



**Politecnico
di Torino**

Politecnico di Torino

Corso di Laurea Magistrale in Ingegneria Energetica e Nucleare

A. a. 2024/2025

Sessione di Laurea Ottobre 2025

Sviluppo di un modello predittivo per l'analisi preliminare della dispersione di idrogeno

Relatore:

Prof. Vittorio Verda

Candidato:

Sofia Castellani

Correlatore:

Ing. Rugiada Scozzari

Abstract

Negli ultimi anni le tecniche di apprendimento automatico si sono diffuse rapidamente in diversi ambiti applicativi, tra cui quello della sicurezza, in quanto consentono di individuare relazioni complesse tra variabili e di sfruttarle per effettuare previsioni. Allo stesso tempo, il crescente interesse verso l'idrogeno nella transizione energetica comporta sfide importanti legate alla gestione del rischio. In questo contesto, la presente tesi illustra lo sviluppo di uno strumento predittivo per l'analisi della dispersione di idrogeno in stazioni di rifornimento per veicoli in aree urbane, valutando in quali casi sia preferibile adottare reti neurali o ricorrere a regressione multivariata, così da affiancare i modelli fisici consolidati con un approccio più rapido e accessibile. A questo scopo, è stato realizzato un codice Python impiegando un database di simulazioni svolte con PHAST, software utilizzato per la modellazione delle conseguenze. Attraverso analisi di sensibilità e considerazioni legate all'applicazione nelle stazioni di rifornimento, sono state definite le variabili di ingresso rilevanti. Sono stati considerati diversi scenari di rilascio accidentale di idrogeno gassoso, per i quali sono stati implementati modelli distinti con l'obiettivo di stimare la portata di rilascio, le distanze ai limiti di infiammabilità e la larghezza massima della nube. L'addestramento delle reti neurali è stato condotto tramite suddivisione 70%-15%-15% dei dati nei set di training, validation e test, selezionando il numero ottimale di neuroni nascosti in base alle prestazioni ottenute. Il contributo principale del lavoro è la definizione di un'architettura modulare, in grado di integrare tecniche diverse a seconda della disponibilità dei dati e dello scenario analizzato, che può essere aggiornata facilmente con nuovi dati per aumentarne la robustezza, in base alle particolari esigenze. La selezione mirata degli input ha semplificato l'addestramento e risulta coerente con l'idea che, per valutazioni preliminari, limitarsi alle variabili più rilevanti sia più conveniente rispetto a dover gestire un numero elevato di parametri. In conclusione, lo strumento proposto non sostituisce i modelli fisici, ma può essere impiegato per uno screening rapido all'interno di domini predefiniti, permettendo di condurre simulazioni più dettagliate solo per scenari critici.

Indice

Elenco delle tabelle	v
Elenco delle figure	vii
1 Introduzione	1
1.1 Idrogeno e analisi del rischio	1
1.2 Stato dell'arte	3
1.3 Obiettivo e struttura del lavoro	4
2 PHAST	6
2.1 Introduzione al software	6
2.2 Scenari di rilascio	8
2.2.1 Modello <i>Orifice</i>	9
2.2.2 Modello <i>Short pipe</i>	10
2.2.3 Modello istantaneo	13
2.3 Simulazioni	13
3 Costruzione del dataset	16
3.1 Selezione dei parametri di ingresso	16
3.1.1 Analisi di sensibilità per il modello <i>Orifice</i>	16
3.1.2 Analisi di sensibilità per il modello <i>Short pipe</i>	24
3.2 Selezione dei risultati	27
3.2.1 Portata di rilascio e distanze	28
3.2.2 Rappresentazione grafica della nube	28
3.3 Dataset	29
3.3.1 Leak	31
3.3.2 Catastrophic rupture	32
3.3.3 Line rupture	33
3.3.4 Disk rupture	33
3.3.5 Relief valve	34

4	Sviluppo del modello	36
4.1	Reti neurali e regressione	36
4.2	Ambiente di sviluppo delle reti e librerie	39
4.2.1	Manipolazione dei dati	40
4.2.2	Visualizzazione	40
4.2.3	Deep Learning	41
4.2.4	Pre-processing e valutazione	44
4.3	Struttura del codice	45
4.4	Implementazione delle reti neurali	49
5	Analisi dei risultati	54
5.1	Metriche	54
5.2	Modelli neurali	56
5.2.1	Leak	56
5.2.2	Catastrophic rupture	59
5.2.3	Line rupture	61
5.3	Modelli regressivi	62
5.3.1	Disk rupture	63
5.3.2	Relief valve	63
5.4	Stima dell'area	64
5.5	Sintesi dei risultati	65
6	Conclusioni	67
6.1	Contributi del lavoro	67
6.2	Discussione e sviluppi futuri	68
A	Dataset di simulazione	70
A.1	Generale	70
A.2	Leak	73
A.3	Catastrophic rupture	75
A.4	Line rupture	77
A.5	Disk rupture	79
A.6	Relief valve	81
B	Prestazioni	83
B.1	Prestazioni globali	83
B.2	Prestazioni per scenario	87
B.3	Approssimazione ellittica dell'area	94
	Bibliografia	97

Elenco delle tabelle

3.1	Valori di ingresso per lo scenario di riferimento.	17
3.2	Variazione percentuale della concentrazione e della distanza massima al variare della temperatura.	18
3.3	Parametri di input per i due scenari di riferimento per l'analisi di sensibilità del modello <i>Short pipe</i>	25
3.4	Variazione percentuale delle distanze raggiunte in funzione della lunghezza della tubazione per lo scenario con DN25.	26
3.5	Variazione percentuale delle distanze raggiunte in funzione della lunghezza della tubazione per lo scenario con DN40.	26
3.6	Variazione percentuale delle distanze raggiunte in funzione dei parametri per lo scenario con DN25.	27
3.7	Variazione percentuale delle distanze raggiunte in funzione dei parametri per lo scenario con DN40.	27
A.1	Riepilogo degli scenari analizzati e dei target considerati.	71
A.2	Range delle variabili di input simulate per lo scenario <i>Leak</i>	73
A.3	Statistiche descrittive per i diversi target per lo scenario <i>Leak</i>	73
A.4	Range delle variabili di input simulate per lo scenario <i>Catastrophic rupture</i>	75
A.5	Statistiche descrittive per i diversi target per lo scenario <i>Catastrophic rupture</i>	75
A.6	Range delle variabili di input simulate per lo scenario <i>Line rupture</i> . . .	77
A.7	Statistiche descrittive per i diversi target per lo scenario <i>Line rupture</i> . .	77
A.8	Range delle variabili di input simulate per lo scenario <i>Disk rupture</i> . . .	79
A.9	Statistiche descrittive per i diversi target per lo scenario <i>Disk rupture</i> . .	79
A.10	Range delle variabili di input simulate per lo scenario <i>Relief valve</i> . . .	81
A.11	Statistiche descrittive per i diversi target per lo scenario <i>Relief valve</i> . .	81
B.1	Metriche di valutazione globali per i diversi target.	83
B.2	Iperparametri per ciascun target dello scenario <i>Leak</i>	87

B.3	Prestazioni sul test set per lo scenario <i>Leak</i>	87
B.4	Valore medio e deviazione standard delle metriche di validazione (10-fold cross validation) per lo scenario <i>Leak</i>	88
B.5	Iperparametri per ciascun target dello scenario <i>Catastrophic rupture</i> .	90
B.6	Prestazioni sul test set per lo scenario <i>Catastrophic rupture</i>	90
B.7	Valore medio e deviazione standard delle metriche di validazione (10-fold cross validation) per lo scenario <i>Catastrophic rupture</i>	90
B.8	Iperparametri per ciascun target dello scenario <i>Line rupture</i>	91
B.9	Prestazioni sul test set per lo scenario <i>Line rupture</i>	91
B.10	Valore medio e deviazione standard delle metriche di validazione (10-fold cross validation) per lo scenario <i>Line rupture</i>	92
B.11	Prestazioni sul test set per lo scenario <i>Disk rupture</i>	93
B.12	Prestazioni sul test set per lo scenario <i>Relief valve</i>	93

Elenco delle figure

1.1	Rappresentazione schematica dell'approccio utilizzato per l'analisi del rischio.	2
2.1	Schema concettuale della sequenza di modellazione di un rilascio di gas in pressione.	7
2.2	<i>Orifice model</i> [13].	9
2.3	Schema di risoluzione del rilascio da foro, rielaborato da documentazione DNV [13].	10
2.4	<i>Short pipe model</i> [13].	11
2.5	Relief valve [13].	11
2.6	Schema di risoluzione del rilascio da breve tratto di tubazione, rielaborato da documentazione DNV [13].	12
2.7	<i>Instantaneous model</i> [13].	13
2.8	Proiezione della nube sul piano orizzontale per diversi rilasci da foro per un serbatoio di idrogeno a 700 bar.	15
3.1	Concentrazione massima raggiunta in funzione della distanza e proiezione della nube sul piano orizzontale per temperatura pari a -10°C rispetto al caso di riferimento.	18
3.2	Concentrazione massima e distanza massima raggiunta per diversi valori della pressione di stoccaggio.	19
3.3	Concentrazione massima e distanza massima raggiunta per diversi valori di diametro del foro.	20
3.4	Visione laterale della nube per le diverse altezze del punto di rilascio.	21
3.5	Visione laterale e proiezione della nube sul piano orizzontale per le diverse velocità del vento.	22
3.6	Concentrazione massima raggiunta in funzione della distanza e proiezione della nube sul piano orizzontale per umidità relativa pari a 90% rispetto al caso di riferimento.	23

3.7	Concentrazione massima raggiunta in funzione della distanza per diversi tipi di terreno.	24
4.1	Rappresentazione schematica delle relazioni tra Intelligenza Artificiale, <i>Machine Learning</i> e <i>Deep Learning</i> (a sinistra) e classificazione delle principali modalità di apprendimento nel <i>Machine Learning</i> (a destra).	37
4.2	Architettura generale di un Multi-layer Perceptron [25].	38
4.3	Visualizzazione dell'andamento della funzione di costo durante il training.	41
4.4	Struttura generale del codice per l'addestramento delle reti neurali.	46
4.5	Struttura generale del codice del modulo centrale.	47
4.6	Grafico relativo a una rottura catastrofica di un serbatoio a 200 bar, con 50 kg di idrogeno stoccato e condizioni meteorologiche 1,5/F.	49
4.7	Confronto distribuzioni per il target <i>Max width</i> dello scenario di rilascio <i>Leak</i>	50
4.8	Illustrazione della procedura alla base della k-fold cross validation.	53
A.1	Numerosità dei dataset per scenario.	72
A.2	Matrice di correlazione con coefficiente di Pearson per lo scenario <i>Leak</i>	74
A.3	Matrice di correlazione con coefficiente di Pearson per lo scenario <i>Catastrophic rupture</i>	76
A.4	Matrice di correlazione con coefficiente di Pearson per lo scenario <i>Line rupture</i>	78
A.5	Matrice di correlazione con coefficiente di Pearson per lo scenario <i>Disk rupture</i>	80
A.6	Matrice di correlazione con coefficiente di Pearson per lo scenario <i>Relief valve</i>	82
B.1	Confronto tra valori reali e predetti per il target <i>Flow rate</i>	84
B.2	Confronto tra valori reali e predetti per il target <i>UFL distance</i>	84
B.3	Confronto tra valori reali e predetti per il target <i>LFL distance</i>	85
B.4	Confronto tra valori reali e predetti per il target <i>1/2 LFL distance</i>	85
B.5	Confronto tra valori reali e predetti per il target <i>Max width</i>	86
B.6	Matrice di confusione per il target <i>UFL distance</i> dello scenario <i>Leak</i>	89
B.7	Matrice di confusione per il target <i>LFL distance</i> dello scenario <i>Leak</i>	89
B.8	Matrice di confusione per il target <i>UFL distance</i> dello scenario <i>Line rupture</i>	92
B.9	Confronto tra i valori reali ed approssimati dell'area per lo scenario <i>Leak</i>	94
B.10	Confronto tra i valori reali ed approssimati dell'area per lo scenario <i>Catastrophic rupture</i>	94

B.11	Confronto tra i valori reali ed approssimati dell'area per lo scenario <i>Line rupture.</i>	95
B.12	Confronto tra i valori reali ed approssimati dell'area per lo scenario <i>Disk rupture.</i>	95
B.13	Confronto tra i valori reali ed approssimati dell'area per lo scenario <i>Relief valve.</i>	96

Capitolo 1

Introduzione

1.1 Idrogeno e analisi del rischio

L'idrogeno è un vettore energetico in grado di offrire un contributo significativo al processo di transizione energetica. Come riportato dall'International Energy Agency (IEA) [1], se prodotto con fonti a basse o nulle emissioni, l'impiego di idrogeno favorisce la decarbonizzazione di diversi settori, tra cui quello dei trasporti. Tuttavia, alcune sue caratteristiche intrinseche, come l'ampio intervallo di infiammabilità di 4-75% in volume in aria, la bassa energia di ignizione, circa un decimo rispetto a quella del metano, e la difficoltà nel rilevamento, comportano sfide importanti in tema di sicurezza [2].

Per promuovere l'utilizzo di idrogeno nel settore dei trasporti, è prevista la realizzazione di stazioni di rifornimento dedicate: in assenza di norme tecniche prescrittive consolidate, è necessario condurre analisi quantitative del rischio (QRA) per dimostrare che sia garantita la sicurezza in questi impianti.

La QRA (*Quantitative Risk Assessment*) è un processo articolato in più fasi che permette di identificare e quantificare i rischi presenti in un determinato contesto, al fine di gestirli e mitigarli. Definito il sistema oggetto di studio, si effettua una prima analisi qualitativa del rischio, volta a identificare i potenziali pericoli attraverso diverse metodologie, come HAZOP (*HAZard and OPerability analysis*) o FMECA (*Failure Mode, Effects and Criticality Analysis*). A ciascun pericolo viene assegnato un punteggio, determinato in funzione della frequenza di accadimento e della gravità del danno, attraverso una matrice di rischio. Gli scenari più critici sono selezionati per un'analisi successiva e costituiscono gli eventi iniziatori da cui possono originarsi sequenze incidentali. Si conduce una valutazione quantitativa del rischio associato a ciascuna sequenza incidentale individuata, effettuando in parallelo un'analisi probabilistica e una stima delle conseguenze, la cui combinazione

restituisce il rischio. Il confronto con criteri di accettabilità può portare a definire il rischio accettabile, tollerabile entro il principio ALARP (*As Low As Reasonably Practicable*), che comporta la valutazione di costi e benefici derivanti da ulteriori interventi di mitigazione, oppure non accettabile; in quest'ultimo caso è necessario adottare opportune misure di sicurezza volte a ridurlo.

Lo schema in Figura 1.1, rielaborato a partire dal materiale del corso *Sicurezza e analisi di rischio* [3] riassume le fasi di un'analisi del rischio quantitativa (QRA).

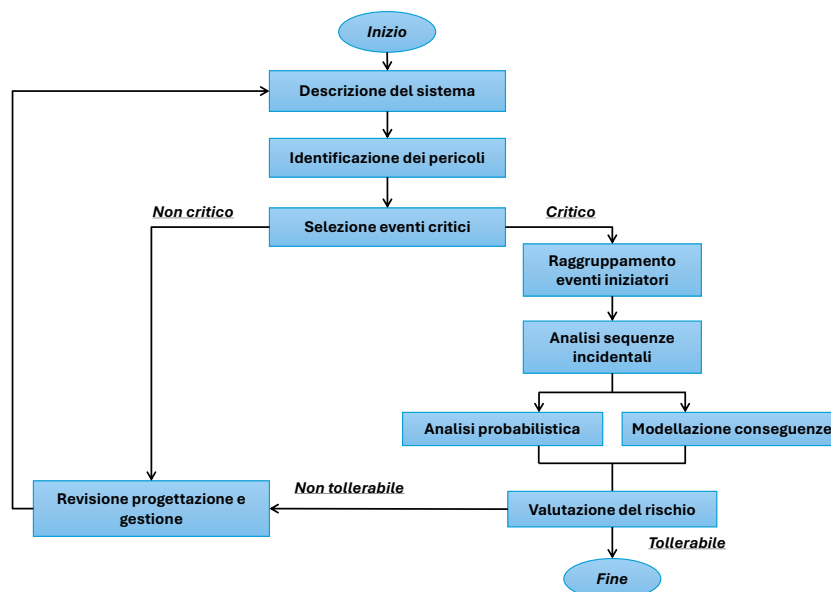


Figura 1.1: Rappresentazione schematica dell'approccio utilizzato per l'analisi del rischio.

La valutazione delle conseguenze è dunque un passaggio chiave previsto dall'approccio QRA e consiste nel quantificare l'estensione delle aree di danno sulla base di criteri distinti per scenario.

In seguito a un rilascio incidentale, l'idrogeno può innescarsi immediatamente dando luogo a un jet fire, quindi una fiamma a getto continua, oppure in un secondo momento, portando ad un'esplosione o ad un flash fire, ossia un incendio che si propaga rapidamente e che coinvolge l'intera nube infiammabile. La differenza tra esplosione e flash fire è legata livelli di sovrappressione generati, trascurabili in quest'ultimo fenomeno. La modellazione della dispersione è necessaria quando l'ignizione non è immediata: la nube, diluendosi progressivamente con l'aria, può raggiungere concentrazioni comprese tra i limiti di infiammabilità e innescarsi anche a distanza dal punto di rilascio [4]. Dal punto di vista fisico, la dispersione è inizialmente governata dalla velocità con cui avviene il rilascio; successivamente, se

la densità della sostanza è inferiore a quella dell'aria, come nel caso dell'idrogeno, la nube tende a sollevarsi e si stabilizza a una certa quota, per poi disperdersi in maniera passiva; ciò viene descritto anche nel report di DNV che illustra la teoria alla base del modello di dispersione implementato in PHAST [5]. Il livello a cui si stabilizza la nube corrisponde al limite superiore del *mixing layer*, che è lo strato di atmosfera in prossimità del suolo dove si verificano moti turbolenti che comportano il mescolamento della sostanza con l'aria. Superiormente la dispersione viene limitata a causa di un'inversione di temperatura e l'altezza di questo strato varia in base alle condizioni meteorologiche. La stabilità atmosferica è un parametro fondamentale per descrivere l'intensità del mescolamento per effetto della turbolenza: per condizioni stabili questa interazione è ridotta, mentre aumenta per condizioni instabili. La classificazione avviene secondo le classi di Pasquill, indicate con lettere che vanno da A, condizioni molto instabili, a F, molto stabili, in funzione della velocità del vento e della radiazione solare. All'aumentare della velocità del vento, inoltre, la nube viene trasportata più rapidamente in direzione orizzontale, mentre le caratteristiche del terreno, come la sua rugosità superficiale, influenzano anch'esse la turbolenza. L'analisi della dispersione è quindi finalizzata a determinare le distanze alle quali vengono raggiunte le concentrazioni di soglia rilevanti ai fini della sicurezza: per le sostanze infiammabili queste corrispondono ai limiti di infiammabilità. Pertanto, è necessario conoscere l'andamento della concentrazione nel tempo e nello spazio attorno al punto di rilascio, mentre per valutare possibili scenari di esplosione e flash fire è necessario stimare anche la quantità di sostanza contenuta nella nube infiammabile.

1.2 Stato dell'arte

La dispersione rappresenta il fenomeno più complesso da simulare, dal momento che ci sono diverse variabili che la influenzano. La complessità dei modelli utilizzati cresce con il numero di variabili considerate, consentendo di rappresentare situazioni più realistiche; al tempo stesso, le simulazioni richiedono più tempo per l'esecuzione. I modelli parametrici sono strumenti semplificati rispetto ai classici CFD (*Computational Fluid Dynamics*) che prevedono la risoluzione numerica delle equazioni di Navier-Stokes; utilizzano, infatti, equazioni più semplici che descrivono il comportamento medio del flusso a seguito di un rilascio. Nella categoria di modelli parametrici e semiempirici rientrano quelli implementati in PHAST (*Process Hazard Analysis Software Tool*), software sviluppato da DNV.

Pouyakian et al. [6] hanno condotto un'analisi della letteratura scientifica prodotta in Iran tra il 2006 e il 2022 evidenziando come, su 40 studi per la modellazione delle conseguenze, 25 abbiano impiegato il software PHAST. La maggior accuratezza nei risultati ottenuti con PHAST, rispetto ad altri software, è stata confermata

anche da Park et al. [7].

Recentemente diversi studi si sono concentrati sull'integrazione di tecniche di machine learning nella modellazione delle conseguenze, molti dei quali hanno previsto la costruzione di database con simulazioni effettuate con PHAST. Nel lavoro di Zhong et al. [8], le simulazioni di diversi scenari di rilascio da serbatoi in pressione sono state utilizzate come base di partenza per lo sviluppo di modelli predittivi. Questi modelli sono finalizzati alla valutazione della distanza alla quale viene raggiunto un certo valore di intensità di radiazione termica in caso di incendio. Sun et al. [9] trattano la predizione delle distanze di soglia per *jet fire*, *early pool fire* e *late pool fire*, dimostrando come sia possibile ottenere modelli accurati per fornire risultati rapidi all'interno del dominio di addestramento.

Sono presenti anche studi che utilizzano un approccio di questo tipo per la modellazione della dispersione. Papadaki et al. [10], ad esempio, propongono quattro modelli per la stima della distanza di sicurezza (SD) associata ad una soglia di tossicità, mentre Lee et al. [11] si soffermano sulla massima distanza raggiunta alla quale si verifica 100% LFL (*Lower Flammability Limit*) e 25% LFL. Jiao et al. [12], infine, propongono la valutazione della distanza massima, minima e della larghezza della nube infiammabile, entro 100% LFL e 50% LFL.

1.3 Obiettivo e struttura del lavoro

Nonostante PHAST sia uno strumento molto efficiente ed affidabile per la valutazione delle conseguenze, potrebbe risultare dispersivo se l'obiettivo è condurre analisi preliminari in cui la rapidità con cui vengono forniti i risultati prevale sulla necessità di elevata precisione. Sulla base di ciò, il presente elaborato descrive lo sviluppo di un modello predittivo che possa essere utilizzato per la stima della dispersione di idrogeno in seguito a rilasci incidentali. Nello specifico, si tratta di un applicativo realizzato con linguaggio Python che richiede un numero limitato di variabili di ingresso significative per la modellazione e restituisce previsioni rapide e chiare.

Partendo da precedenti lavori finalizzati a valutare l'integrazione di tecniche machine learning nell'analisi delle conseguenze, si è cercato di estendere il modello comprendendo diversi scenari di rilascio e più target utili per la valutazione delle aree di danno. Pertanto, nel lavoro è stato condotto uno studio per definire la fattibilità dell'inclusione di alcuni scenari, individuando le modalità più appropriate per lo sviluppo del modello e valutando le eventuali semplificazioni necessarie a garantire la coerenza complessiva del sistema.

L'elaborato è articolato in sei capitoli. Nel Capitolo 2 viene riportata una descrizione sintetica del software PHAST utilizzato per le simulazioni che costituiscono il database del modello, distinguendo tra gli scenari di rilascio disponibili e mostrando

la procedura per avviarle. Nel Capitolo 3 viene descritto l'approccio per la costruzione del database, la selezione dei parametri di ingresso, con relativi intervalli definiti per ciascuno di essi sulla base di analisi di sensibilità ed il tipo di risultati considerati. Nel Capitolo 4 viene proposta una breve introduzione alla teoria dei modelli implementati, seguita dalla presentazione della loro architettura e delle procedure di addestramento e valutazione degli stessi. Nel Capitolo 5 vengono discussi i risultati ottenuti per valutare le prestazioni del modello sviluppato. Infine, nel Capitolo 6 vengono tratte le conclusioni in merito all'impiego dello strumento per la valutazione delle conseguenze e proposti possibili miglioramenti.

Capitolo 2

PHAST

2.1 Introduzione al software

PHAST è un software sviluppato da DNV che si basa su diversi modelli parametrici e semiempirici per la modellazione delle conseguenze derivanti da rilasci incidentali di sostanze pericolose. Come riportato da Park et al. [7], si tratta di uno strumento ampiamente utilizzato nell'analisi del rischio.

L'architettura interna di PHAST è particolarmente complessa: il software impiega modelli diversi in base alla sostanza analizzata, alle sue proprietà fisiche e alle condizioni di rilascio. Per un rilascio di gas da un serbatoio in pressione la sequenza dei fenomeni può essere descritta attraverso lo schema concettuale riportato in Figura 2.1.

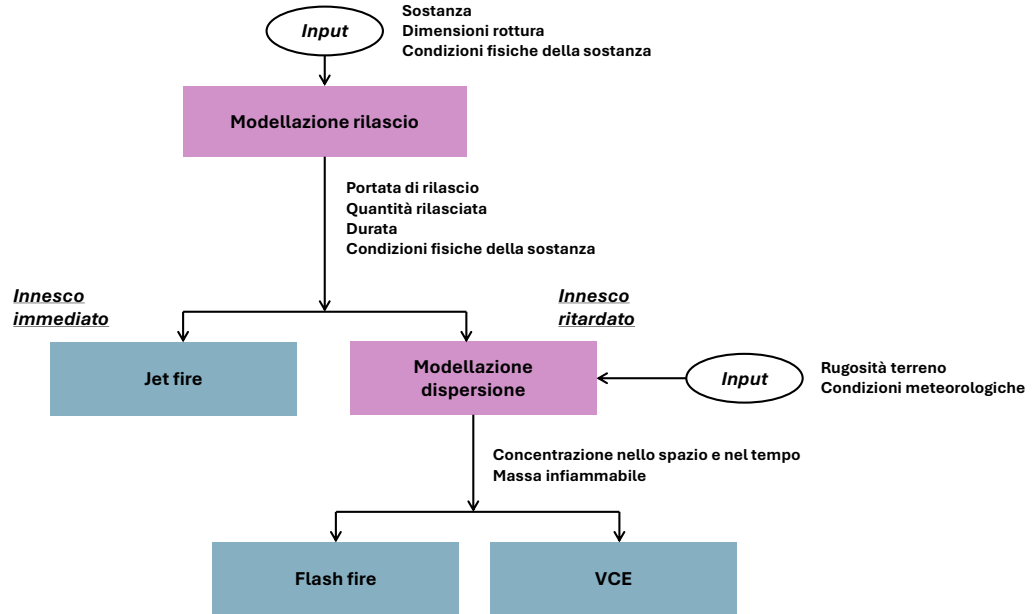


Figura 2.1: Schema concettuale della sequenza di modellazione di un rilascio di gas in pressione.

PHAST propone alcuni elementi dai quali la sostanza può essere rilasciata, come serbatoi in pressione o a pressione atmosferica. A livello di componente, possono essere esplicitate informazioni come il tipo di sostanza e le sue proprietà fisiche. Successivamente, è possibile selezionare lo scenario di rilascio scegliendo, ad esempio, tra modelli stazionari o transitori e tra configurazioni di rilascio da un foro sul serbatoio o da un breve tratto di tubazione. Per i diversi scenari è necessario definire alcuni parametri, come il diametro del foro di perdita.

Configurato lo scenario di interesse, il modello DISC (*Discharge Model*) gestisce il rilascio iniziale, generando in uscita informazioni come la portata di rilascio, la quantità totale rilasciata e la durata. La successiva espansione fino alle condizioni atmosferiche è gestita dal modello ATEX (*Atmospheric Expansion Model*), generando le condizioni di ingresso per il modello di dispersione UDM (*Unified Dispersion Model*), che simula invece la dispersione atmosferica, tenendo conto di altri parametri come stabilità atmosferica e rugosità del terreno.

PHAST simula contemporaneamente *jet fire*, che può avvenire in caso di innesco immediato, *flash fire* e *VCE* (*Vapor Cloud Explosion*), che sono invece le possibili conseguenze di un innesco ritardato della nube formatasi. I risultati vengono forniti sia in forma tabellare riportando, ad esempio, le distanze massime alle quali si

verificano concentrazioni infiammabili o diversi livelli di radiazione termica, sia in forma grafica, con rappresentazioni spaziali e temporali delle soglie di effetto.

2.2 Scenari di rilascio

PHAST impiega diversi modelli di rilascio per simulare incidenti da serbatoi in pressione. Le possibili configurazioni, descritte nella documentazione ufficiale [13], sono:

- Rilascio da foro sul serbatoio (*“Orifice model”*)
- Rilascio da breve tratto di tubazione connessa al serbatoio (*“Short pipe model”*)
- Rilascio dovuto a rottura catastrofica del serbatoio (*“Instantaneous release”*)
- Rilascio da sfiato di vapore durante le operazioni di riempimento (*“Venting from vapour space”*)

Dal momento che il presente lavoro analizza le conseguenze associate esclusivamente a rilasci di idrogeno gassoso, vengono utilizzati solo i primi tre modelli. Per i modelli *Orifice* e *Short pipe* è inoltre possibile specificare se il rilascio è stazionario o variabile nel tempo.

2.2.1 Modello *Orifice*

Il modello di rilascio da foro simula l'espansione della sostanza dalle condizioni di stoccaggio fino alle condizioni di uscita, a valle dell'orifizio.

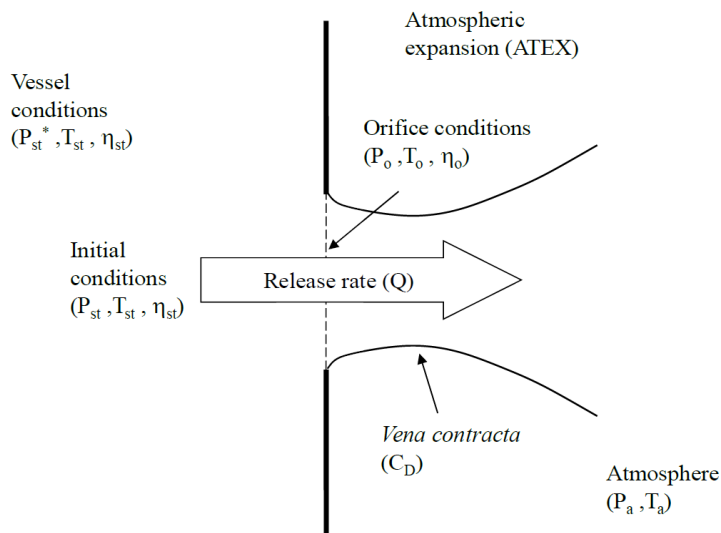


Figura 2.2: *Orifice model* [13].

Si propone in Figura 2.3 una rielaborazione schematica del modello, che riassume il metodo di risoluzione riportato nella documentazione ufficiale di DNV [13].

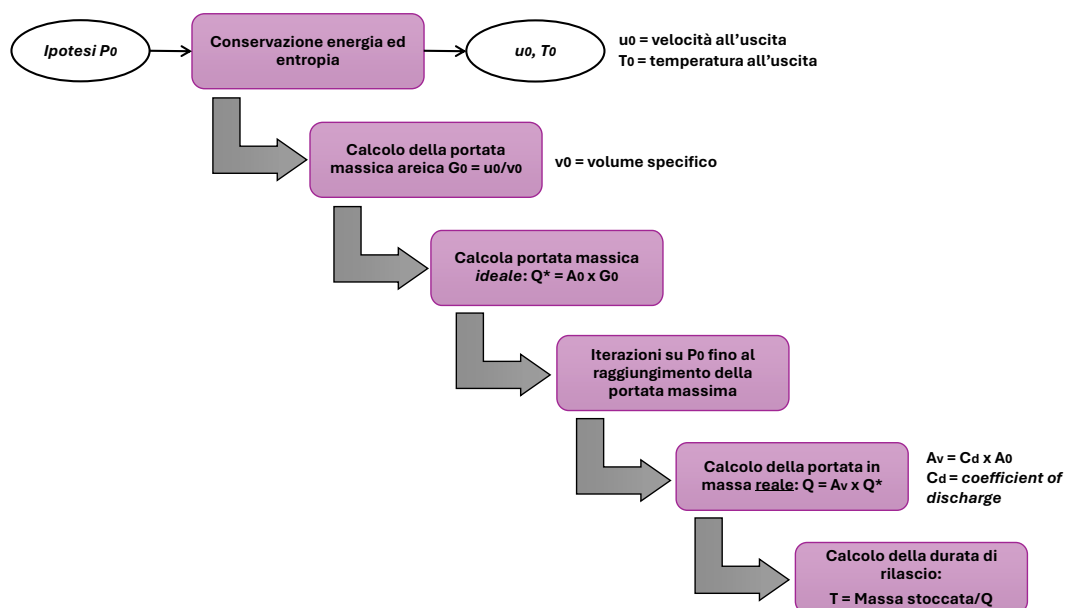


Figura 2.3: Schema di risoluzione del rilascio da foro, rielaborato da documentazione DNV [13].

I parametri esplicitabili per questo scenario sono limitati: diametro del foro, direzione ed altezza di rilascio.

2.2.2 Modello *Short pipe*

Il modello *Short pipe* permette di simulare il rilascio che può avvenire da un breve tratto di tubazione connessa al serbatoio e può essere utilizzato per rappresentare tre differenti scenari:

- Rottura completa della tubazione
- Rilascio da valvola di sicurezza (*Relief valve*)
- Rilascio da disco di rottura (*Bursting disk*)

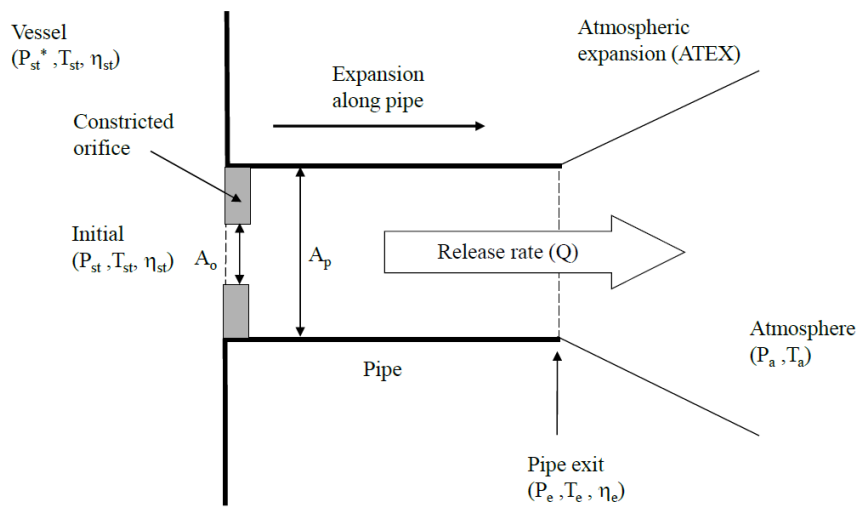


Figura 2.4: *Short pipe model* [13].

Il metodo di risoluzione è lo stesso indipendentemente dallo scenario scelto, con la differenza che nel caso della valvola di sicurezza viene considerata una sezione di passaggio ridotta rispetto a quella della tubazione (Figura 2.5).

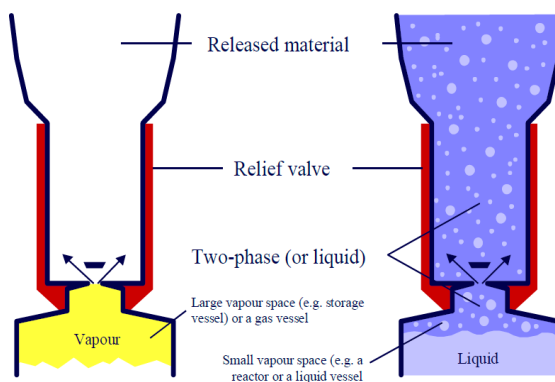


Figura 2.5: Relief valve [13].

In Figura 2.6 è riportata una rappresentazione schematica del metodo di risoluzione.

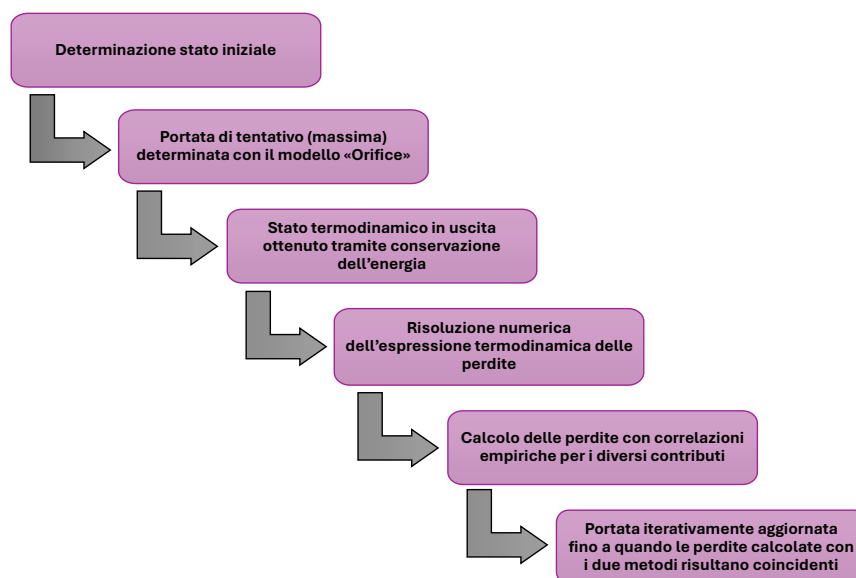


Figura 2.6: Schema di risoluzione del rilascio da breve tratto di tubazione, rielaborato da documentazione DNV [13].

Il modello *Short pipe*, a differenza del modello *Orifice*, richiede un maggior numero di parametri, come la lunghezza e la rugosità superficiale del tratto di tubazione, il numero di valvole o di curve a gomito per metro; infatti, tiene conto delle perdite localizzate e distribuite presenti tra il serbatoio e il punto di rilascio così da fornire una stima più realistica delle conseguenze.

2.2.3 Modello istantaneo

Il modello istantaneo simula la rottura catastrofica del serbatoio e si limita a descrivere l'espansione dalle condizioni di stoccaggio fino a quelle atmosferiche.

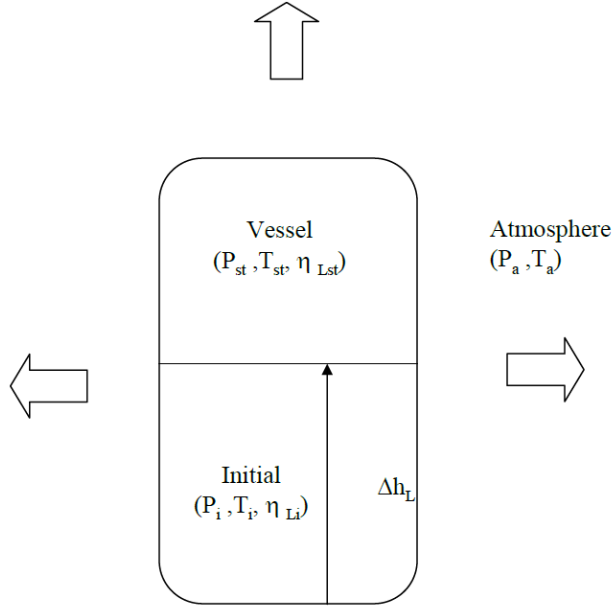


Figura 2.7: *Instantaneous model* [13].

Per questo scenario non è necessario impostare parametri aggiuntivi rispetto a quelli che definiscono le proprietà della sostanza contenuta nel serbatoio.

2.3 Simulazioni

Le simulazioni sono state condotte distinguendo per scenario, analizzando separatamente le diverse tipologie di rilascio e definendo per ciascuno di essi un insieme di combinazioni dei parametri di ingresso. PHAST presenta una struttura gerarchica articolata su diversi livelli per gestire le simulazioni, come riportato nel tutorial fornito da DNV [14]:

- *Workspace*: è il primo livello, in cui è possibile agire sulle impostazioni che influenzano il comportamento del programma.
- *Study*: è il livello successivo, nel quale è possibile definire impostazioni comuni a più simulazioni come, ad esempio, le condizioni meteorologiche.

- *Equipment item*: rappresenta il componente dell'impianto da cui avviene il rilascio che si vuole analizzare. PHAST permette di scegliere tra serbatoi in pressione, serbatoi a pressione atmosferica, tubazioni, moduli per incendi in ambienti confinati e, infine, moduli *standalone*, utilizzati per modellare specifiche dinamiche. In questo lavoro, è stato utilizzato per tutte le simulazioni il serbatoio in pressione.
- *Scenario*: per ogni componente è necessario definire lo scenario di rilascio. Sono stati considerati *Leak*, *Short pipe* e *Catastrophic rupture*. All'interno dello scenario *Short pipe* è possibile selezionare il tipo di rilascio specifico (*Line rupture*, *Relief valve* o *Disk rupture*) tramite un menù a tendina.

In un primo momento, sono state configurate le condizioni meteorologiche, come temperatura e stabilità atmosferica, e la rugosità del terreno. Per la caratterizzazione della dispersione, le soglie di concentrazione corrispondono ai livelli di infiammabilità per l'idrogeno, già presenti nel database delle sostanze di PHAST. Per organizzare le simulazioni in maniera efficiente, sono stati definiti diversi serbatoi per rappresentare i livelli di pressione considerati. A ciascuno di essi, sono stati poi associati più scenari, differenziati per altri parametri caratteristici; ad esempio, per *Leak*, sono stati distinti in base a diametro del foro e altezza di rilascio.

I risultati numerici ottenuti dalle simulazioni vengono prodotti da PHAST in forma tabellare, distinguendo per scenario incidentale e riportando, ad esempio, le distanze massime alle quali vengono raggiunti i valori di soglia; questi, sono stati esportati in formato Excel.

I risultati grafici prodotti sono anch'essi distinti per scenario incidentale e possono essere esportati in formato PNG. Tra le immagini prodotte solo la massima impronta a terra della nube è stata considerata: è stata utilizzata per ottenere il valore della larghezza massima della nube e dell'estensione superficiale della stessa.

Tutti i risultati dipendono dall'altezza di riferimento che è stata posta pari a 1 metro.

Si riporta in Figura 2.8 un esempio dei grafici utilizzati.

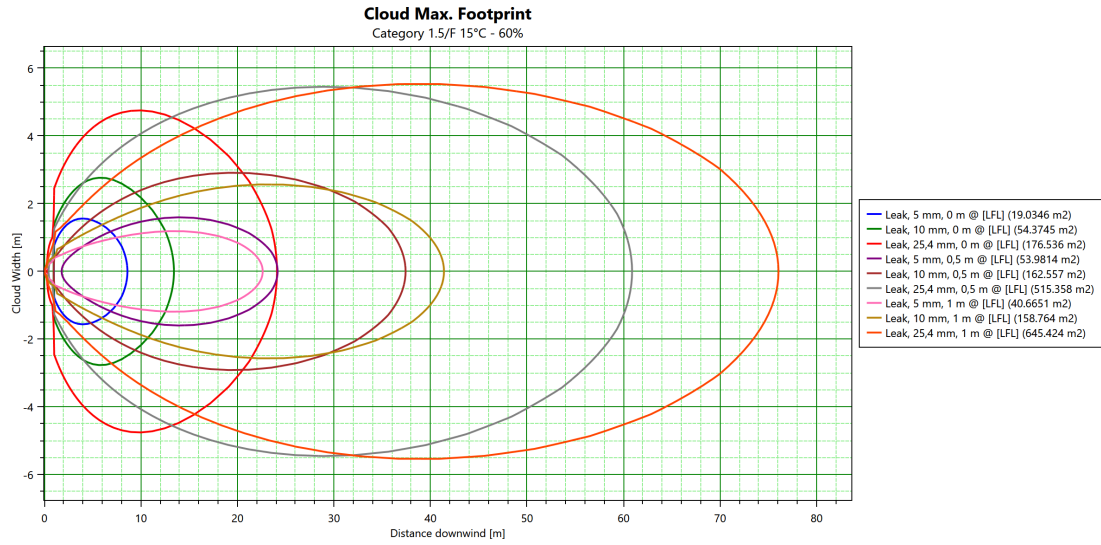


Figura 2.8: Proiezione della nube sul piano orizzontale per diversi rilasci da foro per un serbatoio di idrogeno a 700 bar.

Capitolo 3

Costruzione del dataset

3.1 Selezione dei parametri di input

I dataset utilizzati per l'addestramento delle reti sono stati generati con simulazioni condotte con il software PHAST. Dal momento che non è possibile coprire ogni condizione operativa, è stato necessario individuare un numero limitato di variabili di ingresso, selezionate in funzione dell'influenza di ciascuna sul fenomeno di dispersione.

Per l'idrogeno gassoso, oggetto di questo studio, gli scenari di rilascio considerati riguardano la perdita da un foro sul serbatoio, la fuoriuscita da un breve tratto di tubazione e la rottura catastrofica del serbatoio. Per ciascuno di essi, il software richiede di impostare specifici parametri caratteristici.

Fatta eccezione per la rottura catastrofica che dipende principalmente dalle condizioni di stoccaggio, per i modelli *Orifice* e *Short pipe* sono state svolte delle analisi di sensibilità. L'obiettivo di queste analisi è limitare le variabili di ingresso alle sole grandezze effettivamente impattanti sui risultati, per contenere il numero di simulazioni necessarie per la modellazione di ciascuno scenario.

3.1.1 Analisi di sensibilità per il modello *Orifice*

Nel loro studio, Lee et al. propongono alcune variabili da analizzare [11]. Nello specifico, vengono selezionati 8 parametri di ingresso a partire da quelli considerati in precedenti studi di sensibilità:

- Temperatura di esercizio
- Pressione di esercizio

- Diametro del foro di perdita
- Altezza del punto di rilascio
- Velocità del vento
- Temperatura ambientale
- Umidità relativa
- Rugosità superficiale

A partire da questi, si è cercato di ridurre ulteriormente le variabili di ingresso valutando l'influenza di ciascuna di esse sui risultati, variandone il valore all'interno di un intervallo definito e mantenendo i restanti parametri costanti. Lo scenario di riferimento è il seguente:

Temperatura di esercizio	15°C
Pressione di esercizio	40 bar
Diametro del foro di perdita	1"
Altezza del punto di rilascio	1 m
Velocità del vento	1,5 m/s
Temperatura ambientale	15°C
Umidità relativa	60%
Rugosità superficiale	<i>Land</i>

Tabella 3.1: Valori di ingresso per lo scenario di riferimento.

Per quantificare gli scostamenti rispetto al caso di riferimento sono state valutate la concentrazione massima raggiunta e la massima distanza alla quale si verifica una concentrazione pari alla metà del limite inferiore di infiammabilità. I valori riportati si riferiscono tutti al livello del suolo.

Temperatura di esercizio e temperatura ambientale

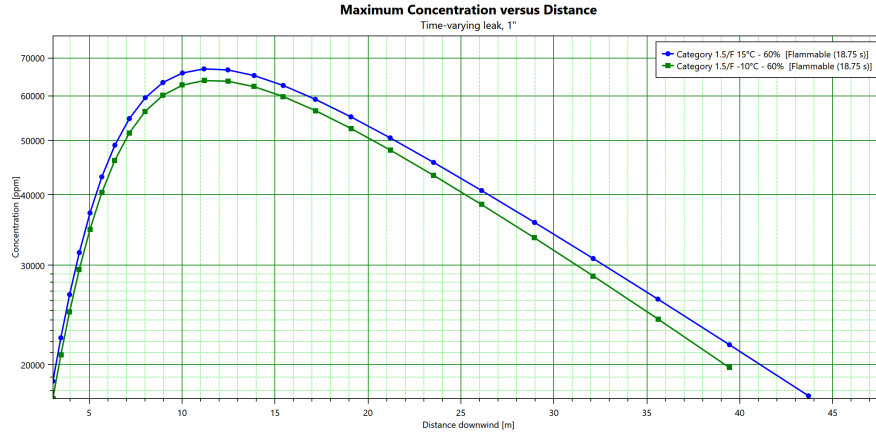
Come affermato da Utgikar e Thiesen in [15], l'idrogeno gassoso può essere stoccato a temperatura ambiente; dunque, è possibile assumere che la temperatura di stoccaggio e quella ambientale costituiscano una sola variabile.

È stata simulata la dispersione con temperatura pari a -10°C e 50°C, per confrontare i risultati con quelli dello scenario di riferimento a 15°C; rispetto a questo, gli scostamenti sono riportati in Tabella 3.2

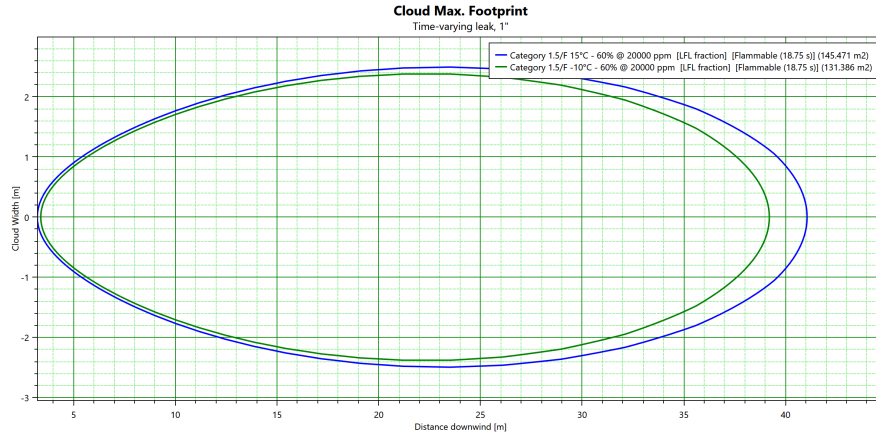
Grandezza	T = 50°C	T = -10°C
Concentrazione massima	+5%	-4%
Distanza massima raggiunta	+7%	-5%

Tabella 3.2: Variazione percentuale della concentrazione e della distanza massima al variare della temperatura.

Dal momento che l'influenza della temperatura sui risultati è contenuta, viene ritenuto opportuno mantenere costante questo parametro per le simulazioni.



(a)

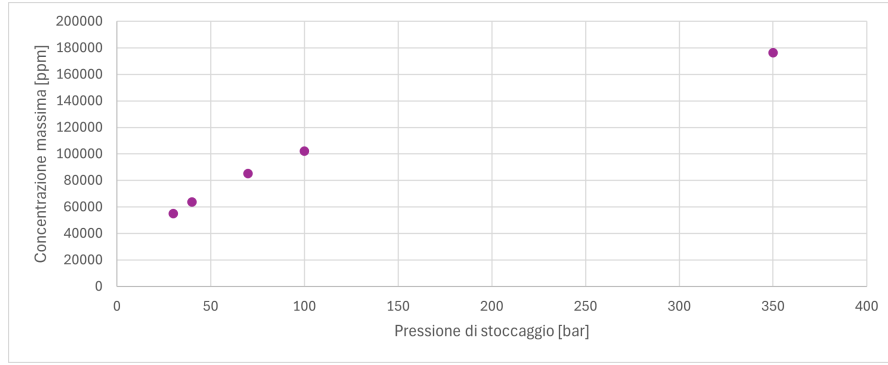


(b)

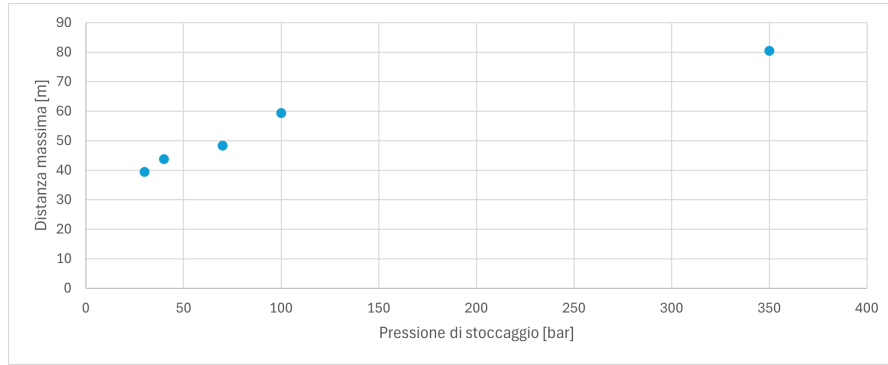
Figura 3.1: Concentrazione massima raggiunta in funzione della distanza e proiezione della nube sul piano orizzontale per temperatura pari a -10°C rispetto al caso di riferimento.

Pressione di esercizio

La portata uscente dal foro di perdita è direttamente proporzionale alla pressione di stoccaggio [16]. La stretta relazione fra queste due grandezze suggerisce che la pressione incida in maniera significativa sull'entità delle conseguenze associate ad un rilascio di idrogeno e ciò viene confermato dai risultati ottenuti.



(a)



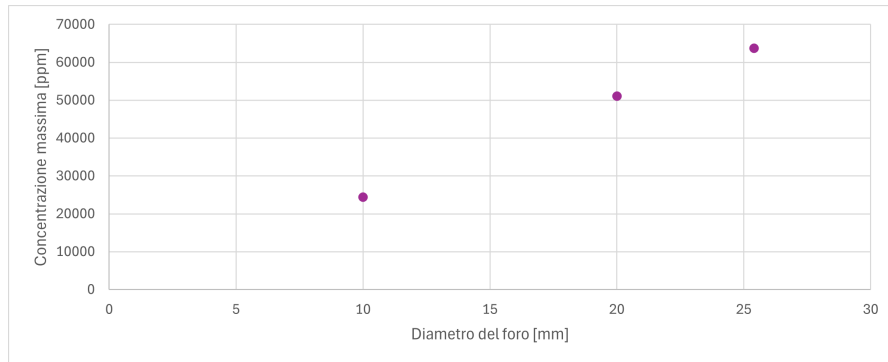
(b)

Figura 3.2: Concentrazione massima e distanza massima raggiunta per diversi valori della pressione di stoccaggio.

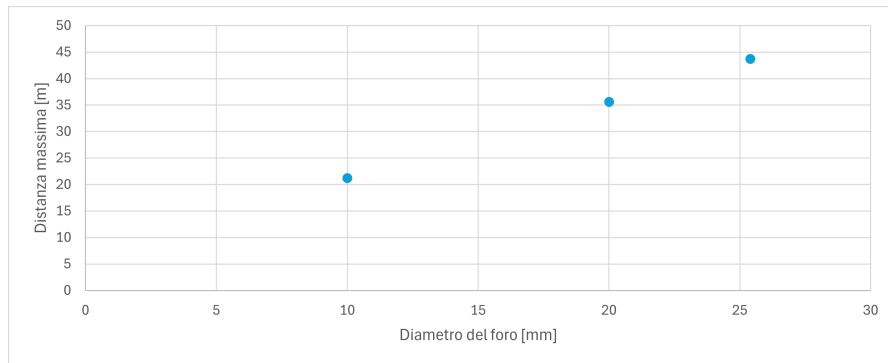
La pressione di stoccaggio viene variata in un ampio intervallo di valori per condurre le simulazioni, cercando di coprire i livelli di pressione normalmente previsti negli impianti che utilizzano idrogeno.

Diametro del foro di perdita

La portata uscente dal foro di perdita è direttamente proporzionale al suo diametro [16]. Pertanto, anche per questa grandezza, a valori maggiori corrispondono concentrazioni massime e distanze massime raggiunte più elevate.



(a)



(b)

Figura 3.3: Concentrazione massima e distanza massima raggiunta per diversi valori di diametro del foro.

Il diametro del foro di perdita è, quindi, un parametro da variare per la costruzione del database.

Altezza del punto di rilascio

L'altezza del punto di rilascio è stata posta pari a 0, 0,5 e 2,5 metri, per permettere un confronto con quella dello scenario di riferimento. È stato possibile osservare come questa grandezza abbia un effetto rilevante sulla forma della nube (Figura 3.4).

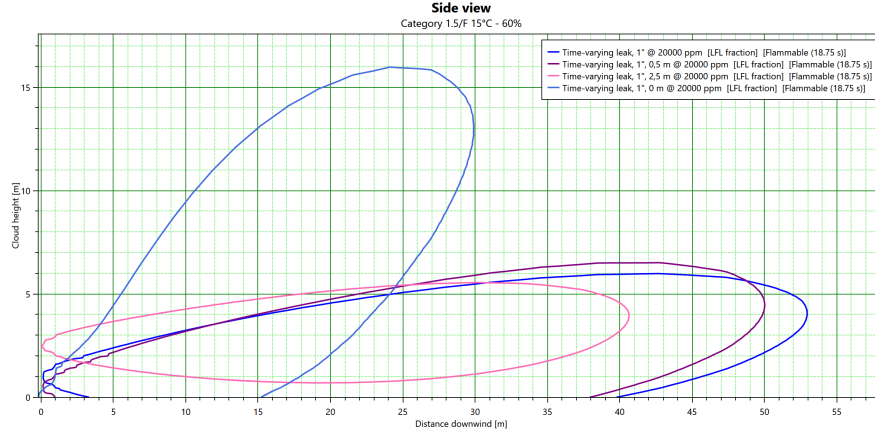
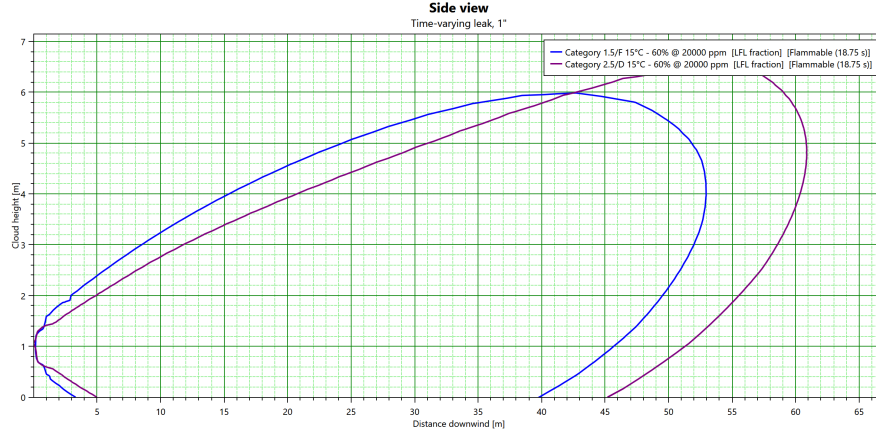


Figura 3.4: Visione laterale della nube per le diverse altezze del punto di rilascio.

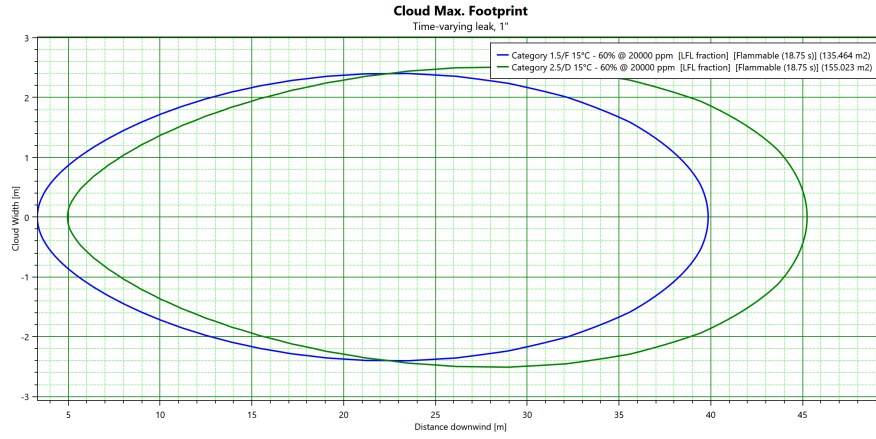
Dal momento che l'altezza di riferimento per i risultati è costante, i valori di concentrazione massima e distanza massima raggiunta variano per una diversa altezza del punto di rilascio, ed è quindi necessario condurre simulazioni a diverse quote. Un rilascio sopraelevato comporta, infatti, una riduzione delle concentrazioni al suolo rispetto a rilasci a quote minori.

Velocità del vento

Nel Purple Book [17] viene suggerito di coprire almeno sei combinazioni di condizioni meteorologiche per l'analisi del rischio, così da rappresentare non solo condizioni conservative, corrispondenti a vento debole e stabilità elevata, ma anche situazioni più realistiche. Tuttavia, nel presente studio l'analisi è stata limitata a due sole condizioni per ridurre il numero di simulazioni: classe di stabilità F e D, associando due velocità del vento rispettivamente pari a 1,5 e 2,5 m/s. È possibile osservare che a velocità del vento maggiori corrispondono una concentrazione massima più bassa (-4%) ma una più elevata distanza massima raggiunta (+10%).



(a)



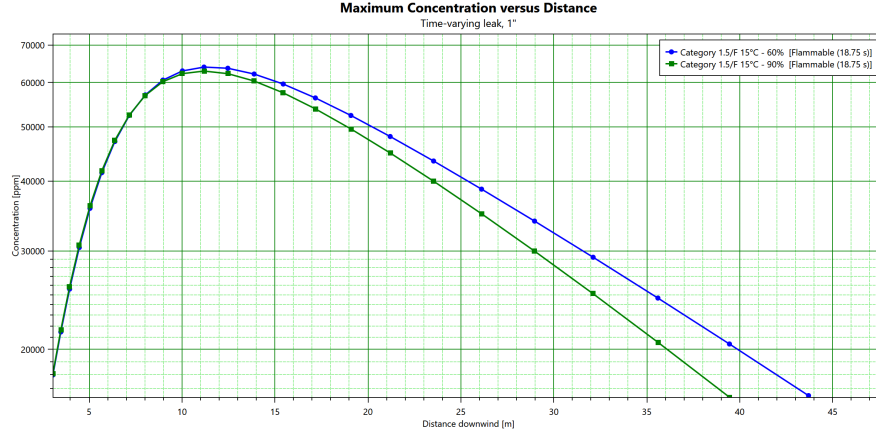
(b)

Figura 3.5: Visione laterale e proiezione della nube sul piano orizzontale per le diverse velocità del vento.

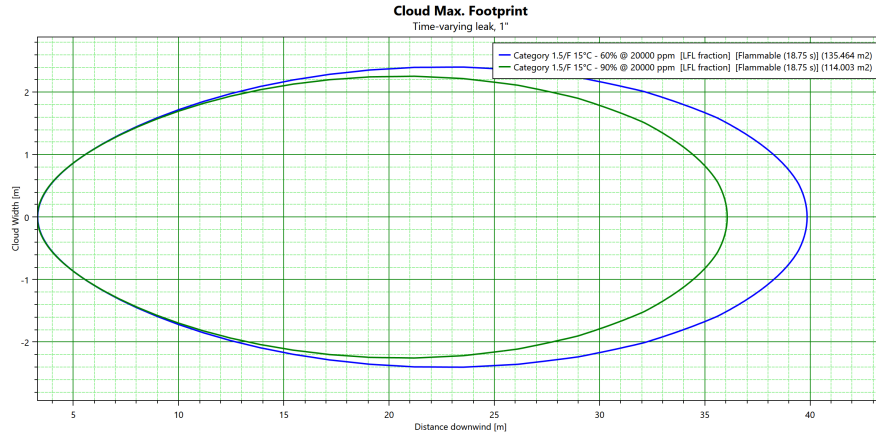
Un vento più intenso favorisce la diluizione del gas, con conseguente riduzione della concentrazione massima. Questo parametro viene fatto variare assumendo solo questi due valori per la definizione del database.

Umidità relativa

Per un'umidità relativa pari a 90%, rispetto allo scenario di riferimento (60%), si osservano riduzioni contenute sia per quanto riguarda la concentrazione massima (-2%) che per la distanza massima raggiunta (-10%).



(a)



(b)

Figura 3.6: Concentrazione massima raggiunta in funzione della distanza e proiezione della nube sul piano orizzontale per umidità relativa pari a 90% rispetto al caso di riferimento.

Pertanto, si ritiene adeguato mantenere costante questo parametro nelle simulazioni pari a 60%.

Rugosità superficiale

Il tipo di terreno predefinito *land* è stato confrontato con le opzioni *city* e *suburb* disponibili in PHAST. Nello specifico, l'opzione *land* rappresenta terreni rurali aperti per i quali il Purple Book [17] suggerisce una lunghezza di rugosità superficiale pari a 0,03 m; *suburb* corrisponde ad aree suburbane con strutture medio-grandi e

lunghezza di rugosità superficiale di 1,0 m; infine, *city* descrive un centro urbano con edifici alti associando lunghezza di rugosità superficiale di 3,0 m.

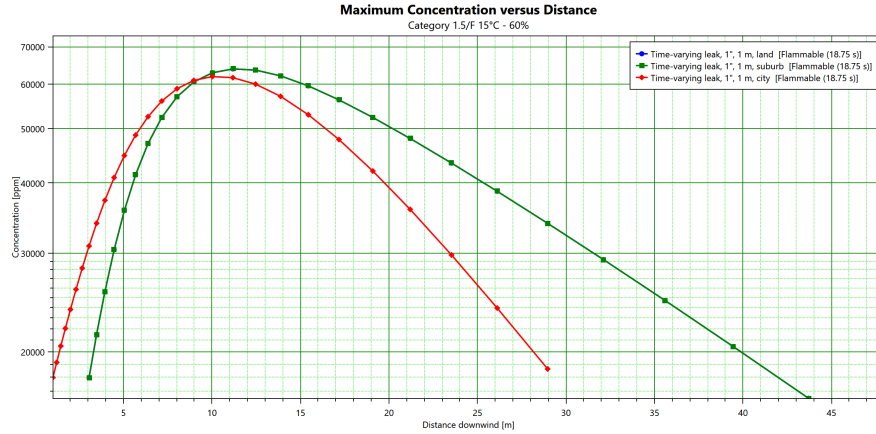


Figura 3.7: Concentrazione massima raggiunta in funzione della distanza per diversi tipi di terreno.

Come è possibile vedere in Figura 3.7, il grafico per tipo di terreno *suburb* si sovrappone a quello dello scenario di riferimento. Per il tipo di terreno *city* si assiste a una riduzione importante della distanza massima raggiunta (-30%) e una variazione marginale per la concentrazione massima (-3%). Tuttavia, si tratta di un parametro legato al contesto fisico del sito: se l'applicazione è nota o si dispone di informazioni sull'ambiente di rilascio, può essere mantenuto costante nelle simulazioni. In questo contesto, si utilizza una lunghezza di rugosità superficiale di 0,5 m, che si ritiene essere un valore rappresentativo per stazioni di rifornimento.

3.1.2 Analisi di sensibilità per il modello *Short pipe*

Prima di procedere con lo sviluppo di una rete neurale per la modellazione dei rilasci da tubazioni con il modello *Short pipe*, è stata condotta un'analisi di sensibilità per verificare che fosse effettivamente necessaria. La scelta di impiegare un modello dedicato, in questo caso, potrebbe essere motivata da stime eccessivamente conservative delle aree di danno prodotte dal modello *Orifice*.

Il modello *Short pipe* richiede alcuni parametri aggiuntivi rispetto a *Orifice*:

- Rugosità della tubazione
- Numero di giunti di accoppiamento per metro
- Numero di curve per metro
- Numero di diramazioni per metro

- Numero di valvole per metro
- Lunghezza della tubazione

Per quanto riguarda il numero di valvole, PHAST permette di considerare tre diverse tipologie e di associare a ciascuna di esse le perdite di carico caratteristiche. Per condurre quest'analisi, sono stati scelti due casi di riferimento, simulati entrambi con il modello *Orifice*, ponendo il coefficiente di efflusso pari a 1. I parametri di questi scenari sono riportati in Tabella 3.3.

Parametro	Unità di misura	Scenario 1	Scenario 2
Pressione	bar	350	350
Altezza rilascio	m	1	1
Velocità del vento	m/s	1,5	1,5
Classe di stabilità	-	F	F
Diametro	mm	26,7	40,9
DN	mm	25	40

Tabella 3.3: Parametri di input per i due scenari di riferimento per l'analisi di sensibilità del modello *Short pipe*.

Per capire quali parametri introdurre per costituire il database degli scenari simulabili con questo modello, vengono variati i valori di ciascuno di essi entro intervalli plausibili per una stazione di rifornimento. Le assunzioni fatte sono le seguenti:

- Rugosità della tubazione: dal momento che i materiali impiegati sono standardizzati e caratterizzati da valori di rugosità relativamente uniformi, viene assunta costante.
- Numero di giunti di accoppiamento per metro: viene posto pari a zero, coerentemente con la prassi progettuale secondo cui le linee per il trasporto di idrogeno vengono realizzate quasi esclusivamente tramite saldatura.
- Numero di curve per metro: si analizza un intervallo di valori compreso tra $0,2 \text{ m}^{-1}$ e $1,5 \text{ m}^{-1}$, tenendo presente che si tratta di un parametro non facilmente generalizzabile e fortemente dipendente dal layout.
- Numero di diramazioni per metro: in base ai P&ID proposti per una stazione di rifornimento [18] viene posto pari a $0,1 \text{ m}^{-1}$, ammettendo la presenza di 1–2 diramazioni, limitate ai punti di discontinuità in corrispondenza delle interfacce tra i vari elementi dell'impianto.
- Numero di valvole per metro: sempre in base ai P&ID [18], è possibile ricavare il numero di valvole per ciascuna tipologia. Nello specifico, sono presenti diverse

valvole a sfera e a solenoide le quali, nella configurazione operativa normale, risultano completamente aperte, comportando perdite di carico trascurabili. Le valvole di non ritorno sono invece meno frequenti, perciò si considera un numero pari a $0,1 \text{ m}^{-1}$. Per questo tipo di valvole, invece, l'elemento mobile presente introduce sempre una resistenza aggiuntiva. Per l'idrogeno possono essere usate valvole di non ritorno a sfera per le quali viene suggerito valore di perdita di $1,5$ [19]. In via cautelativa, si assume un valore di *velocity head losses* pari a 2.

- Lunghezza della tubazione: si considerano valori indicativi dei tratti di tubazione che collegano i diversi elementi nel layout, pari a 0,5, 2 e 6 metri.

Per valutare l'influenza del parametro relativo alla lunghezza della tubazione, gli scenari di riferimento sono modellati con *Orifice*; tuttavia, per l'analisi dei restanti parametri è necessario definire uno scenario di riferimento modellato con *Short pipe*, definendo la lunghezza del tratto di tubazione, richiesta se si vuole simulare la presenza di altri elementi.

Lunghezza [m]	Variazione percentuale (%)		
	UFL dist.	LFL dist.	1/2 LFL dist.
0,5	0	-5,3	-4,6
2	0	-11,0	-9,8
6	0	-20,2	-18,0

Tabella 3.4: Variazione percentuale delle distanze raggiunte in funzione della lunghezza della tubazione per lo scenario con DN25.

Lunghezza [m]	Variazione percentuale (%)		
	UFL dist.	LFL dist.	1/2 LFL dist.
0,5	0	-5,1	-4,6
2	0	-7,5	-7,0
6	0	-15,4	-14,3

Tabella 3.5: Variazione percentuale delle distanze raggiunte in funzione della lunghezza della tubazione per lo scenario con DN40.

Parametro	N. [m^{-1}]	Variazione percentuale (%)		
		UFL dist.	LFL dist.	1/2 LFL dist.
Curve	0,2	0	-0,5	-0,4
Curve	1,5	0	-3,0	-2,7
Diramazioni	0,1	0	-1,0	-0,8
Valvole	0,1	0	-1,9	-1,7

Tabella 3.6: Variazione percentuale delle distanze raggiunte in funzione dei parametri per lo scenario con DN25.

Parametro	N. [m^{-1}]	Variazione percentuale (%)		
		UFL dist.	LFL dist.	1/2 LFL dist.
Curve	0,2	0	-0,5	-0,5
Curve	1,5	0	-3,3	-3,0
Diramazioni	0,1	0	-1,3	-1,2
Valvole	0,1	0	-2,5	-2,3

Tabella 3.7: Variazione percentuale delle distanze raggiunte in funzione dei parametri per lo scenario con DN40.

I risultati mostrano che la lunghezza del tratto di tubazione può influenzare significativamente le caratteristiche della dispersione se assume valori elevati. Al contrario, la presenza di altri elementi in grado di generare perdite localizzate ha un impatto contenuto sulle distanze. Inoltre, va sottolineato che il numero di questi componenti non è facilmente generalizzabile, dal momento che dipende fortemente dal layout specifico dell'impianto.

Per la costruzione del database verrà dunque considerata la lunghezza del tratto di tubazione come parametro di ingresso.

3.2 Selezione dei risultati

L'obiettivo del presente lavoro è fornire uno strumento per valutazioni preliminari, utile per comprendere l'influenza di alcune grandezze chiave nella modellazione delle conseguenze. La scelta dei risultati che le reti neurali sviluppate devono restituire è quindi fondamentale per fornire una stima delle zone coinvolte dagli scenari di rilascio.

Nel presente lavoro, alcuni risultati vengono ricavati dal report in formato tabellare prodotto da PHAST e altri, invece, dai grafici ottenuti.

3.2.1 Portata di rilascio e distanze

I valori della portata di rilascio e delle distanze massime raggiunte ai limiti di infiammabilità vengono estratti dal report tabellare prodotto da PHAST.

La portata di rilascio è quella massima che si verifica all'istante iniziale; tuttavia, per condizioni stazionarie, viene assunta costante per l'intera durata dell'evento. Rappresenta l'output principale nella modellazione del rilascio e costituisce l'input necessario per la successiva simulazione della dispersione. In questo contesto, la stima della portata di rilascio è importante perché può essere utilizzata anche come parametro di ingresso per la stima delle conseguenze con altri software. Per lo scenario di rottura catastrofica la portata di rilascio non viene considerata dato che viene assunto un rilascio istantaneo di tutto il contenuto del serbatoio.

Le distanze massime raggiunte sono risultati tipici della modellazione delle conseguenze, poiché permettono di determinare l'estensione spaziale dell'area pericolosa. Delimitano la zona entro la quale può verificarsi l'innescò della sostanza infiammabile, che può comportare incendio o esplosione, oltre a fornire indicazioni utili per valutare potenziali effetti domino, vale a dire la possibilità che un evento incidentale primario vada a coinvolgere altre apparecchiature dell'impianto. Infine, permettono di pianificare le distanze di separazione e altre misure di protezione. Nello specifico, le distanze massime individuano, lungo la direzione di rilascio, le posizioni in cui vengono raggiunti i valori di soglia. Nonostante l'informazione relativa alla distanza massima corrispondente al limite inferiore di infiammabilità LFL (*Lower Flammability Limit*) possa essere sufficiente a caratterizzare la zona interessata, per completezza vengono calcolate anche quelle relative al limite superiore UFL (*Upper Flammability Limit*) e a metà del limite inferiore di infiammabilità, $1/2$ LFL. In molti casi il limite superiore non viene raggiunto all'altezza di riferimento, comportando la presenza di zeri nel dataset. Per gli scenari che presentano un'elevata percentuale di distanze nulle non viene allenata una rete per la predizione di questo target, a causa di dati insufficienti. La distanza massima corrispondente a metà del limite inferiore di infiammabilità è, invece, un output comune per la modellazione della dispersione poiché fornisce un margine conservativo ed è previsto come limite predefinito in PHAST per la visualizzazione grafica delle zone di effetto.

3.2.2 Rappresentazione grafica della nube

Tra i grafici prodotti da PHAST per la caratterizzazione della dispersione si possono distinguere i seguenti:

- *Side view*: è la visione laterale della nube, rappresentata con contorni chiusi che delimitano l'area all'interno della quale viene superata la soglia di concentrazione definita.

- *Maximum concentration vs distance*: rappresenta un inviluppo delle concentrazioni massime registrate in ogni posizione lungo la direzione di rilascio.
- *Cloud maximum footprint*: rappresenta la massima proiezione dell'estensione della nube sul piano orizzontale all'altezza di riferimento in cui la concentrazione di soglia viene superata anche solo per un istante.

Il grafico considerato più rilevante per la definizione delle aree di danno è la proiezione della nube sul piano orizzontale. In questo caso, l'altezza di riferimento per i risultati è stata posta pari a 1 metro, come indicazione delle linee guida [17], e la concentrazione di soglia utilizzata è il limite inferiore di infiammabilità.

PHAST offre due opzioni di visualizzazione per la proiezione: *shape* o *effect zone*. Nel primo caso si ottiene la forma reale rispetto alla direzione di rilascio, mentre nel secondo un'area semplificata ottenuta dall'inviluppo delle impronte per tutte le direzioni di rilascio. Dalla visualizzazione *shape*, vengono estratti i valori relativi alla larghezza massima della nube e l'area interessata, da integrare con i risultati ricavati dal report. A partire dalla predizione della distanza massima raggiunta dal limite inferiore di infiammabilità e della larghezza massima viene costruita un'ellissi, di cui si calcola l'area che può essere confrontata con quella fornita da PHAST per valutare la possibilità di approssimare la nube con una forma ellittica. Se è possibile ricostruire l'*effect zone* tracciando un'area circolare di raggio pari alle distanze massime raggiunte, risulta invece più complesso ottenere la forma reale della nube rispetto a una determinata direzione di rilascio, dal momento che le reti neurali sviluppate gestiscono solo parametri numerici. Il limitato numero di simulazioni a disposizione non consente, infatti, lo sviluppo di reti che possano produrre come output grafici o immagini.

3.3 Dataset

Gli scenari di rilascio analizzati sono denominati come segue:

- *Leak*: rilascio da foro sul serbatoio (*Orifice model*)
- *Catastrophic rupture*: rilascio dovuto a rottura catastrofica del serbatoio (*Instantaneous release*)
- *Line rupture*: rilascio dovuto a rottura completa della tubazione connessa al serbatoio (*Short pipe model*)
- *Relief valve*: rilascio da valvola di sicurezza (*Short pipe model*)
- *Disk rupture*: rilascio da disco di rottura (*Short pipe model*)

Per tutti gli scenari, fatta eccezione per la rottura catastrofica, è possibile scegliere di modellare il rilascio come stazionario o transitorio. La modellazione in condizioni stazionarie risulta conservativa e, dunque, viene ritenuta una scelta adeguata nell'ambito dell'analisi del rischio. Inoltre, questa scelta comporta che la quantità di idrogeno contenuta nel serbatoio, espressa come massa o volume, non influisca sui risultati, come verificato empiricamente, e sia pertanto considerata come parametro di ingresso solo nel caso di rottura catastrofica.

Per ciascuna grandezza da variare per la costruzione dei dataset vengono definiti valori plausibili per le stazioni di rifornimento di idrogeno: le simulazioni saranno caratterizzate da diverse combinazioni di questi. Ogni simulazione viene quindi descritta tramite i valori delle variabili di ingresso e corrisponde ad una riga di una tabella realizzata con Excel, così da poter aggiungere in seguito i corrispondenti risultati ottenuti.

La pressione di esercizio risulta uno dei parametri più importanti nelle simulazioni. Nella definizione dei valori da assegnare ad essa, si è cercato di coprire l'intervallo di livelli riscontrabili nelle stazioni di rifornimento. Come riportato da Genovese et al. [20] un intervallo di pressioni tra 20 e 900 bar permette di considerare, rispettivamente, la pressione dell'idrogeno in arrivo alla stazione e il valore massimo che viene raggiunto dopo la compressione.

Le condizioni meteorologiche utilizzate per la costruzione dei dataset sono 1,5/F e 2,5/D, selezionate in seguito all'analisi di sensibilità.

Per quanto riguarda l'altezza del punto di rilascio, l'analisi di sensibilità condotta per lo scenario *Leak* ha mostrato che si tratta di un parametro abbastanza influente poiché modifica la geometria della nube. Tuttavia, è stata variata solo per questo scenario, mantenendola costante negli altri casi, al fine di ridurre la complessità dei modelli.

La direzione di rilascio adottata per le simulazioni è sempre orizzontale; le linee guida [17] suggeriscono si tratti di un'ipotesi conservativa quando la direzione di rilascio reale non è nota. Inoltre, è necessario specificare che la direzione di rilascio è parallela a quella del vento, dal momento che PHAST non ne consente la variazione.

Costruiti i dataset per ogni scenario e riportati i risultati di ogni simulazione, è stata calcolata la matrice di correlazione: una tabella quadrata che riporta l'elenco delle variabili come intestazioni di righe e colonne e in ciascuna cella è presente il coefficiente di correlazione di Pearson della coppia corrispondente. Sulla diagonale principale, che rappresenta la correlazione di una variabile con se stessa, il valore è pari a 1.

Il coefficiente di correlazione di Pearson è un indicatore che descrive l'intensità e la direzione della relazione lineare tra due variabili e che può assumere un valore compreso tra -1 e +1:

- Un valore pari a +1 indica una correlazione lineare positiva perfetta, vale a dire che al crescere di una variabile cresce in proporzione anche l'altra.
- Un valore pari a -1 indica correlazione lineare negativa perfetta, cioè al crescere di una variabile l'altra decresce in modo proporzionale.
- Un valore prossimo a 0 indica che non c'è alcuna correlazione lineare, sebbene possano esistere forme di dipendenza non lineare che in questo contesto non vengono considerate.

La matrice di correlazione è uno strumento utile per individuare variabili fortemente correlate, per riconoscere l'assenza di legame tra input e output e supportare quindi la selezione delle feature e l'interpretazione dei dati.

Le *heatmap* corrispondenti alle matrici di correlazione consentono di visualizzare l'intensità del legame tra le variabili e sono riportate in Appendice A per ogni scenario.

3.3.1 Leak

La modellazione del rilascio da un foro di perdita su un serbatoio viene effettuata con lo scenario *Leak* del modello *Orifice*. Il coefficiente di efflusso è stato posto pari a 0,62 [17] e l'altezza di riferimento per i risultati pari a 1 m.

Grazie all'analisi di sensibilità condotta preliminarmente, sono stati individuati i seguenti parametri da variare, insieme ai valori utilizzati, per la costruzione del dataset relativo a questo scenario di rilascio:

- Pressione di esercizio: 30, 50, 70, 100, 350, 500, 700 e 800 bar.
- Velocità del vento: $1,5/F$ e $2,5/D$.
- Diametro del foro di perdita: 5, 10 e 25,4 mm.
- Altezza di rilascio: 0, 0,5 e 1 m.

Le simulazioni da effettuare sono state definite combinando tutti i parametri sopra elencati, ottenendo un dataset costituito da 144 casi. Le caratteristiche del dataset di questo scenario sono descritte in Appendice A.

Dalla matrice di correlazione, presente in Appendice A, emerge che la portata di rilascio è la variabile chiave che collega le condizioni di stoccaggio agli effetti della dispersione.

Il diametro del foro di perdita, essendo ben correlato con la portata di rilascio, lo sarà anche con gli altri parametri di uscita.

L'altezza di rilascio, che per questo scenario è stata variata, mostra correlazioni modeste e appare più legata alla distanza massima a cui viene raggiunto il limite

superiore di infiammabilità. Inoltre, se l'altezza di rilascio coincide con l'altezza di riferimento per la visualizzazione dei risultati, nel punto di rilascio il modello riporta concentrazione pari al 100%; il software, in questo caso, restituisce la distanza alla quale la concentrazione è scesa fino al limite superiore di infiammabilità.

Si osserva una correlazione modesta della velocità del vento con le distanze raggiunte mentre risulta pressoché nulla con la larghezza massima. Quest'ultimo risultato è coerente con la limitazione del modello di PHAST che assume la direzione del vento allineata a quella di rilascio.

3.3.2 Catastrophic rupture

La rottura catastrofica di un serbatoio in pressione viene simulata con il modello di rilascio istantaneo, che richiede un numero limitato di parametri di input rispetto agli altri modelli.

L'altezza del punto di rilascio, in questo caso, ha poca influenza sui risultati e non è quindi stata considerata.

Per questo scenario è necessario specificare la quantità di idrogeno compresso contenuta nel serbatoio. Questa, è stata espressa in termini di massa, prendendo come esempio alcuni valori riportati in [21], [22] e altri valori tipici di massa stoccata, coerenti con la pressione considerata. L'altezza di riferimento per i risultati è pari a 1 m.

A differenza dello scenario *Leak*, non sono state effettuate le simulazioni per tutte le combinazioni dei valori assunti per i diversi parametri; di seguito sono quindi riportati gli intervalli e non i singoli valori utilizzati:

- Pressione di esercizio: 20 - 900 bar.
- Massa stoccata: 20-1000 kg.
- Velocità del vento: 1,5/F e 2,5/D.

Il dataset ottenuto per questo scenario è costituito da 48 simulazioni. Le caratteristiche del dataset di questo scenario sono descritte in Appendice A.

La matrice di correlazione mostra una correlazione positiva tra la pressione di stoccaggio e la massa contenuta all'interno del serbatoio: questo risultato riflette il dominio di simulazione adottato, per il quale configurazioni con masse elevate e pressioni molto basse non sono state considerate poiché non realistiche dal punto di vista operativo. Inoltre, le grandezze target presentano correlazioni più elevate con la massa stoccata che con la pressione.

3.3.3 Line rupture

La modellazione del rilascio dovuto alla rottura di un tratto di tubazione viene effettuata con lo scenario *Line rupture* del modello *Short pipe*. Per evitare un numero eccessivo di simulazioni necessarie, non è stata variata l'altezza di rilascio, posta costante e pari a 1 m. L'altezza di riferimento per i risultati è anch'essa pari a 1 m.

Grazie all'analisi di sensibilità condotta preliminarmente, volta a valutare la necessità di una rete dedicata distinta da quella dello scenario *Leak*, sono stati definiti i seguenti parametri da variare insieme ai valori utilizzati, per la costruzione del dataset relativo a questo scenario di rilascio:

- Pressione di esercizio: 40, 100, 350 e 700 bar.
- Velocità del vento: 1,5/F e 2,5/D.
- Diametro della tubazione: 26,7, 40,9 e 52,5 mm, corrispondenti rispettivamente ai diametri interni per DN 25, DN 40 e DN 50.
- Lunghezza del tratto di tubazione: 0,5, 2 e 6 m.

Per i tre diametri caratteristici della tubazione si è fatto riferimento al diametro interno, determinato sulla base delle dimensioni reali delle tubazioni in Schedule 40.

Le simulazioni da effettuare sono state definite combinando tutti i parametri sopra elencati, ottenendo un dataset costituito da 72 casi. Le caratteristiche del dataset di questo scenario sono descritte in Appendice A.

La matrice di correlazione per questo scenario evidenzia come la pressione sia il parametro di ingresso dominante sui risultati. Anche il diametro della tubazione ha un effetto rilevante mentre la lunghezza del tratto sembra non influire in modo significativo in termini lineari, nonostante sia associata a maggiori perdite di carico e quindi a una minore portata di rilascio.

In questo caso, la velocità del vento presenta elevata correlazione con la distanza massima alla quale viene raggiunto il limite superiore di infiammabilità poiché l'altezza di rilascio coincide con quella di riferimento dei risultati, rendendo l'effetto del vento più marcato.

3.3.4 Disk rupture

Lo scenario di rilascio da disco di rottura è simulato con il modello *Short pipe*. Per questo scenario non è stata variata l'altezza di rilascio, posta costante e pari a 1 m. Sebbene il modello *Short pipe* permetta di variare la lunghezza del tratto di tubazione che collega il serbatoio al disco di rottura, in questo caso la lunghezza

viene considerata pari a 0,1 m, cioè il valore minimo, ipotizzando che, come le valvole di sicurezza, il disco di rottura venga installato praticamente a ridosso del serbatoio. Si assume, inoltre, che il diametro del disco sia pari a quello delle tubazioni sulle quali viene installato. L'altezza di riferimento per i risultati è pari a 1 m.

Si possono definire i seguenti parametri da variare, insieme ai valori utilizzati, per la costruzione del dataset relativo a questo scenario di rilascio:

- Pressione di esercizio: 40, 100, 350 e 700 bar.
- Velocità del vento: 1,5/F e 2,5/D.
- Diametro del disco di rottura: 26,7, 40,9 e 52,5 mm, corrispondenti rispettivamente ai diametri interni per DN 25, DN 40 e DN 50.

Le simulazioni da effettuare sono state definite combinando tutti i parametri sopra elencati, ottenendo un dataset costituito da 24 simulazioni. Una di esse è stata però esclusa poiché i risultati riportati non erano corretti. Le caratteristiche del dataset di questo scenario sono descritte in Appendice A.

La matrice di correlazione indica che la pressione è il parametro di ingresso dominante sui risultati anche per questo scenario, seguita dal diametro del disco.

Non è stata considerata nello sviluppo della rete la distanza al limite superiore di infiammabilità UFL, dal momento che per tutte le simulazioni del dataset questa variabile risulta nulla.

3.3.5 Relief valve

Lo scenario di rilascio da valvola di sicurezza è simulato con il modello *Short pipe*. Per questo scenario non è stata variata l'altezza di rilascio, posta costante e pari a 1 m.

Sebbene il modello *Short pipe* permetta di variare la lunghezza del tratto di tubazione che collega il serbatoio alla valvola, in questo caso la lunghezza viene posta pari a 0,1 m, cioè il valore minimo, ipotizzando che le valvole di sicurezza vengano installate a ridosso del serbatoio [23].

Come unico parametro geometrico rilevante è stato considerato il diametro dell'orificio della valvola di sicurezza, trascurando l'influenza del diametro della tubazione a monte, essendo ridotta la lunghezza del tratto e quindi di influenza limitata sui risultati. Per la definizione dei diametri sono state considerate delle taglie API [24]. L'altezza di riferimento per i risultati è pari a 1 m.

Si possono definire i seguenti parametri da variare, insieme ai valori utilizzati, per la costruzione del dataset relativo a questo scenario di rilascio:

- Pressione di esercizio: 40, 100, 350 e 700 bar.

- Velocità del vento: 1,5/F e 2,5/D.
- Diametro dell'orifizio: 9,5, 12,7 e 15,9 mm corrispondenti rispettivamente alle taglie API D, E e F.

Le simulazioni da effettuare sono state definite combinando tutti i parametri sopra elencati, ottenendo un dataset costituito da 24 simulazioni.

La matrice di correlazione è simile a quella dello scenario Disk rupture.

Il diametro dell'orifizio mostra una correlazione più elevata con i risultati rispetto al diametro del disco di rottura.

Non è stata considerata nello sviluppo della rete la distanza al limite superiore di infiammabilità UFL, dal momento che per tutte le simulazioni del dataset questa variabile risulta nulla.

Capitolo 4

Sviluppo del modello

4.1 Reti neurali e regressione

Per la definizione di un modello predittivo è necessario individuare le relazioni tra le variabili di input e quelle di output, così da poterle utilizzare per fare previsioni su nuovi dati. A tale scopo si possono adottare metodologie di diversa complessità, come approcci lineari o non lineari.

Il *machine learning* è un sottoinsieme dell'intelligenza artificiale che comprende algoritmi in grado di apprendere dai dati senza essere programmati in maniera esplicita. Questi algoritmi vengono spesso definiti “*black-box*”, dal momento che le relazioni tra le variabili di input e quelle di output non sono direttamente interpretabili dall'utente.

Le tecniche di *machine learning* possono essere classificate in base al tipo di apprendimento:

- **Supervisionato:** i dati utilizzati sono etichettati, in modo che a ciascun input corrisponda un output noto.
- **Non supervisionato:** i dati non sono etichettati e l'algoritmo individua autonomamente relazioni e strutture.
- **Per rinforzo:** l'algoritmo apprende interagendo con un ambiente, procedendo per tentativi sulla base di ricompense o penalità.

Il *deep learning* rappresenta a sua volta un sottoinsieme del *machine learning* che utilizza reti neurali artificiali con più strati per estrarre conoscenza.

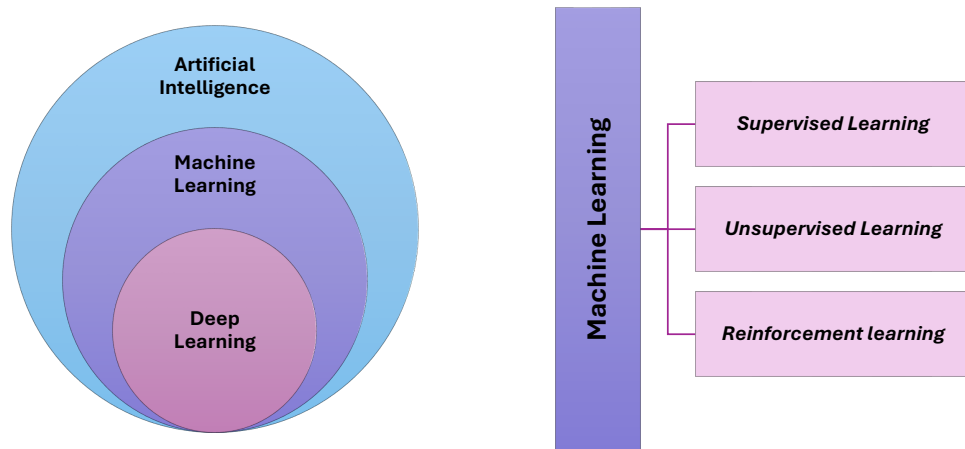


Figura 4.1: Rappresentazione schematica delle relazioni tra Intelligenza Artificiale, *Machine Learning* e *Deep Learning* (a sinistra) e classificazione delle principali modalità di apprendimento nel *Machine Learning* (a destra).

L'unità base di una rete neurale è il neurone. Ciascun neurone riceve in ingresso un certo numero di dati, a ciascuno dei quali viene associato un peso, che rappresenta l'importanza relativa di quella variabile nel calcolo. Alla combinazione lineare degli input ponderati viene aggiunto solitamente un bias, vale a dire un termine costante che consente di aumentare la flessibilità del modello. Infine, il risultato viene trasformato impiegando una funzione di attivazione opportuna che introduce non linearità all'interno del modello: in assenza di tale elemento, la rete si ridurrebbe a una semplice regressione lineare.

Il modello **Multi-layer Perceptron**, rappresentato in Figura 4.2 è un tipo di rete neurale costituita da diversi strati (*layers*) di neuroni e diverse connessioni tra di essi. Si distinguono tre tipi di strati:

- *Input layer*: dove vengono introdotti i dati e il cui numero di neuroni è pari alla dimensione di questi.
- *Hidden layers*: strati interni dove vengono effettuati i calcoli per processare le informazioni in ingresso.
- *Output layer*: strato finale che restituisce le predizioni.

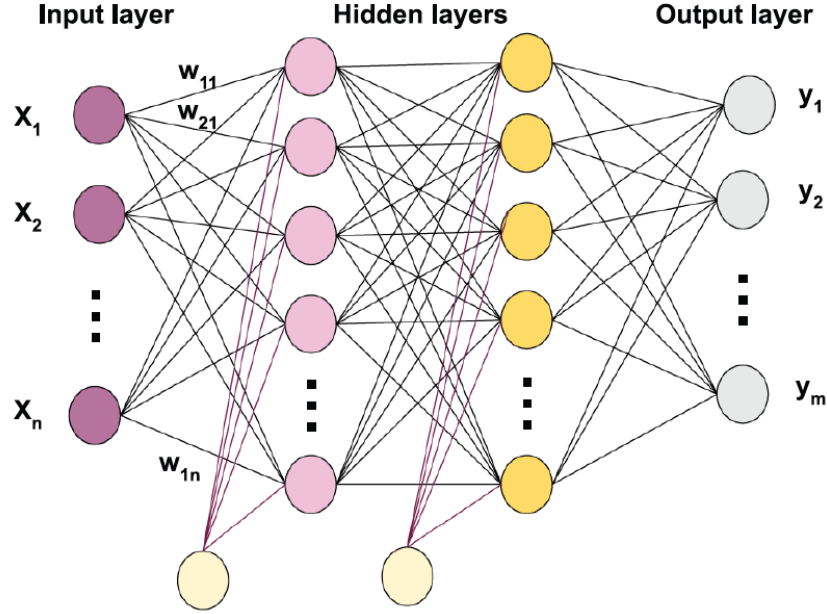


Figura 4.2: Architettura generale di un Multi-layer Perceptron [25].

Il numero di strati e il numero di neuroni per ciascuno di essi, sono iperparametri del modello. Gli iperparametri sono variabili di configurazione che non variano durante il processo di addestramento ma la cui definizione influenza le prestazioni del modello. In questa categoria rientrano anche la dimensione del *batch*, ossia il numero di dati che contiene ciascuna porzione in cui viene suddiviso il dataset utilizzato per l'addestramento, e il numero di epoche, cioè il numero di passaggi completi di questo attraverso la rete.

Fissati gli iperparametri, la rete deve apprendere i parametri del modello sulla base dei dati ad essa forniti. In questa categoria rientrano pesi e bias, che vengono inizializzati in maniera randomica e aggiornati di volta in volta durante il processo di addestramento.

L'allenamento del modello si basa sulla stima dell'errore della predizione rispetto al valore atteso, che viene quantificato con un'opportuna funzione di costo (*loss function*) e il cui monitoraggio durante il processo permette di capire se la rete sta imparando dai dati forniti. L'obiettivo è trovare un minimo globale della funzione di costo e ciò avviene impiegando un ottimizzatore che individua la direzione per raggiungerlo il prima possibile. I parametri sono quindi aggiornati con la seguente regola:

$$w_i^{n+1} = w_i^n - \alpha \frac{\partial L}{\partial w_i} \quad (4.1)$$

dove L è il valore della funzione di costo e α rappresenta un altro iperparametro che controlla la velocità di apprendimento, definendo l'entità della modifica dei parametri ad ogni iterazione.

Il calcolo dei gradienti avviene mediante *backpropagation*, ossia propagando l'errore dall'uscita verso l'ingresso e aggiornando in questo modo i parametri lungo il percorso. Ottenuti i parametri ottimizzati, la rete può essere testata su un'altra serie di dati disponibili per valutarne la capacità di generalizzazione.

Nel presente lavoro, prima di procedere con lo sviluppo di reti neurali, è stato definito per ciascuno scenario un modello di riferimento (*baseline*) mediante regressione lineare. Si tratta di un metodo largamente impiegato, sia in statistica sia nel *machine learning* [26], per determinare la relazione lineare tra variabili di ingresso e target di interesse. L'ipotesi di base è che una combinazione lineare delle prime possa stimare un valore continuo del secondo. Il vantaggio che presenta questo approccio, rispetto a modelli più complessi, è quello di essere interpretabile. La relazione stimata si esprime come:

$$Y = w_0 + w_1x_1 + w_2x_2 + \dots + w_nx_n \quad (4.2)$$

dove w_0 è l'intercetta e w_i sono i coefficienti da moltiplicare per ciascuna variabile di input.

Il modello di regressione è stato sviluppato in Excel ed è stato confrontato con le reti neurali implementate in Python. La decisione di ricorrere a modelli più complessi è stata motivata dall'analisi delle metriche di errore: per gli scenari con dataset più ampi, le reti neurali hanno consentito di cogliere relazioni non lineari, comportando un miglioramento dell'accuratezza predittiva.

4.2 Ambiente di sviluppo delle reti e librerie

Le reti neurali sono state sviluppate con Python, linguaggio di programmazione ampiamente utilizzato nel campo del *machine learning*. In particolare, per la scrittura e l'esecuzione del codice sono stati predisposti dei file Jupyter Notebook, un'applicazione *open source* che consente di organizzare il codice in celle, combinarlo con testi descrittivi e visualizzare i risultati. Per lavorare con Jupyter Notebook è stato utilizzato PyCharm, un ambiente di sviluppo integrato (IDE) utile per la gestione dei file, sviluppato da JetBrains.

Inizialmente è stato creato un ambiente virtuale Virtualenv, opzione predefinita di PyCharm: ciò consente di realizzare un ambiente virtuale isolato specifico per il progetto che consiste in un *interpreter* e i pacchetti installati, così da gestire le impostazioni in maniera indipendente da altri progetti [27]. L'interprete utilizzato per questo ambiente virtuale è Python versione 3.11 e vengono installate le librerie

utili in questo contesto, necessarie per il pre-processing, l'addestramento e la valutazione delle reti.

4.2.1 Manipolazione dei dati

Per la manipolazione dei dati sono state installate le librerie Pandas e Numpy, entrambe indispensabili per i calcoli e l'analisi dei dati. Le versioni utilizzate sono:

- Pandas 2.2.3
- Numpy 2.1.3

In particolare, Pandas è stato necessario per la creazione di DataFrame, una struttura bidimensionale di dati organizzata in righe e colonne, a partire dai dataset forniti sotto forma di tabelle Excel.

```
1 import pandas as pd
2 df = pd.read_excel("leak.xlsx")
```

Listato 4.1: Caricamento del dataset.

Numpy è stato utilizzato, ad esempio, per applicare trasformazioni matematiche ai dati.

```
1 import numpy as np
2 if apply_log:
3     y_real = np.expm1(y_real)
4     y_pred_real = np.expm1(y_pred_real)
```

Listato 4.2: Applicazione della trasformazione inversa logaritmica per i target interessati.

4.2.2 Visualizzazione

Per la generazione di grafici utili durante lo sviluppo delle reti, ad esempio con scopi di diagnostica e valutazione delle prestazioni, è stata impiegata la libreria Matplotlib, versione 3.9.3. Come riporta la documentazione [28] questa libreria presenta un'interfaccia ispirata a MATLAB per creare grafici in maniera semplice ed interattiva. Si riporta in Figura 4.3 il grafico generato dal codice in Listato 4.3.

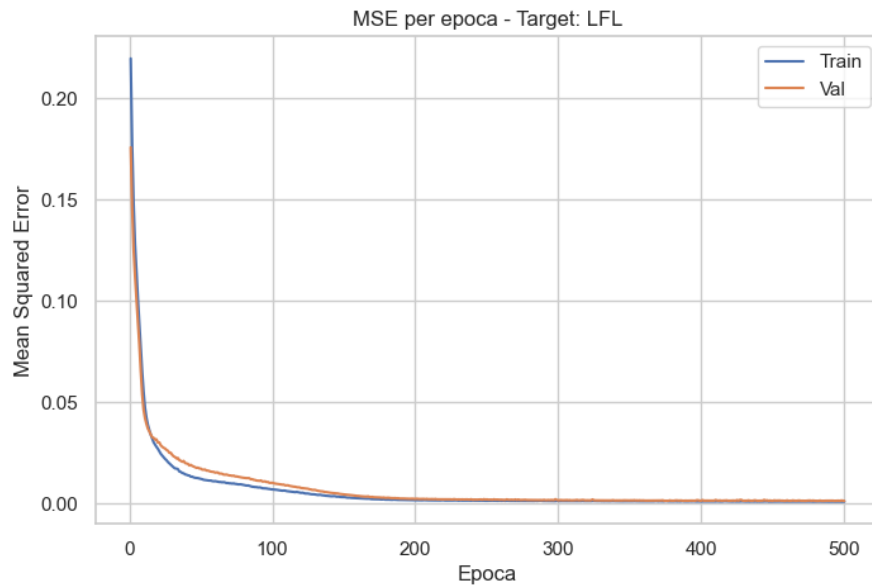


Figura 4.3: Visualizzazione dell'andamento della funzione di costo durante il training.

```

1 plt.figure(figsize=(8,5))
2 plt.plot(range(1, epochs+1), train_losses, label="Train")
3 plt.plot(range(1, epochs+1), val_losses, label="Val")
4 plt.title(f'MSE per epoca - Target: {target_col}')
5 plt.xlabel('Epoca')
6 plt.ylabel('Mean Squared Error')
7 plt.legend()
8 plt.grid(True)
9 plt.show()

```

Listato 4.3: Creazione del grafico dell'andamento della funzione di costo.

4.2.3 Deep Learning

Per lo sviluppo vero e proprio delle reti si è deciso di impiegare la libreria PyTorch, versione 2.5.1. Come riportato in [29] si tratta, infatti, di una libreria che coniuga efficienza e semplicità di utilizzo, ed è quindi stata preferita a TensorFlow, anch'essa utilizzata per il deep learning.

La costruzione, l'addestramento e infine la valutazione delle reti neurali, sono state effettuate importando i moduli di PyTorch `torch.nn`, `torch.optim` e `torch.utils.data`.

I modelli sono stati definiti creando una classe che eredita le regole di `nn.Module`, la classe base di PyTorch, per tutti i moduli di rete neurale. La definizione del modello è riportata in Listato 4.4.

```
1 # Definizione del modello generale
2 class Model(nn.Module):
3
4     def __init__(self, input_size, hidden_sizes, output_size,
5         target_type="default"):
6
7         super(Model, self).__init__()
8
9         # Primo hidden layer
10        self.hidden1 = nn.Linear(input_size, hidden_sizes[0])
11
12        # Secondo hidden layer
13        self.hidden2 = nn.Linear(hidden_sizes[0], hidden_sizes[1])
14
15        # Output layer (lineare)
16        self.output = nn.Linear(hidden_sizes[1], output_size)
17
18        self.target_type = target_type
19        if self.target_type == "flow":
20            self.activation_out = nn.Softplus()
21        else:
22            self.activation_out = nn.Identity()
23
24    def forward(self, x):
25        x = torch.tanh(self.hidden1(x))
26        x = torch.relu(self.hidden2(x))
27        x = self.output(x)
28        x = self.activation_out(x)
29        return x
```

Listato 4.4: Definizione del modello generale.

La struttura generale presenta l’inizializzazione del modello, in cui vengono definiti layers e relative dimensioni e una funzione per descrivere il passaggio dei dati attraverso la rete (“*forward*”). Nel codice il modello viene impostato in due modalità differenti: *training* ed *evaluation*.

Nel ciclo di allenamento si impiega il modello nella configurazione `model.train()` e ciò influenza il comportamento di layer speciali adibiti a:

- *Dropout*: nella fase di allenamento si disattivano casualmente alcuni neuroni per evitare *overfitting* del modello.
- *Batch normalization*: per rendere l'allenamento più stabile, i dati in input al layer successivo vengono normalizzati per ogni *batch*.

Nel momento in cui è necessario validare o testare il modello, si impiega invece la configurazione `model.eval()`: il *dropout* viene disabilitato e vengono impiegate come media e varianza per la normalizzazione dei dati i valori calcolati e salvati durante la fase di training.

Definito il modello, è necessario scegliere un'opportuna funzione di costo per ottimizzare i suoi parametri [30]. In questo lavoro è stata impiegata `nn.MSELoss()`, che calcola l'errore medio quadratico tra output e target. In questo contesto, va definita la regola utilizzata per aggiornare i pesi, ossia l'ottimizzatore. Nel codice viene impiegato `optim.Adam()`, dal modulo `torch.optim`, basato sull'algoritmo Adam e spesso impiegato nel deep learning con PyTorch. In Listato 4.5 si riporta l'inizializzazione del modello.

```
1 if target_col == "Flow rate":
2     model = Model(
3         input_size=X.shape[1],
4         hidden_sizes=hidden_sizes,
5         output_size=1,
6         target_type="flow")
7 else:
8     model = Model(
9         input_size=X.shape[1],
10        hidden_sizes=hidden_sizes,
11        output_size=1)
12
13 optimizer = optim.Adam(model.parameters(), lr=learning_rate)
14
15 criterion = nn.MSELoss()
```

Listato 4.5: Inizializzazione del modello.

All'interno del ciclo di allenamento si fa uso di altre funzionalità di PyTorch (Listato 4.6).

```
1 # Forward pass
2 output = model(inputs)
3
4 # Calcolo loss
5 loss = criterion(output, targets)
6
7 # Backward pass e ottimizzazione
8 optimizer.zero_grad()
9 loss.backward()
10 optimizer.step()
```

Listato 4.6: Funzionalità di PyTorch utilizzate nella fase di addestramento.

Una volta ottenute le predizioni a partire dai valori di input, viene applicata la funzione di costo scelta. Per l'ottimizzazione è necessario azzerare i gradienti accumulati nei tensori dei pesi del modello con `optimizer.zero_grad()`; in seguito, si calcola il gradiente della funzione di costo rispetto a ciascun parametro con `loss.backward()` e, infine, con `optimizer.step()` si aggiornano i parametri sulla base dei gradienti calcolati e in base all'ottimizzatore scelto.

Un modulo fondamentale di PyTorch è `torch.utils.data`. Dopo aver trasformato i valori di input e di output, sia di training che di testing, in tensori, è stata utilizzata la funzione `DataLoader()` che costruisce sopra il dataset fornito una logica più complessa e utile per il training, ad esempio, mescolando i batch ad ogni epoca se richiesto [31].

4.2.4 Pre-processing e valutazione

La libreria Scikit-learn offre diverse funzionalità per il machine learning. In questo lavoro, è stata utilizzata la versione 1.5.2 nella fase di pre-processing dei dati e di valutazione dei modelli.

In una prima fase, si è provveduto a scalare i dati da utilizzare per la rete, attraverso la funzione `MinMaxScaler()` del modulo `sklearn.preprocessing`. In seguito, con `train_test_split()` di `sklearn.model_selection`, il dataset è stato suddiviso in tre insiemi distinti: 70% per l'addestramento, 15% per la validazione e 15% per il test; ciò ha permesso di disporre di un set di validazione per la definizione degli iperparametri e di un set indipendente per valutare la capacità di generalizzazione del modello, una volta scelti gli iperparametri.

Infine, da `sklearn.metrics`, sono state derivate le funzioni necessarie per il calcolo di alcune metriche comuni come `mean_squared_error()`, `r2_score()` e `mean_absolute_percentage_error()`, per valutare quantitativamente le prestazioni delle reti.

4.3 Struttura del codice

Lo strumento consiste in un codice Python nel quale vengono importati i modelli predittivi. Per ciascun target di ogni scenario sono stati sviluppati modelli distinti, così da semplificare l'architettura riducendo al tempo stesso il numero di simulazioni necessarie per l'allenamento.

Regressione lineare

Per gli scenari con un numero ridotto di campioni, ossia *Disk rupture* e *Relief valve*, la regressione multivariata è stata adottata come modello finale. In presenza di dataset limitati, infatti, l'impiego di reti neurali o di modelli più complessi rischierebbe di produrre metriche apparentemente buone, che in realtà riflettono un sovradattamento ai dati di training piuttosto che una reale capacità predittiva. I modelli regressivi sono stati implementati in Excel attraverso lo strumento *Analisi dati* e *Regressione*, incluso nei componenti aggiuntivi del programma. Con questa funzione possono essere stimati i coefficienti di una relazione lineare secondo il metodo dei minimi quadrati e vengono fornite alcune statistiche di supporto, tra cui R^2 che, insieme a MSE e $MAPE$, è stato impiegato come indicatore dell'accuratezza del modello. I coefficienti sono stati riportati in file JSON, così da integrare i modelli regressivi nel framework generale.

Reti neurali

Le singole reti neurali sono state addestrate con script distinti per ciascuno scenario. Il loro sviluppo è stato riservato agli scenari con dataset più ampi, per i quali il rischio di sovradattamento è ridotto e l'utilizzo di modelli più complessi risulta effettivamente vantaggioso per migliorare la capacità predittiva. La Figura 4.4 mostra lo schema che descrive l'organizzazione dello script per l'addestramento delle reti.

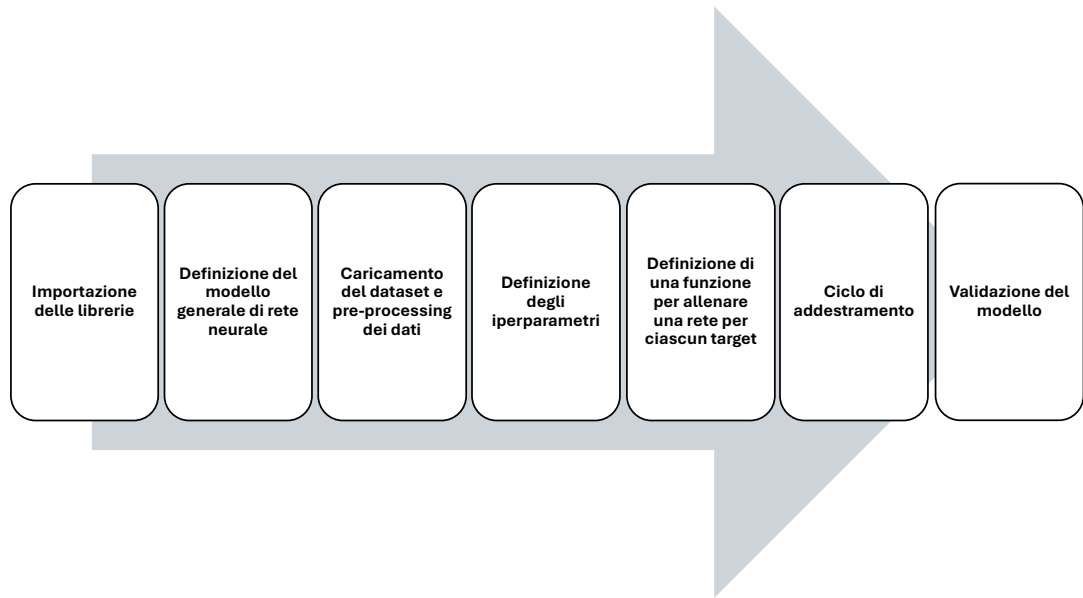


Figura 4.4: Struttura generale del codice per l'addestramento delle reti neurali.

La funzione utilizzata per l'allenamento delle reti è organizzata nel seguente modo:

1. Preparazione dei dati: gli input e gli output vengono normalizzati; per i target può inoltre essere applicata una trasformazione logaritmica.
2. Suddivisione in train, validation e test set: il dataset viene suddiviso in tre sottoinsiemi (70% train, 15% validation, 15% test) per disporre di un set dedicato alla definizione degli iperparametri adeguati ed uno indipendente per la valutazione finale, oltre a quello necessario per l'addestramento.
3. Conversione dei dati in tensori e creazione del Dataloader: i dati vengono trasformati e organizzati in strutture necessarie per l'addestramento dei modelli in PyTorch.
4. Inizializzazione del modello, scelta dell'ottimizzatore e della funzione di costo.
5. Ciclo di allenamento: viene calcolata la perdita media per epoca sul training set; parallelamente si monitora anche quella di validazione per diagnosticare fenomeni di *underfitting* o *overfitting*. I valori medi vengono riportati in un grafico per seguire l'andamento dell'apprendimento.
6. Salvataggio dei parametri ottimizzati e dei valori utilizzati per la normalizzazione.

7. Inferenza finale: vengono calcolate le predizioni sui tre set (train, validation, test).
8. Creazione Dataframe generale: vengono raccolti input, set di appartenenza, risultati reali e predetti per poter effettuare analisi successive.

Il salvataggio dei parametri ottimizzati e dei valori utilizzati per la normalizzazione è fondamentale, dal momento che questi vengono importati nel modulo centrale. La Figura 4.5 mostra lo schema che descrive l'organizzazione del codice del modulo centrale.

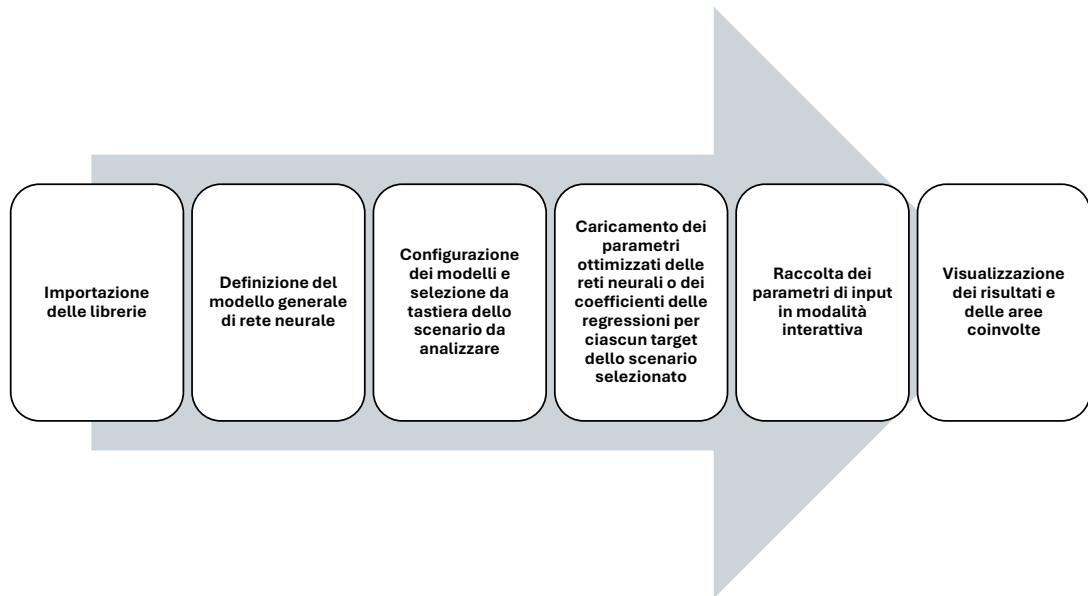


Figura 4.5: Struttura generale del codice del modulo centrale.

Il modello generale deve poter gestire l'importazione dei parametri ottimizzati sia delle reti neurali che dei modelli di regressione. Per la sua configurazione è stata definita una struttura basata su proposizioni condizionali del tipo *if-else*: in base allo scenario specificato dall'utente tramite input da tastiera, viene caricato il modello appropriato.

Per *Disk rupture* e *Relief valve* vengono importati i coefficienti della regressione associati a ciascun target. Per *Leak*, *Catastrophic rupture* e *Line rupture* vengono caricati i dataset utilizzati per l'addestramento delle reti, viene effettuato pre-processing dei dati, vengono definite le variabili di input e di output e vengono impostati gli iperparametri, in maniera analoga a quanto fatto nei singoli script. Per associare i parametri ottimizzati ed eventualmente i valori utilizzati per la normalizzazione, si costruisce un dizionario ordinato, ossia una struttura di Python

che organizza le informazioni utili per ciascun target.

Il caricamento dei modelli è gestito tramite un ciclo che itera sul dizionario di configurazione dello scenario selezionato. In questo modo, vengono caricate le informazioni relative al modello di ciascun target.

Per la raccolta degli input dell'utente, viene inizializzata una lista contenente tutte le variabili richieste dalle reti. Un ciclo *while* permette l'inserimento dei parametri richiesti, così da poter valutare diversi rilasci simultaneamente. Al suo interno, viene resettato il dizionario delle *features* immesse, viene effettuata la codifica della classe di stabilità immessa in forma letterale e vengono memorizzate le nuove variabili di ingresso. Inoltre, nel codice sono stati esplicitati gli intervalli associati ad ogni target e scenario: in generale sono stati fissati un valore minimo e uno massimo, mentre per l'altezza del punto di rilascio, la velocità del vento e la stabilità atmosferica sono stati previsti dei valori discreti, in modo che l'input dell'utente venga ricondotto a quello più vicino. Per le distanze è stata impostata una soglia minima per discriminare il raggiungimento o meno dei limiti di infiammabilità.

La fase di predizione prevede la costruzione di un vettore ordinato delle *features*, così da rispettare l'ordine richiesto dal modello specifico, la normalizzazione dei parametri e l'inferenza sui dati forniti, con eventuale trasformazione inversa logaritmica. Le distanze vengono salvate in un dizionario, inizializzato prima della predizione, contenente anche le coordinate del punto di rilascio.

Definiti i livelli corrispondenti ai limiti di infiammabilità, è stato predisposto un plot delle zone di effetto generate da ciascun rilascio, riportando in una griglia il punto di rilascio e le aree ad esso associate. Le zone sono rappresentate con cerchi che evidenziano l'involuppo delle impronte al suolo della nube per ciascuna direzione di rilascio.

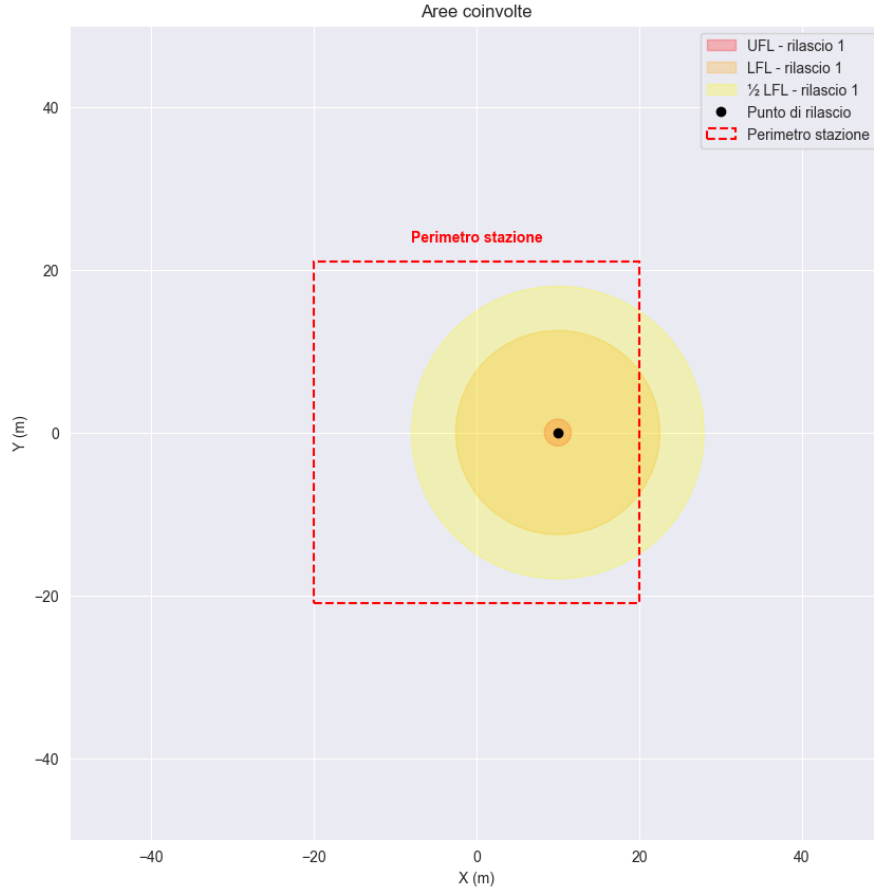


Figura 4.6: Grafico relativo a una rottura catastrofica di un serbatoio a 200 bar, con 50 kg di idrogeno stoccato e condizioni meteorologiche 1,5/F.

4.4 Implementazione delle reti neurali

La costruzione di reti neurali distinte per ciascun target di ogni scenario di rilascio si è rivelata una scelta efficace, in quanto ha permesso di semplificare l'architettura del modello e la fase di addestramento. Dal momento che i target presentavano scale differenti, infatti, l'impiego di una rete neurale multi-target avrebbe comportato prestazioni peggiori per alcune variabili a causa di conflitti in fase di ottimizzazione. Le reti neurali sono sensibili all'ordine di grandezza delle *features* di ingresso; infatti, nel momento in cui viene calcolata la funzione di costo confrontando i valori reali e quelli predetti, le variabili con magnitudo più elevata finiscono per dominare la perdita complessiva, sbilanciando l'ottimizzazione. La normalizzazione dei dati, perciò, è fondamentale nelle fasi preliminari di sviluppo di algoritmi di questo tipo.

La tecnica di normalizzazione impiegata è la Min-Max che consiste nel riscalare i valori all'interno di un intervallo predefinito con la seguente formula:

$$x_{scaled} = \frac{x - x_{min}}{x_{max} - x_{min}} \quad (4.3)$$

Se l'intervallo è $[0,1]$ al valore minimo corrisponde 0 mentre a quello massimo 1. I valori massimi e minimi utilizzati per la normalizzazione di ciascuna rete sono parametri molto importanti per la riproducibilità del modello e, per questo, vengono opportunamente salvati assieme ai parametri ottimizzati.

Le portate di rilascio sono state trasformate in scala logaritmica e, per lo scenario *Leak*, la trasformazione è stata applicata anche al target *Max width*. La trasformazione logaritmica è una tecnica matematica utilizzata nell'analisi dei dati che risulta utile quando la distribuzione dei dati è asimmetrica ed include diversi ordini di grandezza. Consiste nella sostituzione di ogni valore nel seguente modo:

$$x \rightarrow \ln(1 + x) \quad (4.4)$$

Il risultato che si ottiene è una distribuzione dei dati più vicina ad una distribuzione normale, in modo tale che gli errori vengano penalizzati in maniera bilanciata. In Figura 4.7 viene riportata la distribuzione originale e quella ottenuta a seguito della trasformazione logaritmica per il target *Max width* dello scenario di rilascio *Leak*.

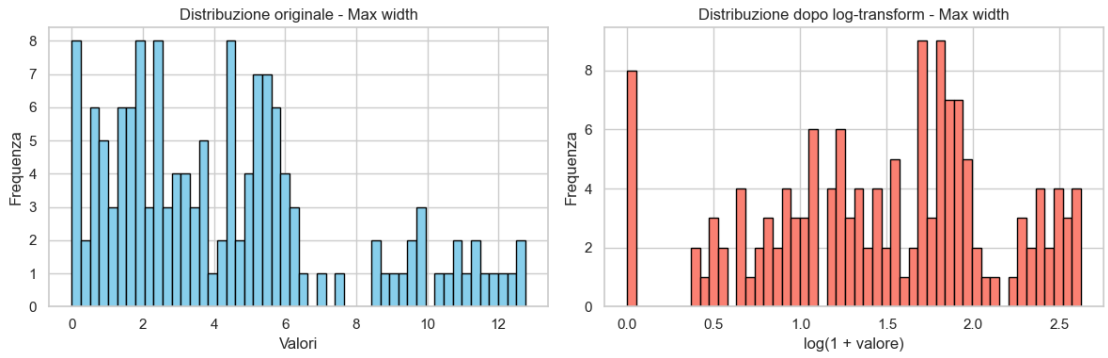


Figura 4.7: Confronto distribuzioni per il target *Max width* dello scenario di rilascio *Leak*.

Per la configurazione dei modelli è necessario definire gli iperparametri, vale a dire quelle variabili che non variano durante il processo di addestramento delle reti neurali ma che ne influenzano le prestazioni. Gli iperparametri da definire sono:

- Numero di strati interni (*hidden layers*)

- Numero di neuroni per ciascuno strato
- Dimensione del *batch*
- Tasso di apprendimento (*learning rate*)
- Numero di epoche

Per ridurre la complessità, in generale gli iperparametri sono stati settati in base a valori tipici [11], riportati in Appendice B, mentre per definire il numero di neuroni di ciascuna rete è stata condotta un'analisi di sensibilità, al fine di individuare il valore ottimale. In generale i modelli prevedono due strati interni, mentre in casi particolari ne è stato definito solo uno in base alle caratteristiche del target e dello scenario; ad esempio, per *Flow rate* è stato utilizzato un singolo strato interno per ridurre la complessità del modello quando non necessaria. Inizialmente sono state definite architetture semplici, con numero ridotto di neuroni per ciascuno strato interno, incrementandolo progressivamente. La scelta finale è stata guidata dalle metriche di prestazione ottenute per ciascuna configurazione, privilegiando la soluzione più semplice in grado di garantire un livello di accuratezza soddisfacente. In particolare, sono stati valutati congiuntamente i valori di R^2 , MSE e $MAPE$ per evitare che il modello risultasse ottimizzato solo rispetto a una metrica.

Per valutare la presenza o meno di fenomeni di *underfitting* o *overfitting*, durante la scelta dell'architettura è stato analizzato l'andamento della funzione di perdita al variare delle epoche per training e validation. Si parla di *underfitting* quando la rete non riesce ad apprendere efficacemente dai dati forniti: in questo caso la perdita si mantiene elevata sia per il training sia per la validation, senza diminuire significativamente con il numero di epoche. L'*overfitting*, invece, si manifesta quando il modello non è in grado di generalizzare per dati mai visti prima: la perdita di training continua a diminuire mentre quella di validation tende a stabilizzarsi o peggiorare, segnalando che il modello si è adattato eccessivamente ai dati di addestramento.

Alle connessioni del primo strato interno è stata applicata in generale come funzione di attivazione la tangente iperbolica (*tanh*), che consente di rappresentare valori sia positivi sia negativi mappandoli in un intervallo compreso tra -1 e 1, mentre al secondo è stata impiegata l'unità lineare rettificata (*ReLU*), che restituisce il valore di ingresso invariato, se positivo, o zero, se negativo. Quest'ultima è stata impiegata anche al primo strato interno per *LFL distance* e $1/2$ *LFL distance* per *Catastrophic rupture*. Per l'output layer in generale viene impiegata la funzione di PyTorch `nn.Identity()`, in modo che la rete restituisca direttamente la combinazione lineare dei pesi e dei bias, mentre per il target *Flow rate* è stata introdotta la funzione `nn.Softplus()`, così da garantire predizioni non negative. Questa scelta è stata motivata dal fatto che, in presenza di valori prossimi allo zero nel dataset, per questo target il modello tendeva in alcuni casi a restituire valori negativi, non

fisicamente plausibili.

La funzione di costo impiegata è MSE, che calcola l'errore medio quadratico tra l'output della rete e il target. Si tratta di una funzione spesso utilizzata in questo contesto, caratterizzata da semplicità di interpretazione e di implementazione [32]. Nel caso in esame non comporta particolari limitazioni: sebbene sia nota la sua sensibilità ad eventuali valori anomali (*outliers*), questo aspetto non costituisce un problema dal momento che i dati utilizzati sono stati generati da un software deterministico. L'algoritmo Adam consente l'ottimizzazione dei parametri del modello. Come descritto in [33] e riportato anche nella documentazione ufficiale di PyTorch, si tratta di un ottimizzatore largamente adottato nel *machine learning* in quanto è in grado di adattare dinamicamente il tasso di apprendimento per ciascun peso della rete.

La suddivisione in training, validation e test è stata effettuata in maniera randomica. Le reti sono state addestrate impiegando il set di training, monitorando al tempo stesso le prestazioni con il validation set. Una volta definita la configurazione ottimale, le prestazioni del modello sono state valutate attraverso il test set, che consiste in dati che il modello non ha mai visto prima, e rispetto al quale sono state calcolate le metriche R^2 , MSE e $MAPE$.

È stata inoltre effettuata una validazione più robusta tramite k-fold cross validation. Data la scarsità di dati a disposizione e quindi l'elevato rischio di *overfitting*, questo tipo di validazione ha permesso di verificare se i modelli implementati mostrassero segni di sovradattamento ai dati di addestramento.

La k-fold cross validation è una tecnica comunemente impiegata sia per la scelta dell'architettura del modello sia per la valutazione delle sue prestazioni. Il parametro k, che dà il nome alla tecnica, rappresenta il numero di gruppi in cui viene suddiviso il dataset, a seguito di mescolamento dei dati: a turno, uno di questi viene impiegato per la validazione mentre i restanti k-1 gruppi per l'allenamento del modello.

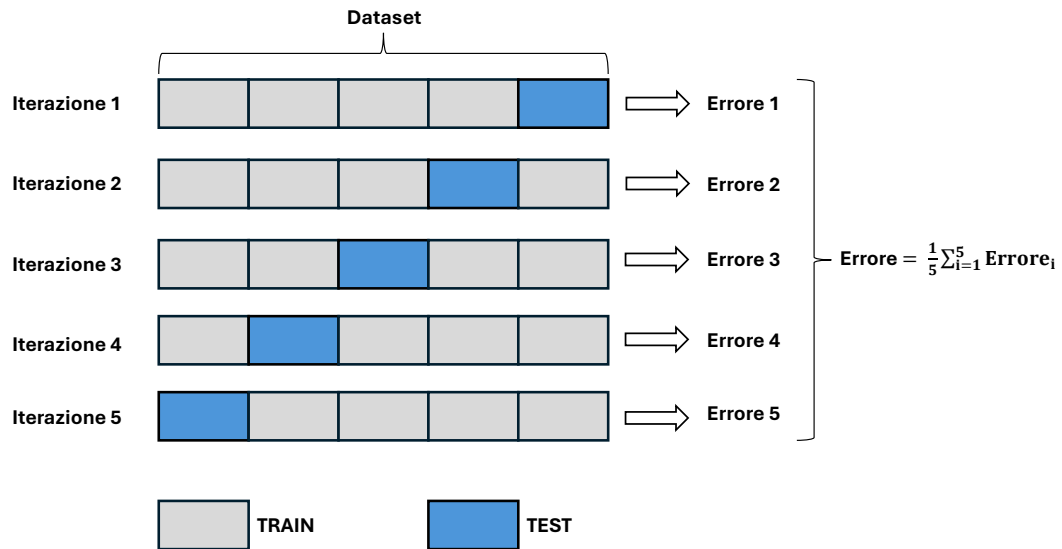


Figura 4.8: Illustrazione della procedura alla base della k-fold cross validation.

Con una validazione di questo tipo, ogni dato può far parte del set di testing e la performance complessiva si ottiene attraverso la media aritmetica delle k metriche ottenute per ogni iterazione sul validation set. Il valore medio delle metriche è accompagnato dalla deviazione standard che indica se il modello si comporta o meno in maniera consistente su tutti i sottogruppi analizzati. Risulta fondamentale la scelta del valore di k adeguato: in questo caso è stato impiegato un k pari a 10, pratica comune.

Capitolo 5

Analisi dei risultati

5.1 Metriche

La valutazione delle prestazioni dei modelli sviluppati si basa sul confronto tra i valori reali, ottenuti con PHAST, e quelli predetti dalle reti. A tale scopo sono state calcolate alcune metriche per quantificare l'accuratezza delle reti:

- *Mean Absolute Error*: misura la media delle differenze assolute tra i valori predetti e quelli reali. È poco sensibile a valori anomali dal momento che penalizza gli errori in maniera lineare.

$$MAE(f) = \frac{1}{N} \sum_{i=1}^N |f(x_i) - y_i| \quad (5.1)$$

- *Mean Squared Error*: la differenza tra i valori predetti e reali viene elevata al quadrato, perciò risulta più sensibile a valori anomali rispetto al MAE. Applicando la radice quadrata al valore di MSE si ottiene RMSE (*Root Mean Squared Error*), che ha la stessa unità di misura del target.

$$MSE(f) = \frac{1}{N} \sum_{i=1}^N (f(x_i) - y_i)^2 \quad (5.2)$$

- *Mean Absolute Percentage Error*: misura l'errore percentuale delle predizioni rispetto ai valori reali. Consente il confronto con modelli con scale diverse di

valori, risultando però molto sensibile a valori prossimi a zero.

$$MAPE(f) = \frac{1}{N} \sum_{i=1}^N \left| \frac{f(x_i) - y_i}{y_i} \right| \times 100 \quad (5.3)$$

- *Coefficient of determination*: misura la capacità del modello di spiegare la variabilità del target rispetto a un modello di riferimento che si limita a predire la media dei valori osservati. Un valore pari a 1 indica una predizione perfetta.

$$R^2 = 1 - \frac{\sum_{i=1}^n (f(x_i) - y_i)^2}{\sum_{i=1}^n (\bar{y} - y_i)^2} \quad (5.4)$$

Le metriche MAE , MSE e $MAPE$ hanno come limite inferiore lo zero, che corrisponde alla condizione ideale di un *fit* perfetto. All'aumentare del loro valore, la qualità delle predizioni risulta peggiore.

Il coefficiente di determinazione R^2 può assumere valori negativi quando il modello fornisce previsioni peggiori rispetto a quelle ottenute utilizzando semplicemente la media dei dati. Un valore pari a zero indica che il modello non spiega in alcun modo la variabilità del target rispetto alla sua media.

D. Chicco et al. in [34] presentano uno studio finalizzato a individuare le metriche più adeguate per confrontare modelli di regressione e uniformare la valutazione dei risultati, concludendo che il coefficiente R^2 rappresenta la misura più significativa. Pertanto, in questo lavoro è stato assunto come metrica di riferimento per valutare l'accettabilità dei modelli, affiancato comunque dagli altri indicatori al fine di fornire un quadro più completo delle prestazioni.

In alcuni casi, i limiti di infiammabilità non vengono raggiunti, perciò è stato definito un valore di soglia per poter assegnare l'etichetta “*not reached*” alle distanze corrispondenti.

Per valutare l'accuratezza del modello nel distinguere i casi di classificazione corretta da quelli errati, si fa ricorso alla matrice di confusione.

La matrice di confusione consiste in una matrice 2x2 nella quale si riportano:

- *True Positive* (TP): numero di campioni positivi classificati correttamente dal modello.
- *True Negative* (TN): numero di campioni negativi classificati correttamente dal modello.
- *False Positive* (FP): numero di campioni negativi classificati erroneamente come positivi, definiti anche "errori di tipo I".

- *False Negative* (FN): numero di campioni positivi classificati erroneamente come negativi, definiti anche "errori di tipo II".

L'accuratezza nella classificazione binaria viene quindi calcolata come:

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \quad (5.5)$$

Sono stati realizzati grafici di confronto tra valori reali e predetti per ogni target, così da valutare visivamente l'accuratezza delle predizioni nei diversi scenari. I punti corrispondenti a ciascuno scenario sono stati rappresentati con colori distinti.

5.2 Modelli neurali

Per gli scenari *Leak*, *Catastrophic rupture* e *Line rupture*, che presentavano dataset più numerosi, sono state confrontate le prestazioni ottenibili con una regressione multivariata rispetto all'impiego di reti neurali, in particolare considerando il valore di R^2 e delle metriche di errore per il set di test. Avendo appurato la convenienza di utilizzare reti neurali rispetto alla regressione multivariata, si è proceduto effettuando un'analisi di sensibilità sul numero di neuroni degli strati interni. L'obiettivo è determinare l'architettura più adatta che rappresenti il miglior compromesso tra accuratezza e complessità, confrontando le prestazioni in termini di R^2 , MSE (e $RMSE$) e $MAPE$ sui set di training, validation e test.

5.2.1 Leak

Le reti analizzate presentano due strati interni in generale, fatta eccezione per il target *Flow rate* in cui ne viene utilizzato uno solo, e un numero di neuroni variabile con l'obiettivo di identificare l'architettura più adeguata.

- **Flow rate** La regressione lineare fornisce R^2 circa pari a 0,63, che arriva a 0,84 nel caso in cui si effettui trasformazione logaritmica, con $RMSE$ che passa da 2,151 kg/s a 1,397 kg/s. Tuttavia, le reti neurali mostrano valori superiori a 0,98 per R^2 già con neuroni ridotti e consentono di ottenere $RMSE$ inferiori a 0,500 kg/s. È stata selezionata la rete con 6 neuroni dal momento che minimizza MSE mantenendo valori di $MAPE$ comparabili agli altri casi. Quest'ultimo resta elevato (circa 50%), dovuto principalmente alla presenza di diversi valori prossimi a zero. Ad esempio, per 30 bar, foro di diametro 25,4 mm, altezza del punto di rilascio 0 m e condizioni 1,5/F, la portata determinata da PHAST è 0,611 kg/s mentre la rete predice 1,080 kg/s, con conseguente errore percentuale elevato (77%).

- UFL distance** La regressione mostra prestazioni molto basse in termini di R^2 , pari a 0,34, e di $MAPE$, di circa il 40%. Le reti neurali sembrano migliorare notevolmente i risultati, elevando R^2 a valori superiori a 0,9. La configurazione [6, 3] rappresenta il miglior compromesso tra semplicità e prestazioni, con R^2 pari a 0,98 e $MAPE$ inferiore al 10% in test set. Dal momento che il dataset presenta molti valori prossimi a zero, è stata introdotta una soglia di 0,1 m per distinguere i casi in cui il limite superiore di infiammabilità non viene raggiunto (“not reached”). Ciò ha consentito di ottenere nella maggior parte dei casi la predizione corretta, corrispondente al limite non raggiunto, mentre rimangono alcune sovrastime, come nel caso di rilascio da serbatoio di 800 bar, foro da 25,4 mm, altezza 0,5 m e condizioni 2,5/D, per cui la rete predice 0,19 m a fronte di un valore reale nullo. Dal momento che i valori elevati per questo target sono scarsi, in alcuni casi il modello tende a sottostimare, come per 800 bar, foro da 25,4 mm, altezza 1 m e condizioni 2,5/D, dove la distanza reale sarebbe 1,87 m mentre la rete restituisce 1,01 m.
- LFL distance** La regressione restituisce R^2 pari a 0,75 e $MAPE$ di quasi il 60% in test, mentre le reti raggiungono valori di 0,92-0,93 e 15-30% a seconda del numero di neuroni. Per questo target è stato necessario incrementare il numero di neuroni, adottando la configurazione [10, 8] che ha consentito di ottenere R^2 tra 0,95 e 0,99 e $MAPE$ inferiori o prossimi al 15%. In alcuni casi le predizioni sono molto accurate, come per 30 bar, foro di diametro 25,4 mm, altezza del punto di rilascio 0 m e condizioni 2,5/D, per cui il valore reale è 21,26 m e quello predetto 21,20 m; in altri casi, si verificano errori significativi, soprattutto per valori più bassi, come la previsione di 4,69 m in corrispondenza di un valore reale nullo.
- 1/2 LFL distance** La regressione in questo caso mostra R^2 più bassi di *LFL distance*, con R^2 intorno a 0,66 e $MAPE$ del 60% circa. Le reti migliorano le prestazioni arrivando a R^2 compresi tra 0,81 e 0,88. Rispetto a *LFL distance* si impiega una rete più complessa, [12, 6], che pur presentando R^2 elevati in training e validation, con valori vicini a 0,97, raggiunge 0,88 in test. Inoltre, i valori di $MAPE$ sono più elevati (fino al 25%). Le predizioni possono essere molto accurate, con differenze inferiori a 1 m, mentre in altri casi si assiste a evidente sottostima, come per 700 bar, foro 5 mm, 0 m, condizioni 2,5/D per cui 63,80 m è il valore reale e 35,86 m quello predetto, o sovrastima, per 500 bar, foro 10 mm, 0 m, condizioni 2,5/D per cui viene stimato 44,29 m a fronte di 30,44 m di PHAST.
- Max width** La regressione fornisce già prestazioni buone, con R^2 di 0,86 e $MAPE$ inferiori al 30%, mentre le reti elevano tali coefficienti a valori compresi tra 0,94 e 0,95 e 10-20% rispettivamente. È sufficiente un’architettura semplice,

[6, 3], che consente di ottenere un netto miglioramento rispetto alla regressione, con *MAPE* inferiori al 20% e R^2 compresi tra 0,94 e 0,98. Le difficoltà maggiori si osservano per valori molto bassi, come nel caso di 30 bar, foro di diametro 10 mm, altezza 0,5 m, condizioni 2,5/D, per cui il valore stimato da PHAST è 0,62 m e quello predetto 1,43 m.

Nel complesso, le reti neurali per questo scenario migliorano sensibilmente la qualità delle predizioni rispetto ad una semplice regressione multivariata, in particolare per i target *Flow rate* e *UFL distance*. Per gli altri target i miglioramenti risultano più contenuti e per *Max width* la regressione si dimostra competitiva.

Dall'analisi delle predizioni ottenute dalla rete sul dataset utilizzato emergono alcune considerazioni aggiuntive. In generale la rete mostra predizioni più accurate per valori di pressione e diametri del foro elevati, ai quali sono associati target di entità maggiore. Per contro, per valori più bassi presenta maggiori difficoltà, in special modo per *Flow rate* e *Max width*. In questi casi gli errori percentuali tendono ad aumentare, nonostante gli scostamenti assoluti rimangano contenuti. Ad esempio, per 30 bar, altezza 0 m e diametro del foro di 25,4 mm, per condizioni meteorologiche 1,5/F, per il quale la predizione del *Flow rate* è di 1,080 kg/s contro 0,611 kg/s reali.

Un'ulteriore criticità riscontrata riguarda il mancato rispetto del vincolo fisico fra la distanza LFL e $\frac{1}{2}$ LFL: in alcuni casi la predizione di $\frac{1}{2}$ LFL distance è minore di quella di LFL distance, a causa dell'incertezza dei modelli. È possibile osservare come ciò si verifichi nei casi in cui l'altezza di rilascio è pari a 0 m, dal momento che i due valori reali sono tra loro molto vicini. Per lo stesso rilascio a 30 bar, la distanza al limite inferiore di infiammabilità predetta è sottostimata, 7,04 m rispetto a 8,69 m di PHAST, mentre quella a metà del limite inferiore è stimata a 3,42 m mentre sarebbe 9,73 m. Situazioni analoghe si verificano per il caso a 50 bar, con errori minori per *Flow rate* ma LFL distance che anche in questo caso non rispetta il vincolo fisico con $\frac{1}{2}$ LFL distance. Questo inconveniente potrebbe essere evitato impiegando reti multi-target nelle quali può essere prevista l'integrazione di tale vincolo nella funzione di perdita; tuttavia, avendo sviluppato modelli indipendenti, questo non è risolvibile a priori: in tali casi viene restituito solo il valore di LFL distance dalla rete.

Per quanto riguarda *UFL distance*, viste le difficoltà derivanti dalla scarsità di valori diversi da zero, per poter esprimere un giudizio sulla capacità del modello si adotta una classificazione binaria per determinare il numero di volte in cui il modello riporta correttamente un valore finito o “not reached”. Dalla matrice di confusione relativa a *UFL distance* emerge che il modello riesce a identificare correttamente tutti i casi in cui tale livello viene raggiunto, evitando falsi negativi; tuttavia, compaiono falsi positivi, 15 casi, in cui viene predetta una distanza finita

quando invece sarebbe “*not reached*”: si tratta delle sovrastime che fa il modello in corrispondenza di valori nulli e che superano il limite di soglia. Per *LFL distance*, invece, non si verificano falsi negativi e i falsi positivi sono solo 8.

Per *Flow rate* il numero di campioni effettivi è inferiore rispetto agli altri target, dal momento che la portata dipende unicamente da due feature, ossia pressione e diametro. Nei grafici si osserva una tendenza sistematica a sovrastimare i valori per portate più basse, mentre per valori più elevati le predizioni si avvicinano maggiormente ai valori reali. Risulta un target più difficile da predire, a conferma di quanto visto dall’analisi delle singole predizioni. Per le distanze UFL e LFL le maggiori criticità sono in corrispondenza dei valori nulli, come già affermato. Per valori di distanze più elevati, invece, i punti si dispongono più vicini alla bisettrice, perciò le predizioni risultano più affidabili. Le distanze a $\frac{1}{2}$ LFL invece non presentano il problema dei valori “*not reached*” essendo sempre finite, perciò, salvo casi specifici, il comportamento della rete per questo target sembra più stabile. Il target *Max width* è un parametro legato anch’esso al raggiungimento del limite inferiore di infiammabilità, dunque può presentare valori nulli; tuttavia, l’andamento complessivo delle predizioni è più omogeneo e le deviazioni sono distribuite in modo uniforme lungo l’intervallo dei valori; perciò, la rete riesce a catturare l’andamento di tale parametro.

La k-fold cross validation, impiegata a posteriori, evidenzia come *UFL distance* e *Max width* siano i target più critici per questo scenario: per il primo, si osserva un valor medio di R^2 pari a 0,52 circa, con deviazione standard molto simile, e *MAPE* medio del 26% circa con deviazione standard del 22%, mentre per il secondo R^2 medio è di circa 0,81 e la deviazione standard di 0,39, con *MAPE* medio del 22% circa e deviazione del 16%. Gli altri target invece mostrano R^2 tra 0,93 e 0,98 con deviazioni tra 0,02 e 0,06. I valori di *MAPE* sono contenuti in generale, sebbene più elevati per *Flow rate* che conferma quanto visto con validazione più semplice, dal momento che sono presenti diversi valori prossimi a zero.

5.2.2 Catastrophic rupture

Le reti analizzate presentano due strati interni in generale e uno solo per *Max width*, con un numero di neuroni variabile, con l’obiettivo di identificare l’architettura più adeguata.

- **UFL distance** Con la regressione si ottiene R^2 pari a circa 0,69 e *MAPE* del 16% in test; le reti migliorano leggermente le prestazioni, portando il valore di R^2 fino 0,92 e *MAPE* inferiori al 10%, con miglioramenti più evidenti in training. La configurazione [8,4] è risultata quella più adeguata, poiché il miglioramento marginale ottenibile passando a quella immediatamente successiva, come R^2 da 0,87 a 0,89, non giustifica l’aumento di complessità.

Per questo scenario non ci sono valori nulli, conseguentemente il target viene predetto con maggior precisione rispetto agli scenari *Leak* e *Line rupture*. Lo scostamento massimo si verifica per 100 bar, 20 kg stoccati e condizioni 2,5/D, per cui la predizione è di 1,44 m rispetto a 1,10 m stimati da PHAST.

- **LFL distance** La regressione fornisce già buoni risultati, con R^2 pari a 0,89 e $MAPE$ del 12% sul set di test, mentre con le reti si arriva a valori compresi tra 0,95 e 0,97 per il primo e inferiori al 10% per il secondo, a seconda dell'architettura. La configurazione [6,3] garantisce prestazioni ottimali, con $RMSE < 1,00$ m e $MAPE < 7\%$. Le predizioni sono buone, con scostamenti contenuti, generalmente entro 2-3 m.
- **½ LFL distance** La regressione presenta risultati soddisfacenti, mentre ulteriori miglioramenti possono essere ottenuti con le reti. L'architettura [6,3] risulta ottimale, con R^2 di circa 0,94 e $MAPE < 10\%$. Le predizioni sono molto precise, in alcuni casi persino migliori di *LFL distance*. Ad esempio, l'errore massimo si ottiene per 100 bar, 20 kg stoccati e condizioni 2,5/D, con valore predetto di 21,19 m rispetto a quello di PHAST di 16,64 m.
- **Max width** La regressione garantisce R^2 di 0,90 e $MAPE$ intorno al 10% sul test set mentre le reti consentono di arrivare a 0,96-0,99 per il primo e a valori intorno al 6% per il secondo. Si adotta l'architettura [6] che restituisce per tutti e tre i set $MAPE < 6\%$. Questo target presenta leggere difficoltà aggiuntive rispetto a *Leak* e *Line rupture* ma gli scostamenti rimangono contenuti, come per 100 bar, 20 kg stoccati e condizioni 2,5/D, per cui il valore predetto è di 22,16 m rispetto mentre quello di PHAST 18,40 m.

Nel complesso i miglioramenti introdotti dalle reti neurali in questo scenario risultano più contenuti rispetto a *Leak* e *Line rupture*, probabilmente perché il caso è più semplice dal punto di vista della modellazione. Si è deciso comunque di adottare le reti per le metriche globalmente migliori.

Dai grafici di confronto tra valori reali e predetti si osserva un andamento complessivamente omogeneo. Sottostime e sovrastime si distribuiscono in modo bilanciato e le predizioni risultano vicine alla bisettrice, confermando una buona accuratezza del modello.

Per questo scenario la k-fold cross validation mostra che il modello di *UFL distance* è accurato e abbastanza stabile, con R^2 medio di 0,91 circa e deviazione standard di 0,18 e ancor più stabile quello di *Max width* con R^2 medio di 0,95 e deviazione standard di circa 0,10. In questo caso appaiono critici i target *LFL distance* e *1/2 LFL distance*: sebbene il valor medio delle metriche possa risultare accettabile, la variabilità tra i fold comporta deviazioni standard della stessa entità dei valori medi, evidenziando instabilità dei modelli. Ad esempio, si osserva un valore di R^2

medio compreso di 0,80-0,82 e deviazione 0,50-0,29 rispettivamente, mentre per *MAPE* un valore medio di 7-12% e deviazione di 4-17% circa.

5.2.3 Line rupture

Le reti analizzate presentano due strati interni in generale, fatta eccezione per il target *Flow rate* in cui ne viene utilizzato uno solo, e un numero di neuroni variabile con l'obiettivo di identificare l'architettura più adeguata.

- **Flow rate** La regressione lineare presenta un R^2 in test di 0,68 circa che sale a 0,86 applicando la trasformazione logaritmica, mentre *RMSE* passa da circa 8 kg/s a 5 kg/s. Le reti neurali, tuttavia, consentono di ottenere valori di R^2 tra 0,91 e 0,97, con *RMSE* inferiori a 4 kg/s. È stata selezionata la configurazione [8], che garantisce R^2 superiori a 0,96, riducendo significativamente *MSE* rispetto alle configurazioni più semplici. Le maggiori difficoltà si osservano per la predizione di portate ridotte con sovrastime marcate, come il caso a 40 bar, diametro della tubazione di 26,7 mm, lunghezza di 6 metri, per cui la rete restituisce 1,461 kg/s mentre PHAST 0,744 kg/s.
- **UFL distance** La regressione restituisce un R^2 in test piuttosto basso, di circa 0,45, e *MAPE* del 10% circa. Con la configurazione di rete [4, 2] si raggiunge R^2 di 0,99 e *MAPE* inferiori al 3%, con prestazioni migliori rispetto a reti più complesse. Rispetto allo scenario *Leak*, le sovrastime sono meno frequenti e la rete riesce quasi sempre ad evitare di restituire predizioni finite quando il limite non è raggiunto.
- **LFL distance** La regressione permette di ottenere R^2 di 0,91 e *MAPE* intorno al 10%, mentre le reti, a seconda della complessità, consentono di salire fino a valori di R^2 compresi tra 0,91 e 0,99 e *MAPE* inferiori al 10% nella maggior parte dei casi. La scelta ricade sulla configurazione [8,6] che garantisce un netto miglioramento di R^2 rispetto a soluzioni più semplici, soprattutto in test, con R^2 di 0,99 e *MAPE* inferiori al 5% in tutti i set. In generale le predizioni sono buone, salvo rari casi come 40 bar, diametro della tubazione di 52,5 mm, 0,5 m di lunghezza e condizioni 1,5/F, per cui PHAST stima 53,12 m mentre viene restituito dalla rete 60,82 m.
- **½ LFL distance:** Si verifica un comportamento analogo a quello di *LFL distance*, con R^2 di circa 0,92 per la regressione e valori prossimi a 0,99 per alcune reti, con *MAPE* che passa da valori intorno al 9% a valori inferiori al 5%. Prestazioni paragonabili a quelle ottenute per *LFL distance* si ottengono con la configurazione [8,4], con *MSE* leggermente più elevati a causa della diversa scala del target. In alcuni casi le predizioni sono meno precise, come per

40 bar, diametro della tubazione di 26,7 mm, lunghezza di 0,5 m e condizioni 1,5/F per cui il valore predetto è 66,38 mentre quello di PHAST è 47,33 m.

- **Max width** Con regressione si ottiene R^2 di 0,90 e $MAPE$ vicino al 15%, mentre con le reti valori del primo compresi tra 0,97 e 0,99 e del secondo tra 5-10%. La configurazione [8,4] è stata selezionata perché assicura maggiore uniformità dei risultati fra i vari set; non è stata necessaria la trasformazione logaritmica visto i valori più elevati di questo target rispetto a quelli dello scenario *Leak*.

Anche per questo scenario l'impiego di reti neurali consente di migliorare la qualità delle predizioni rispetto alla regressione multivariata, in particolare per i target *Flow rate* e *UFL distance*, mentre per gli altri i miglioramenti risultano più contenuti. Rispetto allo scenario *Leak*, per *Line rupture* si osservano predizioni complessivamente migliori, grazie anche al fatto che i valori dei target sono in media più elevati. Il *Flow rate* viene predetto con maggiore accuratezza rispetto allo scenario *Leak*, mostrando solo una leggera tendenza alla sovrastima per portate più basse. I casi di *UFL distance* nulli sono meno frequenti e quindi le predizioni appaiono più precise, anche se restano alcuni scostamenti in corrispondenza dei valori nulli. *LFL distance*, $\frac{1}{2}$ *LFL distance* e *Max width* presentano distribuzioni più omogenee rispetto alla bisettrice nei grafici di confronto tra valori reali e predetti, a conferma di un miglior adattamento della rete. Per questo scenario solo *UFL distance* presenta alcuni valori “*not reached*”, motivo per cui si riporta solo la matrice di confusione associata ad essa, nella quale si osserva un solo caso falso positivo.

I risultati della k-fold cross validation evidenziano come *UFL distance* sia un target critico anche per questo scenario, con R^2 medio di circa 0,78 e deviazione standard di 0,40. Per gli altri target, i valori medi di R^2 variano tra 0,96 e 0,98 e le deviazioni standard sono comprese tra 0,01 e 0,04, mentre il valore medio di $MAPE$ per le distanze e la larghezza è di circa il 4%, con deviazioni di 1-2%, mentre per la portata di circa il 17% con deviazione del 5%.

5.3 Modelli regressivi

Gli scenari *Disk rupture* e *Relief valve* disponevano di un numero di campioni limitato, per cui l'uso delle reti neurali non sarebbe giustificato, dato il rischio di sovradattamento ai dati forniti senza garantire un reale vantaggio rispetto a modelli più semplici. In questi casi i dataset sono stati suddivisi in 85% training e 15% test. Sono state calcolate le metriche per ogni target, considerando anche in tal caso la trasformazione logaritmica per *Flow rate*.

5.3.1 Disk rupture

- **Flow rate** Le prestazioni sono scarse: senza trasformazione logaritmica R^2 cala da 0,92 in training a 0,19 in test; migliora leggermente con la trasformazione logaritmica, passando da 0,89 a 0,53 in test mentre MSE e $MAPE$ si mantengono elevati in entrambi i casi. Si osservano sia sovrastime, come per 700 bar, diametro di 26,7 mm e condizioni 1,5/F, per cui la portata predetta è di 33,184 kg/s mentre quella stimata da PHAST è 18,862 kg/s, sia sottostime evidenti, come per 350 bar, diametro 52,5 mm e condizioni 2,5/D, 23,821 kg/s predetta e 38,810 kg/s di PHAST.
- **UFL distance** Non valutata per insufficienza di campioni non nulli.
- **LFL distance** Prestazioni discrete, con R^2 che passa da 0,95 in training a 0,85 in test; $MAPE$ contenuto, inferiore al 10% anche in test, mentre MSE risulta elevato. In alcuni casi si riscontrano sovrastime, come per 700 bar, diametro 26,7 mm e condizioni 1,5/F per cui la predizione di tale distanza è 105,94 m mentre PHAST restituisce 87,22 m, in altri invece sottostime, come per 350 bar, diametro 52,5 mm e condizioni 2,5/D, con predizione di 107,03 m a fronte di 119,89 m.
- $\frac{1}{2}$ **LFL distance** I valori di R^2 sono simili a quelli riscontrati per *LFL distance*, con errori percentuali comparabili. Anche per questo target possono verificarsi sovrastime significative, come per 700 bar, diametro 26,7 mm e condizioni 1,5/F, per cui la predizione è di 129,15 m invece di 108,52 m.
- **Max width** R^2 passa da 0,95 a 0,83 in test mentre $MAPE$ si mantiene intorno al 10% sia in training che test. Lo scostamento massimo che si verifica è una sovrastima di quasi 4 metri per il caso a 700 bar, diametro 26,7 mm e condizioni 1,5/F.

Nel complesso la regressione descrive sufficientemente le distanze e la larghezza massima, risultando meno adatta per la predizione della portata di rilascio che subisce un calo delle prestazioni passando da training a test.

5.3.2 Relief valve

- **Flow rate** Le prestazioni sono scarse: R^2 passa da 0,92 in training a 0,14 in test, mentre il $MAPE$ non è significativo a causa di valori molto bassi del target. La trasformazione logaritmica migliora i risultati, con R^2 che scende da 0,96 in training a 0,81 in test e $MAPE$ che rimane inferiore al 30%. Gli scostamenti assoluti restano sotto i 2 kg/s con il caso peggiore rappresentato

da 350 bar, diametro 20,3 mm e 2,5/D, per cui la predizione è di 5,342 kg/s rispetto ai 6,963 kg/s di PHAST.

- **UFL distance** Non valutata per insufficienza di campioni significativi.
- **LFL distance** Prestazioni buone e stabili, con R^2 che passa da 0,95 in training a 0,92 in test e $MAPE$ intorno al 10%. Le differenze tra predizioni e valori reali sono inferiori ai 10 m, ad esempio per 350 bar, diametro 15,9 mm e condizioni 2,5/D per cui viene stimato un valore di 54,41 m mentre PHAST restituisce 62,97 m.
- **½ LFL distance** Stesso andamento osservato per *LFL distance*, con $MAPE$ leggermente migliore, inferiore al 10% anche in test. In valore assoluto gli errori sono più grandi, ma in percentuale più bassi: ad esempio per 350 bar, diametro 15,9 mm e condizioni 2,5/D, per cui il valore predetto è di 72,71 m rispetto a 82,09 m.
- **Max width** R^2 stabile, con un valore che passa da 0,96 a 0,92, e $MAPE$ inferiore al 15%. Gli scostamenti più rilevanti in termini percentuali si osservano per valori bassi, come per 40 bar, diametro 15,9 mm e condizioni 1,5/F, per cui la predizione è di 3,56 m rispetto ai 2,64 m di PHAST.

Anche per questo scenario la regressione multivariata risulta adeguata per la predizione di distanze e larghezza massima, mentre la portata di rilascio si conferma un target critico, migliorando solo parzialmente con trasformazione logaritmica.

5.4 Stima dell'area

Dal momento che l'area reale della proiezione della nube sul piano orizzontale non risultava predetta in maniera accurata dai modelli, si è proceduto con una stima alternativa basata su un'ellissi costruita a partire da *LFL distance* e *Max width*, individuando così l'area all'interno della quale la concentrazione supera 40000 ppm (LFL). L'accuratezza di questa approssimazione è stata valutata confrontando i valori reali con quelli ellittici rispetto alla bisettrice e calcolando MAE e $MAPE$ come metriche di errore.

- **Leak** Il valore di MAE è il più basso tra gli scenari (14,53 m^2), mentre il $MAPE$ è elevato (28,03%), a causa delle aree di ridotta entità presenti in questo caso, che perciò fanno corrispondere anche ad errori assoluti contenuti degli errori percentuali relativamente alti. I punti si distribuiscono attorno alla bisettrice, presentando sia sovrastime che sottostime.

- **Catastrophic rupture** Il MAE si attesta a $119,33 \text{ m}^2$ mentre il $MAPE$ al 13,91%. L'approssimazione comporta sia sovrastime che sottostime con una deviazione più evidente per le aree maggiori, dove emerge una tendenza a sovrastima.
- **Line rupture** Per questo scenario il MAE è di $35,57 \text{ m}^2$ mentre il $MAPE$ è inferiore al 10% (6,60%). L'approssimazione risulta più accurata rispetto ad altri scenari, con punti ben distribuiti lungo la bisettrice e solo una lieve tendenza a sovrastima per i valori più elevati.
- **Disk rupture** Nonostante le aree siano di entità simile a quelle dello scenario *Line rupture*, gli errori sono più elevati: MAE di $134,30 \text{ m}^2$ e $MAPE$ del 21,75%. Gli errori crescono per valori più grandi di area, dove l'approssimazione peggiora.
- **Relief valve** Il MAE è pari a $43,04 \text{ m}^2$ mentre il $MAPE$ al 28,21%. Sovrastime e sottostime si bilanciano in questo caso.

In generale l'approssimazione ellittica consiste in un approccio semplice per la stima dell'estensione della nube al limite inferiore di infiammabilità. Tuttavia, essendo calcolata a partire dalle predizioni di due target, la sua accuratezza dipende direttamente dalla precisione con cui il modello è in grado di stimare tali parametri. La qualità della predizione varia quindi in funzione dello scenario: gli errori risultano contenuti per *Line rupture*, che mostra le prestazioni migliori, mentre diventano più significativi per gli scenari con dataset ridotti, come *Disk rupture* e *Relief valve*, che non garantiscono predizioni altrettanto affidabili.

5.5 Sintesi dei risultati

Dai valori delle metriche ottenute si osserva come sia vantaggioso impiegare modelli più complessi nei casi in cui si disponga di un numero di dati adeguato. In particolare, questi hanno permesso di ottenere prestazioni migliori per alcuni target per i quali la regressione non forniva stime soddisfacenti, come la portata di rilascio. Per altri target invece i miglioramenti osservati sono più contenuti e in alcuni casi la regressione può essere competitiva se si vuole prediligere semplicità. Negli scenari con dataset ridotti è stata esclusa a priori l'implementazione di modelli più complessi, limitandosi a una regressione multivariata che però ha mostrato limiti nella predizione della portata di rilascio specialmente.

Sebbene lo sviluppo di modelli separati necessiti di un minor numero di simulazioni e sia più semplice da implementare rispetto a una rete multi-target, questo non permette di inglobare nella funzione di perdita eventuali vincoli fisici importanti per avere risultati congruenti.

L'analisi k-fold cross validation conferma l'accuratezza dei modelli sviluppati, già evidenziata dalle metriche ottenute con la suddivisione in set di training, validation e test. Per diversi target i modelli non sono solo accurati ma anche stabili, con metriche che risultano quindi indipendenti dalla particolare suddivisione dei dati. Tuttavia, alcuni target risultano critici, presentando una maggiore variabilità delle prestazioni, dovuta principalmente alla limitata disponibilità di dati.

Infine, l'approssimazione dell'area di estensione della nube sul piano orizzontale mediante ellissi si è dimostrata una soluzione semplice per la sua stima, pur risentendo dell'accuratezza dei modelli sviluppati per i singoli scenari. In questo modo lo strumento proposto è in grado di fornire un'informazione aggiuntiva utile per valutazioni di massima.

Capitolo 6

Conclusioni

6.1 Contributi del lavoro

L'obiettivo di questa tesi era sviluppare uno strumento predittivo per la stima di parametri caratteristici di dispersione, per rilasci incidentali di idrogeno gassoso in stazioni di rifornimento. Per la sua costruzione sono state impiegate simulazioni condotte con PHAST, software largamente utilizzato per la modellazione delle conseguenze.

Dopo una rassegna della letteratura, che ha confermato il diffuso utilizzo di PHAST e la crescente integrazione di tecniche di machine learning nell'analisi delle conseguenze, sono stati definiti gli scenari di rilascio da considerare per le simulazioni. Per ciascuno di essi sono stati valutati i parametri di ingresso rilevanti sull'entità dei risultati di interesse, da dover variare per la costruzione del dataset. La selezione è stata guidata da analisi di sensibilità per le grandezze fisiche generali comuni a più scenari e da considerazioni specifiche per parametri caratteristici di ciascuno. In questo modo, si è voluta bilanciare la rappresentatività e la semplicità del modello da ottenere. Tale scelta risulta coerente con l'obiettivo di un modello che in base a poche informazioni sia in grado di fornire stime preliminari. Un numero maggiore di variabili di ingresso può aumentare le informazioni disponibili e la variabilità catturata ma è necessario correlare questo aspetto al numero di dati di cui dispone, per evitare fenomeni di *overfitting* [26]. Per ciascuno scenario è stata definita una matrice di correlazione delle variabili con coefficiente di Pearson. Tuttavia, questo indicatore riesce a interpretare esclusivamente relazioni di tipo lineare tra le variabili, mentre potrebbe risultare opportuno adottare anche altri indici.

Le simulazioni ottenute sono complessivamente 311, distinte in seguito per scenario per costituire il dataset di ognuno. In una prima fase sono stati definiti dei modelli di riferimento con regressione multivariata, implementando poi reti neurali per gli

scenari con più campioni, analizzando in particolare il numero di neuroni ottimale per gli strati interni. I risultati hanno mostrato che con dataset sufficientemente numerosi le reti presentano prestazioni migliori rispetto a semplici modelli regressivi. Nei casi in cui il numero di campioni è ridotto, invece, non è consigliabile adottare modelli più complessi per il rischio di sovradattamento degli stessi ai dati utilizzati per l'addestramento; per questi scenari si è scelto di mantenere i modelli regressivi che, pur non garantendo prestazioni particolarmente elevate, si ritiene possano essere migliorati da eventuali simulazioni aggiuntive.

L'ampliamento del dataset è auspicabile anche per gli scenari che attualmente presentano più simulazioni: ciò consentirebbe di irrobustire la rete, che beneficia di un maggior numero di dati da cui può apprendere, e permetterebbe di rappresentare ulteriori condizioni, ad esempio includendo scenari con situazioni meteorologiche differenti in linea con le indicazioni del Purple Book [17]. Questo risulta semplice da realizzare: in base alle esigenze, possono essere condotte nuove simulazioni in PHAST, i cui parametri di ingresso e risultati vengono riportati in file Excel; l'esecuzione degli script di addestramento dei singoli modelli consente di aggiornare automaticamente il codice principale.

A differenza di molte applicazioni di machine learning, in cui i dati a disposizione risultano eterogenei e spesso caratterizzati da rumore o valori mancanti, per questo lavoro si disponeva di dati omogenei provenienti da un software deterministico; questo, oltre a ridurre i tempi di preprocessing, ha compensato in parte la limitata numerosità dei dati con una maggiore qualità degli stessi.

Lo strumento è pensato per ottenere stime di massima: non sostituisce i software di simulazione consolidati ma può affiancarli riducendo i tempi necessari nelle fasi preliminari. Il campo di applicazione è definito in base agli intervalli dei parametri simulati, segnalando all'utente quando i dati di ingresso non rientrano in questi.

Il principale contributo di questo lavoro è l'implementazione di una struttura modulare facilmente ampliabile con nuovi dati, che si ritiene possa essere utile per screening degli scenari incidentali, così da riservare approfondimenti con software dedicati solo per gli scenari più critici.

6.2 Discussione e sviluppi futuri

Nonostante per alcuni scenari fosse disponibile un numero maggiore di simulazioni, la dimensione complessiva dei dataset è risultata comunque contenuta rispetto alle quantità di dati normalmente richieste per modelli data-driven, come le reti neurali. Le prestazioni ottenute sono risultate soddisfacenti, probabilmente grazie alla semplicità del problema e alla coerenza dei dati generati con un software deterministico. È opportuno sottolineare che tali risultati si ritengono validi all'interno del dominio di addestramento dei modelli, mentre eventuali estensioni richiedono

verifiche aggiuntive.

L'affidabilità del modello predittivo può essere ulteriormente rafforzata con la validazione su dati sperimentali, così da verificarne l'utilità fuori da condizioni controllate. In scenari più complessi potrebbe emergere la necessità di considerare variabili trascurate inizialmente, che però sono impattanti sui risultati. Una delle principali limitazioni del modello integrale di dispersione di PHAST riguarda l'impossibilità di rappresentare ostacoli definiti nello spazio; in contesti reali la complessità geometrica del sito influenza significativamente la dispersione, limitando l'applicabilità del modello a scenari in campo aperto.

Per lo sviluppo del modello, inizialmente è stata adottata una regressione multivariata come riferimento, per valutare l'efficacia di approcci più complessi. La scelta del modello dipende, infatti, sia dalla disponibilità dei dati che dalla fisica del fenomeno oggetto di studio. In questo lavoro, la quantità di dati generati con PHAST è risultata sufficiente per testare reti neurali per alcuni scenari, ma il numero di campioni potrebbe non essere stato comunque ottimale per valorizzare tale approccio. I modelli di regressione lineare in alcuni casi hanno fornito risultati coerenti con la natura del fenomeno osservato, come per le distanze che tendono effettivamente ad avere una dipendenza quasi lineare dai dati di ingresso entro certi intervalli.

Un'ulteriore verifica della robustezza dei modelli neurali sviluppati è stata effettuata mediante k-fold cross validation, che ha permesso di evidenziare non solo l'accuratezza dei modelli ma anche la loro stabilità. Le metriche tra i fold sono stabili per la maggior parte dei modelli, indicando assenza di *overfitting* evidente. Alcuni target hanno invece mostrato maggiore variabilità delle prestazioni tra i fold, dovuto sia alla natura del fenomeno che al numero limitato di dati: in questi casi i modelli sono meno affidabili e i risultati vanno interpretati con cautela.

Lo strumento può essere esteso alla stima dei parametri caratteristici di scenari come jet fire, flash fire ed esplosioni, prevedendo ulteriori grandezze come la distanza massima di radiazione termica e la quantità di sostanza contenuta nella nube che può innescarsi. In questo modo, si può ottenere un modello semplice e rapido da utilizzare, utile come supporto per la progettazione e la sicurezza delle stazioni di rifornimento di idrogeno, considerando anche i diversi scenari incidentali possibili.

Appendice A

Dataset di simulazione

A.1 Generale

Target

- *Flow rate*: portata massica di idrogeno rilasciato, espressa in kg/s. Il valore è quello di picco, che si verifica solo all'istante iniziale ma che nelle ipotesi di rilascio stazionario si mantiene costante per tutta la durata del rilascio.
- *UFL distance*: distanza massima lungo la direzione del vento fino alla quale la concentrazione di idrogeno risulta superiore al limite superiore di infiammabilità (Upper Flammability Limit, UFL), corrispondente a 750000 ppm.
- *LFL distance*: distanza massima lungo la direzione del vento fino alla quale la concentrazione di idrogeno risulta superiore al limite inferiore di infiammabilità (Lower Flammability Limit, LFL) corrispondente a 40000 ppm.
- *1/2 LFL distance*: distanza massima lungo la direzione del vento fino alla quale la concentrazione di idrogeno risulta superiore a metà del limite inferiore di infiammabilità, corrispondente a 20000 ppm.
- *Max width*: larghezza massima della proiezione orizzontale della nube all'altezza di riferimento che delimita la zona in cui la concentrazione di idrogeno risulta superiore al limite inferiore di infiammabilità (LFL). Fornisce un'indicazione dell'estensione laterale della zona infiammabile.

Scenario	Modello	N. simulazioni	Target analizzati
Leak	Orifice Model	144	Flow rate UFL distance LFL distance 1/2 LFL distance Max width
Catastrophic rupture	Instantaneous Model	48	UFL distance LFL distance 1/2 LFL distance Max width
Line rupture	Short Pipe Model	72	Flow rate UFL distance LFL distance 1/2 LFL distance Max width
Disk rupture	Short Pipe Model	23	Flow rate LFL distance 1/2 LFL distance Max width
Relief valve	Short Pipe Model	24	Flow rate LFL distance 1/2 LFL distance Max width

Tabella A.1: Riepilogo degli scenari analizzati e dei target considerati.

Parametri mantenuti costanti

- Temperatura di esercizio posta pari a quella ambientale e fissata a 15°C.
- Umidità relativa pari al 60%, corrispondente a condizioni medie.
- Lunghezza di rugosità superficiale posta pari a 0,5 m, come compromesso tra ambienti completamente aperti e zone densamente edificate, plausibile per stazioni di rifornimento.

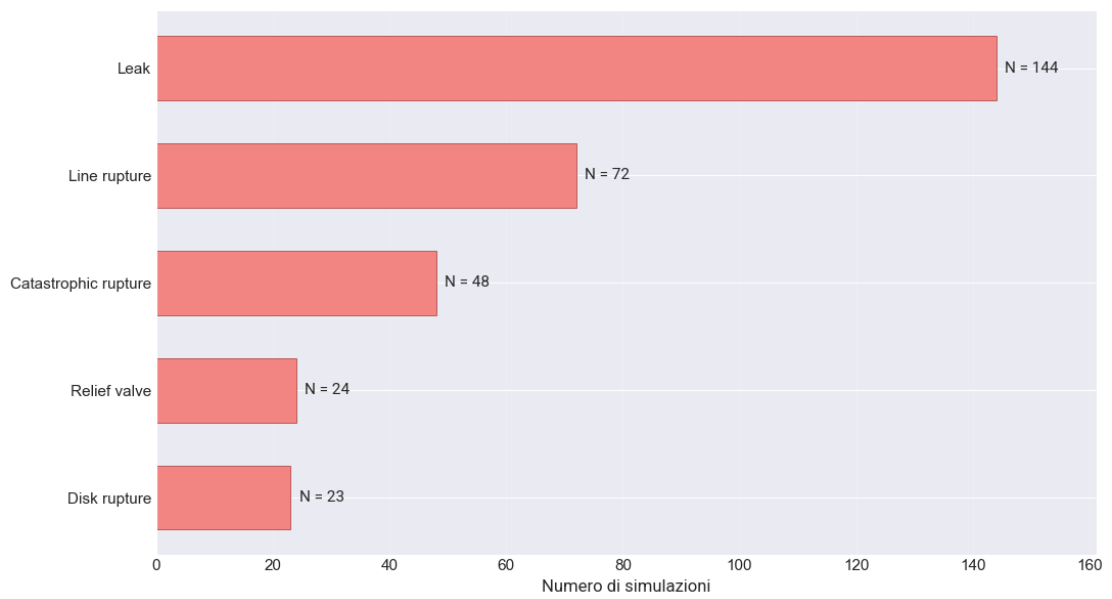


Figura A.1: Numerosità dei dataset per scenario.

A.2 Leak

Variabile	Tipo	Range/Valori	Unità
Pressure	Continuo	30–800	bar
Height	Discreto	0; 0,5; 1	m
Hole diameter	Continuo	5–25,4	mm
Wind speed	Discreto	1,5; 2,5	m/s
Stability class	Discreto	D; F	–

Tabella A.2: Range delle variabili di input simulate per lo scenario *Leak*.

	Flow rate [kg/s]	UFL distance [m]	LFL distance [m]	1/2 LFL distance [m]	Max width [m]
Mediana	0,580	0,57	24,92	32,87	3,74
Range	14,153	1,63	89,36	112,57	12,76

Tabella A.3: Statistiche descrittive per i diversi target per lo scenario *Leak*.

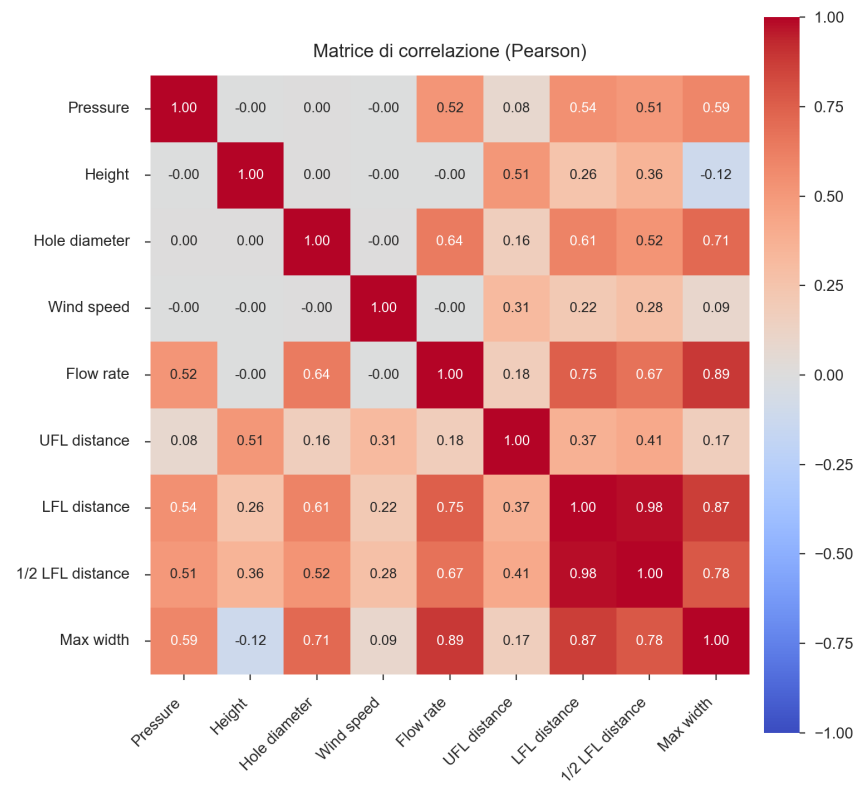


Figura A.2: Matrice di correlazione con coefficiente di Pearson per lo scenario *Leak*.

A.3 Catastrophic rupture

Variabile	Tipo	Range/Valori	Unità
Pressure	Continuo	20–900	bar
Mass	Continuo	20–1000	kg
Wind speed	Discreto	1,5; 2,5	m/s
Stability class	Discreto	D; F	–

Tabella A.4: Range delle variabili di input simulate per lo scenario *Catastrophic rupture*.

	UFL distance [m]	LFL distance [m]	1/2 LFL distance [m]	Max width [m]
Mediana	2,34	16,45	26,92	31,54
Range	3,47	29,03	51,97	49,48

Tabella A.5: Statistiche descrittive per i diversi target per lo scenario *Catastrophic rupture*.

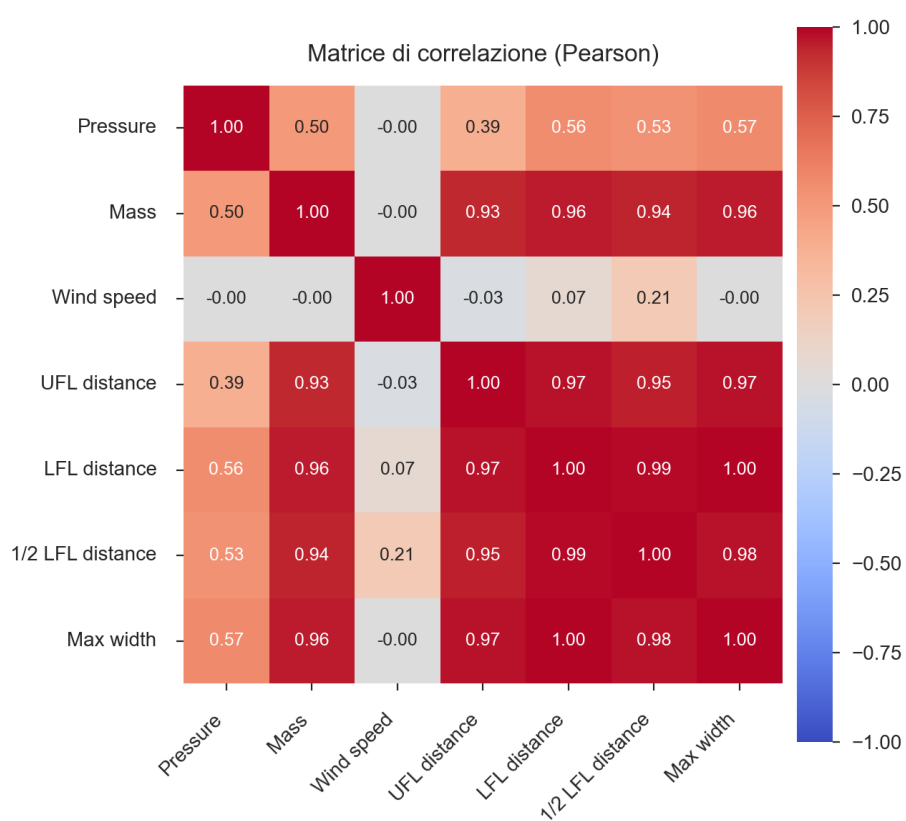


Figura A.3: Matrice di correlazione con coefficiente di Pearson per lo scenario *Catastrophic rupture*.

A.4 Line rupture

Variabile	Tipo	Range/Valori	Unità
Pressure	Continuo	40–700	bar
Pipe length	Continuo	0,5–6	m
Pipe diameter	Continuo	26,7–52,5	mm
Wind speed	Discreto	1,5; 2,5	m/s
Stability class	Discreto	D; F	–

Tabella A.6: Range delle variabili di input simulate per lo scenario *Line rupture*.

	Flow rate [kg/s]	UFL distance [m]	LFL distance [m]	1/2 LFL distance [m]	Max width [m]
Mediana	9,586	1,60	70,35	88,52	10,24
Range	72,180	2,46	114,33	127,56	19,68

Tabella A.7: Statistiche descrittive per i diversi target per lo scenario *Line rupture*.

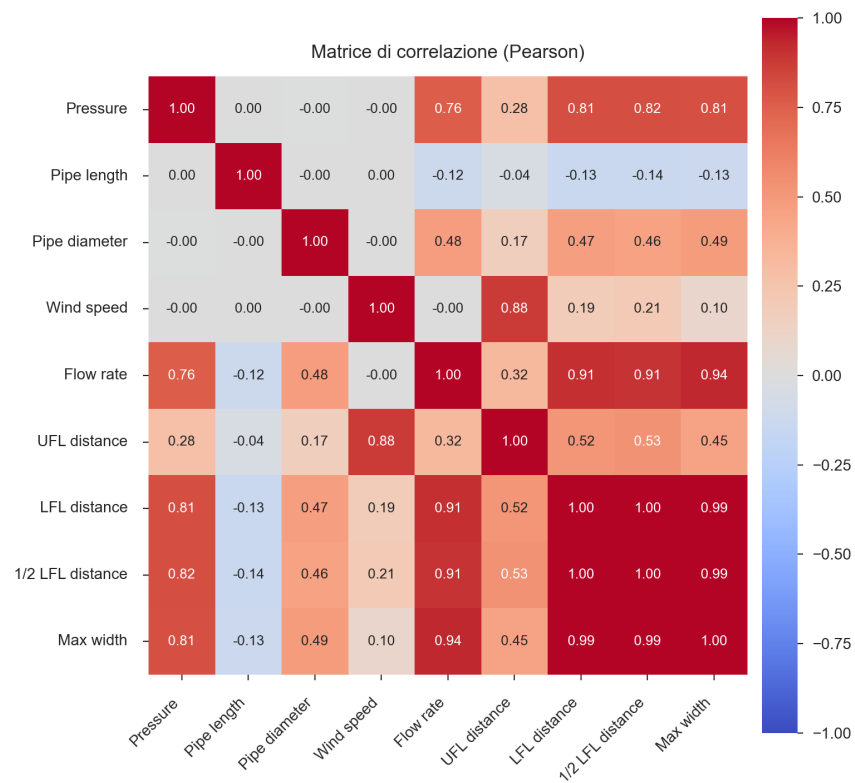


Figura A.4: Matrice di correlazione con coefficiente di Pearson per lo scenario *Line rupture*.

A.5 Disk rupture

Variabile	Tipo	Range/Valori	Unità
Pressure	Continuo	40–700	bar
Disk diameter	Continuo	26,7–52,5	mm
Wind speed	Discreto	1,5; 2,5	m/s

Tabella A.8: Range delle variabili di input simulate per lo scenario *Disk rupture*.

	Flow rate [kg/s]	LFL distance [m]	1/2 LFL distance [m]	Max width [m]
Mediana	10,038	73,21	93,20	10,78
Range	71,686	109,81	121,51	19,08

Tabella A.9: Statistiche descrittive per i diversi target per lo scenario *Disk rupture*.

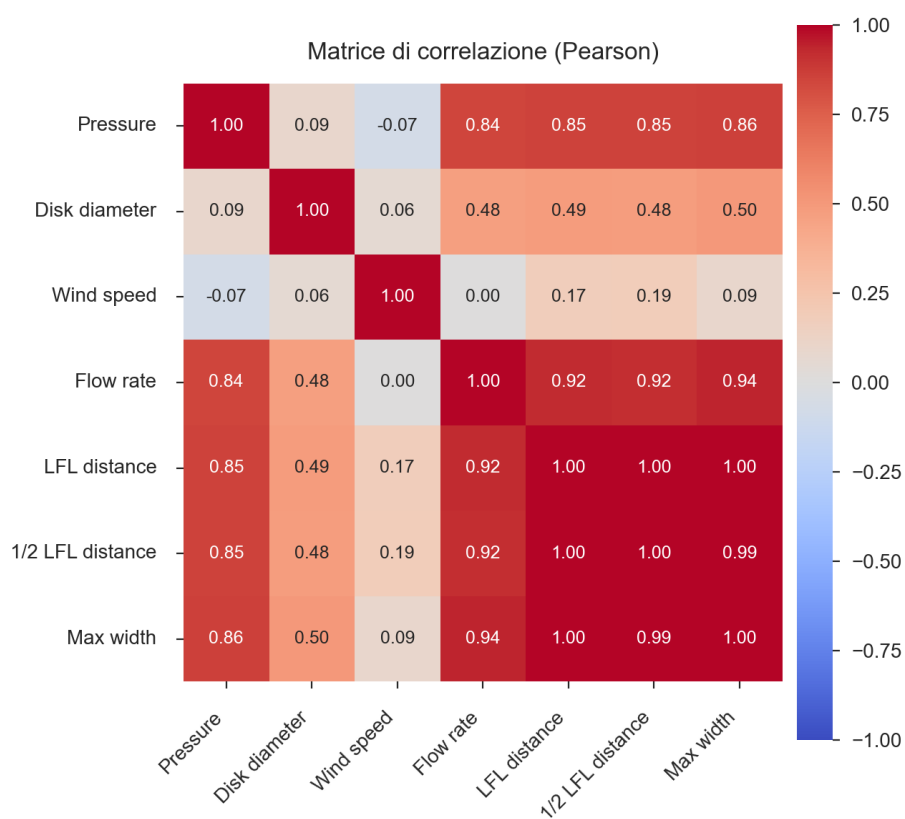


Figura A.5: Matrice di correlazione con coefficiente di Pearson per lo scenario *Disk rupture*.

A.6 Relief valve

Variabile	Tipo	Range/Valori	Unità
Pressure	Continuo	40–700	bar
Orifice diameter	Continuo	9,5–20,3	mm
Wind speed	Discreto	1,5; 2,5	m/s

Tabella A.10: Range delle variabili di input simulate per lo scenario *Relief valve*.

	Flow rate [kg/s]	LFL distance [m]	1/2 LFL distance [m]	Max width [m]
Mediana	2,413	47,95	64,85	5,88
Range	12,896	76,59	87,76	11,20

Tabella A.11: Statistiche descrittive per i diversi target per lo scenario *Relief valve*.

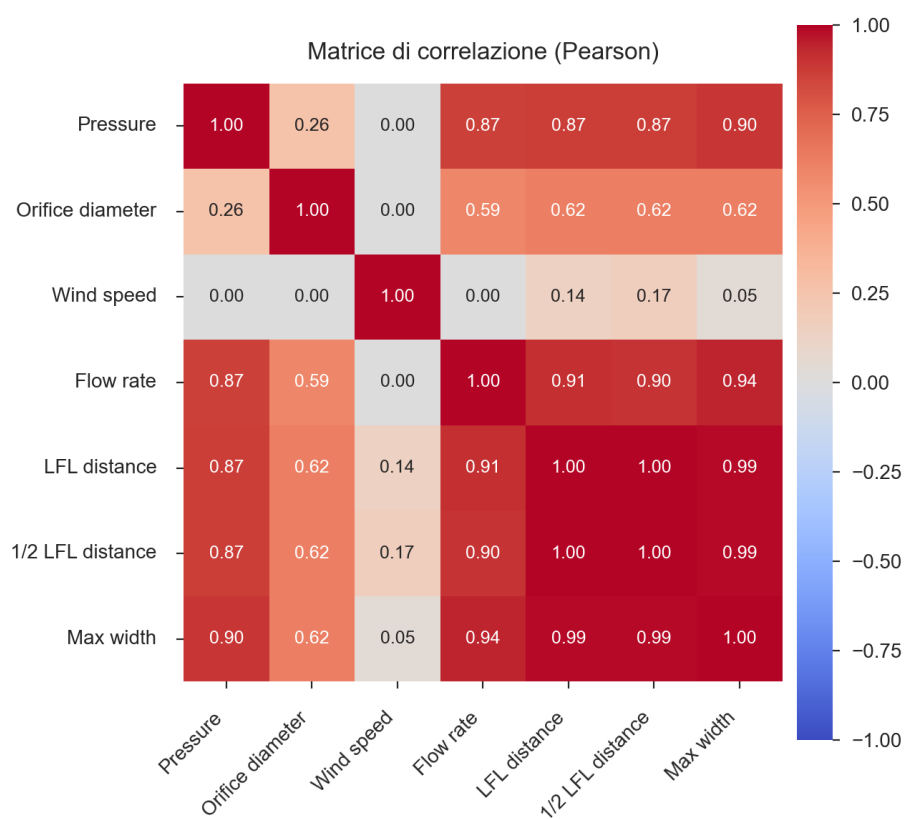


Figura A.6: Matrice di correlazione con coefficiente di Pearson per lo scenario *Relief valve*.

Appendice B

Prestazioni

B.1 Prestazioni globali

Target	MSE	MAPE (%)	R^2
Flow rate	6,814	35,88	0,96784
UFL distance	0,019	13,50	0,98557
LFL distance	11,829	8,77	0,98875
1/2 LFL distance	22,249	9,03	0,98447
Max width	0,994	10,17	0,99461

Tabella B.1: Metriche di valutazione globali per i diversi target.

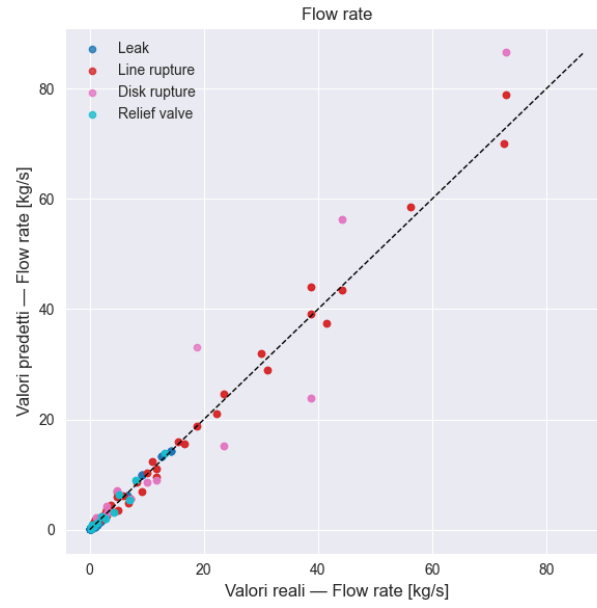


Figura B.1: Confronto tra valori reali e predetti per il target *Flow rate*.

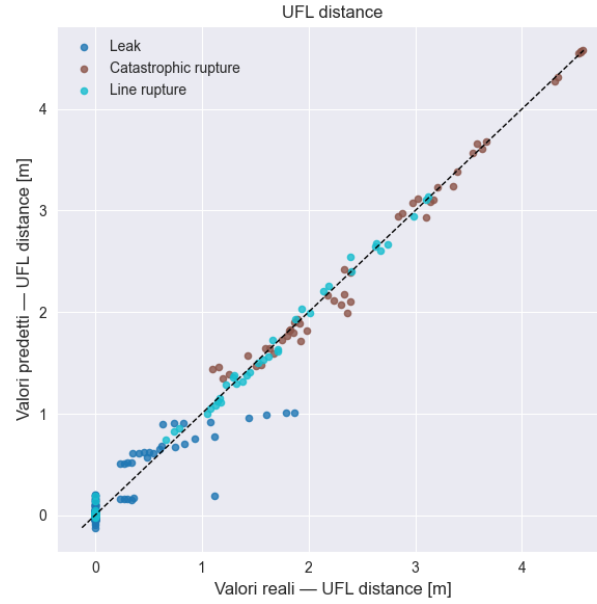


Figura B.2: Confronto tra valori reali e predetti per il target *UFL distance*.

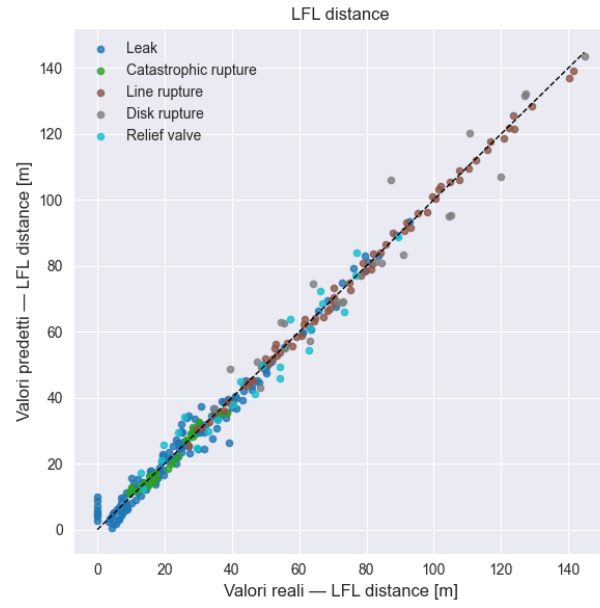


Figura B.3: Confronto tra valori reali e predetti per il target *LFL distance*.

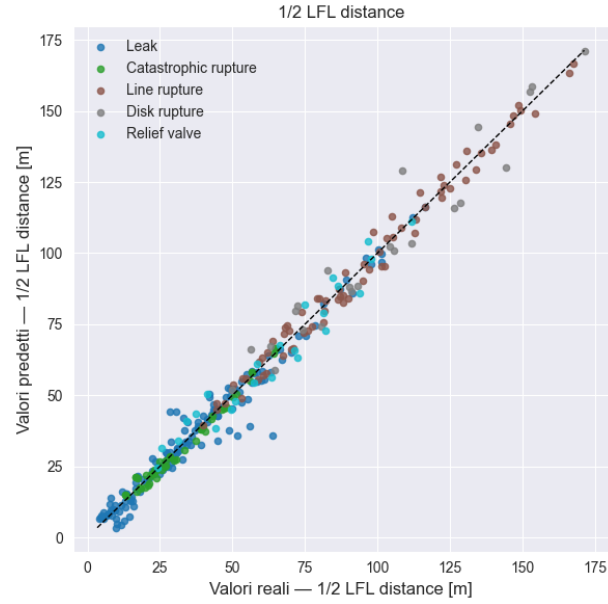


Figura B.4: Confronto tra valori reali e predetti per il target *1/2 LFL distance*.

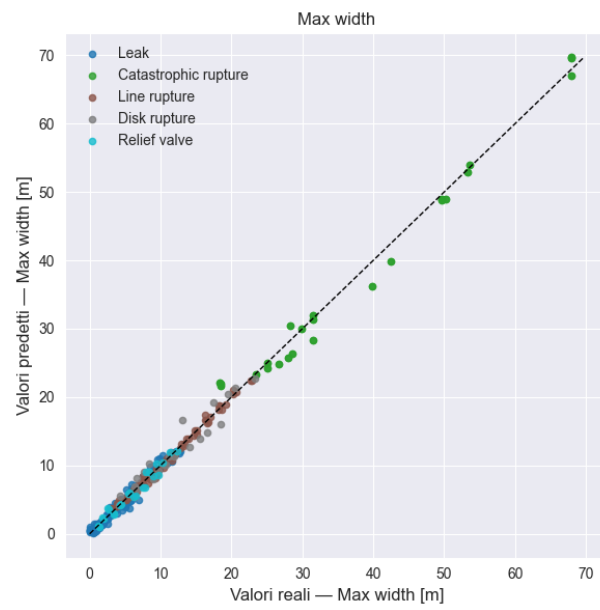


Figura B.5: Confronto tra valori reali e predetti per il target *Max width*.

B.2 Prestazioni per scenario

Target	Hidden sizes	Learning rate	Batch size	Epochs
Flow rate	[6]	0,001	4	500
UFL distance	[6, 3]	0,001	4	500
LFL distance	[10, 8]	0,0005	4	500
1/2 LFL distance	[12, 6]	0,0005	4	500
Max width	[6, 3]	0,001	4	500

Tabella B.2: Iperparametri per ciascun target dello scenario *Leak*.

Target	MSE	MAPE (%)	R^2
Flow rate	0,081	60,54	0,99322
UFL distance	0,002	9,51	0,98494
LFL distance	16,909	15,31	0,95755
1/2 LFL distance	69,412	25,64	0,88652
Max width	0,559	17,59	0,94589

Tabella B.3: Prestazioni sul test set per lo scenario *Leak*

Target	MSE		MAPE (%)		R^2	
	Media	Std	Media	Std	Media	Std
Flow rate	0,067	0,037	38,39	15,95	0,98853	0,02008
UFL distance	0,026	0,031	26,14	22,17	0,51291	0,52122
LFL distance	12,982	7,639	13,98	4,40	0,95667	0,04061
1/2 LFL distance	32,304	20,331	13,91	4,82	0,93150	0,05597
Max width	1,388	2,725	21,63	15,81	0,81281	0,39273

Tabella B.4: Valore medio e deviazione standard delle metriche di validazione (10-fold cross validation) per lo scenario *Leak*.

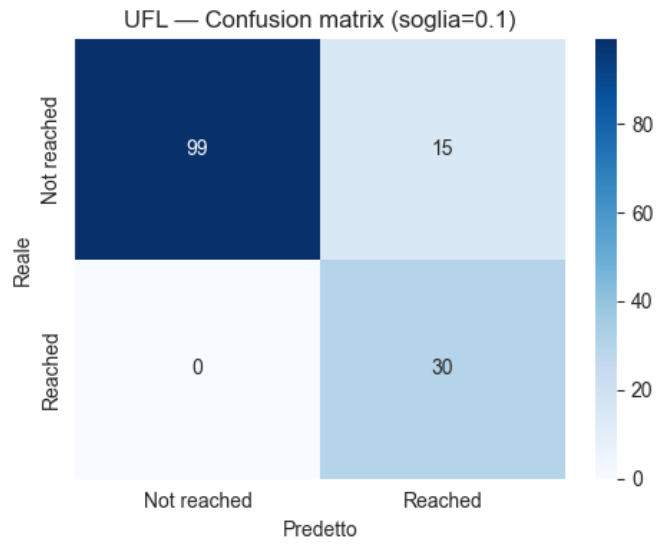


Figura B.6: Matrice di confusione per il target *UFL distance* dello scenario *Leak*.

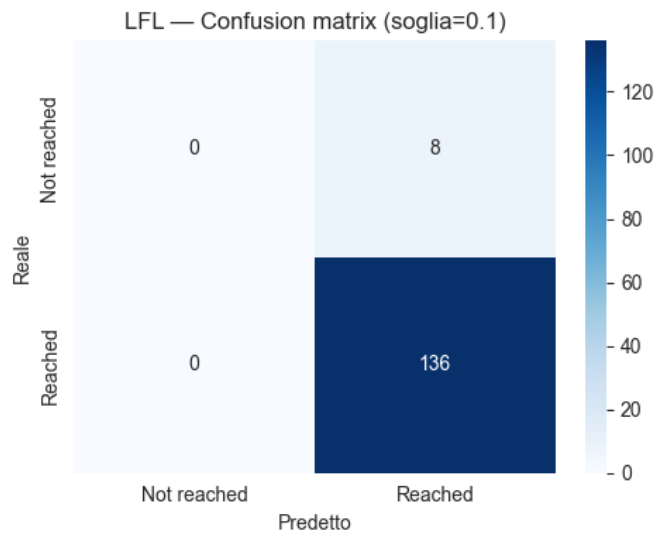


Figura B.7: Matrice di confusione per il target *LFL distance* dello scenario *Leak*.

Target	Hidden sizes	Learning rate	Batch size	Epochs
UFL distance	[8, 4]	0,001	4	500
LFL distance	[6, 3]	0,0005	4	500
1/2 LFL distance	[6, 3]	0,001	4	500
Max width	[4]	0,001	4	500

Tabella B.5: Iperparametri per ciascun target dello scenario *Catastrophic rupture*.

Target	MSE	MAPE (%)	R ²
UFL distance	0,044	9,03	0,87948
LFL distance	1,025	6,51	0,96887
1/2 LFL distance	4,747	8,87	0,95445
Max width	4,195	6,10	0,95720

Tabella B.6: Prestazioni sul test set per lo scenario *Catastrophic rupture*.

Target	MSE		MAPE (%)		R ²	
	Media	Std	Media	Std	Media	Std
UFL distance	0,028	0,033	6,33	4,41	0,91252	0,18059
LFL distance	1,475	1,475	6,23	3,87	0,90712	0,20135
1/2 LFL distance	22,712	61,602	12,05	22,22	0,65496	0,68427
Max width	4,074	2,226	5,39	2,53	0,94566	0,09541

Tabella B.7: Valore medio e deviazione standard delle metriche di validazione (10-fold cross validation) per lo scenario *Catastrophic rupture*.

Target	Hidden sizes	Learning rate	Batch size	Epochs
Flow rate	[8]	0,001	4	500
UFL distance	[8, 6]	0,001	4	500
LFL distance	[8, 6]	0,0005	4	500
1/2 LFL distance	[8, 4]	0,0005	4	500
Max width	[8, 4]	0,001	4	500

Tabella B.8: Iperparametri per ciascun target dello scenario *Line rupture*.

Target	MSE	MAPE (%)	R^2
Flow rate	7,584	30,87	0,96167
UFL distance	0,011	17,19	0,96592
LFL distance	5,344	3,59	0,99219
1/2 LFL distance	2,559	1,91	0,99705
Max width	0,398	8,67	0,98052

Tabella B.9: Prestazioni sul test set per lo scenario *Line rupture*

Target	MSE		MAPE (%)		R^2	
	Media	Std	Media	Std	Media	Std
Flow rate	7,377	7,123	17,22	5,15	0,96760	0,04537
UFL distance	0,088	0,215	11,72	12,08	0,78171	0,39355
LFL distance	8,955	6,247	3,93	1,83	0,98368	0,01724
1/2 LFL distance	17,845	10,026	4,47	2,13	0,97462	0,01907
Max width	0,170	0,138	4,04	1,86	0,98752	0,01596

Tabella B.10: Valore medio e deviazione standard delle metriche di validazione (10-fold cross validation) per lo scenario *Line rupture*.

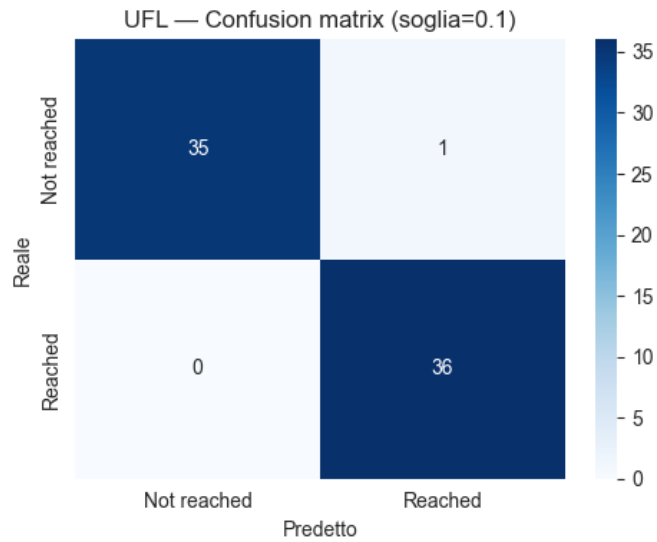


Figura B.8: Matrice di confusione per il target *UFL distance* dello scenario *Line rupture*.

Target	MSE	MAPE (%)	R ²
Flow rate	107,662	48,13	0,52858
LFL distance	110,895	9,29	0,84745
1/2 LFL distance	135,938	8,23	0,84675
Max width	4,045	10,46	0,83117

Tabella B.11: Prestazioni sul test set per lo scenario *Disk rupture*.

Target	MSE	MAPE (%)	R ²
Flow rate	0,696	30,07	0,81182
LFL distance	30,019	9,72	0,91786
1/2 LFL distance	35,600	6,67	0,92874
Max width	0,573	13,91	0,92173

Tabella B.12: Prestazioni sul test set per lo scenario *Relief valve*.

B.3 Approssimazione ellittica dell'area

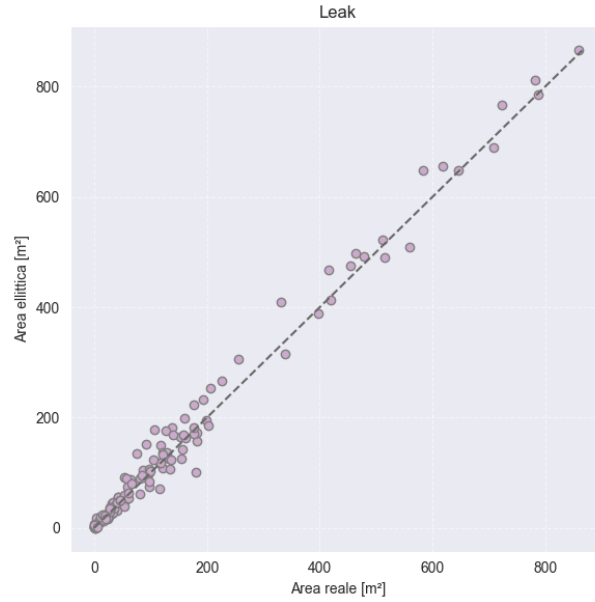


Figura B.9: Confronto tra i valori reali ed approssimati dell'area per lo scenario *Leak*.

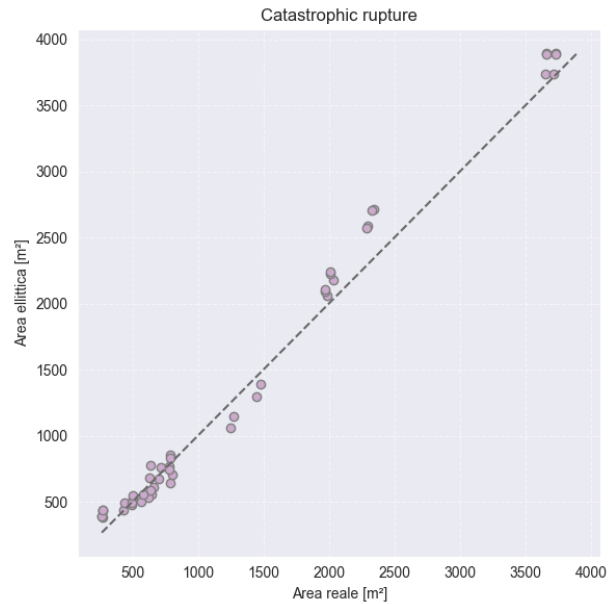


Figura B.10: Confronto tra i valori reali ed approssimati dell'area per lo scenario *Catastrophic rupture*.

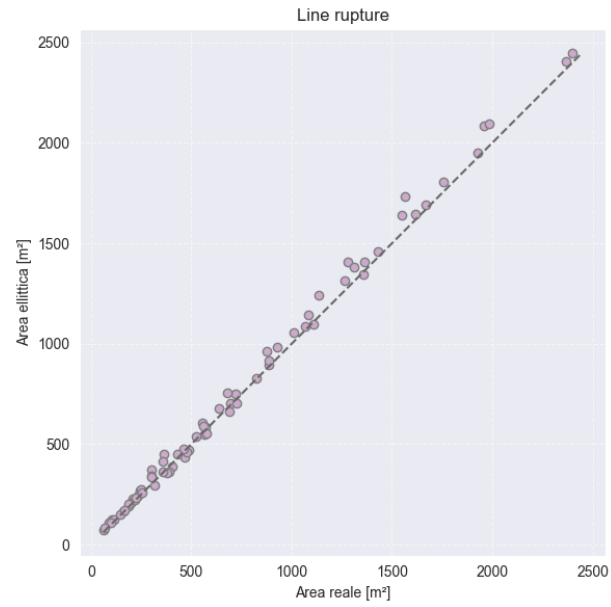


Figura B.11: Confronto tra i valori reali ed approssimati dell'area per lo scenario *Line rupture*.

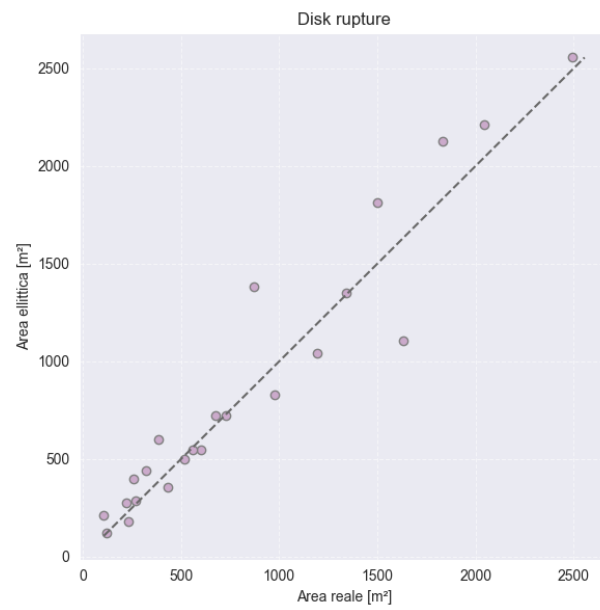


Figura B.12: Confronto tra i valori reali ed approssimati dell'area per lo scenario *Disk rupture*.

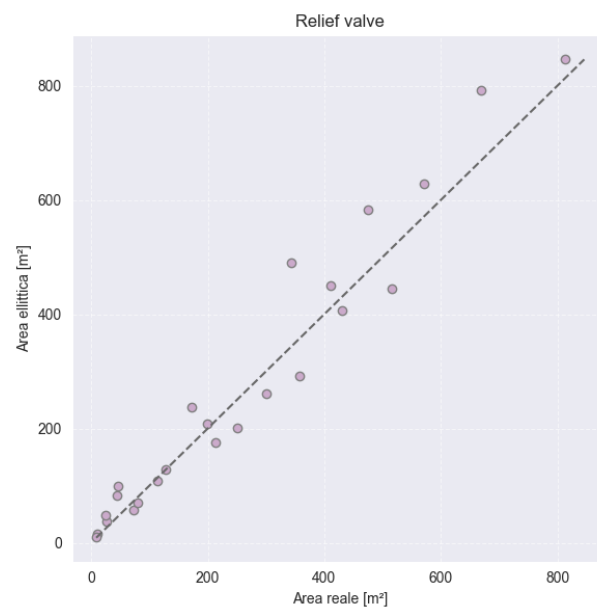


Figura B.13: Confronto tra i valori reali ed approssimati dell'area per lo scenario *Relief valve*.

Bibliografia

- [1] International Energy Agency. *Hydrogen – Low-emission fuels*. 2024. URL: <https://www.iea.org/energy-system/low-emission-fuels/hydrogen> (cit. a p. 1).
- [2] Liejin Guo, Jinzhan Su, Zhiqiang Wang, Jinwen Shi, Xiangjiu Guan, Wen Cao e Zhisong Ou. «Hydrogen safety: An obstacle that must be overcome on the road towards future hydrogen economy». In: *International Journal of Hydrogen Energy* 51.D (2024), pp. 1055–1078. DOI: 10.1016/j.ijhydene.2023.08.248 (cit. a p. 1).
- [3] Andrea Carpignano. «Dispense del corso di Sicurezza e analisi di rischio». Materiale didattico, Politecnico di Torino (cit. a p. 2).
- [4] Alfonso Ibarreta, Achim Wechsung, Ryan Hart, Delmar Trey Morrison e Nicholas Reding. «Blended natural gas/hydrogen fuel gas systems: An evaluation of risk». In: *Process Safety Progress* 44.2 (2025), pp. 232–238. DOI: 10.1002/prs.12677 (cit. a p. 2).
- [5] DNV Digital Solutions. *UDM Theory*. Technical report, DNV, Høvik, Norway. 2023. URL: https://mysoftware.dnv.com/download/public/phast/technical_documentation/05_dispersion/udm/theory/UDM%20Theory.pdf (cit. a p. 3).
- [6] Mostafa Pouyakian, Maryam Ashouri, Shaghayegh Eidani, Rohollah Fallah Madvari e Fereydoon Laal. «A systematic review of consequence modeling studies of the process accidents in Iran from 2006 to 2022». In: *Heliyon* 9.2 (2023). DOI: 10.1016/j.heliyon.2023.e13550 (cit. a p. 3).

-
- [7] Sunhwa Park, Bashir Hashim, Umer Zahid e Junghwan Kim. «Global risk assessment of hydrogen refueling stations: Trends, challenges, and future directions». In: *International Journal of Hydrogen Energy* 106 (2025), pp. 1462–1479. DOI: 10.1016/j.ijhydene.2025.01.438 (cit. alle pp. 4, 6).
- [8] Wei Zhong, Shuangli Wang, Tan Wu, Xiaolei Gao e Tianshui Liang. «Optimized Machine Learning Model for Fire Consequence Prediction». In: *Fire* 7.4 (2024), p. 114. DOI: 10.3390/fire7040114 (cit. a p. 4).
- [9] Yue Sun, Jingyao Wang, Wen Zhu, Shuai Yuan, Yizhi Hong, M. Sam Mannan e Benjamin Wilhite. «Development of Consequent Models for Three Categories of Fire through Artificial Neural Networks». In: *Industrial & Engineering Chemistry Research* 59.1 (2020), pp. 464–474. DOI: 10.1021/acs.iecr.9b05032 (cit. a p. 4).
- [10] Artemis Papadaki, Alba Àgueda e Eulàlia Planas. «Machine learning for safety distances prediction during emergency response of toxic dispersion accidental scenarios». In: *Journal of Loss Prevention in the Process Industries* 95 (2025), pp. 105–604. DOI: 10.1016/j.jlp.2025.105604 (cit. a p. 4).
- [11] Junseo Lee, Sehyeon Oh e Byungchol Ma. «Performance analysis of optimized machine learning models for hydrogen leakage and dispersion prediction via genetic algorithms.» In: *International Journal of Hydrogen Energy* 97 (2025), pp. 1287–1301. DOI: 10.1016/j.ijhydene.2024.10.183 (cit. alle pp. 4, 16, 51).
- [12] Zeren Jiao, Yue Sun, Yizhi Hong, Trent Parker, Pingfan Hu, M. Sam Mannan e Qingsheng Wang. «Development of Flammable Dispersion Quantitative Property-Consequence Relationship Models Using Extreme Gradient Boosting». In: *Industrial & Engineering Chemistry Research* 59 (2020), pp. 15109–15118. DOI: 10.1021/acs.iecr.0c02822 (cit. a p. 4).
- [13] DNV Digital Solutions. *DISC Model Theory*. Technical report, DNV, Høvik, Norway. 2023. URL: <https://mysoftware.dnv.com/knowledge-centre/phast-and-safeti/tech-doc/03discharge> (cit. alle pp. 8–13).
- [14] DNV GL. *PHAST and PHAST LITE Tutorial Manual*. DNV GL Software. 2018. URL: <http://www.dnvgl.com/software> (cit. a p. 13).

- [15] Vivek P. Utgikar e Todd Thiesen. «Safety of compressed hydrogen fuel tanks: Leakage from stationary vehicles». In: *Technology in Society* 27.3 (2005), pp. 315–320. DOI: 10.1016/j.techsoc.2005.04.005 (cit. a p. 17).
- [16] D. Michael Johnson, Robert Crewe, Graham Boaler e John Evans. *Review of the Current Understanding of Hydrogen Jet Fires and the Potential Effect on PFP Performance*. Hazards 31 Conference, Paper 55. 2021. URL: <https://www.icheme.org/media/27744/hazards-31-paper-55-johnson.pdf> (cit. a p. 19).
- [17] Committee for the Prevention of Disasters. *Guidelines for Quantitative Risk Assessment “Purple Book” CPR 18E*. Committee for the Prevention of Disasters, 2005 (cit. alle pp. 21, 23, 29–31, 68).
- [18] C. Ainscough. *H2FIRST Reference Station Design Task: Stand-Alone Piping and Instrumentation Diagrams*. Operated for the U.S. Department of Energy by the Alliance for Sustainable Energy, LLC. 2014 (cit. a p. 25).
- [19] Val-Matic Valve & Mfg. Corp. *Design and Selection of Check Valves*. Rapp. tecn. 2018. URL: <https://www.valmatic.com> (cit. a p. 26).
- [20] Matteo Genovese, Viviana Cigolotti, Elio Jannelli e Petronilla Fragiaco. «Current standards and configurations for the permitting and operation of hydrogen refueling stations». In: *International Journal of Hydrogen Energy* 48.51 (2023), pp. 19357–19371. DOI: 10.1016/j.ijhydene.2023.01.324 (cit. a p. 30).
- [21] Hyunjun Kwak, Minji Kim, Mimi Min, Byoungjik Park e Seungho Jung. «Assessing the Quantitative Risk of Urban Hydrogen Refueling Station in Seoul, South Korea, Using SAFETI Model». In: *Energies* 17.4 (2024), p. 867. DOI: 10.3390/en17040867 (cit. a p. 32).
- [22] Anirudha Joshi, Fereshteh Sattari, Lianne Lefsrud, Modusser Tufail e M.A. Khan. «Mitigating uncertainty: A risk informed approach for deploying hydrogen refueling stations». In: *International Journal of Hydrogen Energy* 74 (2024), pp. 136–150. DOI: 10.1016/j.ijhydene.2024.06.085 (cit. a p. 32).

- [23] *Hydrogen Pipelines: Guide to the Design, Construction and Operation of Hydrogen Pipelines*. Rapp. tecn. IGC DOC 121/14. European Industrial Gases Association (EIGA), 2014. URL: <https://www.eiga.eu/uploads/documents/DOC121.pdf> (cit. a p. 34).
- [24] *API Standard 526: Flanged Steel Pressure Relief Valves*. American Petroleum Institute, 2023 (cit. a p. 34).
- [25] Kit Yan Chan, Bilal Abu-Salih, Raneem Qaddoura, Ala' M. Al-Zoubi, Vasile Palade, Duc-Son Pham, Javier Del Ser e Khan Muhammad. «Deep neural networks in the cloud: Review, applications, challenges and research directions». In: *Neurocomputing* 545 (2023), p. 126327. DOI: 10.1016/j.neucom.2023.126327 (cit. a p. 38).
- [26] Zheng Zhou, Cheng Qiu e Yufan Zhang. «A comparative analysis of linear regression, neural networks and random forest regression for predicting air ozone employing soft sensor models». In: *Scientific Reports* 13 (2023), p. 22420. DOI: 10.1038/s41598-023-49899-0 (cit. alle pp. 39, 67).
- [27] JetBrains. *Creating a virtual environment*. URL: <https://www.jetbrains.com/help/pycharm/creating-virtual-environment.html> (cit. a p. 39).
- [28] Matplotlib Development Team. *Pyplot tutorial*. URL: <https://matplotlib.org/stable/tutorials/pyplot.html> (cit. a p. 40).
- [29] Adam Paszke et al. *PyTorch: An Imperative Style, High-Performance Deep Learning Library*. <https://arxiv.org/abs/1912.01703>. 2019 (cit. a p. 41).
- [30] PyTorch Team. *Neural Networks*. URL: https://docs.pytorch.org/tutorials/beginner/blitz/neural_networks_tutorial.html (cit. a p. 43).
- [31] PyTorch Team. *Datasets & DataLoaders*. URL: https://docs.pytorch.org/tutorials/beginner/basics/data_tutorial.html (cit. a p. 44).
- [32] Juan Terven, Diana-Margarita Cordova-Esparza, Julio-Alejandro Romero-González, Alfonso Ramírez-Pedraza e E. A. Chávez-Urbiola. «A comprehensive survey of loss functions and metrics in deep learning». In: *Artificial Intelligence Review* 58 (2025), p. 195. DOI: 10.1007/s10462-025-11198-7 (cit. a p. 52).
- [33] Diederik P. Kingma e Jimmy Ba. *Adam: A Method for Stochastic Optimization*. <https://arxiv.org/abs/1412.6980>. 2015 (cit. a p. 52).

- [34] Davide Chicco, Matthijs J. Warrens e Giuseppe Jurman. «The coefficient of determination R-squared is more informative than SMAPE, MAE, MAPE, MSE and RMSE in regression analysis evaluation». In: *PeerJ Computer Science* 7 (2021), e623. DOI: 10.7717/peerj-cs.623 (cit. a p. 55).