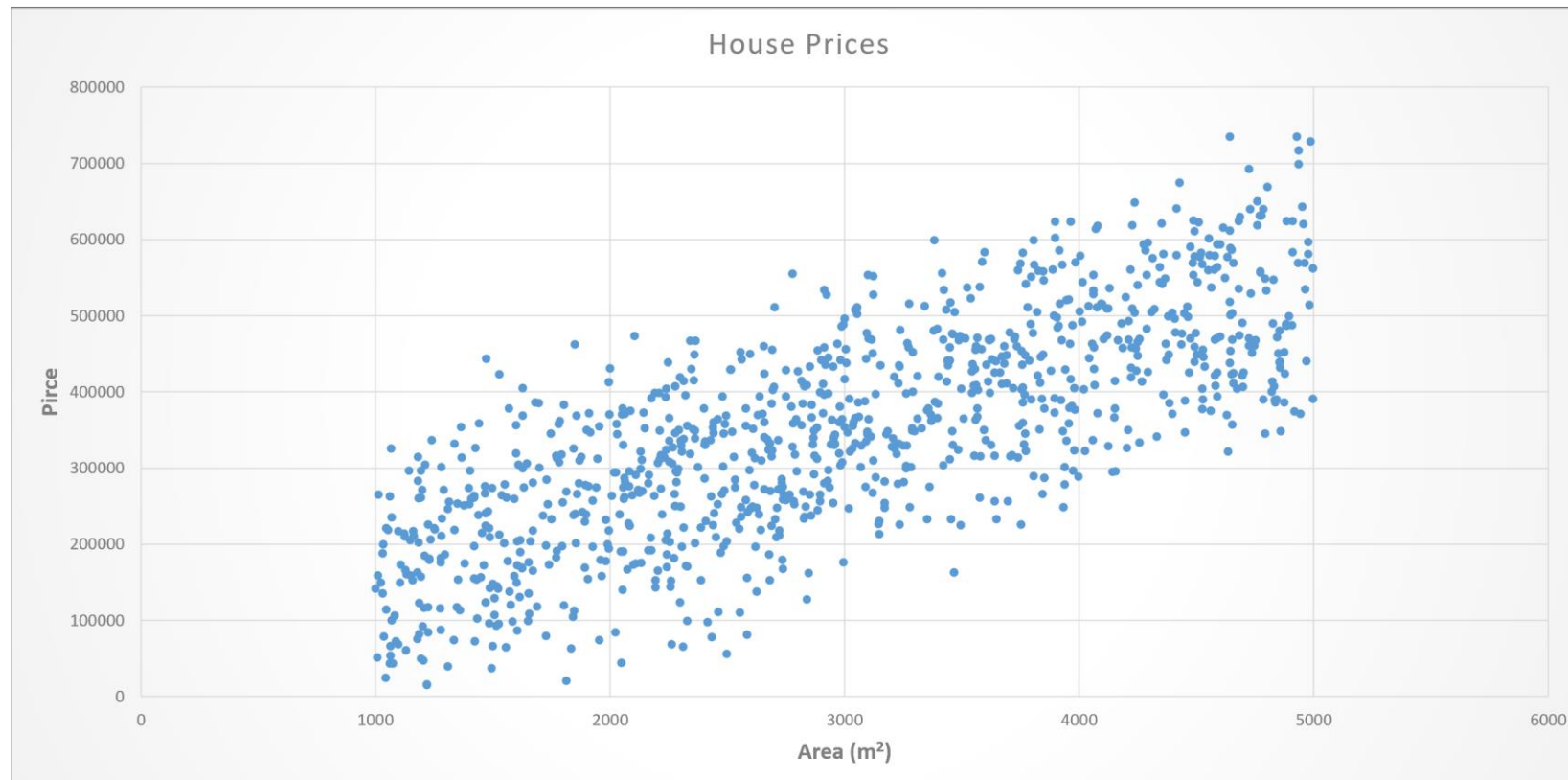
@data
15,9,95,1,Gold
28,8,110,1,Gold
12,10,130,1,Gold
...
```

Supervised Learning – Decision Trees



Supervised Learning – Regressione Lineare

Regressione Lineare



Esempio: previsione del prezzo di un appartamento a partire dalla sua estensione (metri quadrati).

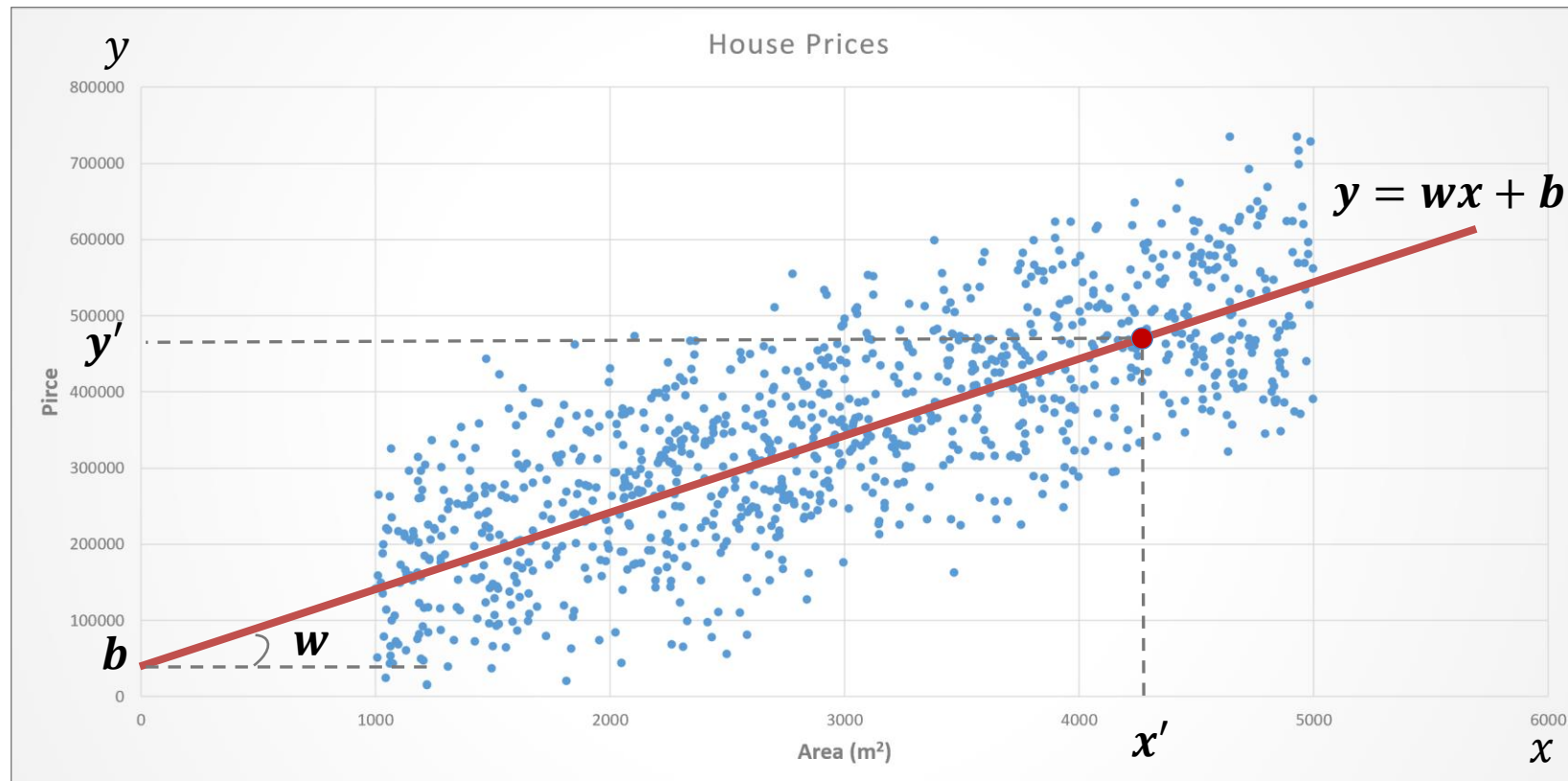
Si parte da un dataset annotato, in cui si hanno a disposizione degli esempi di appartamenti di cui si conosce sia l'estensione in metri quadrati, sia il prezzo.

Si applica un modello predittivo lineare: la **regressione lineare**

Dataset from GitHub: <https://github.com/kshitizrohilla/house-price-prediction-using-simple-linear-regression/blob/main/data/house-prices.csv>

Supervised Learning – Regressione Lineare

Regressione Lineare



L'obiettivo è quindi trovare una funzione lineare (retta: $y = wx + b$) che approssimi al meglio i dati di partenza, in modo da poter usare la funzione trovata per predire nuovi esempi (non appartenenti al dataset di partenza), ovvero il prezzo di un appartamento (y') di cui si conosce solo l'estensione in metri quadri (x').

Dataset from GitHub: <https://github.com/kshitizrohilla/house-price-prediction-using-simple-linear-regression/blob/main/data/house-prices.csv>

Machine Learning: Definizioni

- « Machine Learning is a field of study that gives computers the *ability to learn without being explicitly programmed.* » (A. L. Samuel)
 - ✓ « A computer *can be programmed so that it will learn to play a better game of checkers than can be played by the person who wrote the program* ».
 - ...e...
 - ✓ « *Programming computers to learn from experience* should eventually eliminate the need for much of this detailed programming effort. »

(A. L. Samuel, "Some Studies in Machine Learning Using the Game of Checkers", 1959)
- « A computer program is said to learn from **experience E** with respect to some class of **tasks T** and **performance measure P** , if its performance at tasks in **T** , as measured by **P** , improves with experience **E** . »

(T. M. Mitchell, "Machine Learning", McGraw Hill, 1997)

Supervised Learning – Regressione Lineare

Regressione Lineare

- **Experience E** : Dato un set osservabile di dati annotati $\mathbb{T} = \{(x_i, y_i)\}_{i=1}^N$, dove $\mathbf{x} = x_1, x_2, \dots, x_N$ è un vettore di ***dati in ingresso*** (ad esempio, l'estensione di un appartamento in metri quadri), e $\mathbf{y} = y_1, y_2, \dots, y_N$ un vettore di ***risultati noti in uscita*** (il prezzo dell'appartamento).
- **Task T** : stimare tramite apprendimento la funzione non nota $f: \mathbf{y} = f(\mathbf{x})$ attraverso un modello lineare:

$$h_{\mathbf{w}, \mathbf{b}}: \hat{\mathbf{y}} = \mathbf{w}\mathbf{x} + \mathbf{b}$$

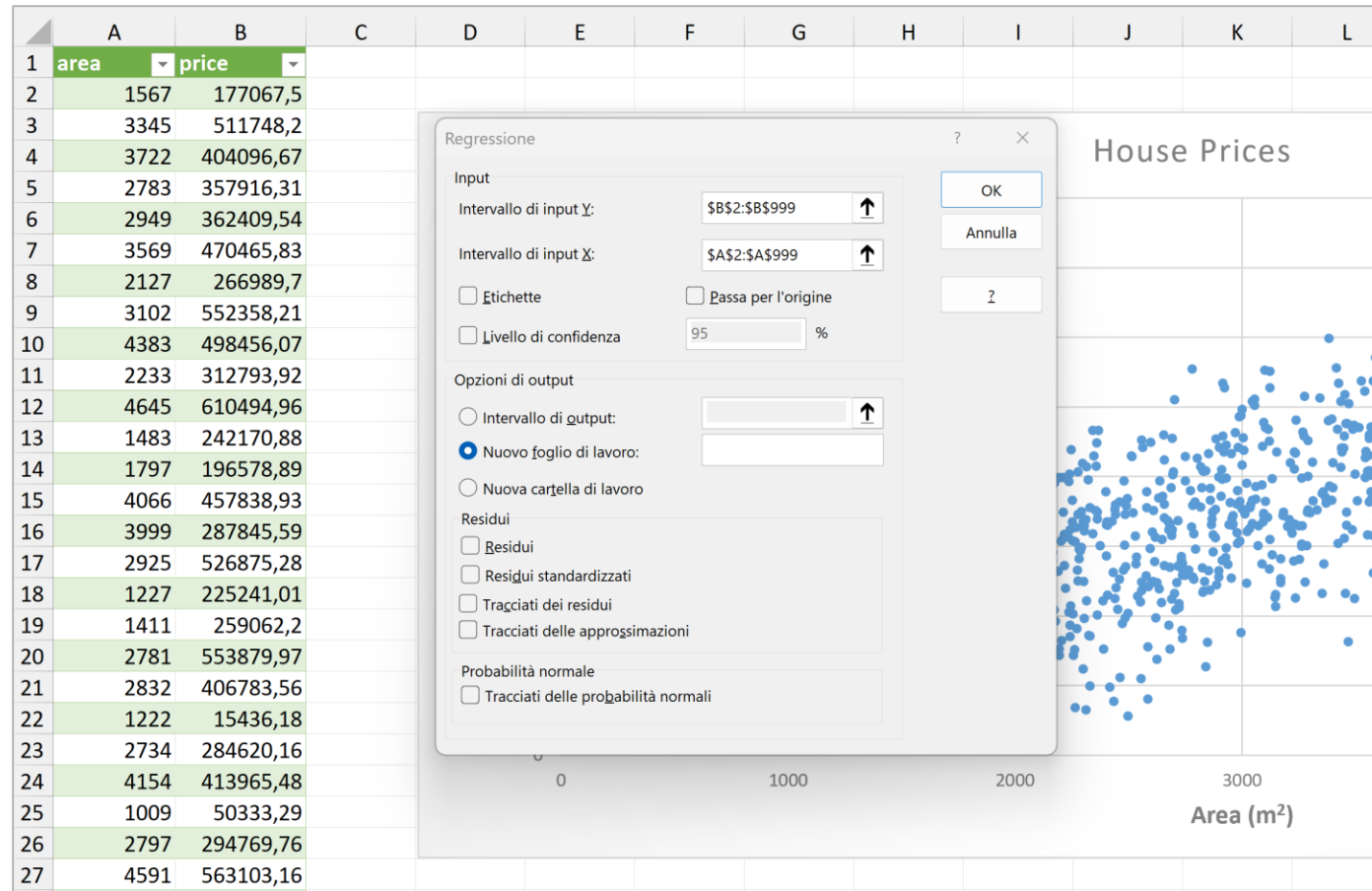
che approssimi f , per predire valori non noti $\hat{\mathbf{y}}$ partendo da dati non precedentemente annotati;

- **Performance P** : minimizzare l'errore di predizione (errore quadratico medio, **MSE**):

$$MSE = \underset{\substack{\text{Cost Function} \\ \downarrow}}{cf} = \frac{1}{N} \sum_{i=1}^N (\hat{y}_i - y_i)^2 = \frac{1}{N} \sum_{i=1}^N (\hat{y}_i - (\mathbf{w}x_i + \mathbf{b}))^2$$

Supervised Learning – Regressione Lineare con Excel

Dati → Analisi Dati → Regressione



Supervised Learning – Regressione Lineare con Excel

OUTPUT RIEPILOGO

Statistica della regressione

R multiplo	0,765397969
R al quadrato	0,585834051
R al quadrato corretto	0,585418221
Errore standard	91007,29115
Osservazioni	998

R è il coefficiente di Correlazione di Pearson. Misura la forza e la direzione di una relazione lineare tra variabili: un valore pari a 1 indica massima correlazione positiva, 0 nessuna correlazione, -1 massima correlazione negativa.

R quadro è una misura che fornisce informazioni sulla bontà di adattamento di un modello di regressione lineare. Più è vicino a 1, più le variazioni della variabile dipendente sono spiegate dalla variabile indipendente.

ANALISI VARIANZA

	<i>gdl</i>	<i>SQ</i>	<i>MQ</i>	<i>F</i>	<i>Significatività F</i>
Regressione	1	1,16684E+13	1,16684E+13	1408,833138	7,4362E-193
Residuo	996	8,2492E+12	8282327043		
Totale	997	1,99176E+13			

	<i>Coefficienti</i>	<i>Errore standard</i>	<i>Stat t</i>	<i>Valore di significatività</i>	<i>Inferiore 95%</i>	<i>Superiore 95%</i>	<i>Inferiore 95,0%</i>	<i>Superiore 95,0%</i>
Intercetta	62920,7389	8144,569225	7,725483959	2,71015E-14	46938,25463	78903,22316	46938,25463	78903,22316
Variabile X 1	96,0560245	2,559144621	37,53442604	7,4362E-193	91,03409056	101,0779584	91,03409056	101,0779584

$$y = 96.0560 x + 62920.74$$

w *b*

Limiti del Machine Learning

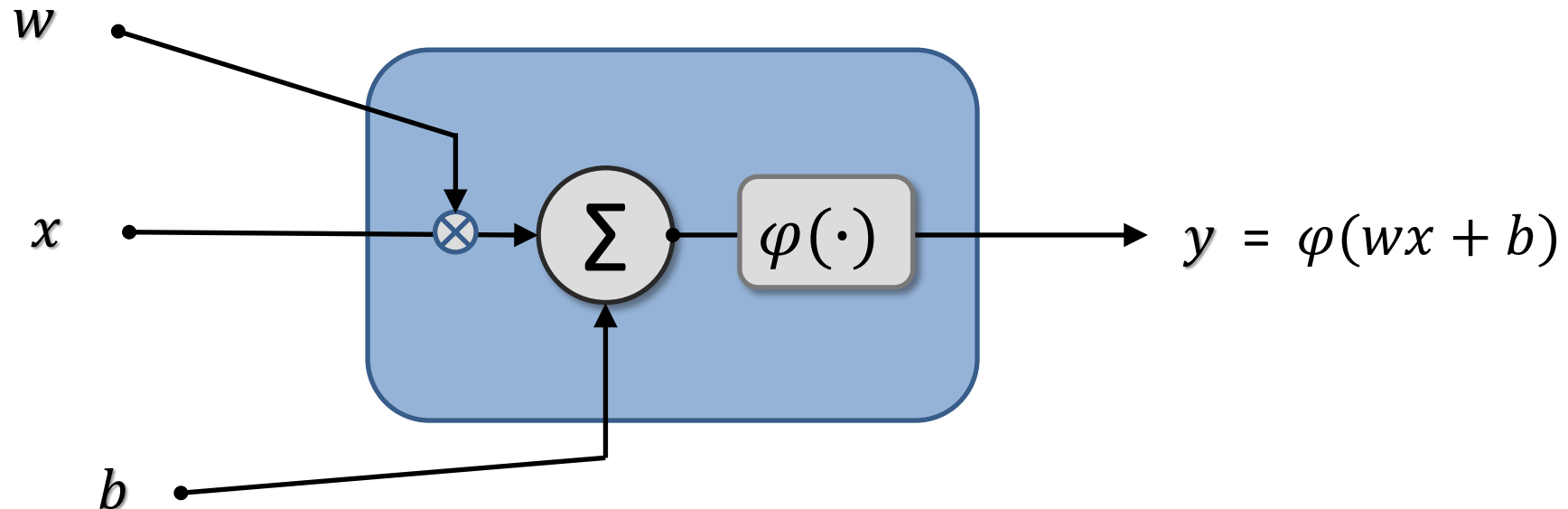
- **Feature engineering manuale:** bisogna dire al modello «cosa guardare»
- **Difficoltà con dati non strutturati:** testo, immagini, audio, ...
- **Relazioni non lineari complesse:** le regole con cui istruire l'algoritmo diventano troppo intricate
- **Esempio:** riconoscere un gatto in un dataset di foto di vari animali
 - Machine Learning classico: necessario definire esplicitamente le features: "orecchie", "baffi", "coda"
- **Problema:** e se il gatto è di profilo? Se è parzialmente nascosto?...

Deep Learning: Introduzione

- Il *Deep Learning* è una modalità di apprendimento basata su reti neurali (*Artificial Neural Networks – ANN*), le quali posso ricercare pattern, effettuare predizioni, esplicitare relazioni ed associazioni esistenti tra un insieme di coppie di dati di input e di output.
- La terminologia «*Deep Learning*» si basa sul concetto di profondità di una rete neurale, che è associato al numero di strati interconnessi (*layers*) di celle neurali di cui essa è composta. Tale parametro è rappresentativo quindi della complessità della rete neurale.

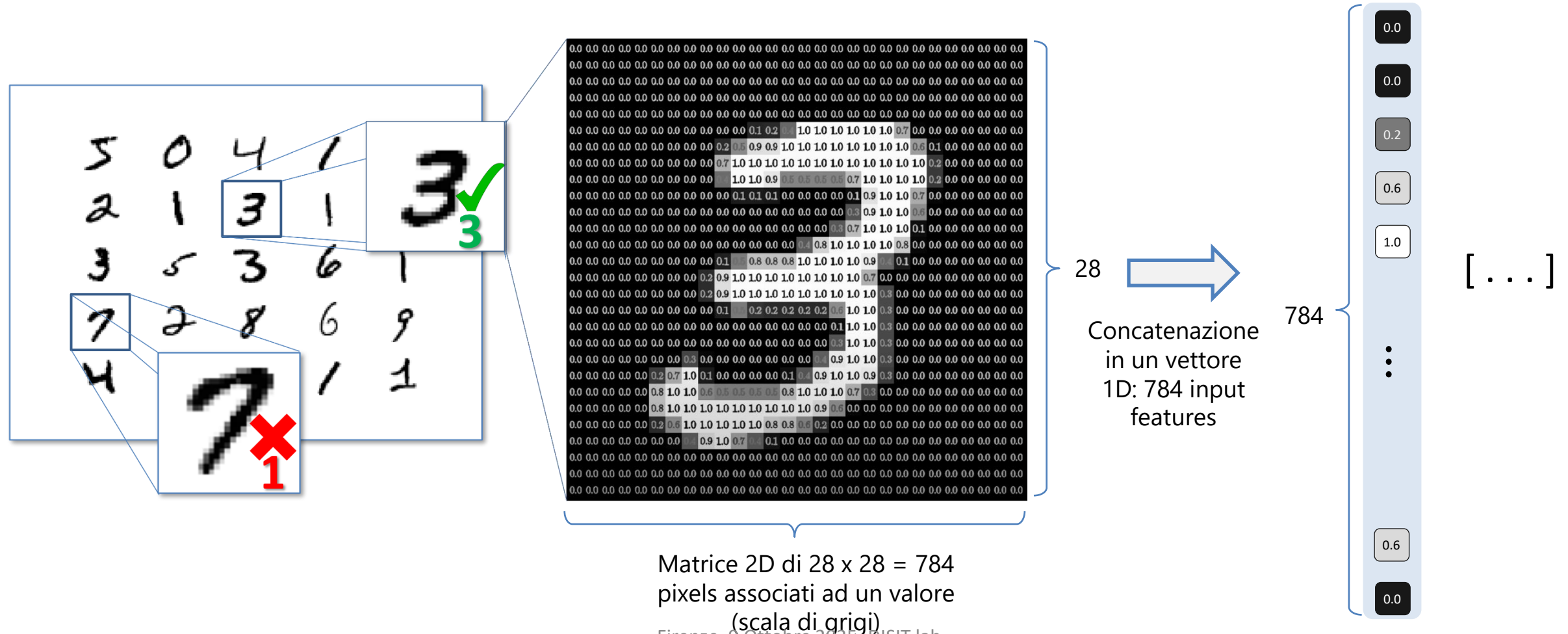
Deep Learning – Modello del Neurone: il Perceptron

- L'unità base di information processing è il *perceptron*. Una rete neurale è formata generalmente da più strati (layers) di neuroni, interconnessi tramite collegamenti caratterizzati da pesi (*synaptic weights*).
- Ciascun neurone riceve in ingresso degli input (x), li combina linearmente (moltiplicandoli per i pesi w e sommandoli a dei termini b detti *bias*) e li propaga attraverso una funzione di attivazione φ .



Esempio Input features: Riconoscimento Cifre Scritte a Mano

➤ Handwritten Characters Recognition



Deep Learning – Domini Applicativi

- **Classificazione / Predizione di serie temporali** – Reti Neurali Ricorrenti (RNN, LSTM)
- **Riconoscimento / Classificazione di immagini / Tracking** – Reti Convoluzionali (CNN)
- **Text-to-Speech / Speech-to-Text**
- **Traduzione multi-lingua** – Modelli Encoder-Decoder, Sequence-to-Sequence (seq2seq), Transformers
- **Comprensione e generazione di testo in linguaggio naturale** – Transformers, Large Language Models (LLM)
- **Sistemi di guida automatizzata / Veicoli autonomi** – Deep Learning + Reinforcement Learning
- **Robotica e Automazione**

Agenda

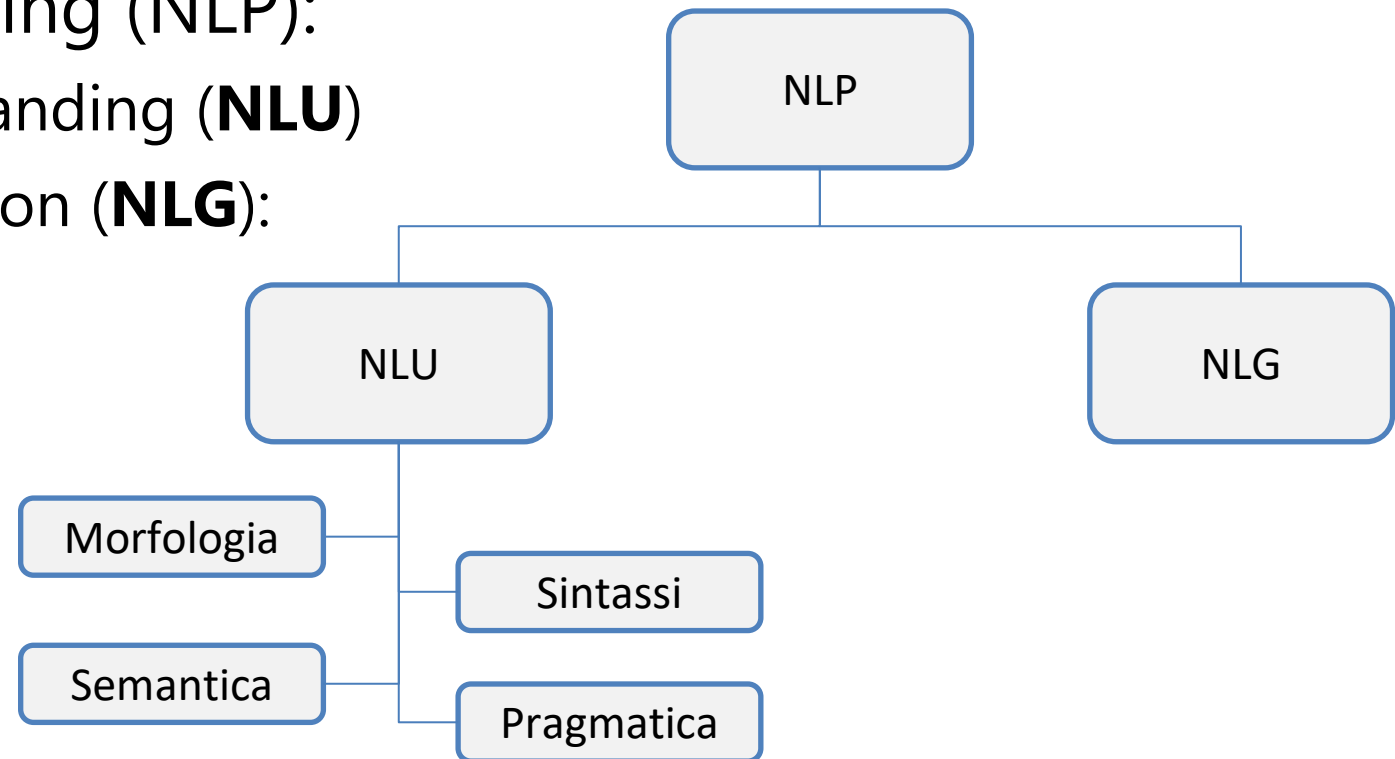
- Ragionamento semantico e simbolico
- Dalla statistica al machine learning e all'intelligenza artificiale
- Apprendimento supervisionato e non supervisionato
- **Dall'elaborazione del linguaggio naturale ai modelli linguistici di grandi dimensioni (LLM)**
- Intelligenza artificiale per previsioni, rilevamento di anomalie, classificazione
- Applicazioni industriali: raccomandazioni
- Esercitazione su LLM
- Test

Come siamo arrivati al contesto attuale dell'AI

- **Per molti anni** l'approccio dell'IA alla comprensione del testo si è basata su modelli di NLP, **Natural Language Processing**
 - Tramite analisi lessico-grammaticali, modelli statistici, etc.
 - Per l'identificazione di pattern (di frasi) con sistemi a regole grammaticali, fondate su conoscenza codificata delle regole del linguaggio.
 - Per esempio, codificate tramite ontologie, tassonomie...
 - Per riconoscere intenzioni, contesti, luoghi, disambiguare, etc.
 - Sono nati i primi Virtual Assistant:
 - *Watson, Siri, Cortana, Alexa, Paval (UNIFI), etc.*
- **Recentemente sono arrivati modelli di apprendimento: Deep Learning, come Recurrent Neural Networks (RNN):**
 - Questi con un numero elevato di strati di Neuroni possono apprendere sequenze complesse di simboli/parole, codificate, etc.

NLP – Prima dell'avvento delle Reti Neurali

- Natural Language Processing (NLP):
 - Natural Language Understanding (**NLU**)
 - Natural Language Generation (**NLG**):



NLP – Prima dell'avvento delle Reti Neurali

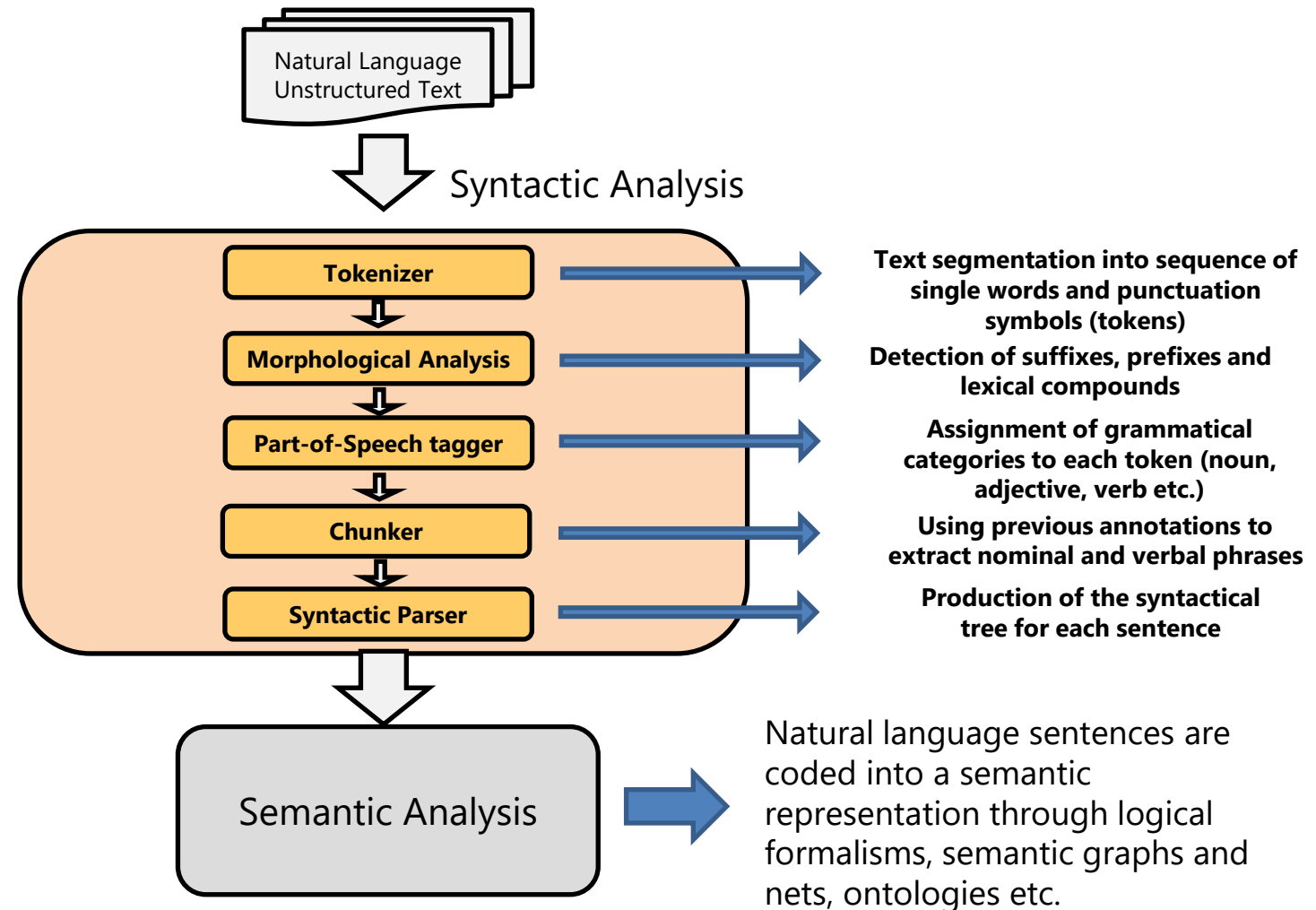
- Principali sfide del Natural Language Processing:
 - IL linguaggio naturale presenta spesso ambiguità in ciascuna delle attività del NLU.
 - 4 livelli di ambiguità:
 - *Lessicale/Morfologica*: « amare » (verbo, aggettivo)
 - *Sintattica*: « Ho visto un uomo con il binocolo »
 - *Semantica*: « acuto » (intelligente, ma anche riferito a forte dolore, e anche in geometria per indicare gli angoli minori di 90°...)
 - *Pragmatica*: « La porta è aperta » (l'intento della comunicazione può riferirsi ad un invito ad entrare, ma potrebbe anche esprimere il desiderio di chiudere la porta perché si è create una fastidiosa corrente d'aria)



Tutte queste ambiguità rendono NLP un insieme di task molto complessi

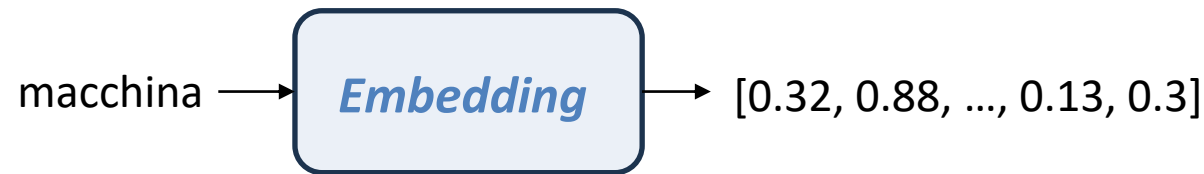


NLP – Prima dell'avvento delle Reti Neurali

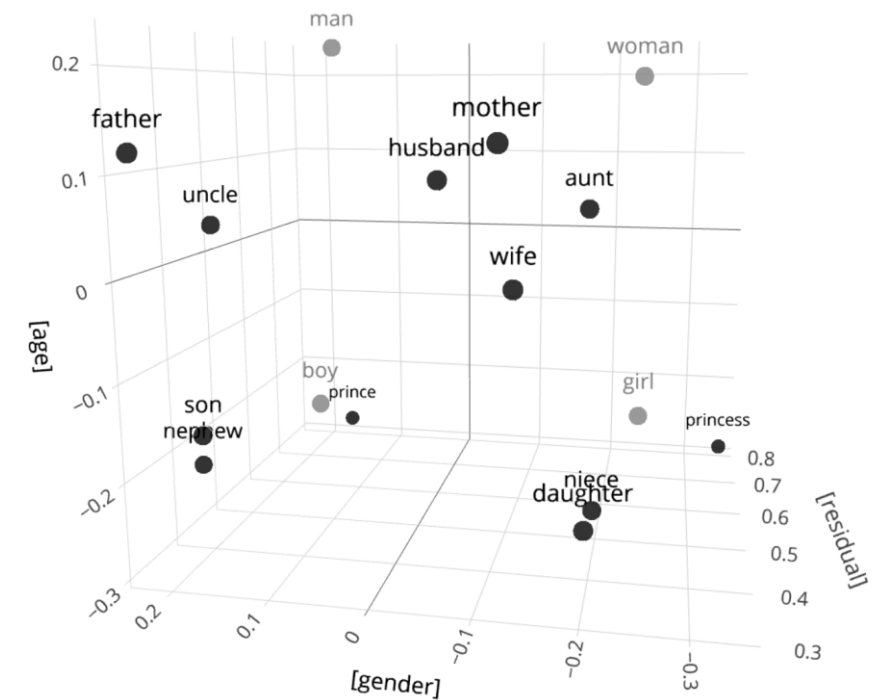


NLP e Deep Learning

- I modelli di linguaggio non comprendono le parole, ma solo sequenze numeriche.
- Per questo motivo, ogni parola in ingresso ad un modello di linguaggio viene tradotta in vettore di numeri: questo processo si chiama «**Word Embedding**».



- Parole con significati simili hanno rappresentazioni numeriche simili (cioè si trovano «vicine» nello spazio di rappresentazione); parole semanticamente molto diverse stanno lontane.



Source: Word Embedding DEMO
<https://www.cs.cmu.edu/~dst/WordEmbeddingDemo/index.html>

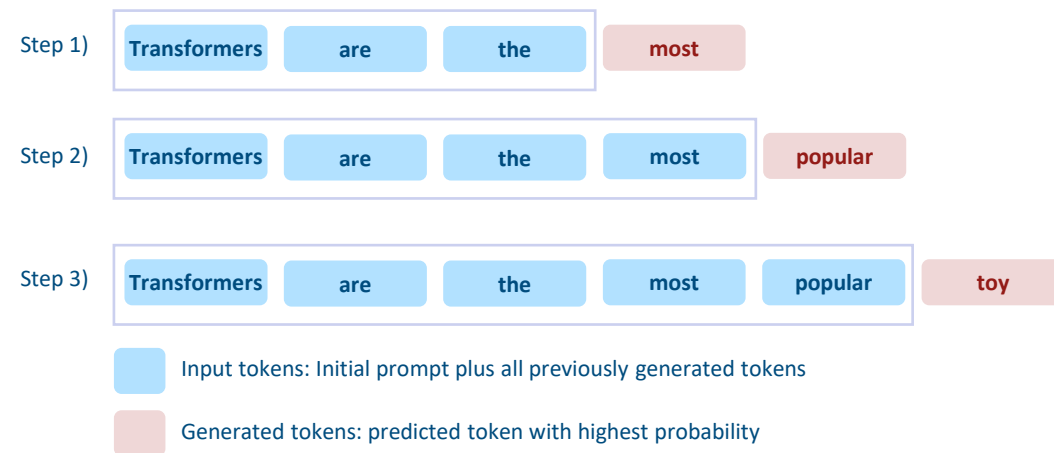
La rivoluzione dei Transformers (2017)

- Hanno sostituito le RNN andando a introdurre il meccanismo di «Attention» per pesare la rilevanza delle parole nelle frasi
- **BERT (Bidirectional Encoder Representations from Transformers)** (2018),
 - BERT: uno dei primi modelli basati su Transformers a ottenere risultati straordinari per i classici compiti di NLP:
 - riconoscimento, classificazione, disambiguazione, etc.
- **GPT (Generative Pretrained Transformer)**: introdotto da OpenAI
 - tra i primi a passare a modelli di auto-regressione per generare testo.
 - **Generazione di risposte in linguaggio naturale**
 - La generazione di risposte, *una sorta di predizione a partire da una sequenza data di parole, il prompt, la richiesta.*

Da BERT ai Large Language Models (LLM)

- I LLM, Large Language Model sono l'evoluzione dei **GPT** (e.g., chat **GPT-4**, **GPT-5**), e si distinguono per la loro scala in termini di:
 - **Numero di parametri**: i LLM hanno miliardi di parametri, per catturare informazioni molto più complesse e dettagliate sul linguaggio.
 - **Dimensione del dataset**: i LLM sono addestrati su dataset enormi, spesso contenenti centinaia di Gbyte-TByte di testo provenienti da varie fonti (libri, articoli, pagine web).
 - **Capacità zero-shot e few-shot**: i LLM, grazie alla loro scala, possono eseguire compiti complessi:
 - Zero-shot: produzione di risposte sulla base della loro conoscenza di base
 - Few-shot: produzione di risposte sulla base di conoscenza aggiuntiva
 - **Alcune architetture di LLM continuano a imparare** ma magari non è sempre cosa buona, anche perché si rischia di virare il modello, e dimenticare...

Modelli di Linguaggio Generativi (GPT)



- L'obiettivo del modello consiste nello stimare il token successivo più probabile in base al prompt e a tutti i token generati in precedenza.
- Dopo aver generato il token successivo, questo viene aggiunto al prompt, estendendo così la sequenza di input per prevedere il token successivo.

- Questo processo iterativo continua fino a quando non viene generato uno speciale token di fine sequenza (`<|endoftext|>`) o quando viene raggiunta una lunghezza massima predefinita.
- Questo processo è chiamato **generazione di testo condizionale**. Ampia gamma di applicazioni: dal completamento di frasi alla generazione di contenuti creativi basati su input specifici dell'utente.

Large Language Models: LLM

➤ Intuizione a partire dal modello **GPT-2**:

- I modelli supervisionati tradizionali si basano su grandi set di dati etichettati, limitandone l'adattabilità
- Relazione tra approcci supervisionati e non supervisionati: i task non supervisionati possono allinearsi quelli supervisionati imparando a prevedere le sequenze di testo.
- I modelli linguistici generativi autoregressivi di grandi dimensioni non supervisionati come GPT-2 (***self-supervised*** con 1,5 miliardi di parametri) possono acquisire caratteristiche linguistiche complesse e generalizzare tra le attività senza la necessità di dati specifici etichettati.
- Training su vasti corpora di testi:
- Addestrato sul set di dati WebText (milioni di pagine Web curate).
- Consente al modello di comprendere task complessi come la traduzione, produrre risposte alle domande da dati di testo non elaborati.

➤ Zero-Shot Learning:

- Il modello apprende le istruzioni dell'attività in modo implicito dai dati di testo naturale.
- Apprendimento Zero-Shot: esegue un'attività senza esempi o addestramento su quell'attività specifica (***fine-tuning***) , basandosi esclusivamente sul contesto o sulle istruzioni fornite durante l'inferenza.
- GPT-2 può eseguire task senza fine-tuning su set di dati specifici.

Large Language Models: LLM

- Il modello **GPT-3** ha ulteriormente sviluppato e migliorato il concetto GPT-2:
- 175 miliardi di parametri.
 - Addestrato su 300 miliardi di token da vari set di dati: Common Crawl (una vasta raccolta di pagine web), WebText2, Wikipedia...
 - Miglioramenti alle attività di apprendimento contestuale, zero-shot, one-shot e few-shot.
 - Genera testi coerenti e contestualmente rilevanti, che riflettono la comprensione di vari domini e stili.
 - Capace di eseguire ragionamenti aritmetici e di utilizzare correttamente il nuovo vocabolario nel contesto, mostrando abilità cognitive avanzate.

Large Language Models: LLM

➤ Novità introdotte con i successivi modelli **GPT-4** e **GPT-5**:

- **Multimodalità**: il modello accetta **testo e immagini** (sempre generando un output testuale)
- Capacità di ragionamento
- Chain-of-Thoughts (CoT) avanzato: invece di fornire una risposta immediata e diretta, il modello risolve un problema complesso scomponendolo in una sequenza di passaggi logici intermedi. Questo processo graduale permette di arrivare a una risposta finale più precisa e spiegabile.
- Significativi progressi su generazione di codice in vari linguaggi di programmazione, Instruction-Following e gestione di contesti molto lunghi (centinaia di migliaia / milioni di token)
- Primi sviluppi di capacità «Agentic» (GPT-5): abilità del modello di pianificare e svolgere sequenze di azioni (eseguire applicazioni e codice, interrogare basi di dati, fare chiamate web, interazione con pagine web ecc.)

➤ **Sviluppi Futuri:**

- Migliorare capacità agentiche
- Multimodalità profonda (testo-immagini-audio-video) in tempo reale

Retrieval Augmented Generation: RAG

➤ Principali limiti e problemi dei LLM:

- Generazione di informazione errata (*Hallucination*)
- Informazioni non sempre aggiornate
- Costoso e complesso addestrarli (*Fine-Tuning*) su basi di conoscenza specifiche (documentazione tecnica, manualistica, documenti legali, contratti ecc.)
- Mancanza di conoscenza in domini specifici

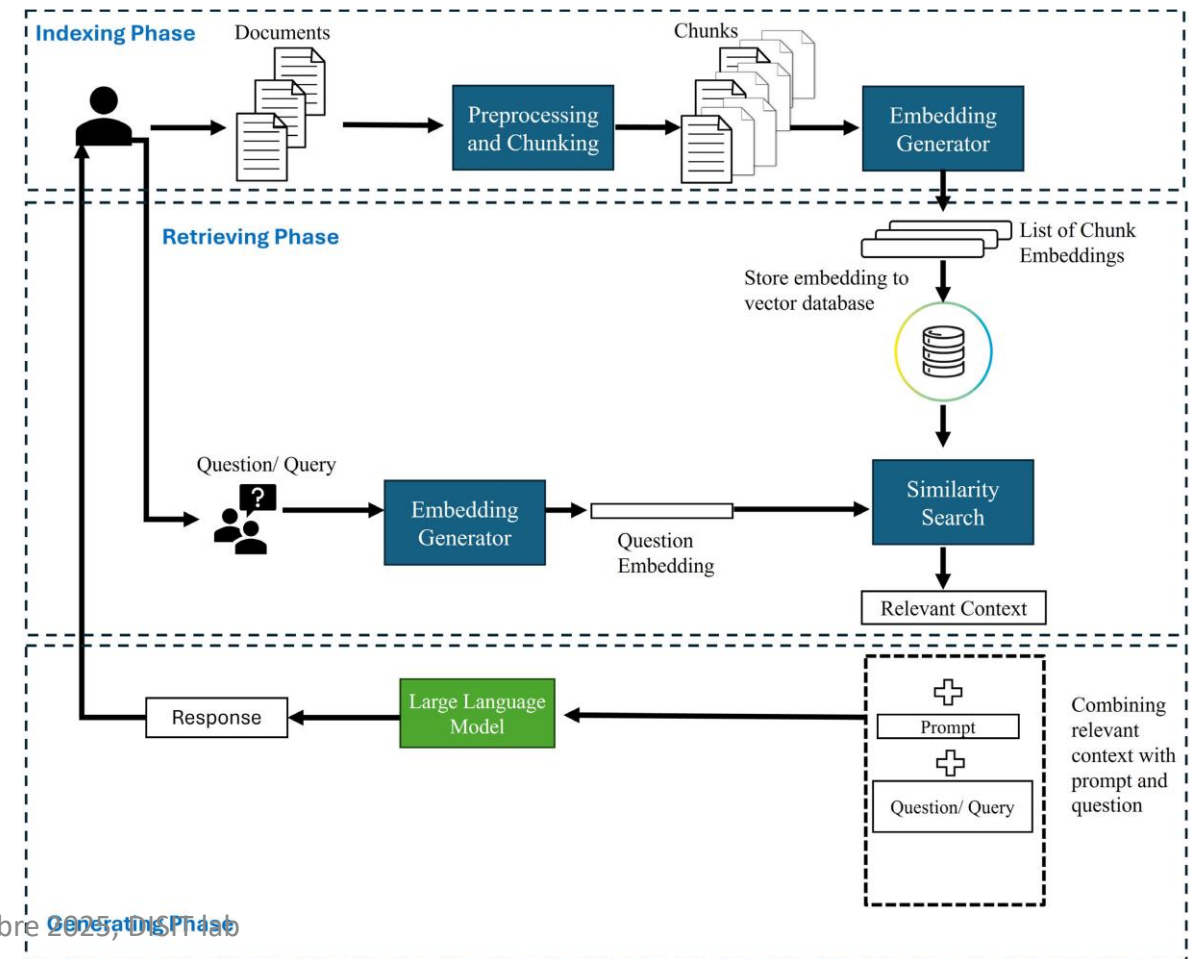
➤ Ulteriori requisiti che applicazioni di AI devono soddisfare in ambiti specifici:

- Produrre risposte accurate e complete su basi di conoscenza specifiche
- Gestire aggiornamenti anche frequenti delle basi di conoscenza
- Spiegabilità (Explainability) e tracciabilità delle risposte
- Riduzione dei costi di Fine-Tuning

Retrieval Augmented Generation: RAG

➤ Approccio RAG:

- Una base di conoscenza specifica viene suddivisa in frammenti (*chunks*), in genere consistenti dal punto di vista semantico (*semantic chunking*);
- Per ogni chunk vengono generati i corrispondenti embeddings e indicizzati in un *Vector Database*;
- **Retrieval**: viene generato l'embedding della domanda dell'utente (*query*) e viene fatta una ricerca di similarità nel Vector DB per recuperare i chunks più rilevanti rispetto alla domanda iniziale;
- **Augmented Generation**: I chunks più rilevanti vengono inclusi come contesto insieme alla query iniziale (*augmented prompt*) ed inviati al LLM, che in questo modo è in grado di produrre una risposta più accurata e centrata sulle parti rilevanti della base di conoscenza specifica.



Agenda

- Ragionamento semantico e simbolico
- Dalla statistica al machine learning e all'intelligenza artificiale
- Apprendimento supervisionato e non supervisionato
- Dall'elaborazione del linguaggio naturale ai modelli linguistici di grandi dimensioni (LLM)
- **Intelligenza artificiale per previsioni, rilevamento di anomalie, classificazione**
- Applicazioni industriali: raccomandazioni
- Esercitazione su LLM
- Test

Predictions

- **Computing predictions**

- **Why?**
- **They can be always computed?**
 - Time series, time trends, seasonality, etc.
- **Which data are needed?**
- **Precision needed and precision which can be obtained?**
- **Computational costs?**



- **Technically:**

- **Time range**, in most cases they are defined such as:

• Short:	5-15 Minutes;	Mid:	30-45 minutes;
• Long:	1 day, week;	very long:	weeks / months / years
- **Computational Model needed ?**

Management

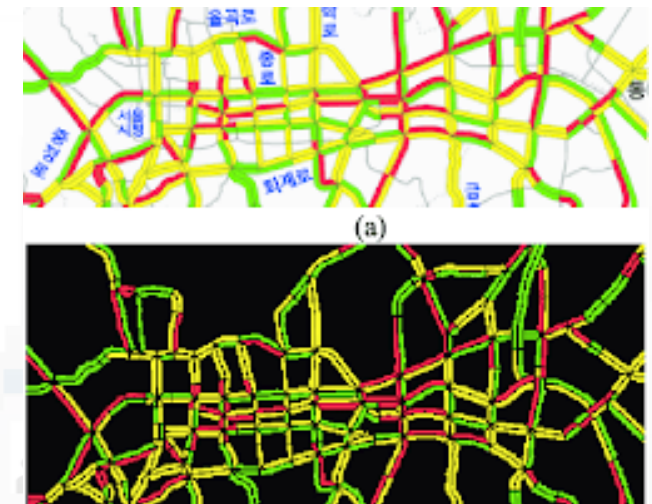
Tactics/strategy

Why Computing Predictions

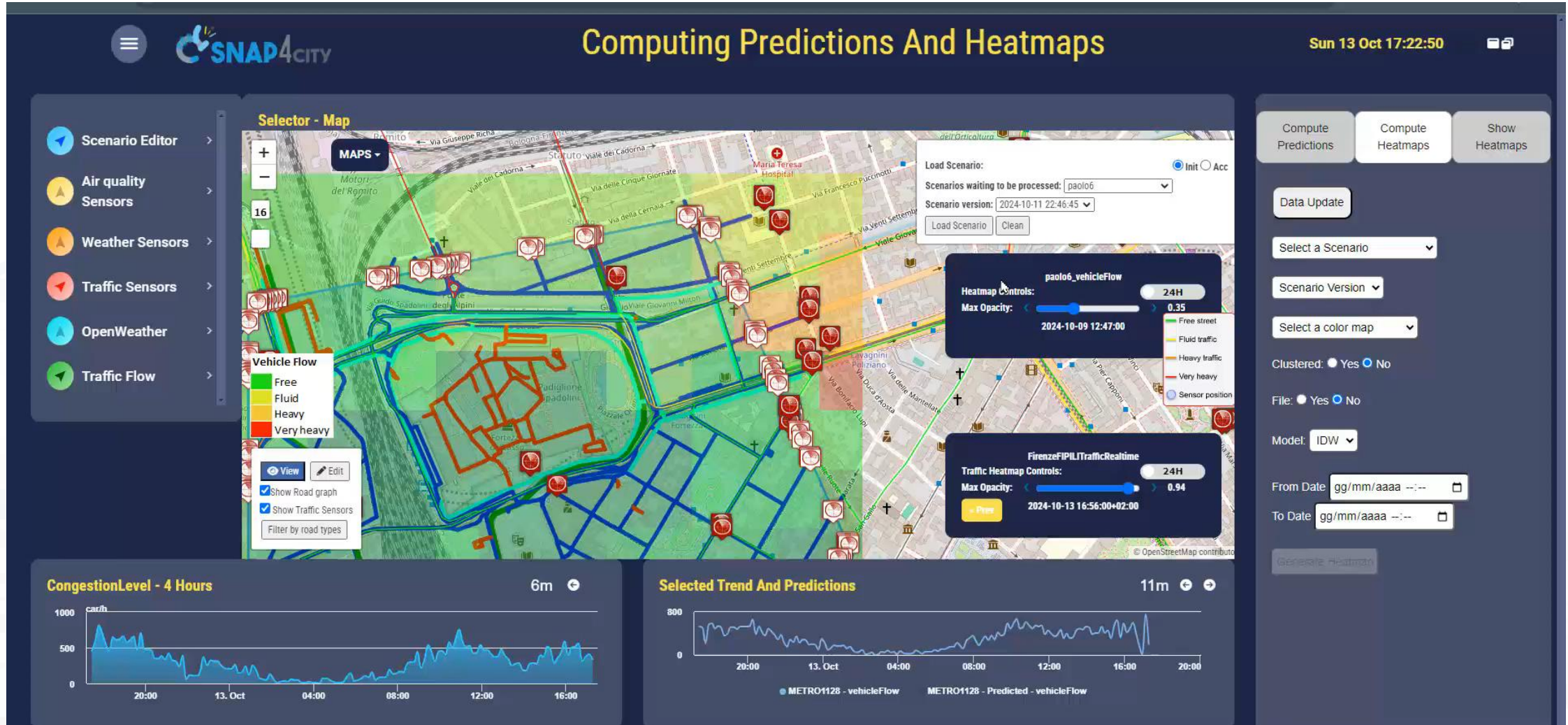
- if I know how many people will attend an event
 - I can **detect anomalies** earlier if an unexpected event will occur, intervene
 - I can **organize** better services, cleaning and preventive security
 - I can **inform, mitigate, plan, save money and time**, etc.
- Other Cases:
 - **Traffic** → pollutant, luminaries, city plan, be prepared critical conditions
 - **Parking** → inform in advance the users, save money and time,
 - **Energy** → be prepared for critical conditions
 - **Pollutant** → to avoid taking taxes, planning trips, etc.
 - **Waste** → save money and time,

Predictions

- **For Cases:**
 - Free parking slots
 - Free bikes, and free slots on bike racks
 - Pollutant: NOX, NO2, CO2
 - Land Slide
 - People behavior
 - Energy consumption
 - Waste production for collection
 - Etc.
- **→ Anomaly Detections**



Predictions and Heatmaps in Real Time



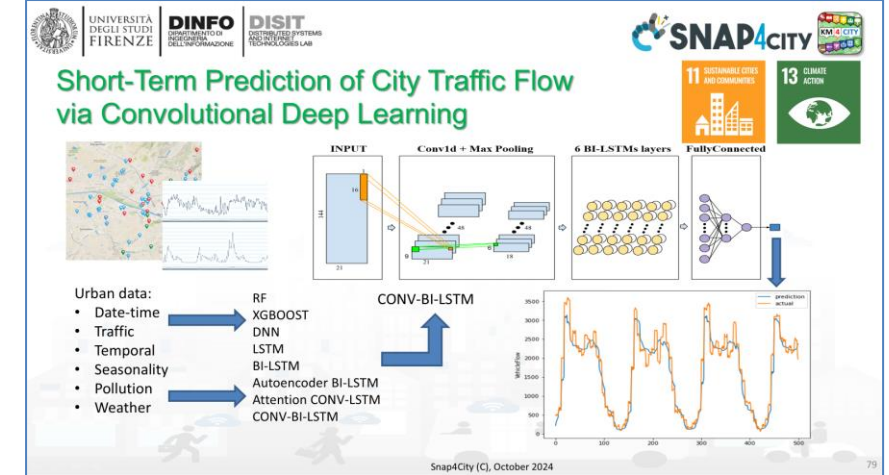
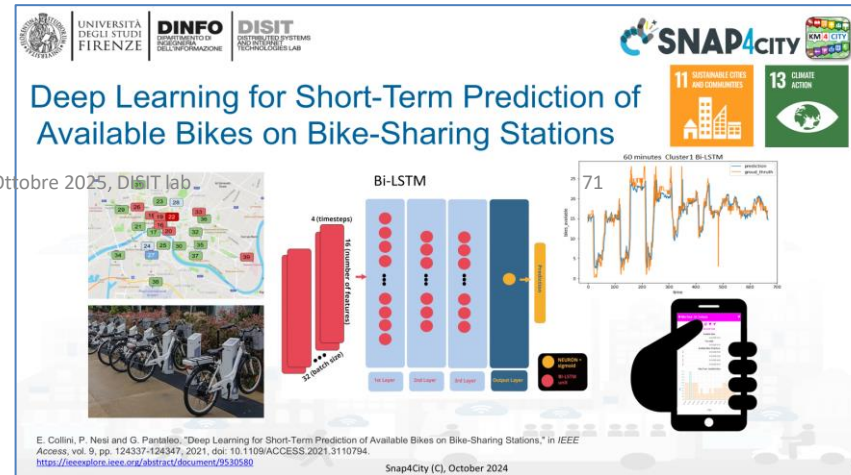
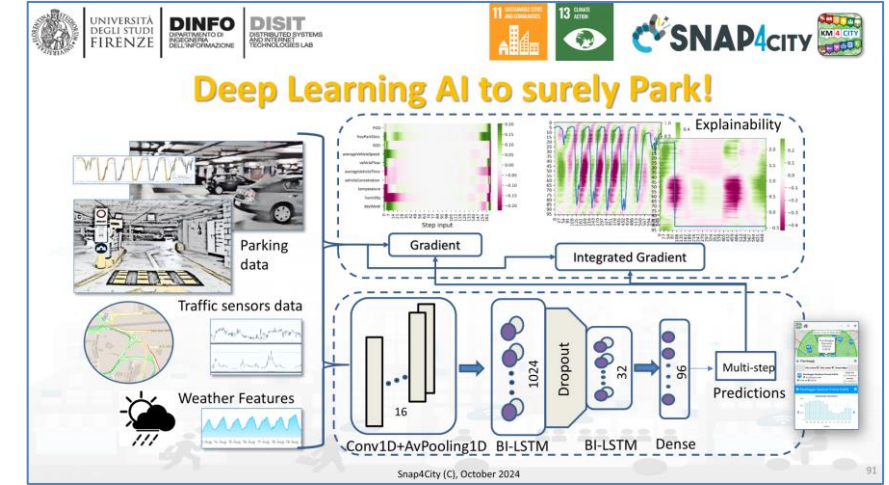
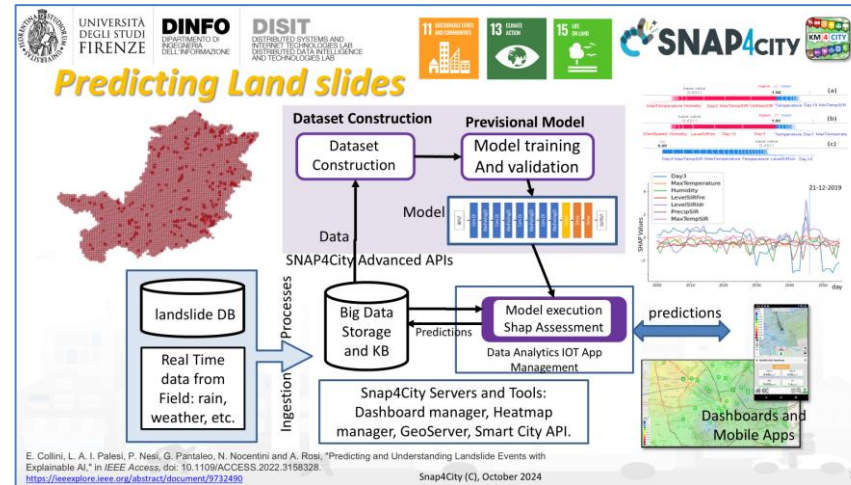
Predictions onselected.....

Aim in operation:

- Short Term predictions
- Reduction of costs, prevention
- Early warning of critical conditions
- Reduction of emissions

AI inject:

- Deep Learning
- Data integration
- XAI addition for suggestions



Smart Parking: *#free slots predictions*

11 SUSTAINABLE CITIES
AND COMMUNITIES



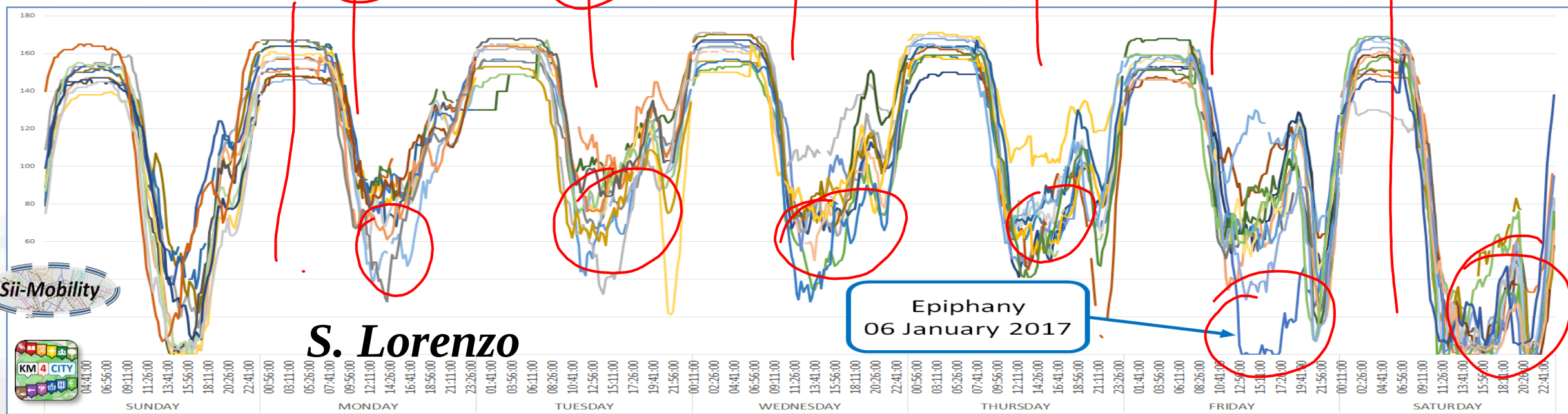
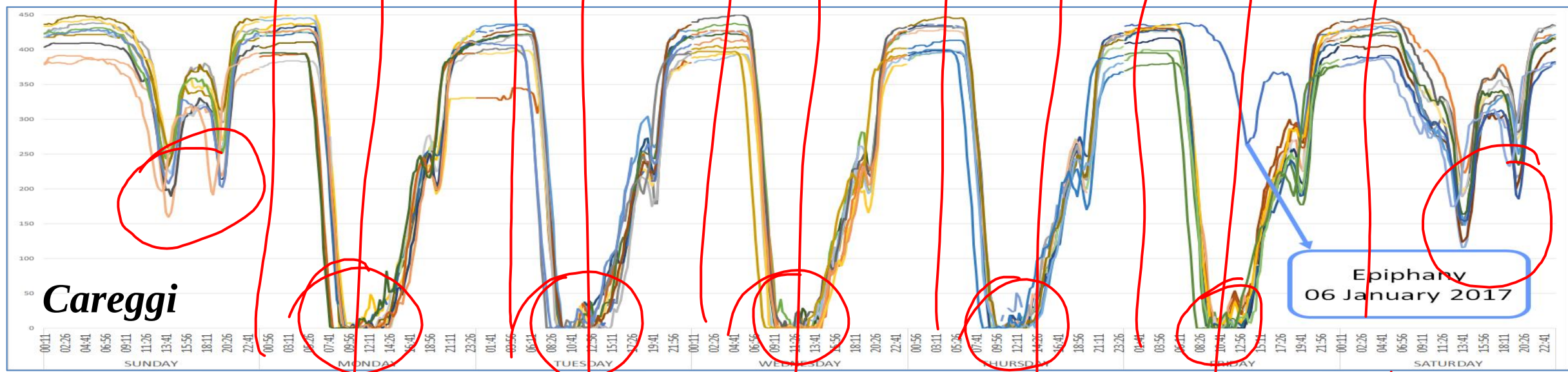
13 CLIMATE
ACTION



Data Analytic



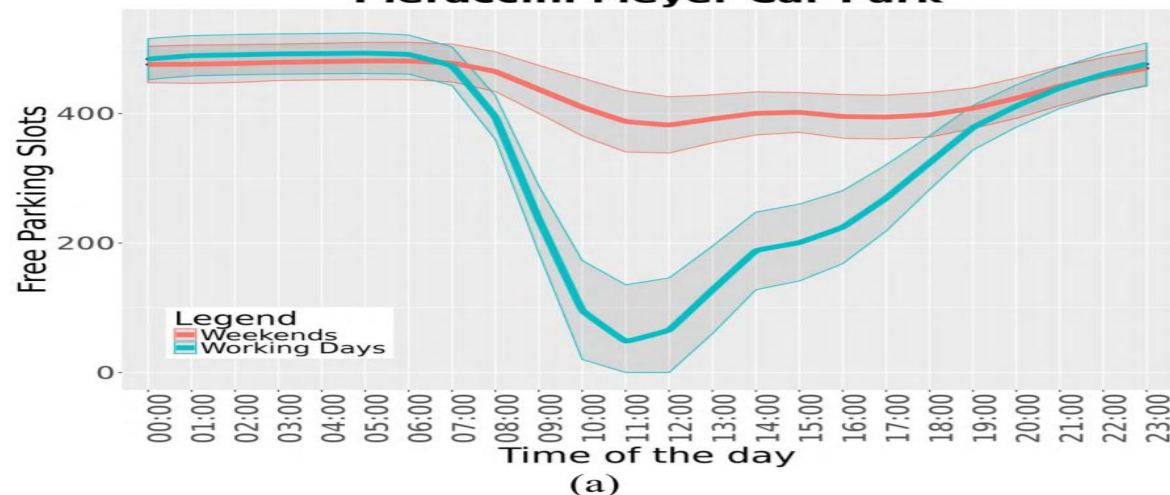
Free Parking space trends



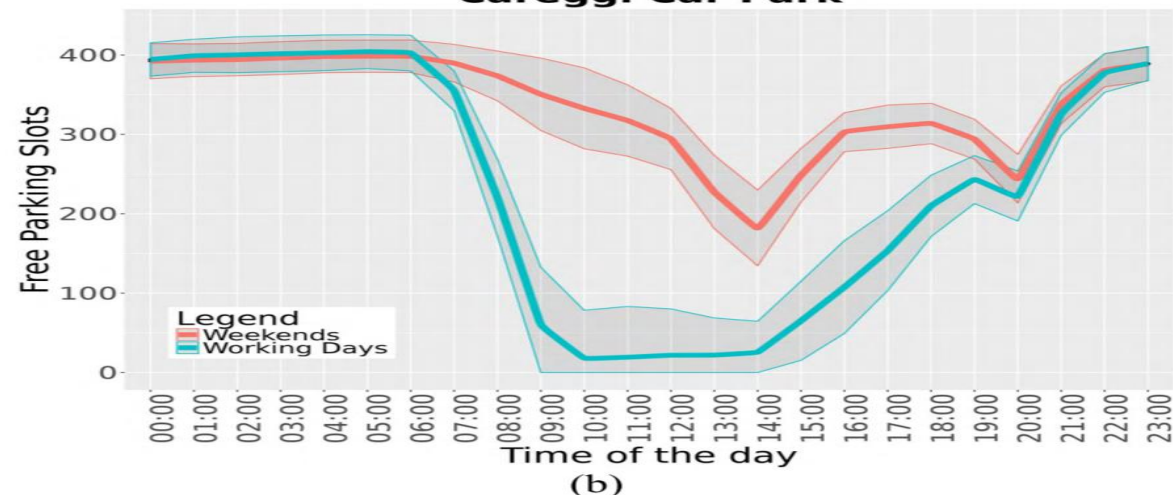
Free Parking space trends



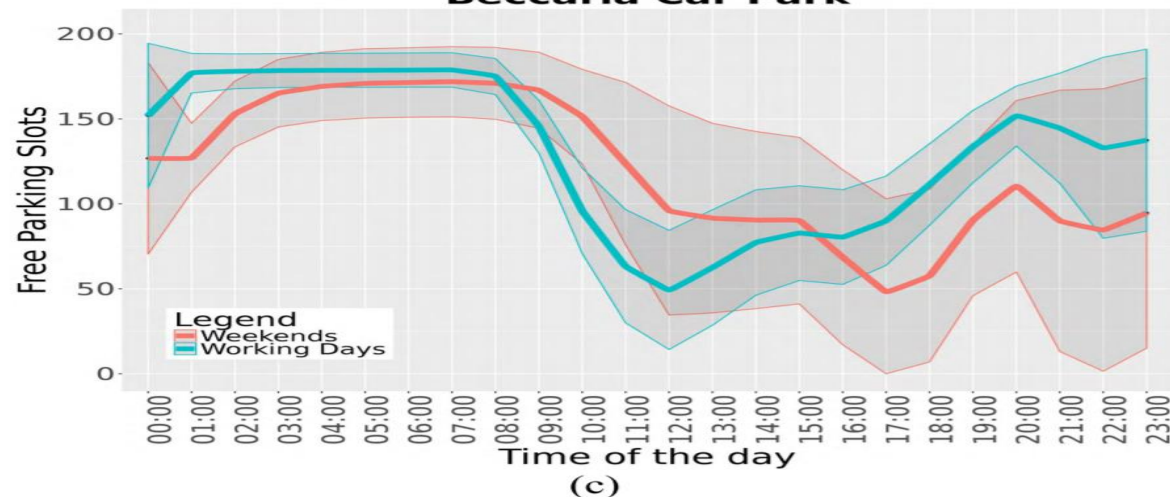
Pieraccini Meyer Car Park



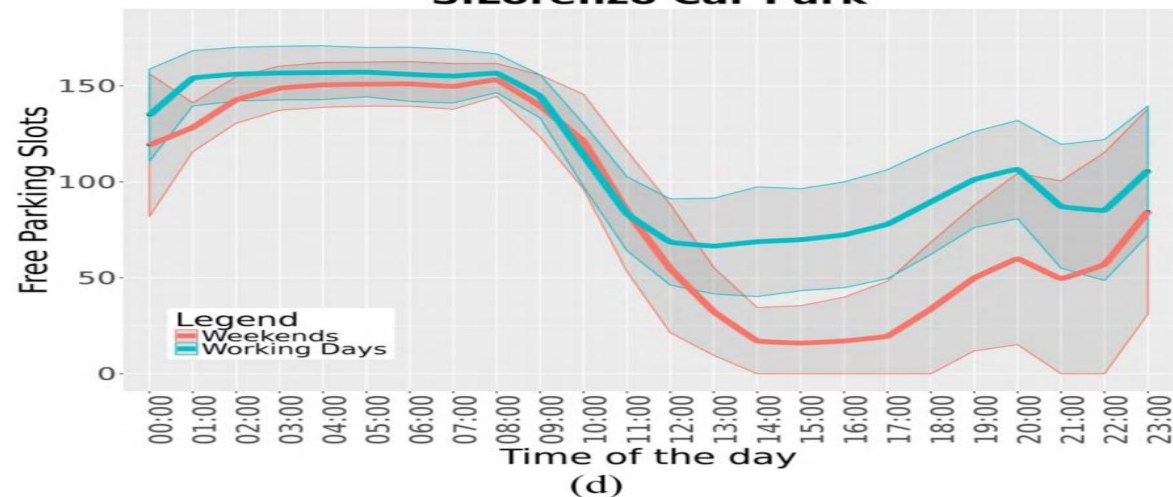
Careggi Car Park



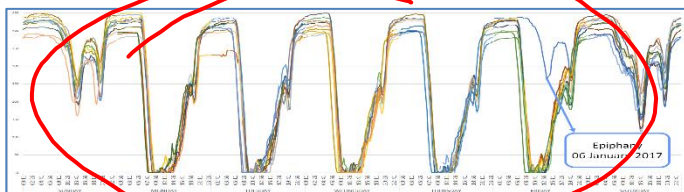
Beccaria Car Park



S.Lorenzo Car Park



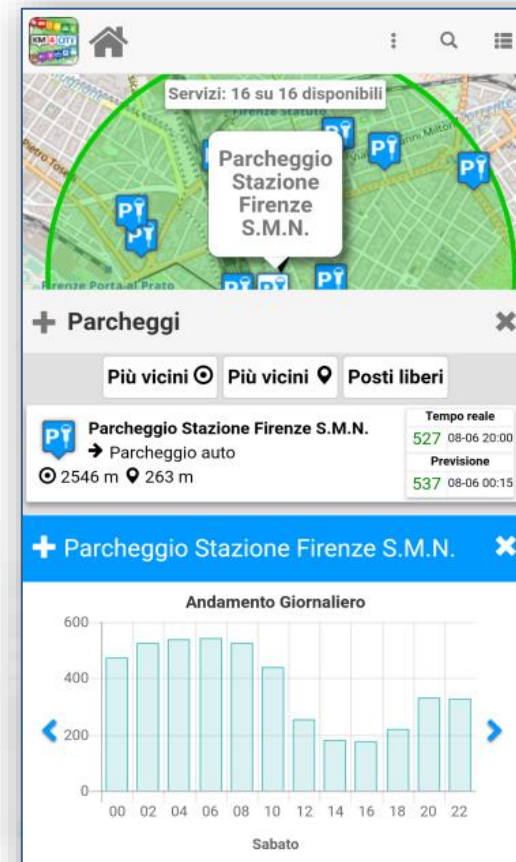
I would arrive to surely Park in 45 Minutes??



Cat ory	Features	Description of features variable
Baseline features of free slot data	Free parking slots	Real number of available slots recorded every 15 minutes
	Time	Hours and minutes
	Month	Month of the year (1-12)
	Day	Day of the month (1-31)
	Day week	Day of the week (0-6)
	Weekend	0 for working days, 1 else
	Previous observation's difference (POD)	Difference between the number of free spaces at time i and number of free spaces at time $(i - 15 \text{ minutes})$ recorded in the previous week
Weather features	Subsequent observation's difference (SOD)	Difference between the number of free spaces at time i , and the number of free spaces at time $(i + 15 \text{ minutes})$ recorded in the previous week
	Temperature	City temperature measured one hour earlier than Time ($^{\circ}\text{C}$)
	Humidity	City humidity measured one hour earlier than Time (%)
Traffic Sensors features	Rainfall	City rainfall measured one hour earlier than Time (mm)
	Average Vehicle Speed	Average speed of vehicles on the road being closest to the parking, over one-hour period (km/h)
	Vehicle Flow	Number of vehicles passing by closest to the parking, over one-hour period
	Average Vehicle Time	Average of distance between vehicles, over one-hour period
	Vehicle Concentration	Number of vehicles per kilometer, over one-hour period

Artificial Intelligence
Predictions

97% of precision



11 SUSTAINABLE CITIES
AND COMMUNITIES



13 CLIMATE
ACTION

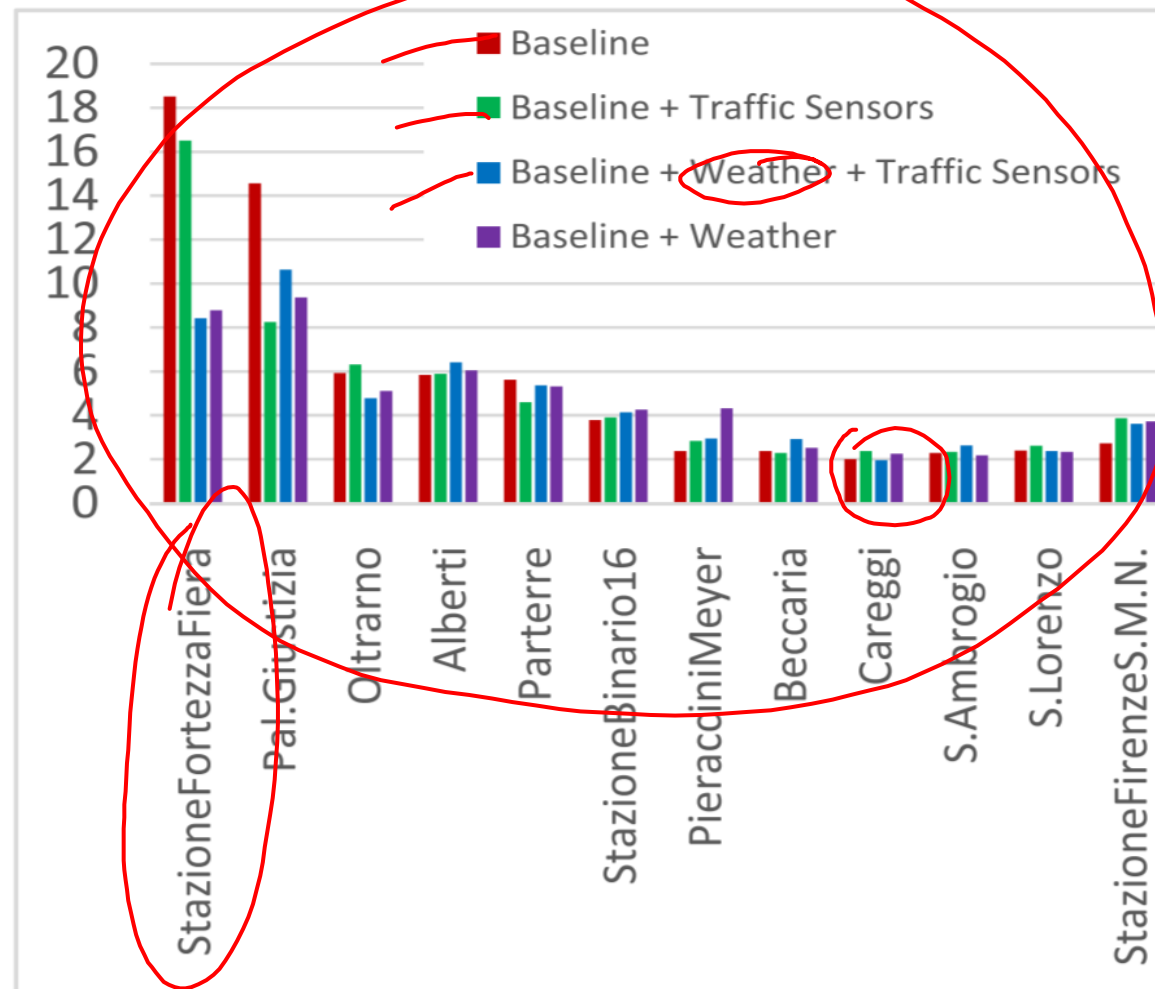


Free Parking PREDICTIONS



C. Badii, P. Nesi, I. Paoli, "Predicting available parking slots on critical and regular services exploiting a range of open data", IEEE Access, preprint, 2018, <https://ieeexplore.ieee.org/abstract/document/8430514/>

Comparison Error	Forecasting Techniques		
	BRANN	SVR	RNN
Careggi car park			
MASE Night	34.85	16.29	20.01
MASE Morning	0.76	1.42	2.82
MASE Afternoon	1.89	4.34	3.66
MASE Evening	1.99	1.51	2.33
MASE	1.87	2.34	3.16
Pieraccini Meyer car park			
MASE Night	6.08	12.83	10.03
MASE Morning	0.86	1.27	4.90
MASE Afternoon	1.87	2.91	6.75
MASE Evening	1.36	1.57	10.23
MASE	1.37	2.06	6.67
S. Lorenzo car park			
MASE Night	10.33	11.81	18.34
MASE Morning	2.13	1.91	3.93
MASE Afternoon	2.70	3.15	2.37
MASE Evening	2.15	3.09	3.82
MASE	2.72	3.21	4.19
Beccaria car park			
MASE Night	9.32	7.80	12.47
MASE Morning	0.95	1.25	4.87
MASE Afternoon	2.49	2.14	2.45
MASE Evening	2.96	4.75	5.91
MASE	2.13	2.67	4.85



The best selected models for the purpose have been:

– BRNN/BRANN:

- Bayesian Regularized Artificial Neural Network

– SVR:

- Support Vector Regression

– ARIMA

- Autoregressive Integrated Moving Average

– RNN

- Recurrent neural networks



Free Parking Predictions



Careggi car park

Model features

BRNN model results

R-squared

RMSE

MASE

Baseline

0.974

24

1.87

Baseline + Weather

0.975

24

1.75

Baseline + Traffic sensors

0.975

24

2.04

Baseline + Weather + Traffic sensors

0.975

24

1.87

Active on Mobile Apps as:

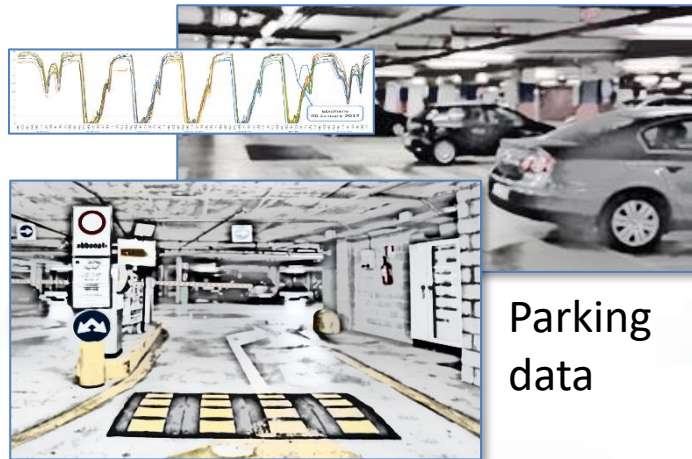
- «Firenze dove cosa»
- «Toscana dove cosa»

Best compromise

Precision: 97,5%



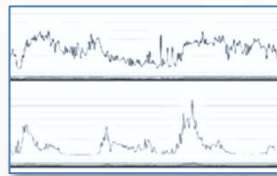
Deep Learning AI to surely Park!



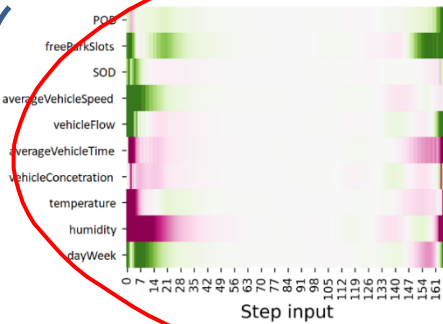
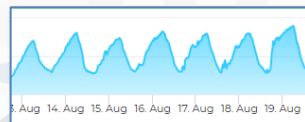
Parking data



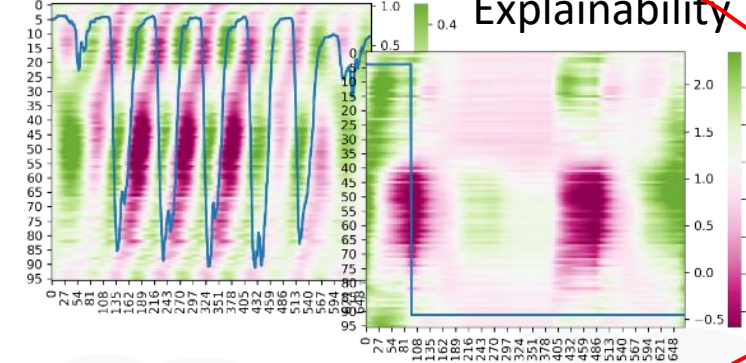
Traffic sensors data



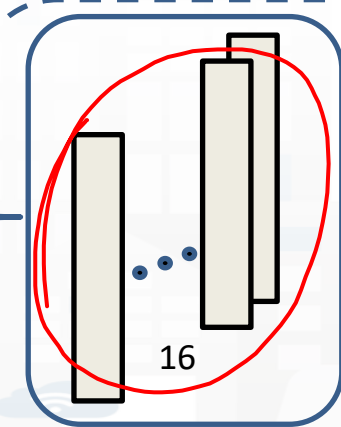
Weather Features



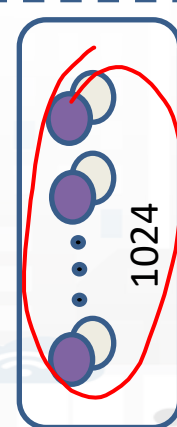
Gradient



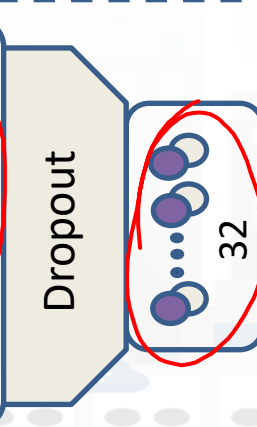
Integrated Gradient



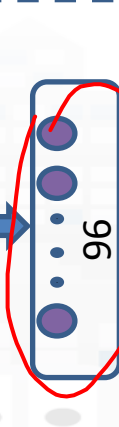
Conv1D+AvPooling1D



BI-LSTM



BI-LSTM

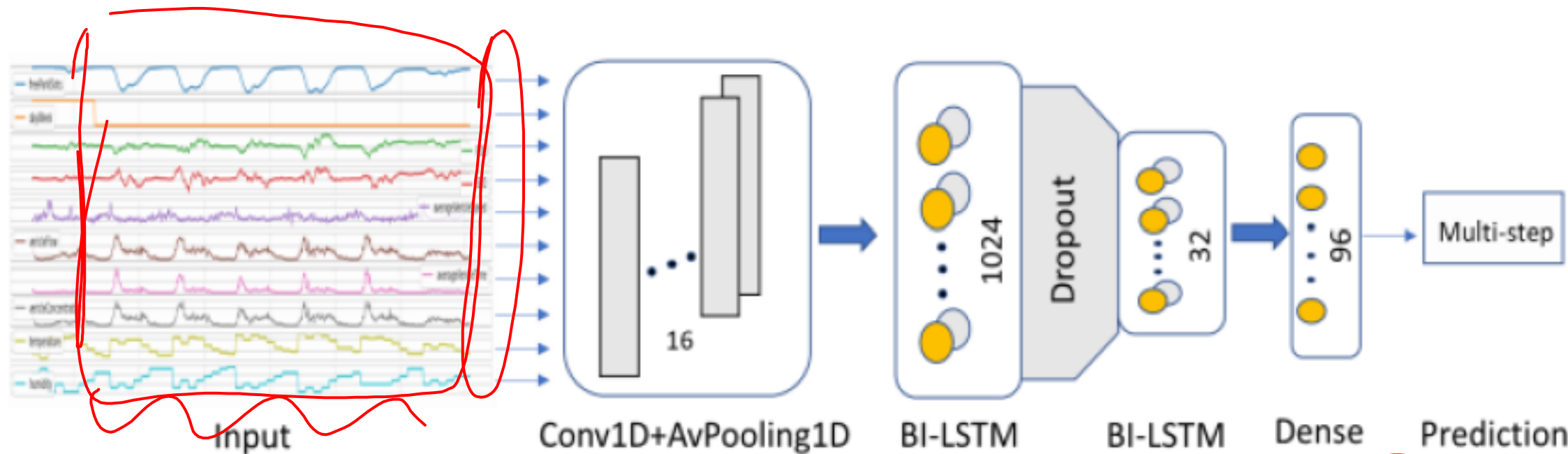


Dense

Multi-step
Predictions
95%



CNN-BI-LSTM model architecture



1-week data
observations sampled each hour
(168 samples / timestamps)

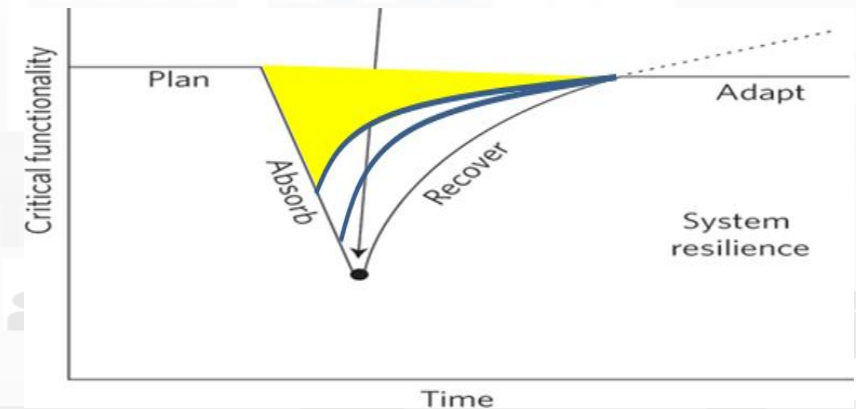
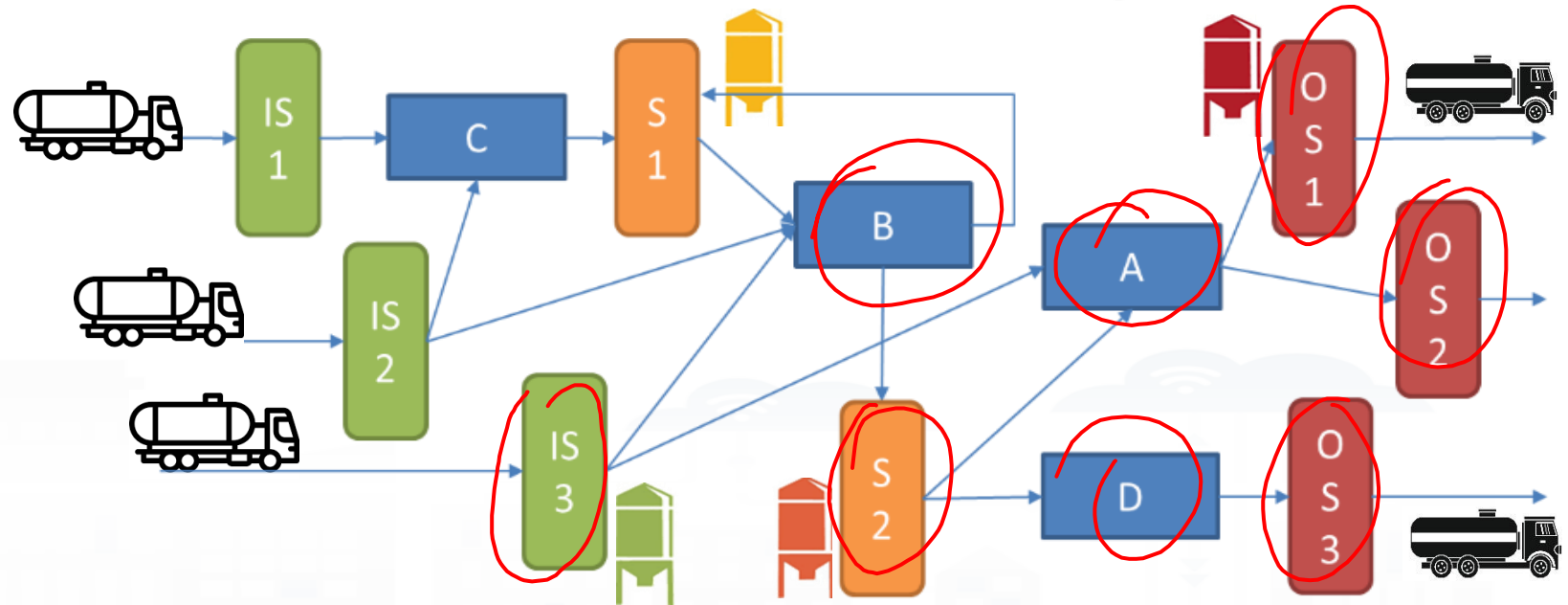
96 timestamps per
day (24 hours and
samples every 15
minutes)

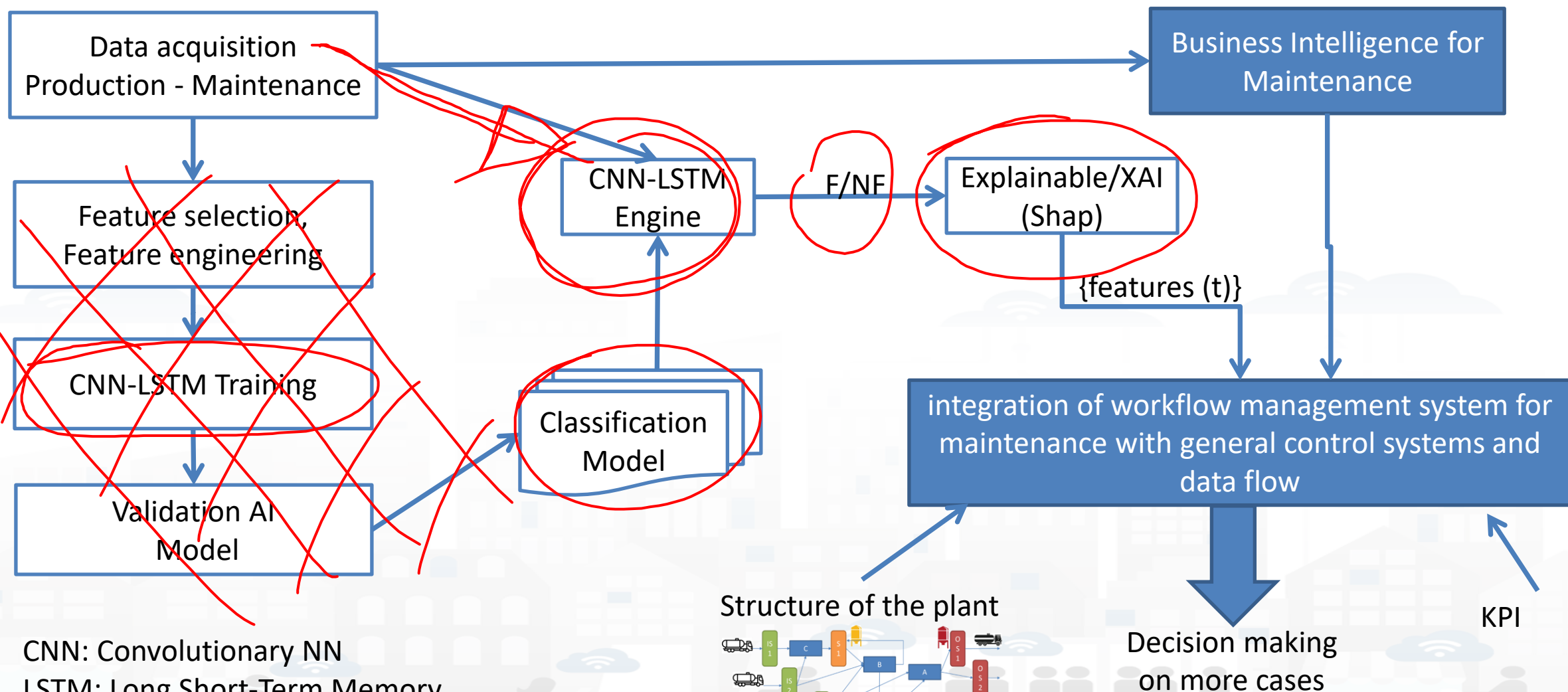
Agenda

- Ragionamento semantico e simbolico
- Dalla statistica al machine learning e all'intelligenza artificiale
- Apprendimento supervisionato e non supervisionato
- Dall'elaborazione del linguaggio naturale ai modelli linguistici di grandi dimensioni (LLM)
- **Intelligenza artificiale per previsioni, rilevamento di anomalie, classificazione**
- Applicazioni industriali: raccomandazioni
- Esercitazione su LLM
- Test

Complex cause-effect relationships

- **Elements:**
 - Machines: A...C
 - Storage: silos...
 - Flows:...
- **Dependencies**
 - Cascade effects
- **Early warning**
 - Reduction of costs
 - Recovering from failure is more expensive than correcting in advance
 - Possible advanced replan and reschedule: secondary solutions





CNN: Convolutionary NN
LSTM: Long Short-Term Memory
(deep learning model)

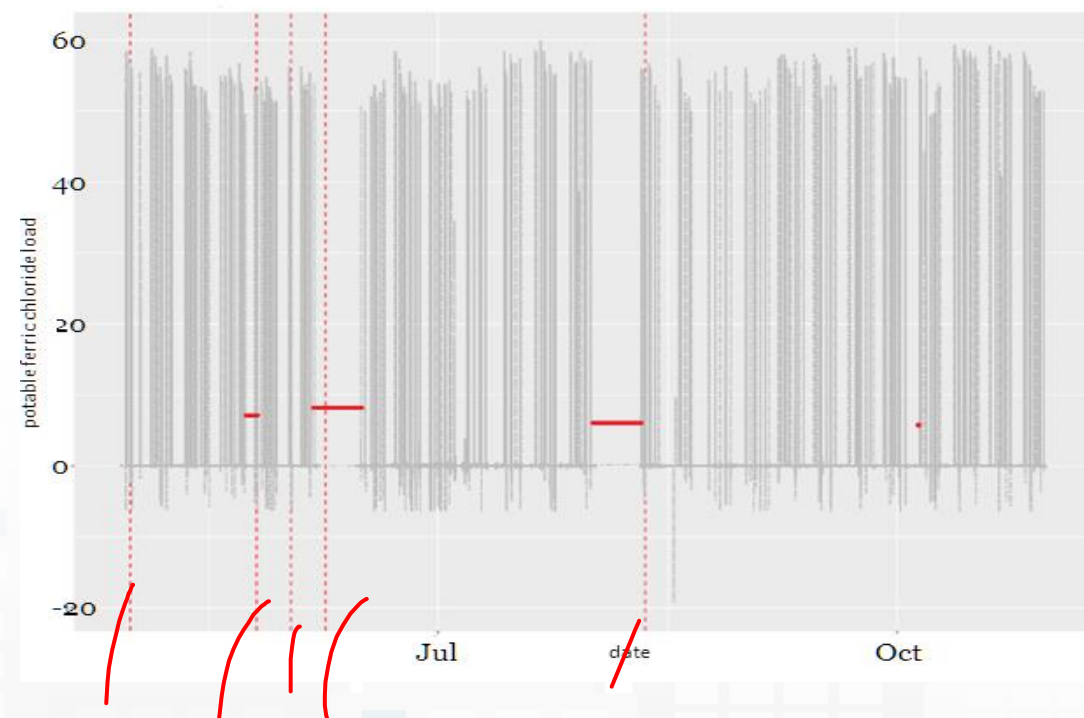
Production:

- 1-minute observation from 2020-04-28 to 2021-01-04
- 343.183 observations for 147 features/variables
- production, storage, status, several temperatures of elements, gear plants, process/safety parameters, chemicals compounds produced

Fault:

- List all the details: event datetime, Permission List, Plant, Signature, Specialty, Status, Job Type, Air Temperature, air humidity and rain
- Ticket and stop classification as "GENERAL PLANT STOP", "ORDINARY", "PLANT STOP" and "EMERGENCY "

Example of a failure



Overview Features

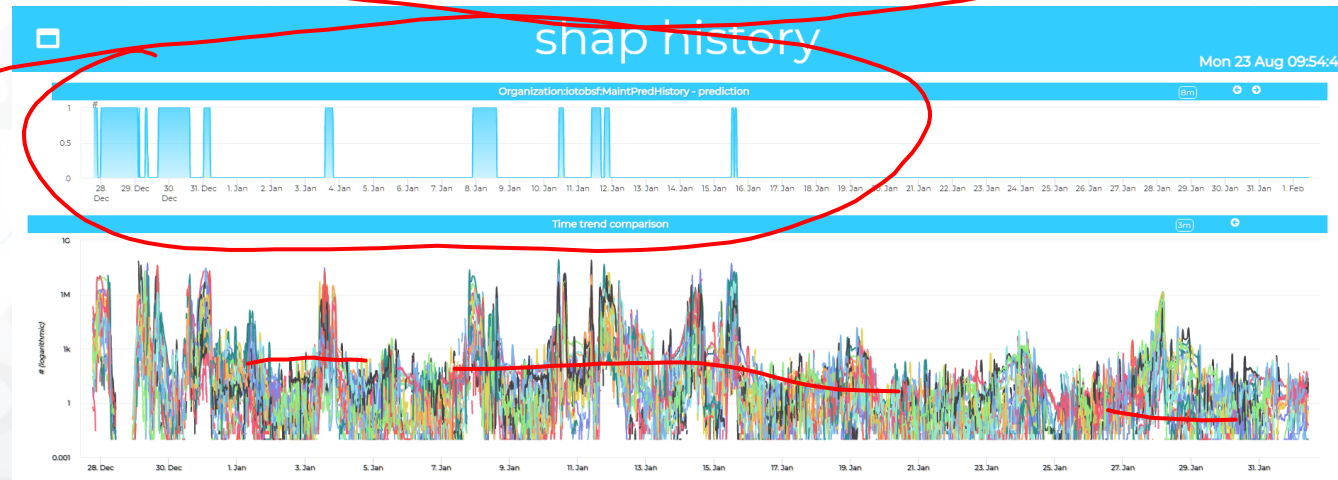
Feature	Plant	Description	Unit of measure
TempreactoreR4001 - TempreactoreR4002 - TempreactorR4003	chlorine paraffins (CPS)	reactor temperature indication	°C
S904A - S904B - S904C	Potable Ferric std	Storage level indication	%
S4304	chlorine paraffins (CPS)	Storage level indication	%
standardFerric Chloride	Potable Ferric std	flow rate measurement and totalization	m3
potFerricChloride	Potable Ferric Chloride	flow rate measurement and totalization	m3
S904E - S904D	Potable Ferric Chloride	Storage level indication	%
QuantNaOHperBatchNaClO - QuantNaOHBatchNaClO_2	NaOH KOH	flow rate measure and totalization	lt – m3
ConversionNaOH - ConversionKOHlinea1	NaOH KOH	electrolysis load adjustment (production)	kA
KOH_1_charge - KOH_2_charge	NaOH KOH	flow rate measure and totalization	m3
S487 - S484 - S5104	NaOH KOH	Storage level indication	%
hypo sodium	sodium hypochlorite	quantity of material produced	m3
S851 - S852 - S854 - S856 - S857	sodium hypochlorite	Storage level indication	%
S871	HCl	Storage level indication	%
RedoxFeCl3Pot	Ferric Chloride std	potential measure redox Ferric Chloride	mV

Explainable/XAI - CNN-LSTM (SHAP)

Explanation of prediction generated by model for fault



Explanation of prediction generated by model for normality



Agenda

- Ragionamento semantico e simbolico
- Dalla statistica al machine learning e all'intelligenza artificiale
- Apprendimento supervisionato e non supervisionato
- Dall'elaborazione del linguaggio naturale ai modelli linguistici di grandi dimensioni (LLM)
- Intelligenza artificiale per previsioni, rilevamento di anomalie, classificazione
- **Applicazioni industriali: raccomandazioni**
- Esercitazione su LLM
- Test

- **Feedback Project:**

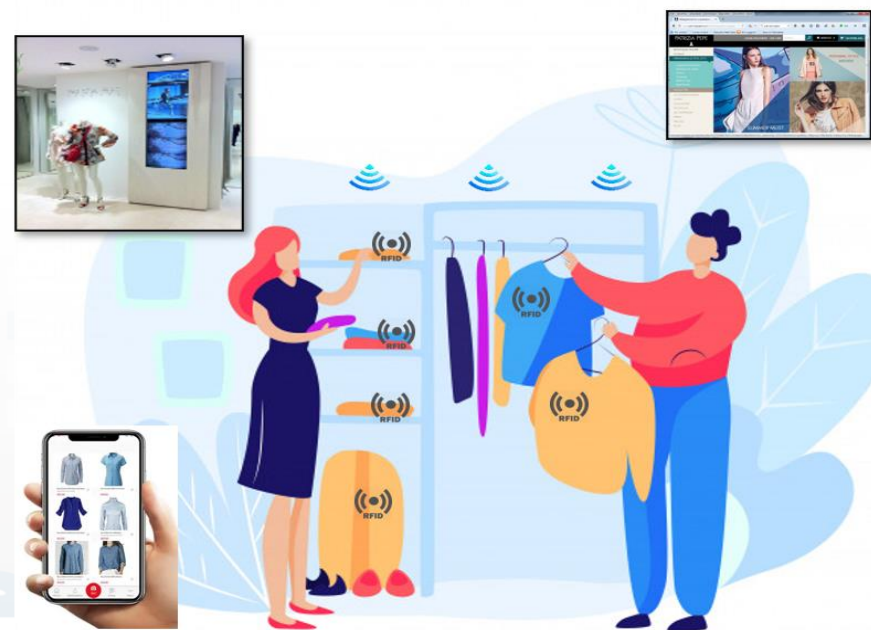
- Flexible Advanced Engagement Exploiting User Profiles and Product/Production Knowledge
- VAR, PatriziaPepe (Tessilform), DISIT Effective Knowledge, SICE
- Keywords: retail, GDO, ...

- **Goals and drivers:**

- adaptive user engagement, customer experience
- Advanced user profiling, user behaviour analysis
- IOT and instrumentation
- Predictive models for engagement
- Integrated in city customer experience

- **Aiming to solve current State of the Art issues:**

- Cold start problems in generating recommendations for new users, also addressing seasonality of products and items
- GDPR compliance



FeedBack Admin Tools

FeedBack Engagement Tools

Intranet

Tool Admin

Tool Engager

Recommender

Totem

Surfaces

Web site - App

Checkout

Sensor
Manager

Network
Manager



Sensors and Markers

Apps and Tags

feedback

Suggestions

customer similarity for each customer cluster the most representative items are suggested;

item similarity: considering the last items purchased by the customer according to the information contained into its profile, and randomly selecting items in the same item clusters;

item complementary: considering items that may complement the last items that have been bought by the customer according to a table of complementary items;

item associated: in order to improve a customer's purchase frequency, we generated suggestions for customers who purchased an item in the last three months;

suggestions for serendipity: randomly selecting items to be suggested from the whole present collection, taking also into account what is available in the physical shop;

Item Cluster	Complementary Item Clusters				
	cluster	support	confidence	lift	count
1	2	0.26486066	0.6069351	1.106003	12935
	7	0.24864345	0.5697729	1.253423	12143
	3	0.24465057	0.5606231	1.213722	11948
	8	0.24336057	0.5576670	1.277549	11885
	4	0.22298667	0.5109797	1.282096	10890
2	3	0.34351004	0.6259701	1.355196	16776
	7	0.32391425	0.5902612	1.298495	15819
	8	0.31392182	0.5720522	1.310504	15331
	4	0.29840080	0.5437687	1.364367	214573
3	2	0.34351004	0.7436830	1.355196	16776
	7	0.30397035	0.6580814	1.447690	14845
	8	0.29868747	0.6466442	1.481385	14587
	4	0.27753548	0.6008511	1.507592	13554
	1	0.24465057	0.5296569	1.213722	11948
4	2	0.29840080	0.7487156	1.364367	214573
	3	0.27753548	0.6963625	1.507592	13554
	7	0.26578209	0.6668722	1.467029	12980
	8	0.27260069	0.6839807	1.566918	13313
	1	0.22298667	0.5594945	1.282096	10890

Item selection

1. Item previously not purchased
2. Confidence recommended item. Confidence established with Market Basket Analysis

Customer Segmentation

Customer segmentation is the practice of dividing a customer base into groups of individuals that are similar in specific ways relevant to marketing



Demographic information
GENDER, AGE,
FAMILY
STATUS, ...

Geographical information
COUNTRY,
DISTRICT, CITY,
ZIP CODE, ...

Behavioral data
FIDELITY CARD,
FIRST BUY, LAST
BUY, PREFERRED
POS, ...

Behavioral data
ONLINE (Days
Frequency, Days
Recency), Avg
Amount,
TRANSACTION(D
ays Recency,
Days Frequency)
...

Behavioral data
CUSTOMER Presence in
store (Days Frequency
Days Recency), Interest
Level item (totem, cart,
check out, tried in the
dressing room), ...

Interest item

Level of interest for
cluster,
Number purchased
item into cluster

Dataset:

Customer retail fashion and large scale
retail distribution

1: Items Segmentation

ITEMS	
SKU_TYPE	Type
SKU_CONFIGURATION_CD	Configuration
SKU_PATTERN_CD	Color
SKU_MODEL_CD	Code model
SKU_PACKAGING_TYPE_CD	Type packaging
PRODUCTION_CATEGORY_CD	Category
PRODUCTION_SUBCATEGORY	Sub-category
MERCHANDISE_MCR_TYP_DS	Merchandise category
MERCHANDISE TYPOLOGY_DS	Merchandise sub-category
MERCHANDISE_FAMILY_CD	Merchandise family
MERCHANDISE_GROUP_DS	Merchandise group
MERCHANDISE_SUBGROUP_DS	Merchandise sub-group
MATERIAL_GROUP_DS	Material group
MATERIAL_SUBGROUP_DS	Material sub-group
GENDER_CD	Gender
LINE_CD	Line
BRAND_CD	Brand
STYLE_GROUP_CD	Style
BIRTH_SEASON_CD	Season
BIRTH_COLLECTION_CD	Collection
PERIODICITY_CD	Periodicity
IS_CLOTHING_ITEM_FL	Flag is clothing
NRM_CAT_LVL_1_CD	Code normalized business classification level 1
NRM_CAT_LVL_2_CD	Code normalized business classification level 2
NRM_CAT_LVL_3_CD	Code normalized business classification level 3
PRICE	NUMERIC
132 X HASHTAG	CATEGORICAL

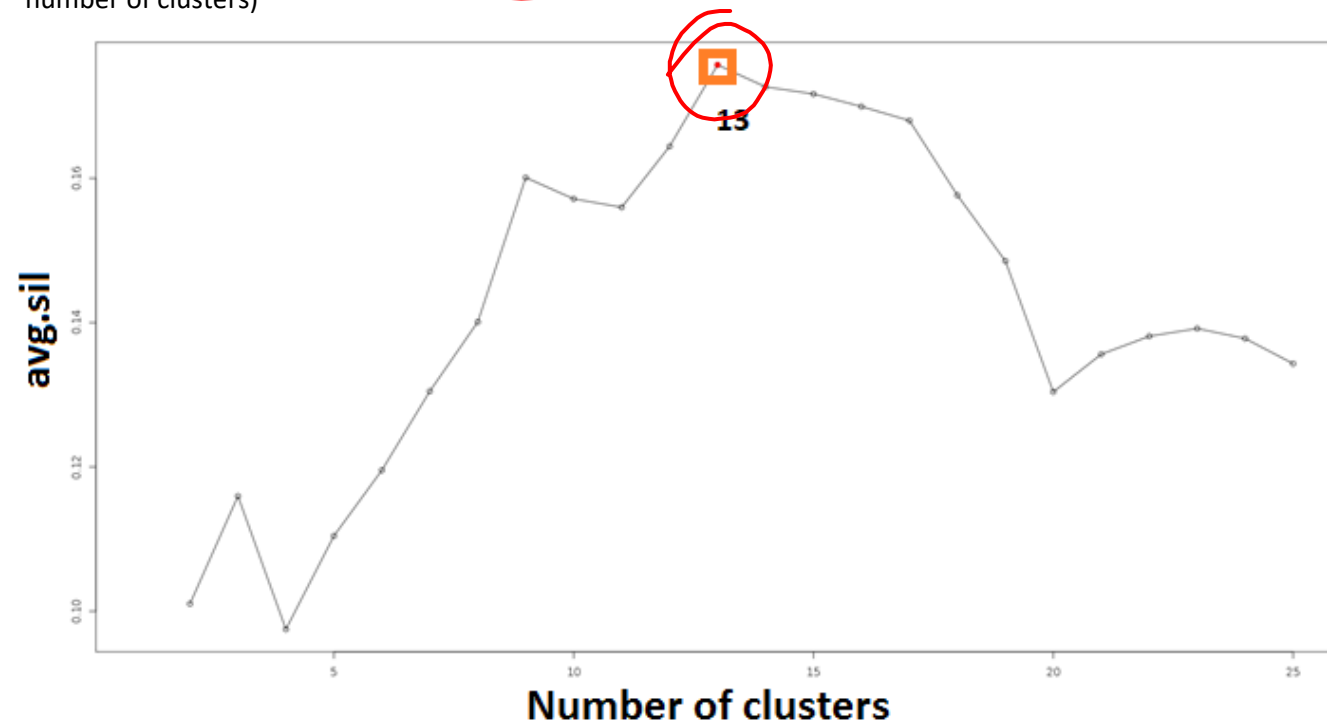
Qualitative data



Methods: K-Medoids

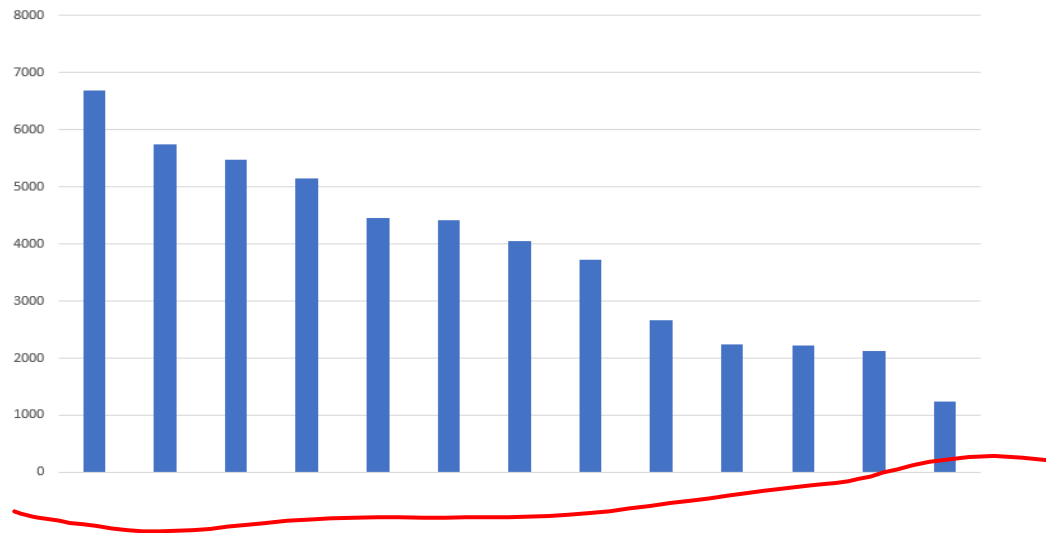
items 50204 X 158 variables

Calculate optimal number of clusters: **Silhouette analysis** (The location of the maximum is considered as the appropriate number of clusters)



1: Items Segmentation

Cluster Distribution



Visualizing the clusters 2D

Method reduction dimensionality (2 dim X/Y): Barnes-Hut Implementation
Of T-Distributed Stochastic Neighbor Embedding.

Cluster	Descrizione sintetica Cluster	# articoli in totale	# articoli in vendita
1	Borsa	6689	1726
2	Abito	5750	807
3	Maglieria A/I	5471	719
4	Pantaloni	5155	708
5	Accessori (cappello, foulard, sciarpa, collana, bracciale, guanti)	4456	620
6	Maglieria P/E	4422	533
7	Camicia	4053	614
8	Cappotto-Giubbotto-Piumino	3726	571
9	Scarpe	2659	408
10	Giacca	2247	320
11	Gonna	2221	254
12	Cintura	2121	352
13	Abbigliamento Bambini	1234	2

2: Add Interest Customer for customer

CUSTOMER TABLE	
GENDER	Gender M or F
AGE	Age
FAMILY_STATUS	Family status Not Married - Married
FIDELITY_CARD_LEVEL	Fidelity card level 1-2-3
ACTUAL_SCORE_NM	Score: integer
ACTUAL_DISCOUNT	Actual discount used
TOTAL_SCORE_NM	Score loyalty: integer
STD_COUNTRY	Country ISO code
STD_CITY_DS	City
FIRST_BUY_DT	Date first buy
FIRST_BUY_POS_CD	Store first buy
LAST_BUY_DT	Date last buy
LAST_BUY_POS_CD	Store last buy
TRN_DaysFrequency	Days frequency to buy
TRN_DaysRecency	Days recency to buy
TRN_AvgAmount	Average purchase amount
PRS_ONLINE_DaysFrequency	Frequency presence in the store online
PRS_ONLINE_DaysRecency	Recency presence in the store online
PRS_ONPREM_DaysFrequency	Frequency presence in the store
PRS_ONPREM_DaysRecency	Recency presence in the store
Preferred_POS_ID	Preferred pos
PercFidelityUsage	Percentage use of loyalty points



Interest Cluster1	Interest Cluster2		Interest Cluster13	# bought cluster1	# bought cluster2		# bought cluster13

Max interest for item in the cluster

Number purchased item in the cluster

0: No interest

1: Observed (Totem, Online, etc.)

2: Tried

3: Purchased item

3: Customer clustering

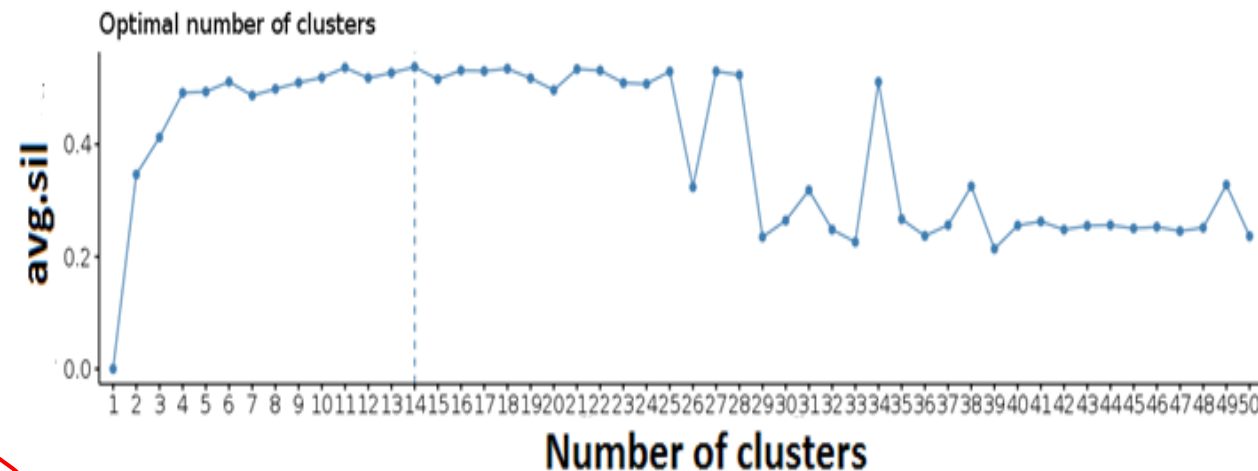
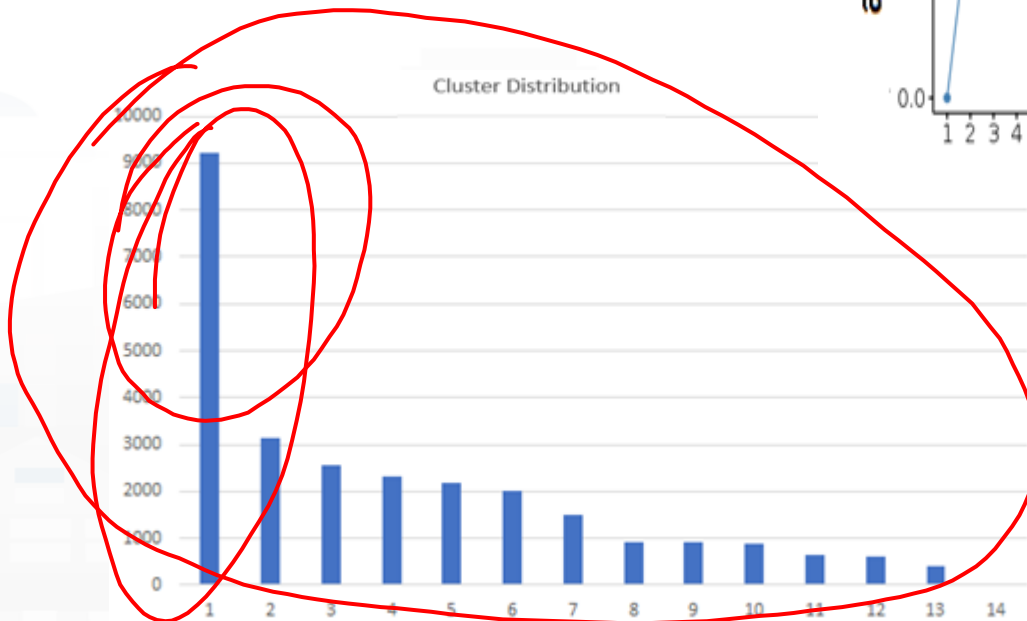
Mixed data →

Methods: K-
Medoids

Calculate optimal number of clusters: **Silhouette analysis** (The location of the maximum is considered as the appropriate number of clusters).

Visualizing the clusters
2D

Method reduction dimensionality (2 dim X/Y): Barnes-Hut
Implementation Of T-Distributed Stochastic Neighbor Embedding.



Features cluster

- Average distance in days between purchases
- Days since last purchase
- Average amount per purchase
- Online
- Onprem
- Fidelity Usage Range
- Most rated cluster
- Clusters with more purchase

4: Making Suggestion

5 X Clustering Item: correlated item

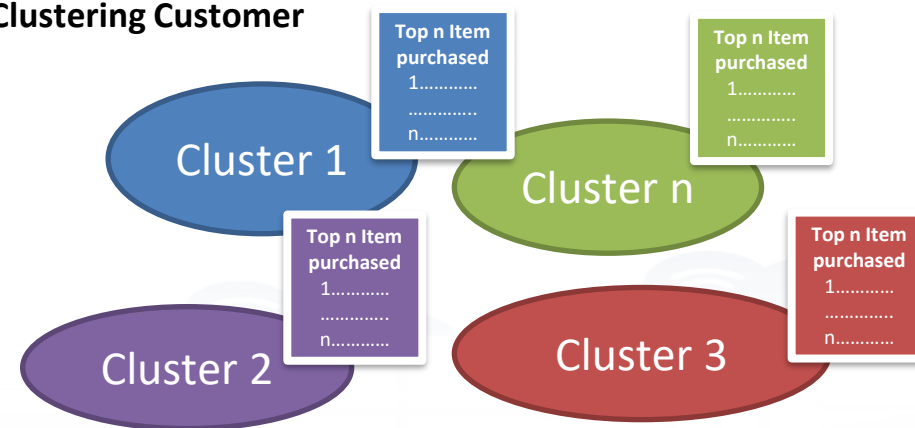
If customer bought into cluster n → System recommends an item into cluster n

5 X Clustering Item: complementary item

If customer bought into cluster n → System recommends an item into complementary cluster i,j

Cluster	Cluster description	Complementary Cluster
1	Abito	4-5-10-11-12-13
2	Maglieria A/I	3-4-5-6-7-10-11-12-13
3	Gonna	2-4-5-6-8-10-11-12-13
4	Giacca	1-2-3-6-10-11-12-13
5	Cappotto-Giubbotto-Piumino	1-2-3-6-7-10-11-12-13
6	Camicia	2-3-4-5-7-10-11-12-13
7	Pantaloni	2-4-5-6-8-10-11-12-13
8	Maglieria P/E	3-7-10-11-12-13
9	Abbigliamento Bambini	
10	Accessori (cappello, foulard, sciarpa, collana, guanti)	1-2-3-4-5-6-7-8
11	Borsa	1-2-3-4-5-6-7-8
12	Cintura	1-2-3-4-5-6-7-8
13	Scarpe	1-2-3-4-5-6-7-8

5 X Clustering Customer



5 X Serendipity

Item selection

1. Item previously not purchased
2. Confidence recommended item. Confidence established with Market Basket Analysis

Agenda

- Ragionamento semantico e simbolico
- Dalla statistica al machine learning e all'intelligenza artificiale
- Apprendimento supervisionato e non supervisionato
- Dall'elaborazione del linguaggio naturale ai modelli linguistici di grandi dimensioni (LLM)
- Intelligenza artificiale per previsioni, rilevamento di anomalie, classificazione
- **Applicazioni industriali: dai commenti degli utenti (email, questionari, post, etc.) alle indicazioni di suggerimenti per le aziende, pubbliche amministrazioni**
- Esercitazione su LLM
- Test

Logged in as:
edoardo_tester

 Your Profile

Logout

SnapAssistant model:

Llama3.3 ▾

New Chat

EdoChat1

EdoChat2

EdoChat3

EdoChat4

prova

chatNuova

chatTest1LLAMA

chatTest2LLAMA

testGianniLLAMA

testTastoInvioPreChanges

ChatTestPostPSWINVIO

<https://www.snap4city.org/1116>

SnapAdvisor

Current chat: debug0

Visualize references

Personal
Advisor LLM



Come creo una dashboard?



Per creare una dashboard, puoi utilizzare la piattaforma Snap4City. La dashboard si compone di un insieme di widget grafici che possono essere impostati separatamente assegnando una serie di parametri, come origine dati, dimensioni, colori, forma, font, allarmi, relazioni con altri, ecc. (SOURCE: Snap4 City per la creazione di Smart City Control Room, SCCR). Inoltre, la piattaforma offre la possibilità di creare vere e proprie applicazioni interattive, multipagina/multidashboard, e non semplici storie monopagina da scrollare. Puoi anche condividere la dashboard con altri utenti anche fuori dalla Smart City Control Room tramite canali autenticati e protetti in HTTPS, in accordo al GDPR (SOURCE: Snap4 City per la creazione di Smart City Control Room, SCCR). Per vedere alcuni esempi di dashboard, puoi visitare il sito <https://www.snap4city.org/dashboardSmartCity/view/index.php?iddasboard=MjU0Mg==> (SOURCE: Snap4 City 1st Overview v4 8 annot). Tuttavia, non ho trovato informazioni dettagliate sulla procedura di creazione di una dashboard. Se hai bisogno di ulteriore assistenza, ti consiglio di contattare direttamente il team di supporto di Snap4City. Non so dirti altro in merito.

Write a message...



SnapAdvisor: AI based Assistants, DSS

- **Sna4CityAdvisor:**

- for developers and promoters knowing all Snap4City and DISIT Lab literature, articles, manuals, slides, codes, etc. The SnapAdvisor with this skill is accessible only for selected Snap4City users. It allows users of saving a lot of time by providing answers to facilitate the development of smart solutions, to the exploitation of Snap4City tools, and recently on data and services accessible on the platform, producing offers, etc.

- **Legal Advisor**

- expert on specific disputes at service of the Legal department of Careggi Hospital of Florence. It allows to save time in recovering precise information from complex legal documents, ordering of events, understanding causes and effects, producing reports, etc.

- **Expert of industrial machines, user manuals, technical manuals, rules,**

- Answering on technical manuals

- **Complains and Questionnaire analysis**

- Answering on trends, via questionnaires collected via QR, blobs, emails, etc.

- **Commercial Advisor:**

- processing orders understanding them and preparing the offer, thus reducing the time to process them.

- **Generative Designs:**

- Multimodal Generative AI supporting designers in producing innovations

- etc.

SnapAdvisor

- **working on your private content** and thus producing answers.
 - Select a subset of documents
- **domain control**: you can decide the knowledge base (internal wikis, PDFs, APIs),
- **explainability** capability, provide references to your documents
- **multilingual via content**
- **modularity**: it is possible to pass from one collection of documents to another, and multiple users can work on the advisor asking for different topic on different collections/domains at the same time, independently as needs change, without any interferences among them.
- **lower hallucinations**

Complains Analysis to Support Decision Making Processes

QRCode Creation TOURISMO

Mon 17 Mar 13:40:47

Infos

Name:

Form's url:

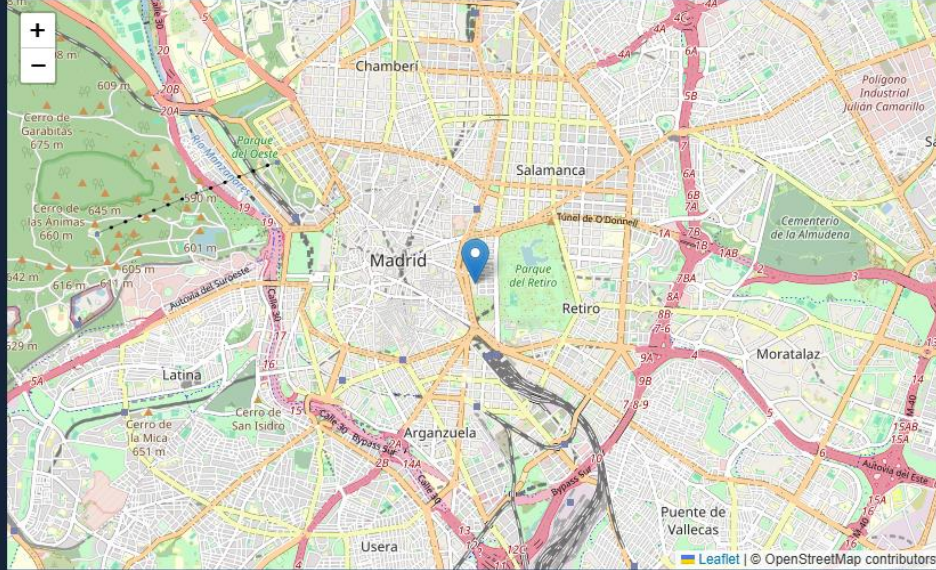
Form's edit url:

Form's stats url:

Form's excel url:

Position

Latitude: Longitude:



QrCode



[Privacy Policy](#)
[Cookies Policy](#)
[Terms and Conditions](#)
[Contact us](#)

Complains Analysis to Support Decision Making Processes

Standard data/form
creation
for reputation analysis:

- Initial Sentiment Analysis/NLP by analysing questionnaires collected from QR code distribution.

Summary of Findings (June–August 2025, 36 responses received - Golden Bay, Malta) produced by a Large Language Model (LLM) based on the responses obtained from questionnaire

Average Ratings (scale -2 to +2)

Overall experience: +0.67 → slightly positive but inconsistent (range -2 to +2)

Toilets: -0.83 → strongly negative, among the worst-rated aspects.

Showers: -0.31 → generally negative.

Bins distribution: -0.33 → slightly negative, with complaints about lack of bins.

Safety & lifeguard services: +0.55 → moderately positive.

Cleanliness: -0.52 → overall negative, many complaints about cigarette butts and plastics.

Natural conditions (dunes, seawater, etc.): -0.40 → negative, issues with seagrass, water quality.

Agenda

- Ragionamento semantico e simbolico
- Dalla statistica al machine learning e all'intelligenza artificiale
- Apprendimento supervisionato e non supervisionato
- Dall'elaborazione del linguaggio naturale ai modelli linguistici di grandi dimensioni (LLM)
- Intelligenza artificiale per previsioni, rilevamento di anomalie, classificazione
- Applicazioni industriali:
- **Esercitazione su LLM**
- Test

Esercitazione: *Bad Prompts VS Good Prompts* (Prompt Engineering)

- Esercitazione su ChatGPT (utilizzando *il* modello base, non è necessario essere registrati). Aprire la pagina web: <https://chatgpt.com/>
- Prompt «vago»:
Cos'è il Digital Marketing?
- Come si può migliorare la risposta, rendendola più completa ed accurata?
 - Assegnare un ruolo specifico all'assistente (ChatGPT) in modo da focalizzarlo nel dominio di competenze richiesto per poter generare una risposta più completa e accurata.
 - Formulare indicazioni e richieste dettagliate, possibilmente organizzate in un elenco numerico o per punti se si vogliono risposte specifiche ad ogni domanda. E' utile anche specificare il background dell'utente che legge la risposta (ad esempio se ha competenze specifiche sull'argomento richiesto o meno...).
 - Aggiungere ulteriori indicazioni su come si vuole far generare la risposta: tono formale/informale, tecnico, didattico, chiaro e semplice («*for dummies*»); indicazioni sullo stile: discorsivo, formattazione tabellare o per punti; chiedere riferimenti sulle fonti, ecc.

Esercitazione: *Bad Prompts VS Good Prompts* (Prompt Engineering)

- Esercitazione su ChatGPT (utilizzando il modello base, non è necessario essere registrati). Aprire la pagina web: <https://chatgpt.com/>

- Prompt «vago»:

Cos'è il Digital Marketing?

- Prompt avanzato:

Ruolo e Contesto

Sei un esperto di Digital Marketing, e ti occupi di formazione per PMI sul digital marketing.

Richieste/Obiettivi

1. Spiega in modo chiaro, pratico e rigoroso cos'è il digital marketing.

2. Fornisci un quadro operativo sintetico, adatto a un'utenza non tecnica.

Altre Indicazioni, Stile, Riferimenti...

Usa bullet e tabelle per massimizzare la leggibilità e fornisci i riferimenti alle fonti da cui prendi le informazioni, citandole nel testo.

Agenda

- Ragionamento semantico e simbolico
- Dalla statistica al machine learning e all'intelligenza artificiale
- Apprendimento supervisionato e non supervisionato
- Dall'elaborazione del linguaggio naturale ai modelli linguistici di grandi dimensioni (LLM)
- Intelligenza artificiale per previsioni, rilevamento di anomalie, classificazione
- Applicazioni industriali:
- Esercitazione su LLM
- **Test**

Test

- Test finale al seguente link:

<https://docs.google.com/forms/d/e/1FAIpQLSfC8BaJxShuUNqe8NcUhwTwapkLmgYd1ky3srg7aOeH8CjonA/viewform>

Grazie per l'attenzione!

? Q&A !

Firenze, 9 Ottobre 2025,
DISIT lab

113



UNIVERSITÀ
DEGLI STUDI
FIRENZE

DINFO
DIPARTIMENTO DI
INGEGNERIA
DELL'INFORMAZIONE

DISIT
DISTRIBUTED SYSTEMS AND
INTERNET TECHNOLOGIES LAB
DISTRIBUTED DATA INTELLIGENCE
AND TECHNOLOGIES LAB



CONTATTI

Paolo Nesi, paolo.nesi@unifi.it

Dipartimento di Ingegneria dell'Informazione, DISIT
Cell: 335-5668674, Lab <https://www.disit.dinfo.unifi.it>

<https://www.snap4city.org>

<https://www.snap4industry.org>



114

Firenze, 9 Ottobre 2025,
DISIT lab



UNIVERSITÀ
DEGLI STUDI
FIRENZE

DINFO
DIPARTIMENTO DI
INGEGNERIA
DELL'INFORMAZIONE

DISIT
DISTRIBUTED SYSTEMS AND
INTERNET TECHNOLOGIES LAB
DISTRIBUTED DATA INTELLIGENCE
AND TECHNOLOGIES LAB

SNAP4CITY

TUSCANY X.O

EDIH | European
Digital Innovation
Hubs Network

Finanziato
dall'Unione europea
NextGenerationEU