

Progetto InSITE: INtelligent energy management of Smartgrids
based on IoT and edge/cloud Technologies

D2.1 SOTA of AI-based techniques for DT training and energy production and demand forecasting

Finanziato a valere sul fondo di finanziamento della Ricerca di Sistema elettrico nazionale (Fondo RdS) Piano Triennale (PT) 2019-2021 – Bando b – progetto “InSITE - INtelligent energy management of Smartgrids based on IoT and edge/cloud Technologies”, riferimento Prog. CSEAB_00320



Contratto	Contratto di ricerca per il Progetto “INtelligent energy management of Smartgrids based on IoT and edge/cloud Technologies” [InSITE], ammesso al finanziamento con decreto del Ministero della Transizione Ecologica, 20 settembre 2021, pubblicato su G.U.R.I. n. 234 del 30 settembre 2021
Titolo	D2.1 Analisi dello Stato dell’Arte di tecniche basate sull'intelligenza artificiale per l’addestramento di DT e la previsione della produzione e della domanda di energia
Title	D2.1 SOTA of AI-based techniques for DT training and energy production and demand forecasting
Progetto	“INtelligent energy management of Smartgrids based on IoT and edge/cloud Technologies” [InSITE]
WP	WP2
Linea di Attività	LA2.1 SOTA of AI-based techniques for DT training and energy production and demand forecasting
Keywords	State of the Art, Artificial Intelligence, Digital Twin, Modelling, Machine Learning, Energy production forecasting, Energy demand forecasting, Electricity price forecasting
N. Pagine	245
Emesso il	28/02/2023
Elaborato da	G. Conte, M. Surico, D. Oresta, P. Serafino, D. Scarafile
Verificato da	M. Surico
Approvato da	G. Conte

Sommario

1	Il concetto di Digital Twin	4
1.1	Classificazione delle tecniche di intelligenza artificiale	9
2	Tecniche di modellazione per il monitoraggio e l'ottimizzazione di asset energetici	12
2.1	Modellazione di carichi	14
2.1.1	Approccio NILM per la disaggregazione delle misure elettriche	29
2.2	Modellazione di generatori	51
2.2.1	Fonti rinnovabili	52
2.2.1.1	Impianti eolici	52
2.2.1.2	Impianti fotovoltaici	64
2.2.2	Fonti non rinnovabili	79
2.2.3	Accumulatori	97
3	Previsione di domanda, produzione e costo dell'energia	112
3.1	Previsione della domanda di energia	113
3.1.1	Short-term load forecast (STLF)	114
3.1.2	Medium-term load forecast (MTLF)	138
3.1.3	Long-term load forecast (LTLF)	155
3.2	Previsione della produzione di energia	167
3.2.1	Fonti rinnovabili	167
3.2.1.1	Energia eolica	168
3.2.1.2	Energia solare	186
3.2.2	Fonti non rinnovabili	200
3.3	Previsione del costo dell'approvvigionamento da reti di fornitura	213
4	Conclusioni	236
	Riferimenti	237

1 Il concetto di Digital Twin

Repliche virtuali di prodotti fisici che forniscono in tempo reale una fotografia dello stato del prodotto, i Digital Twin consentono di prevedere le prestazioni future dell'asset fisico e di sperimentare miglioramenti senza doverli testare sul prodotto stesso. Le potenzialità e i campi di applicazione sono numerose. Ad oggi i Digital Twin vengono utilizzati principalmente nell'industria, ma la tecnologia si sta affermando anche nel campo sanitario e nel retail.

Un Digital Twin, o gemello digitale, è una rappresentazione virtuale di un oggetto o di un sistema, collegato ad esso per tutto il ciclo di vita. Il Digital Twin viene aggiornato in tempo reale dai dati raccolti dai sensori collegati all'asset fisico e usa programmi di simulazione, l'apprendimento automatico e il ragionamento per fornire informazioni utili sull'asset e per elaborare modelli predittivi delle prestazioni future e delle reazioni dell'oggetto a determinate condizioni. In termini più semplici, il gemello digitale è un modello virtuale altamente complesso, che è l'esatta replica del suo corrispettivo fisico. Questo può essere qualsiasi cosa: da un'auto, a un macchinario industriale, a un aereo, un ponte, un edificio e così via.

Anche se la terminologia si è evoluta nel tempo, il concetto base è rimasto lo stesso. Si basa sull'idea che un costrutto informativo digitale che riguarda un sistema fisico può essere creato come un'entità a sé stante. Grazie all'avanzamento delle tecnologie digitali, il Digital Twin permette una gestione più efficiente dei sistemi, permettendo un risparmio di costi e di tempo, sin dalla fase di progettazione.

Il concetto viene per la prima volta espresso da Grieves nel 2002, in una presentazione per la proposta di un Product Lifecycle Management (PLM) center. In quell'occasione, Grieves descrive per la prima volta gli elementi alla base di quello che in seguito verrà chiamato Digital Twin: uno spazio reale, uno spazio virtuale e il collegamento del flusso di dati e di informazioni tra spazio reale allo spazio virtuale. Questo PLM non sarebbe stato un oggetto statico, ma avrebbe seguito tutto il ciclo di vita del sistema reale, dalla creazione, alla realizzazione, all'utilizzo e, infine, al suo smaltimento. Il sistema virtuale era dunque uno specchio di quello reale, e venne perciò definito "Mirrored Spaces Model" e, in un secondo momento, "Information Mirroring Model". Anche in alcune pubblicazioni successive, si fa riferimento al concetto proposto da Grieves con il termine "Information Mirroring

Model". È nel 2011 con il saggio "Virtually Perfect: Driving Innovative and Lean Products through Product Lifecycle Management", che per la prima volta Grieves parla di Digital Twin. Il termine viene utilizzato dal ricercatore per descrivere "un insieme di costrutti informativi virtuali che descrivono completamente un manufatto fisico potenziale o reale, dal livello micro atomico al livello macro geometrico".

Prima di allora, il sistema doveva essere fisicamente creato, inizialmente sotto forma di progetto e, in seguito, come un prototipo. La costruzione di un modello 3D fisico permetteva una migliore comprensione del sistema stesso e del suo comportamento. Tuttavia, se venivano riscontrati problemi, malfunzionamenti e rischi di collisione tra le componenti (quando il sistema prevedeva il movimento di alcune di esse) si doveva tornare al modello in 2D, modificarlo e ricominciare il ciclo da capo. Un processo molto dispendioso, sia in termini di tempo che di costi economici. L'avanzamento delle tecnologie digitali ha semplificato il processo, permettendo di creare modelli 3D virtuali. Questi consentono di fare, virtualmente, tutte le valutazioni per cui era necessario costruire un prototipo, risparmiando così tempo e risorse.

Il gemello digitale di un prodotto è una fonte preziosissima di informazioni per ingegneri e operatori. Informazioni che vengono ricavate attraverso la combinazione di più tecnologie, dal cloud all'Internet of Things, all'Intelligenza Artificiale. Il Digital Twin, infatti, è connesso al prodotto fisico attraverso vari sensori posti su aree vitali alla funzionalità dell'oggetto. Questi sensori producono dati su diversi aspetti delle prestazioni dell'oggetto fisico, dalla temperatura, all'energia utilizzata/prodotta, alle condizioni meteorologiche e così via. L'analisi di questi dati, combinata con altre fonti di informazione, permette di capire non solo il comportamento del prodotto, ma anche di predire come il prodotto si comporterà in futuro.

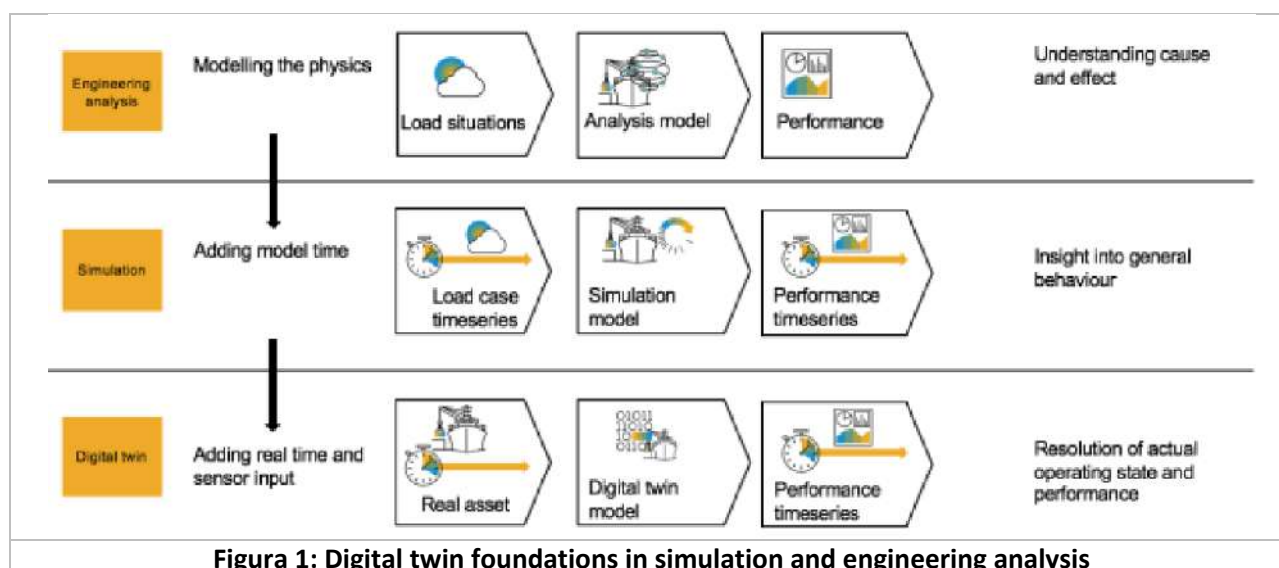
Questo flusso continuo di informazioni permette al gemello digitale di eseguire simulazioni, rilevare e analizzare eventuali problemi di prestazioni del prodotto e studiare possibili miglioramenti. Il flusso di dati tra il prodotto fisico e il suo gemello virtuale funziona a doppio senso: il Digital Twin riceve i dati dai sensori di cui è dotato il prodotto fisico, per poi restituire insight.

Con lo sviluppo delle tecnologie dell'Industria 4.0, sempre più aziende manifatturiere ricorrono ai Digital Twin per ottimizzare i processi e i prodotti. Il gemello digitale di un macchinario, ad esempio, è in grado di fornire informazioni molto preziose sulle prestazioni e lo stato di salute della macchina.

Questo permette all'azienda di rilevare malfunzionamenti, ottimizzare le prestazioni della macchina e programmare i necessari interventi di manutenzione. Il Digital Twin permette dunque di abilitare la manutenzione predittiva, evitando così stop alla produzione dovuti al guasto dei macchinari. Inoltre, analizzando i dati provenienti dai macchinari presenti in un impianto è anche possibile monitorare il consumo energetico e ottimizzare quei processi che comportano sprechi di energia. Allo stesso modo, il gemello digitale di un prodotto è una fonte preziosissima di informazioni per il manufacturer. Potendo seguire le prestazioni del prodotto per tutto il suo ciclo di vita, il produttore può utilizzare i dati per migliorare il prodotto stesso e fornire ai clienti servizi post-vendita customizzati ai loro bisogni.

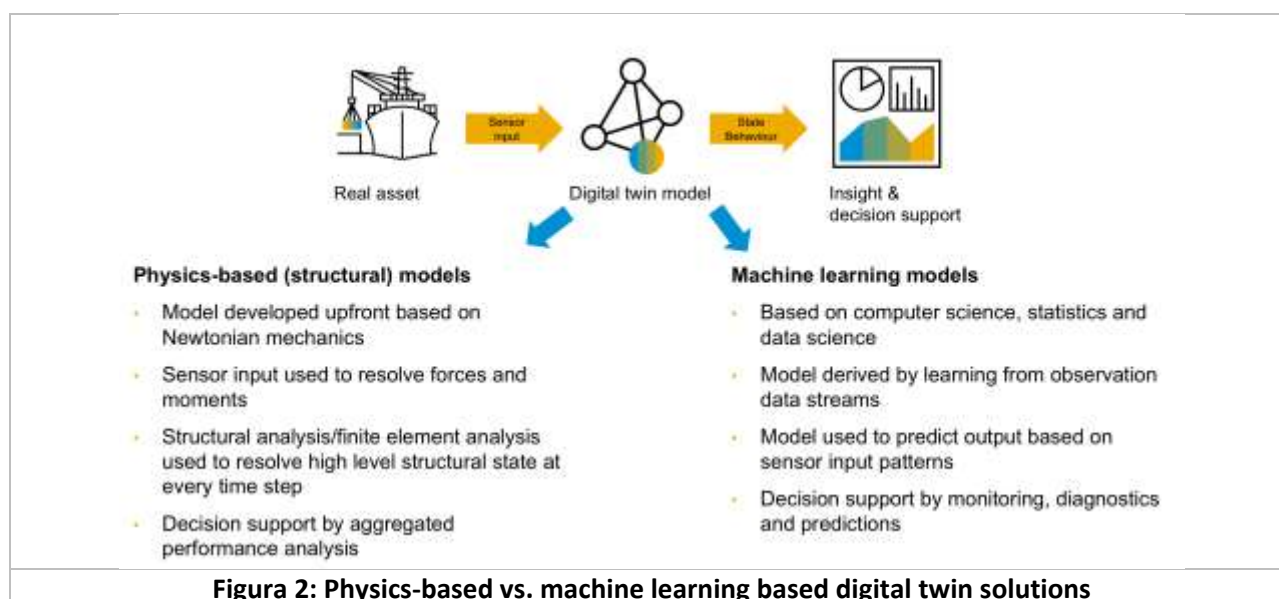
Anche la tecnologia dei Digital Twin si è evoluta. Quando le aziende hanno iniziato ad adottare i Digital Twin, questi venivano apprezzati per la loro capacità di monitorare, simulare, e ottimizzare i dati di diversi dispositivi. Recentemente, le innovazioni in campo di Intelligenza Artificiale e l'aumento dell'adozione hanno permesso di scalare i modelli, fino a creare reti neurali con più layer. Le aziende stanno iniziando a collegare massicce reti di gemelli intelligenti, collegando molti Digital Twin insieme, per creare modelli virtuali di intere fabbriche, cicli di vita dei prodotti, catene di approvvigionamento, porti e città.

I metodi di modellazione di un DT possono essere generalmente raggruppati in due tipi: metodi basati sulla fisica (ad es. modellazione meccanica) e metodi basati sui dati (ad es. deep learning).



Sebbene il concetto di DT basato sulla fisica sia relativamente nuovo [1], le sue fondamenta possono essere fatte risalire a costrutti ben noti che sono stati ampiamente utilizzati nella comunità ingegneristica, come illustrato in Figura 1. Si basa sull'analisi ingegneristica e il suo fondamento è nella fisica newtoniana. Aggiungendo il tempo come dimensione, è possibile valutare il comportamento degli asset, consentendo l'analisi delle operazioni in base ai casi di carico previsti. Da qui, l'ulteriore evoluzione del DT si basa su due contributi fondamentali; passare dai casi di carico (previsti) alle osservazioni basate su sensori come input principale per il modello e passare dal tempo simulato al tempo reale (vicino).

Ancor più dei DT basati sulla fisica, le tecnologie dei big data, come l'apprendimento automatico e l'analisi dei dati, hanno ricevuto una notevole attenzione e sono una parte importante del panorama IoT. Svolgono un ruolo fondamentale nell'estrarre intuizioni e conoscenze dall'enorme quantità di dati che vengono creati dall'ampia installazione di sensori nello scafo, nei sistemi di macchinari e nei macchinari di coperta. Il potenziale utilizzo di modelli di apprendimento automatico come base per soluzioni di digital twin si baserebbe su principi piuttosto diversi rispetto alla soluzione basata sulla fisica qui presentata.



L'idea dei data driven digital twin è motivata dalla crescente quantità di dati disponibili: si pensi ai macchinari in una fabbrica intelligente. L'attenzione è posta sui dati perché sono l'elemento chiave che può essere utilizzato per determinare e monitorare gli stati di sistemi complessi. Lo sfruttamento efficace di questi dati contribuisce allo sviluppo di modelli con un'alta fedeltà per i sistemi di produzione. Gli strumenti di machine learning sono metodi chiave che possono essere utilizzati per scoprire informazioni da tali dati. Il loro ruolo nella simulazione, e quindi nel quadro del digital twin, è quello di migliorare le capacità del modello utilizzando la conoscenza estratta. Questa conoscenza sarà aggiornata online per riflettere qualsiasi cambiamento avvenuto nell'ambiente di produzione reale. Ciò consente di migliorare continuamente i diversi aspetti del sistema reale anche in presenza di cambiamenti.

Nella Tabella che segue sono riassunti i vantaggi e gli svantaggi di questi due approcci.

	Machine learning based	Physics based
PRO	- Modello derivato solo dai dati: non è necessaria la conoscenza del dominio - Generico e flessibile: gestisce flussi di dati eterogenei (anche non fisici) - Il modello migliora nel tempo (apprendimento)	- I modelli catturano la profonda conoscenza esistente basata sulla fisica newtoniana - Le relazioni causali forniscono intuizione e comprensione - Incertezza controllata dall'input e dalla precisione

	per rinforzo) • Bravo a scoprire relazioni e modelli complessi	della modellazione • Il modello ha validità universale – prevedere qualsiasi punto coperto dal modello
CONTRO	• La disponibilità dei dati di addestramento necessari per sviluppare il modello • Correlazioni, non causalità. Blackbox, nessuna spiegazione (in particolare, deep learning) • Metodi di approssimazione, nessuna matematica esatta • Le capacità predittive si deteriorano rapidamente al di fuori dell'ambito del training set • Difficile prevedere condizioni estreme/critiche (poche osservazioni)	• Richiede una vasta conoscenza del dominio (fisica). • Calcolo intensivo, sfida in tempo reale • Presupposti completi su input-output devono essere fatti in anticipo

L'utilizzo dei Physics-based implica una profonda conoscenza dei concetti statistici/probabilistici e le relazioni che simulano il sistema possono diventare molto complesse e difficili da gestire. Di contro i metodi di machine learning e intelligenza artificiale pur rappresentando una sorta di black box possono descrivere, con modelli anche semplici, fenomeni molto complessi; infine l'uso di questa tipologia di metodi garantisce un elevato livello di efficienza e flessibilità, in particolar modo quando vi è l'esigenza di applicare le tecniche di modellazione ad un vasto numero di dispositivi ed impianti dei quali non si ha conoscenza a priori. Per tale ragione nel seguito della presente relazione ci si focalizzerà solo su di essi.

1.1 Classificazione delle tecniche di intelligenza artificiale

All'interno dell'intelligenza artificiale (AI) e dell'apprendimento automatico, esistono due approcci di base [2]: l'apprendimento supervisionato e l'apprendimento non supervisionato. La differenza principale è che uno utilizza dati etichettati per aiutare a prevedere i risultati, mentre l'altro no. Tuttavia, ci sono alcune sfumature tra i due approcci e aree chiave in cui uno supera l'altro.

L'apprendimento supervisionato è un approccio di apprendimento automatico definito dall'uso di set di dati etichettati. Questi set di dati sono progettati per addestrare o "supervisionare" gli

algoritmi nella classificazione dei dati o nella previsione accurata dei risultati. Utilizzando input e output etichettati, il modello può misurare la sua accuratezza e apprendere nel tempo.

L'apprendimento supervisionato può essere separato in due tipi di problemi durante il data mining, classificazione e regressione:

- I problemi di classificazione utilizzano un algoritmo per assegnare accuratamente i dati del test in categorie specifiche, come separare le mele dalle arance. Oppure, nel mondo reale, è possibile utilizzare algoritmi di apprendimento supervisionato per classificare lo spam in una cartella separata dalla posta in arrivo. Classificatori lineari, support vector machines, decisional tree e random forest sono tutti tipi comuni di algoritmi di classificazione.
- La regressione è un altro tipo di metodo di apprendimento supervisionato che utilizza un algoritmo per comprendere la relazione tra variabili dipendenti e indipendenti. I modelli di regressione sono utili per prevedere valori numerici basati su diversi punti dati, come le proiezioni dei ricavi delle vendite per una determinata attività. Alcuni popolari algoritmi di regressione sono la regressione lineare, la regressione logistica e la regressione polinomiale.

L'apprendimento non supervisionato utilizza algoritmi di apprendimento automatico per analizzare e raggruppare set di dati senza etichetta. Questi algoritmi scoprono modelli nascosti nei dati senza la necessità di intervento umano (quindi, sono "non supervisionati"). I modelli di apprendimento non supervisionato vengono utilizzati per tre attività principali: raggruppamento, associazione e riduzione della dimensionalità:

- Il clustering è una tecnica di data mining per raggruppare i dati senza etichetta in base alle loro somiglianze o differenze. Ad esempio, gli algoritmi di clustering K-means assegnano punti di dati simili in gruppi, dove il valore K rappresenta la dimensione del raggruppamento e la granularità. Questa tecnica è utile per la segmentazione del mercato, la compressione delle immagini, ecc.
- L'associazione è un altro tipo di metodo di apprendimento non supervisionato che utilizza regole diverse per trovare relazioni tra variabili in un determinato set di dati. Questi metodi sono spesso utilizzati per l'analisi del paniere di mercato e per i motori di raccomandazione, sulla falsariga delle raccomandazioni "I clienti che hanno acquistato questo articolo hanno

acquistato anche".

- La riduzione della dimensionalità è una tecnica di apprendimento utilizzata quando il numero di caratteristiche (o dimensioni) in un determinato set di dati è troppo elevato. Riduce il numero di input di dati a una dimensione gestibile preservando al tempo stesso l'integrità dei dati. Spesso questa tecnica viene utilizzata nella fase di pre-elaborazione dei dati, ad esempio quando gli autoencoder rimuovono il rumore dai dati visivi per migliorare la qualità dell'immagine.

La principale distinzione tra i due approcci è l'uso di set di dati etichettati. Per dirla semplicemente, l'apprendimento supervisionato utilizza dati di input e output etichettati, mentre un algoritmo di apprendimento non supervisionato no. Nell'apprendimento supervisionato, l'algoritmo "apprende" dal set di dati di addestramento facendo iterativamente previsioni sui dati e adattandosi alla risposta corretta. Sebbene i modelli di apprendimento supervisionato tendano ad essere più accurati dei modelli di apprendimento non supervisionato, richiedono un intervento umano iniziale per etichettare i dati in modo appropriato. Ad esempio, un modello di apprendimento supervisionato può prevedere quanto durerà un tragitto giornaliero in base all'ora del giorno, alle condizioni meteorologiche e così via. Ma prima si dovrà addestrarlo per sapere che il tempo piovoso prolunga il tempo di guida.

I modelli di apprendimento senza supervisione, al contrario, lavorano da soli per scoprire la struttura intrinseca dei dati senza etichetta. Si noti che richiedono ancora un intervento umano per la convalida delle variabili di output. Ad esempio, un modello di apprendimento non supervisionato può identificare che gli acquirenti online spesso acquistano gruppi di prodotti contemporaneamente. Tuttavia, un analista di dati dovrebbe confermare che ha senso per un motore di raccomandazione raggruppare i vestiti per bambini con un ordine di pannolini, salsa di mele e tazze con beccuccio.

Altre differenze chiave tra apprendimento supervisionato e non supervisionato sono riportate di seguito.

- Obiettivi: nell'apprendimento supervisionato, l'obiettivo è prevedere i risultati per i nuovi dati e si conosce in anticipo il tipo di risultati attesi. Con un algoritmo di apprendimento non

supervisionato, l'obiettivo è ottenere informazioni dettagliate da grandi volumi di nuovi dati. L'apprendimento automatico stesso determina ciò che è diverso o interessante dal set di dati.

- Applicazioni: i modelli di apprendimento supervisionato sono ideali per il rilevamento dello spam, l'analisi del sentiment, le previsioni meteorologiche e le previsioni dei prezzi, tra le altre cose. Al contrario, l'apprendimento senza supervisione è perfetto per il rilevamento di anomalie, motori di raccomandazione, caratteristiche dei clienti e imaging medico.
- Complessità: l'apprendimento supervisionato è un metodo semplice per l'apprendimento automatico, tipicamente calcolato attraverso l'uso di programmi come R o Python. Nell'apprendimento non supervisionato, sono necessari strumenti potenti per lavorare con grandi quantità di dati non classificati. I modelli di apprendimento non supervisionato sono complessi dal punto di vista computazionale perché richiedono un ampio set di addestramento per produrre i risultati previsti.
- Svantaggi: i modelli di apprendimento supervisionato possono richiedere molto tempo per l'addestramento e le etichette per le variabili di input e output richiedono esperienza. Invece, i metodi di apprendimento senza supervisione possono avere risultati estremamente imprecisi a meno che non si avesse un intervento umano per convalidare le variabili di output.

2 Tecniche di modellazione per il monitoraggio e l'ottimizzazione di asset energetici

Per sviluppare e impiantare DT, è necessario raccogliere le informazioni generate nell'ambiente fisico. Queste informazioni vengono raccolte da sensori e altri dispositivi che raccolgono dati reali sullo stato del processo o del prodotto. D'altra parte, queste informazioni devono essere trattate ed elaborate correttamente. Per gestire tutto questo volume di dati, i sensori e i dispositivi devono essere collegati in un sistema basato su cloud in grado di ricevere i dati in tempo reale ed elaborarli. Una volta che i dati sono stati elaborati, è possibile generare la replica virtuale, che è un DT. Il DT riceve le informazioni dai sensori in tempo reale e imita ciò che accade nella realtà. Di conseguenza,

consente di rappresentare virtualmente possibili problemi che sorgono durante lo sviluppo o il funzionamento di un processo o di una macchina ed è molto utile per il miglioramento dei comportamenti e dell'efficienza delle macchine.

Pertanto, è necessario che l'intero processo sia coinvolto nella creazione di questi ST, e se si ottengono informazioni da tutte le fasi del processo, questa replica virtuale sarà più rappresentativa del processo reale.

I Digital Twin hanno numerosi vantaggi poiché migliorano il comportamento dei processi e dei prodotti, il che generalmente significa migliorarne l'efficienza.

I Digital Twin possono essere applicati in tutti i tipi di industria e per diverse applicazioni, garantendo i seguenti vantaggi:

- **Ottimizzazione delle prestazioni:** è possibile sviluppare modelli basati su dati reali per prevedere comportamenti che ottimizzino le attività per la generazione di energia analizzando grandi quantità di questi dati. I Digital Twin consentono di modellare le prestazioni delle centrali elettriche rappresentando lo stato e il funzionamento. Possono anche aiutare a identificare la domanda di energia, ridurre i costi di implementazione di nuovi impianti e migliorare il processo decisionale nelle fasi di stoccaggio dell'energia.
- **Manutenzione predittiva:** I Digital Twin possono simulare comportamenti e prevenire possibili guasti durante il funzionamento dei dispositivi energetici (sia carichi che generatori). La manutenzione predittiva aiuta gli ingegneri a determinare esattamente quando le apparecchiature necessitano di manutenzione. Riduce i tempi di inattività e previene i guasti alle apparecchiature consentendo di programmare la manutenzione in base alle effettive necessità piuttosto che a un programma predeterminato.
- **Efficientamento della produzione dei dispositivi:** I DT consentono di simulare i processi di produzione dei dispositivi energetici con tutte le loro fasi. Ciò consente, a sua volta, di ottimizzare l'intera produzione. Può essere applicato direttamente sia nel processo di progettazione dei dispositivi energetici che nella loro effettiva produzione. Consentono di prevedere il funzionamento di un prodotto o di un processo produttivo prima del suo avvio, il

che consente di anticipare possibili errori e risparmiare sui costi.

2.1 Modellazione di carichi

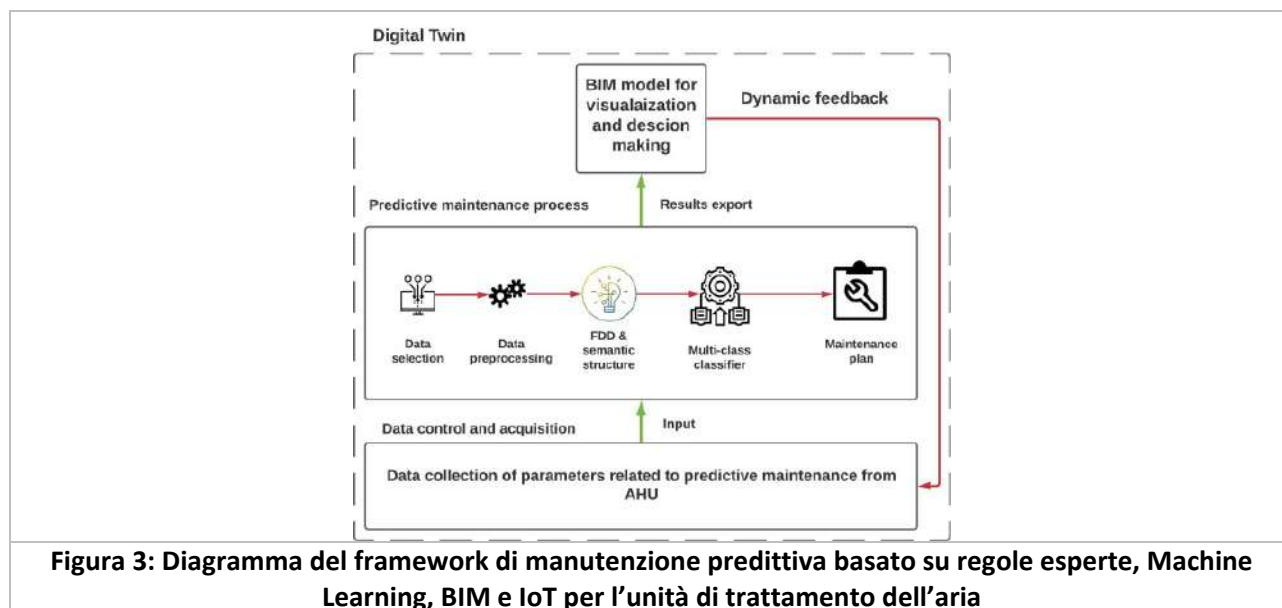
In questa sezione si riportano i più recenti esempi di applicazione di algoritmi di intelligenza artificiale per la modellazione digital twin di carichi in ambito building management e industriale. In particolare, si considereranno le applicazioni per l'ottimizzazione delle prestazioni o il calcolo delle attività manutentiva con approccio predittivo (predictive maintenance). Una sottosezione specifica (2.1.1) tratterà dell'approccio NILM, che ambisce ad modellare singoli carichi partendo da un valore di consumo energetico aggregato.

In [3] viene proposto un framework di manutenzione predittiva, tramite digital twin, di un'unità di trattamento dell'aria (AHU) al fine di superare i limiti dei sistemi di gestione della manutenzione delle strutture (FMM) attualmente in uso negli edifici. Ciò consentirà di diagnosticare i guasti e prevedere la condizione dei componenti dell'edificio, in modo che il personale addetto alla gestione della struttura possa prendere decisioni migliori al momento giusto. La tecnologia digital twin adottata utilizza il Building Information Modeling (BIM), l'Internet of Things (IoT) e le tecnologie semantiche per creare una migliore strategia di manutenzione per le strutture degli edifici.

Il metodo sviluppato dagli autori è stato ottenuto integrando le nuove tecnologie, in particolare BIM, IoT, metadati semantici, expert rules e ML.

Il framework comprende tre fasi principali, ovvero l'acquisizione dei dati, il processo di manutenzione predittiva e il modello BIM per la visualizzazione e il monitoraggio delle informazioni. Le informazioni spaziali possono essere ottenute dal modello BIM. Il modello BIM è stato integrato con i risultati della manutenzione predittiva per supportare il processo decisionale sviluppando un'estensione plug-in per Autodesk Revit utilizzando C sharp, in modo che il team FM possa comprendere facilmente i dati.

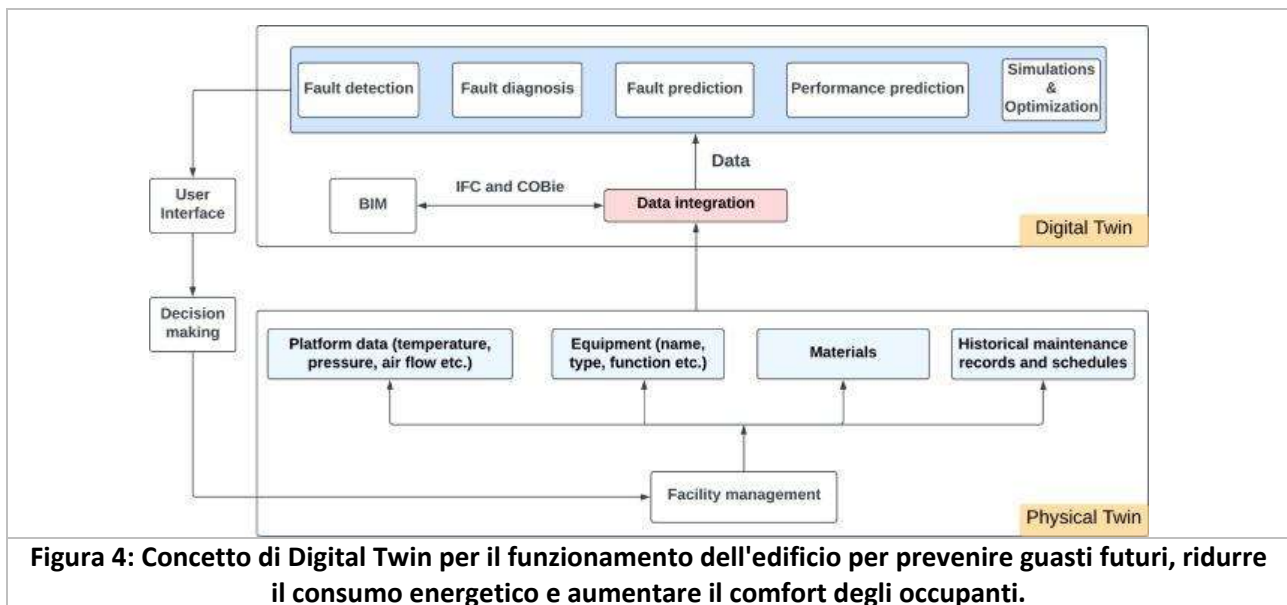
Più in dettaglio, il framework è stato implementato tramite tre moduli che realizzano la manutenzione predittiva: rilevamento dei guasti operativi nell'AHU basato sul metodo APAR (Air Handling Unit Performance Assessment Rules), previsione delle condizioni utilizzando tecniche di apprendimento automatico, **Artificial Neural Network (ANN)**, **Support Vector Machine (SVM)** e **Decision Trees (DT)**, e pianificazione della manutenzione.



I risultati dimostrano che i dati continuamente aggiornati combinati con APAR e algoritmi di apprendimento automatico possono rilevare guasti e prevedere lo stato futuro dei componenti dell'unità di trattamento dell'aria, il che può aiutare nella pianificazione della manutenzione. La rimozione dei guasti operativi rilevati ha comportato un risparmio energetico annuo di diverse migliaia di dollari grazie all'eliminazione dei guasti operativi identificati.

In [4] le prestazioni degli edifici in termini di comfort degli occupanti vengono valutate utilizzando un modello probabilistico basato su **reti bayesiane (BN)**. Il modello BN si è fondato su un'analisi approfondita delle risposte ai sondaggi di soddisfazione e su uno studio approfondito dei parametri prestazionali dell'edificio. Questo studio ha presentato anche una visualizzazione user-friendly compatibile con il BIM per semplificare la raccolta dei dati in due casi di studio dalla Norvegia con

dati dal 2019 al 2022. Nell'articolo viene proposto un nuovo approccio Digital Twin per incorporare il Building Information Modeling (BIM) con sensore che riceve dati in tempo reale, feedback degli occupanti, un modello probabilistico del comfort degli occupanti e rilevamento e previsione dei guasti HVAC che possono influire sul comfort degli occupanti. Vengono inoltre presentati nuovi metodi per l'utilizzo del BIM come piattaforma di visualizzazione, nonché un metodo di manutenzione predittiva per rilevare e anticipare i problemi nel sistema HVAC. Questi metodi aiuteranno i decisori a migliorare le condizioni di comfort degli occupanti negli edifici. Tuttavia, a causa dell'intricata interazione tra numerose apparecchiature e dell'assenza di integrazione dei dati tra i sistemi FM, i dati CMMS, BMS e BIM sono stati integrati in un framework che utilizza i grafici dell'ontologia per generalizzare il framework Digital Twin, in modo che possa essere applicato a molti edifici. I risultati di questo studio possono aiutare i responsabili delle decisioni nel settore del facility management, offrendo informazioni sugli aspetti che influenzano il comfort degli occupanti, accelerando il processo di identificazione dei malfunzionamenti delle apparecchiature e indicando possibili soluzioni.



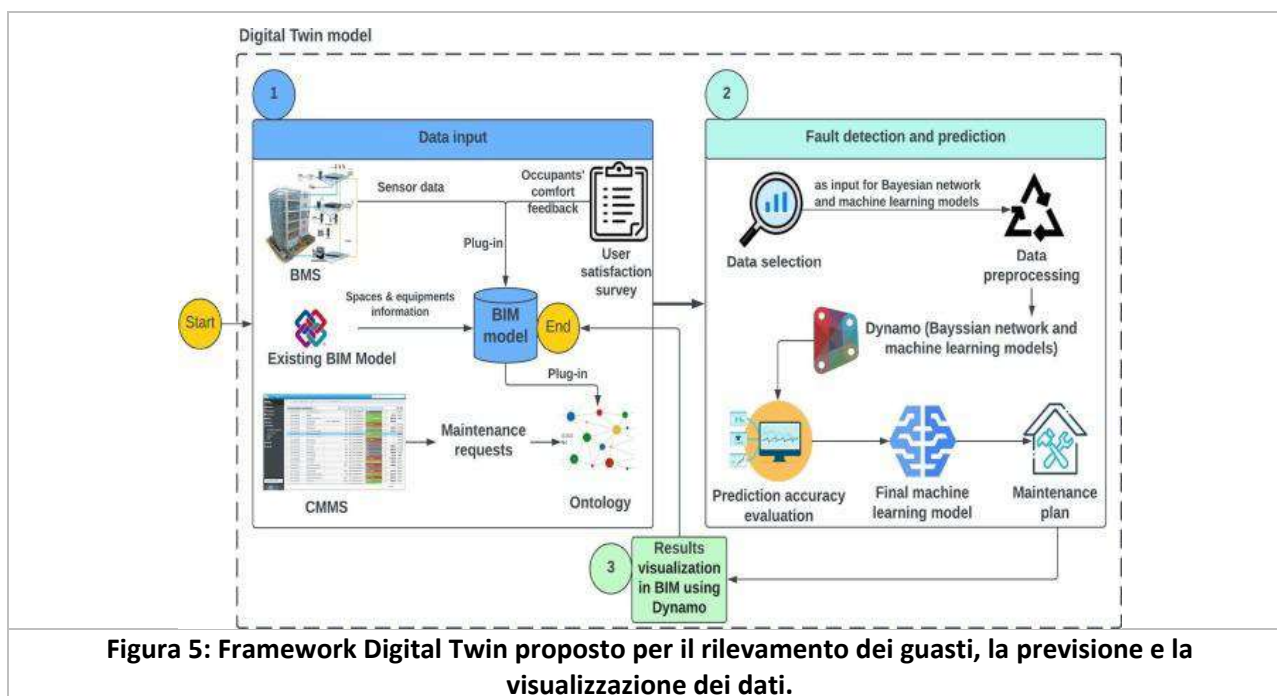


Figura 5: Framework Digital Twin proposto per il rilevamento dei guasti, la previsione e la visualizzazione dei dati.

In conclusione, considerando che l'analisi di diversi fattori ambigui è necessaria per valutare le prestazioni di comfort degli edifici e che l'utilizzo di metodologie convenzionali per quantificare e valutare il comfort degli occupanti negli edifici potrebbe essere difficile sulla base di tali fattori indeterminati, questa ricerca ha dimostrato la creazione di un modello BN per il controllo del comfort termico degli edifici esistenti. Il BN suggerito può descrivere il comfort in un edificio come un processo probabilistico piuttosto che deterministico e quindi fornire livelli di prestazioni di comfort sotto forma di distribuzioni di probabilità. Il vantaggio principale di BN è la sua adattabilità, che gli consente di includere molti tipi di dati e prove, incluso il giudizio di esperti. Pertanto viene introdotto un nuovo approccio Digital Twin che incorpora il feedback degli occupanti (comfort termico, qualità dell'aria interna, comfort visivo, comfort acustico e adeguatezza dello spazio), i dati dei sensori in tempo reale, il modello probabilistico del comfort degli occupanti e la manutenzione predittiva nel BIM, classificati per aspetti di comfort.

I risultati di due casi di studio di due edifici in Norvegia hanno dimostrato che il metodo suggerito potrebbe approfondire la conoscenza degli elementi che influenzano i livelli di disagio degli occupanti e la connessione tra tali fattori, l'ambiente interno e le proprietà fisiche degli edifici. Il framework Digital Twin proposto potrebbe rilevare e diagnosticare più di 17 guasti che il BMS

tradizionale non è in grado di rilevare. Inoltre, con altissima precisione, il framework potrebbe prevedere i guasti che si verificheranno nei prossimi 2 mesi. Inoltre, l'analisi sensibile della qualità dell'aria interna ha dimostrato che la densità di occupazione e i difetti di progettazione HVAC hanno il maggiore impatto sui livelli di comfort. Allo stesso modo, il rapporto finestre/pareti ha la maggiore influenza sulla qualità della luce, suggerendo che un rapporto finestre/pareti da qualche parte nel mezzo, tra il 10 e il 40 per cento, è la scelta migliore.

In [5] viene proposto un modello di pianificazione delle operazioni del sistema di accumulo di energia (ESS) da applicare allo spazio virtuale durante la costruzione di una microrete utilizzando la tecnologia del digital twin. Viene stabilita una programmazione ottimale di carica/scarica ESS per ridurre al minimo le bollette dell'elettricità ed è stata implementata utilizzando tecniche di apprendimento supervisionato come i **Decision Trees (DT)**, i **modelli Nonlinear AutoRegressive network with eXogenous inputs (NARX)** e **MARS** invece delle tecniche di ottimizzazione esistenti. NARX e alberi decisionali sono tecniche di apprendimento automatico. MARS è un modello di regressione non parametrico e la sua applicazione è in aumento. Le sue prestazioni sono state analizzate derivando indicatori di valutazione delle prestazioni per ciascun modello.

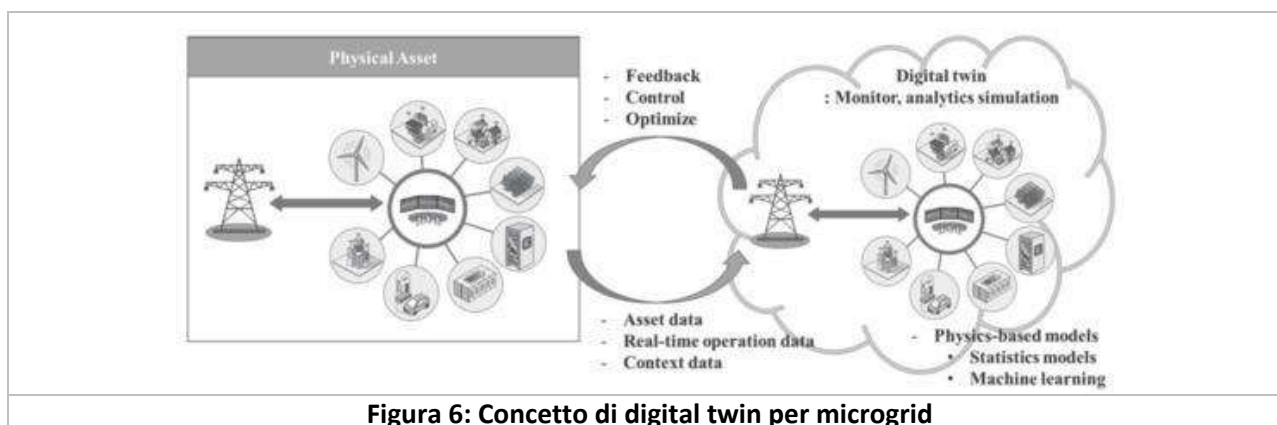


Figura 6: Concetto di digital twin per microgrid

Nello specifico, viene proposto un modello di apprendimento automatico allo scopo di stabilire un programma orario di carica/scarica dell'ESS giornaliero necessario per il funzionamento della microrete, e questo è stato applicato a un ESS effettivo per la valutazione. Attraverso la tecnica di

apprendimento automatico, anche se l'utente non conosce le caratteristiche fisiche dell'ESS, esso può essere modellato in modo simile all'ESS effettivo sulla base di dati passati e poiché viene applicata la tecnica di apprendimento automatico basata sui dati anziché l'algoritmo di ottimizzazione, la difficoltà di stabilire separatamente la funzione obiettivo e i vincoli può essere risolta. È stata valutata l'idoneità di ciascun modello di RMSE e MAE, e questo è stato l'indice di valutazione delle prestazioni del modello di programmazione ottimale ESS. I risparmi nelle bollette dell'elettricità sono stati calcolati e confrontati con i risparmi nelle bollette dell'elettricità ottenuti utilizzando l'ESS effettivo.

Se il problema viene risolto utilizzando un algoritmo basato sull'ottimizzazione, è possibile risolvere il problema derivando la corretta funzione obiettivo piuttosto che i dati passati, ma se il problema viene risolto utilizzando l'apprendimento automatico, è necessario accumulare dati storici per un sufficiente periodo di tempo.

Utilizzando il modello proposto, è emerso in un caso di studio che l'importo del risparmio sulla bolletta dell'elettricità durante il funzionamento dell'ESS è maggiore di quello sostenuto nell'effettivo funzionamento dell'ESS. L'idoneità del modello è stata valutata mediante un'analisi comparativa con il modello di programmazione di carica/scarica ESS basato sull'ottimizzazione.

In [6] viene proposto un nuovo sistema Digital Twin per il riscaldamento, la ventilazione e l'aria condizionata (HVACDT) per ridurre il consumo di energia aumentando il comfort termico. Il framework è stato sviluppato per aiutare i facility manager a comprendere meglio il funzionamento dell'edificio per migliorare la funzione del sistema HVAC. Il framework Digital Twin si basa sul Building Information Modeling (BIM) combinato con un plug-in di nuova creazione per ricevere i dati dei sensori in tempo reale, nonché il comfort termico e il processo di ottimizzazione attraverso la programmazione Matlab. Per determinare se il quadro suggerito è pratico, i dati sono stati raccolti da un edificio per uffici norvegese tra agosto 2019 e ottobre 2021 e utilizzati per testare il quadro. Una **Artificial Neural Network (ANN)** in un modello Simulink e un **algoritmo genetico multiobiettivo (MOGA)** vengono quindi utilizzati per migliorare il sistema HVAC. Il sistema HVAC è composto da

distributori d'aria, unità di raffreddamento, unità di riscaldamento, regolatori di pressione, valvole, porte d'aria e ventilatori, tra gli altri componenti. In questo contesto, diverse caratteristiche, come temperatura, pressione, flusso d'aria, controllo del funzionamento del raffreddamento e del riscaldamento e altri fattori sono considerate variabili decisionali. Per determinare le funzioni obiettivo, vengono calcolati sia la percentuale prevista di insoddisfatti (PPD) sia il consumo di energia HVAC. Di conseguenza, le variabili decisionali di ANN e la funzione obiettivo sono state ben correlate. Inoltre, MOGA presenta diversi fattori progettuali che possono essere utilizzati per ottenere la migliore soluzione possibile in termini di comfort termico e consumo energetico.

Come anticipato, il nuovo metodo HVACDT è stato sviluppato per ridurre il consumo energetico di HVAC e aumentare il comfort termico degli occupanti negli edifici. Il processo del metodo HVACDT inizia estraendo i dati appropriati dal modello BIM e termina inviando i dati delle opzioni di progettazione ottimali al modello BIM.

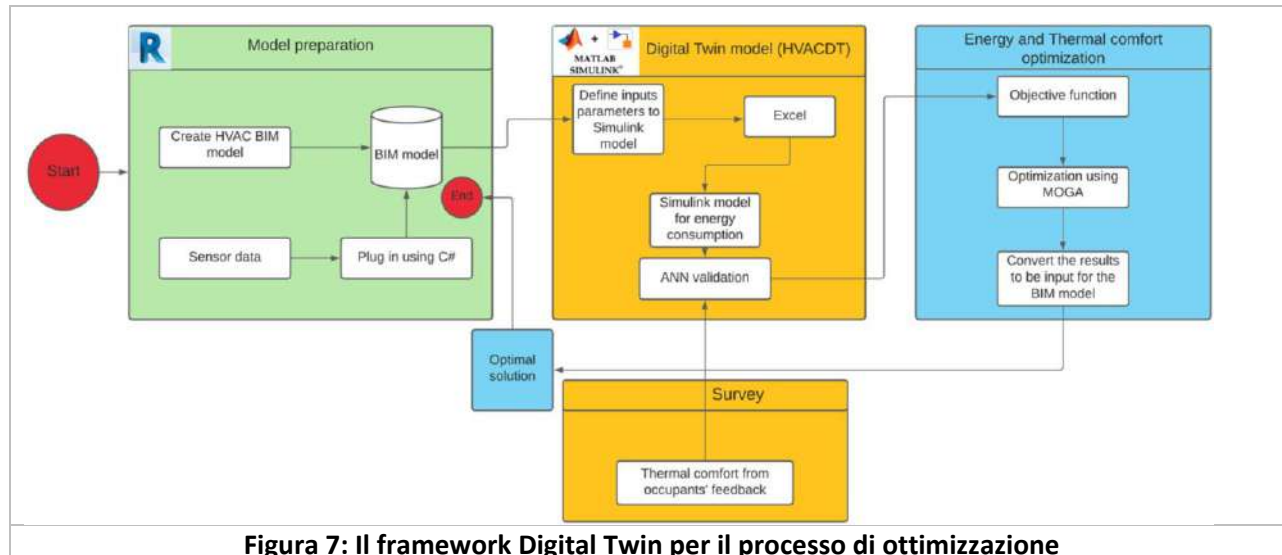
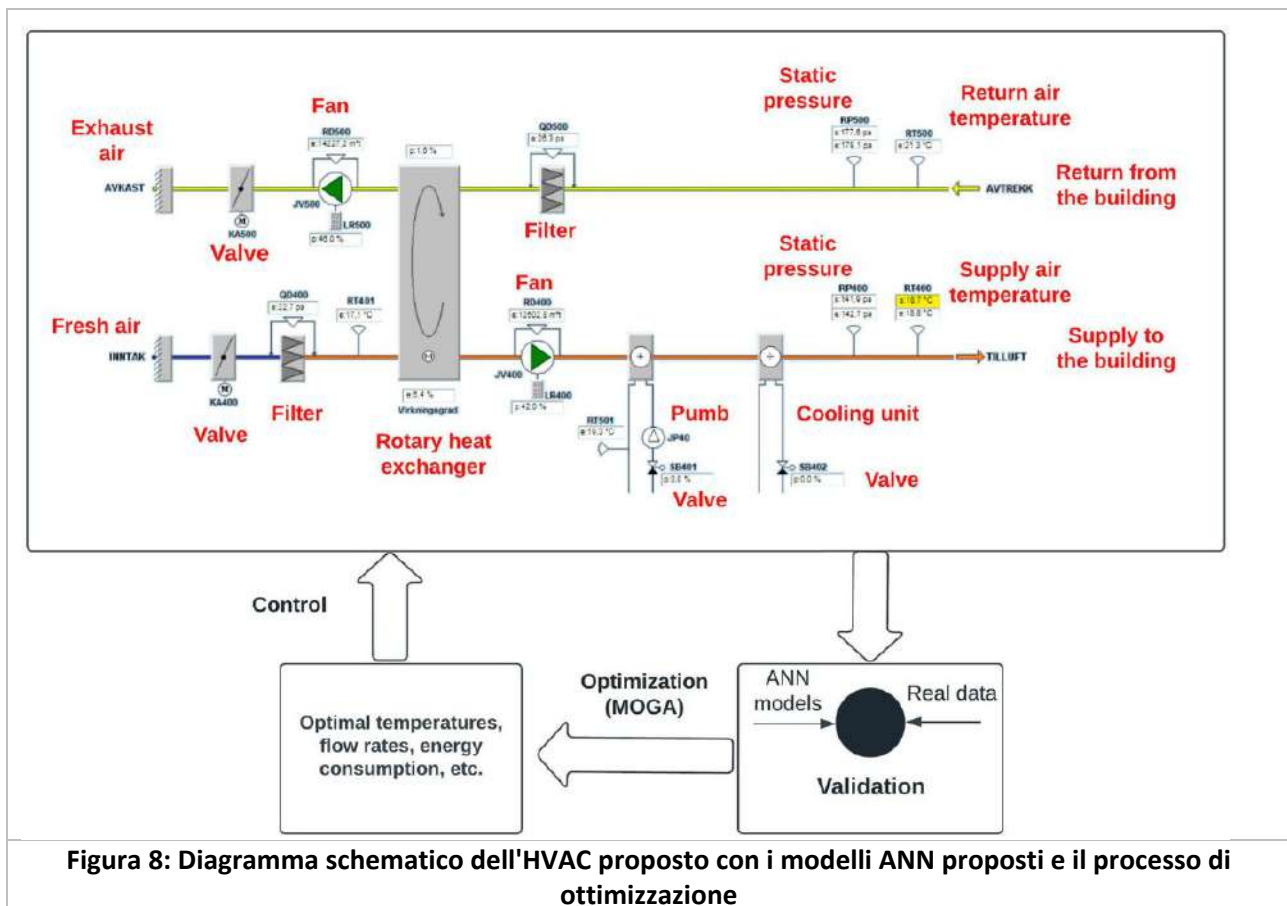


Figura 7: Il framework Digital Twin per il processo di ottimizzazione

La seguente figura mostra, invece, una rappresentazione schematica dell'HVAC. Questo schema illustra i modelli di rete neurale artificiale (ANN) proposti e le loro interazioni con il resto del sistema.



Lo sviluppo del sistema HVACDT ha incluso la creazione e la preparazione del modello BIM per l'estrazione dei dati, la programmazione MATLAB che si concentra sulla personalizzazione di ANN e MOGA per adattarsi al caso di studio e generare soluzioni di ottimizzazione e riportare la soluzione ottimizzata al modello BIM. Un approccio di ottimizzazione multi-obiettivo che combina ANN e MOGA è stato efficacemente impiegato in Matlab per definire l'operazione di costruzione ideale.

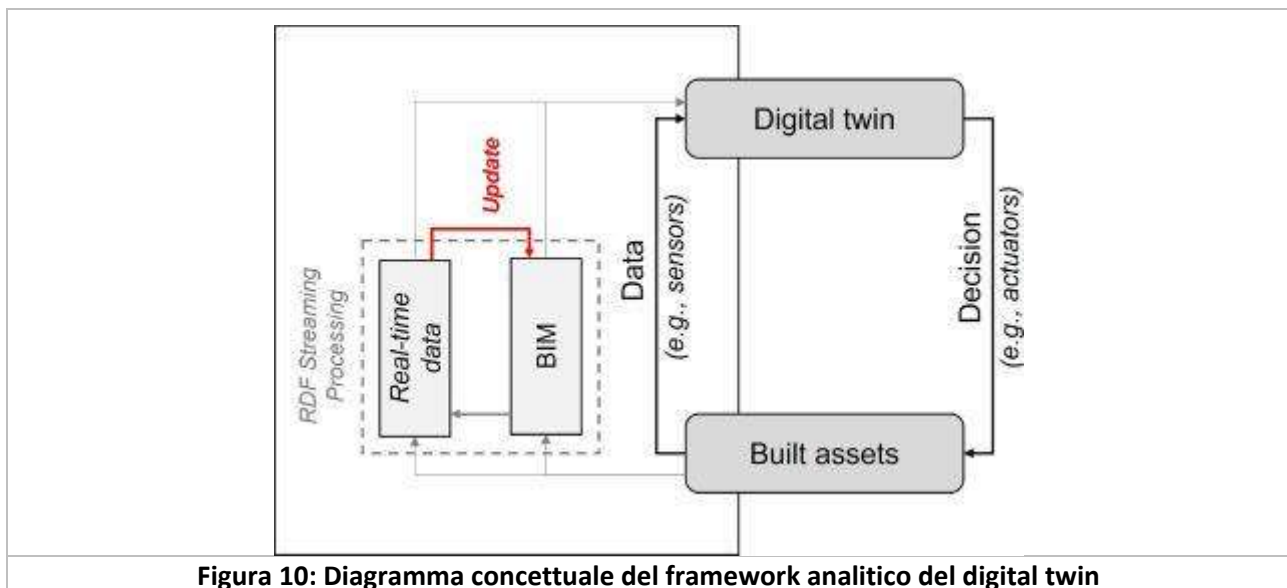
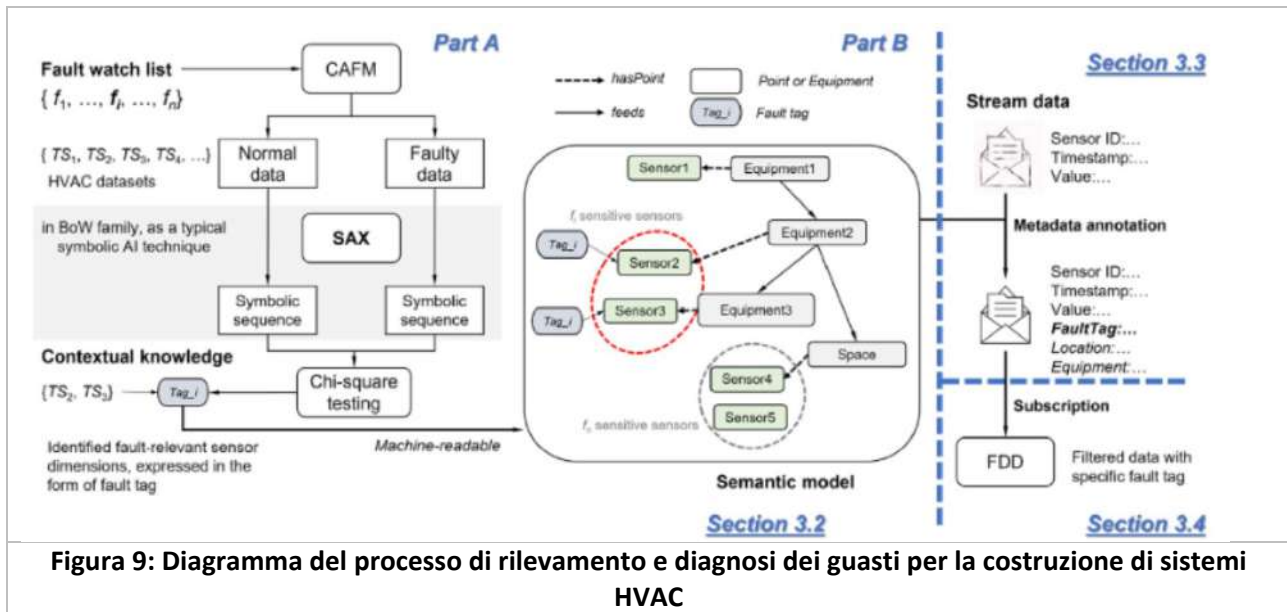
In base ai risultati dell'ottimizzazione, rispetto al progetto effettivo, l'ottimizzazione multi-obiettivo ha migliorato significativamente il funzionamento HVAC per il comfort termico mantenendo un basso consumo energetico. L'integrazione di BIM, la programmazione C# e le tecniche di ottimizzazione multi-obiettivo nel processo di progettazione di HVACDT hanno consentito un esame più approfondito delle configurazioni di progettazione HVAC e un migliore supporto per le decisioni dei progettisti. Inoltre, in questo studio viene fornito un nuovo processo di gestione dei dati basato su Application Programming Interfaces (Revit API) in un ambiente di programmazione (C# e

Windows Presentation Foundation (WPF)) invece di utilizzare il formato di scambio dei dati esistente, come IFC.

Secondo i risultati, il risparmio medio di energia per il raffrescamento per quattro giorni estivi è di circa il 13,2%, 10,8% per i tre mesi estivi (giugno, luglio e agosto) contribuendo a mantenere il PPD inferiore al 10%. Il PPD minimo è del 6,2% per l'inverno e del 3,4% per l'estate, corrispondenti rispettivamente a 46,4 e 42,9 kW.

In questa ricerca, l'HVACDT si è concentrato esclusivamente sul sistema HVAC. Il processo decisionale può essere migliorato in futuro includendo più variabili, come l'indice di utilizzo dell'energia (IUE), l'illuminazione diurna, i costi del ciclo di vita e l'efficienza della ventilazione naturale.

Prendendo come caso il processo di rilevamento e diagnosi dei guasti per la costruzione di sistemi HVAC l'articolo [7] ha adottato una tecnica di **intelligenza artificiale simbolica** per identificare dimensioni sensoriali informative per guasti specifici dell'edificio esplorando la rappresentazione simbolica di serie temporali etichettate. Per preservare questa conoscenza temporale ad hoc nell'ecosistema del digital twin, vengono definiti tag di errore leggibili dalla macchina per etichettare le corrispondenti entità del sensore. Viene sviluppata una piattaforma dati digital twin per annotare i dati in tempo reale con tag di errore e produrre flussi di dati filtrati a bassa latenza associati a un tag specificato per automatizzare questo processo. Nello specifico, viene utilizzato un approccio basato sui digital twin per identificare e raccogliere automaticamente dati informativi per supportare la gestione dinamica delle risorse. Dati abbondanti, in particolare dati in tempo reale, vengono generati in sistemi di costruzione dinamici. Questi pezzi di dati, che contengono una sconcertante quantità di informazioni, devono essere selezionati e filtrati per guidare le corrispondenti funzionalità intelligenti dell'edificio, altrimenti le dimensioni dei dati irrilevanti maschererebbero e cancellerebbero i dati più informativi.



Nel caso di FDD, viene introdotto il metodo di estrazione e selezione delle caratteristiche basato su Bag-of-Words (cioè SAX), come una tipica tecnica di intelligenza artificiale simbolica, per aumentare l'efficienza computazionale. Esso cerca di estrarre le caratteristiche temporali delle serie temporali storiche e identificare le dimensioni sensoriali distintive dai dati etichettati normali e difettosi, il che approfondisce la comprensione di tutti i difetti interessati. I "tag di errore" sono definiti di conseguenza nel modello Brick per etichettare queste entità sensoriali rilevanti, preservando la

conoscenza/comprendimento contestuale appresa in modo leggibile dalla macchina. La piattaforma di dati del digital twin viene adottata per integrare i tag di conoscenza definiti con i dati in tempo reale. Annotando il flusso di dati con tag di errore ausiliari, è possibile estrarre automaticamente flussi di dati in tempo reale a larghezza di banda elevata a bassa latenza aggiunti con il tag di errore specificato per alimentare la funzionalità FDD. Questo informa il modo per abilitare le funzionalità di gestione dinamica delle risorse attraverso il gemellaggio digitale. La strategia "divide et impera" qui utilizzata contribuisce al carattere in tempo reale e alla riduzione del carico computazionale per la fornitura di funzionalità intelligenti di costruzione. Per il futuro dei Digital Twin, il framework proposto consente di superare il problema dei "grandi" dati (alto volume, alta varietà, ma non alta velocità) e filtrare la grande quantità di dati provenienti dall'ecosistema dei Digital Twin. Grazie agli approcci di tagging e annotazione dei metadati, le metodologie sviluppate possono essere replicate e consentono di trarre vantaggio dalla conoscenza ontologica e dalla conoscenza ad-hoc delle tecniche di intelligenza artificiale per supportare funzionalità di costruzione intelligente. L'espressività e la ricchezza della semantica vengono sfruttate per lo sviluppo di funzionalità complesse, consentendo di rappresentare meglio le comprensioni complete sui sistemi dell'edificio e guidare la gestione in tempo reale delle risorse fisiche.

In [8] viene proposto un digital twin multistrato basato sui dati del sistema energetico che mira a rispecchiare il consumo energetico effettivo delle famiglie sotto forma di un gemello digitale domestico (HDT). Se collegato al gemello digitale per la produzione di energia (EDT), l'HDT consente al modello di ottimizzazione energetica incentrato sulla famiglia di raggiungere l'efficienza desiderata nell'uso dell'energia. Il modello intende migliorare l'efficienza della produzione di energia appiattendolo i livelli di fabbisogno energetico giornaliero. Ciò viene fatto riorganizzando in modo collaborativo i modelli di consumo energetico delle abitazioni residenziali per evitare i picchi di domanda, soddisfacendo al contempo le esigenze dei residenti e riducendo i costi energetici. Infatti, il sistema proposto incorpora il primo modello HDT per misurare l'impatto di varie modifiche sulla bolletta energetica domestica e, successivamente, sulla produzione di energia. Il sistema energetico proposto viene applicato a un set di dati IoT del mondo reale che copre oltre due anni e diciassette famiglie.

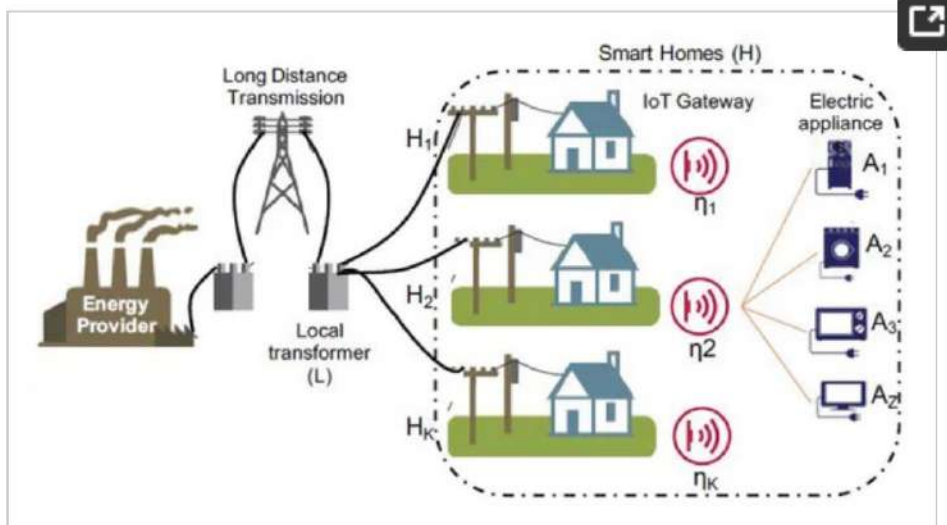


Figura 11: Modello di sistema che collega il produttore di energia alle famiglie residenziali nelle case intelligenti

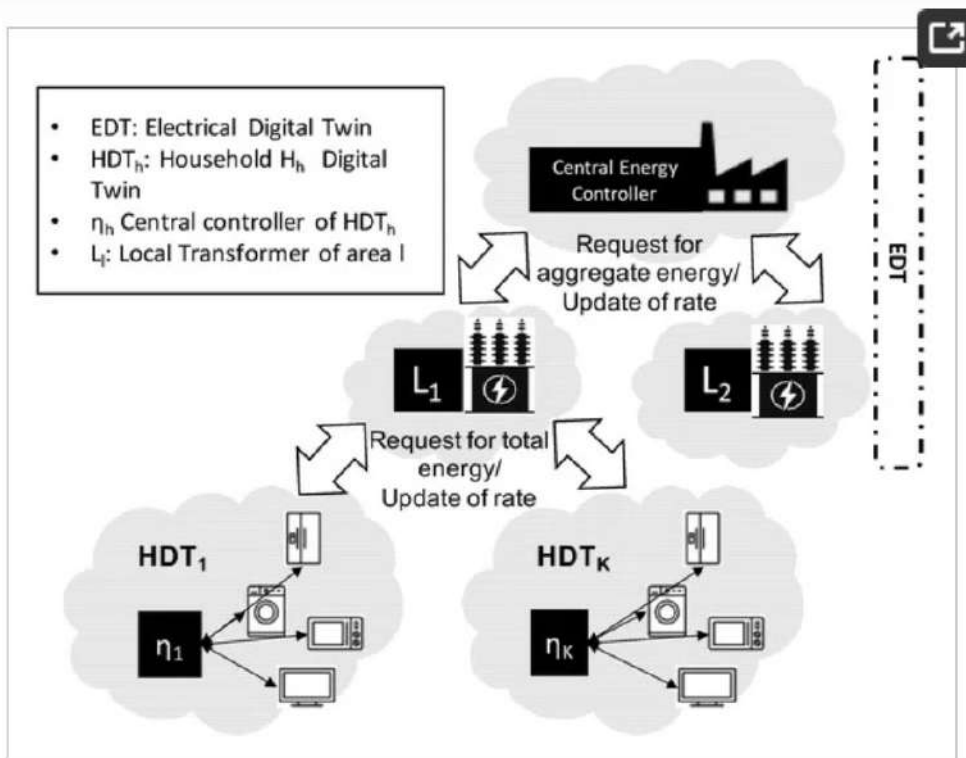


Figura 12: Rappresentazione Digital Twin (DT) a più livelli del modello di sistema e scambio di dati.

Viene proposto un approccio di **Reinforced Learning (RL)** basato sull'edge per riprogrammare intenzionalmente gli elettrodomestici e spingere la domanda energetica collettiva verso un modello

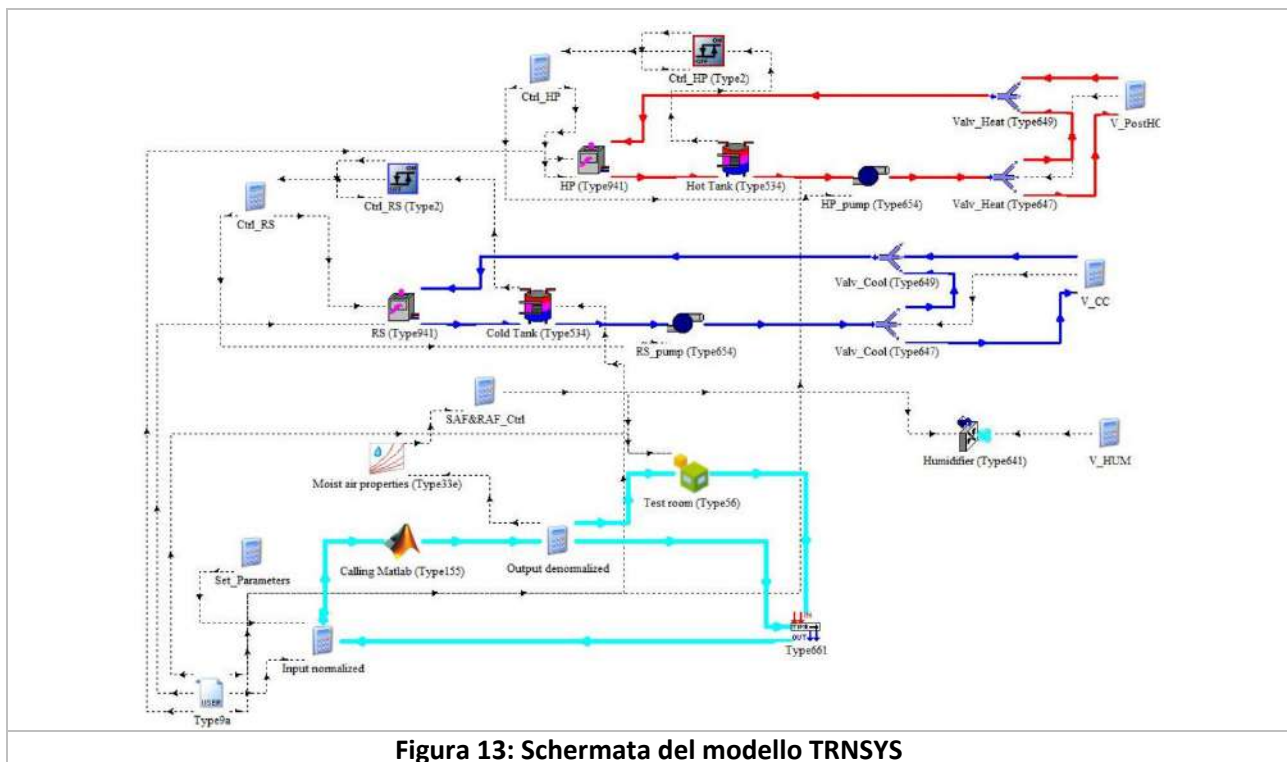
più piatto. La nuova architettura ha protetto la privacy della famiglia ai margini del sistema, ovvero un gateway intelligente IoT installato in ogni famiglia. Il gateway intelligente ha raccolto il consumo energetico orario in tempo reale per tutti gli elettrodomestici di una determinata abitazione. Ha quindi condiviso le informazioni aggregate con l'impianto di produzione di energia senza rivelare dati e comportamenti domestici specifici della casa.

L'approccio proposto per il Reinforced Learning è stato adattivo. Ad esempio, quando si installano nuovi elettrodomestici o si hanno nuovi membri della famiglia, RL può adattarsi in modo efficace e ottenere risultati ottimizzati regolando la programmazione degli elettrodomestici in ogni famiglia per ridurne al minimo il costo energetico. In linea di principio, l'ottimizzazione avviene nella replica virtuale (HDT) e verrebbe applicata agli asset fisici solo se i risultati sono soddisfacenti. Nel complesso, l'obiettivo principale dell'algoritmo è stato ridurre i costi energetici per il settore residenziale massimizzando al contempo il comfort dell'utente. Poiché l'EDT controlla i parametri di fatturazione dell'energia, questi sono stati progettati in modo efficace in modo che il metodo RL basato su edge potesse ottimizzare con successo i modelli di utilizzo collettivo dell'energia ed evitare le richieste di picco di energia.

Gli esperimenti condotti su set di dati sintetici e di case intelligenti del mondo reale hanno mostrato che l'architettura proposta e l'approccio RL autoadattativo hanno ridotto efficacemente la dispersione della domanda collettiva di energia diurna rispettivamente del 20,9% e del 20,4% per i set di dati sintetici e della vita reale. Il metodo proposto ha ridotto con successo il costo energetico per famiglia rispettivamente del 10,7% e del 17,7% per i set di dati sintetici e della vita reale.

In [9] viene indagato sperimentalmente l'impianto di riscaldamento, ventilazione e condizionamento (HVAC) a servizio della sala prove del SENS i-Lab del Dipartimento di Architettura e Disegno Industriale dell'Università della Campania Luigi Vanvitelli (Aversa) attraverso un serie di test eseguiti durante l'estate e l'inverno sia in scenari normali che difettosi. In particolare, cinque guasti tipici distinti sono stati implementati artificialmente nel sistema HVAC e analizzati durante il funzionamento transitorio e stazionario. Un modello ottimale di sistema basato su rete neurale

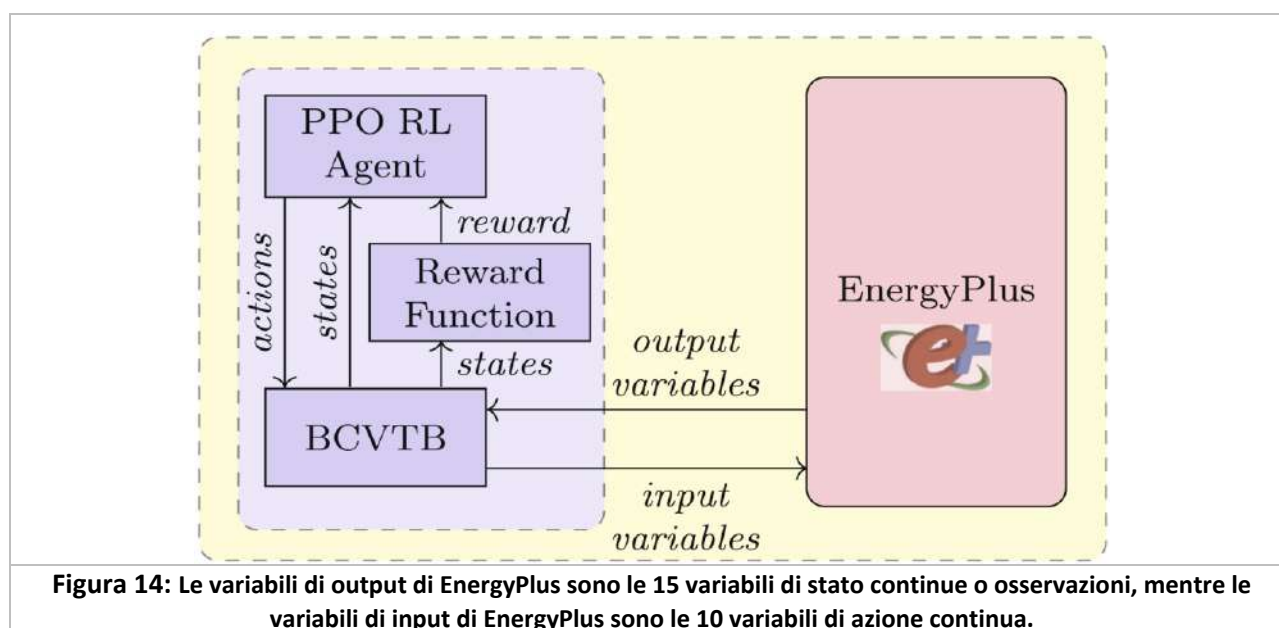
artificiale è stato creato nella piattaforma MATLAB e verificato confrontando i dati sperimentali con le previsioni di ventidue diverse architetture di rete neurale. L'architettura di rete neurale artificiale selezionata è stata accoppiata con un modello di simulazione dinamica sviluppato utilizzando la piattaforma software TRaNsient SYStems (TRNSYS) con gli obiettivi principali di (i) rendere disponibile un set di dati sperimentale caratterizzato da dati etichettati normali e difettosi che coprono un'ampia gamma di condizioni operative e climatiche; (ii) fornire uno strumento di simulazione accurato in grado di generare dati operativi per assistere ulteriori ricerche nel rilevamento dei guasti e nella diagnosi delle unità HVAC; e (iii) valutare l'impatto di guasti selezionati sul comfort termoigrometrico interno degli occupanti, sulle tendenze temporali dei parametri chiave del sistema operativo e sui consumi di energia elettrica.



Il modello basato su **Artificial Neural Network (ANN)** del sistema HVAC (sviluppato nell'ambiente MATLAB e confrontato con i dati misurati) ha evidenziato che è in grado di fornire una caratterizzazione rigorosa delle prestazioni stazionarie e transitorie del sistema HVAC sotto scenari normali oppure difettosi. Più in dettaglio, il modello è stato caratterizzato dai valori medi del coefficiente di determinazione R^2 nella previsione della temperatura dell'aria di mandata,

dell'umidità relativa dell'aria di mandata, della percentuale di apertura della valvola della batteria di postriscaldamento, della percentuale di apertura della valvola della batteria di raffreddamento e della percentuale di apertura della valvola dell'umidificatore molto vicina ai valori massimi e, rispettivamente, pari a 0,95 °C, 0,93%, 0,95%, 0,97% e 0,96%.

In [10] si propone una nuova architettura di **Reinforced Learning (RL)** per la programmazione e il controllo efficienti del sistema di riscaldamento, ventilazione e condizionamento dell'aria (HVAC) in un edificio commerciale, sfruttando al contempo i suoi potenziali di risposta alla domanda (DR). Con i progressi nei sistemi automatizzati di gestione degli edifici, questo può essere ottenuto senza soluzione di continuità da un agente RL autonomo intelligente che intraprende l'azione migliore, ad esempio, una modifica del setpoint della temperatura HVAC, necessaria per modificare il modello di utilizzo dell'elettricità di un edificio in risposta ai segnali di risposta alla domanda e con un minimo impatto sul comfort termico per i clienti. Precedenti ricerche in quest'area hanno affrontato solo singoli aspetti del problema utilizzando RL. In particolare, a causa delle sfide nell'implementazione della risposta alla domanda con modelli di intero edificio, sono stati utilizzati modelli analitici più semplici che catturano male la realtà. E dove vengono applicati i modelli dell'intero edificio, RL viene utilizzato per il controllo HVAC principalmente per raggiungere obiettivi di efficienza energetica trascurando la risposta alla domanda. Così, in questa ricerca, si implementa un quadro olistico progettando un controller RL efficiente per un modello di edificio intero che impara a ottimizzare e controllare il sistema HVAC per migliorare l'efficienza energetica e i livelli di comfort termico oltre a raggiungere gli obiettivi di risposta alla domanda.



I risultati della simulazione hanno dimostrato che è possibile ottenere un notevole risparmio energetico utilizzando l'apprendimento per rinforzo rispetto al controller EnergyPlus artigianale, mantenendo livelli di comfort termico accettabili.

Nello specifico, i risultati della simulazione hanno mostrato che, applicando l'apprendimento per rinforzo per il normale funzionamento HVAC, è possibile ottenere una riduzione energetica massima settimanale fino al 22% rispetto a un controller di base realizzato a mano. Inoltre, utilizzando un controller RL compatibile con DR durante i periodi di risposta alla domanda, è possibile ottenere riduzioni o aumenti di potenza media fino al 50% su base settimanale rispetto al controller RL predefinito, mantenendo i livelli di comfort termico degli occupanti entro limiti accettabili.

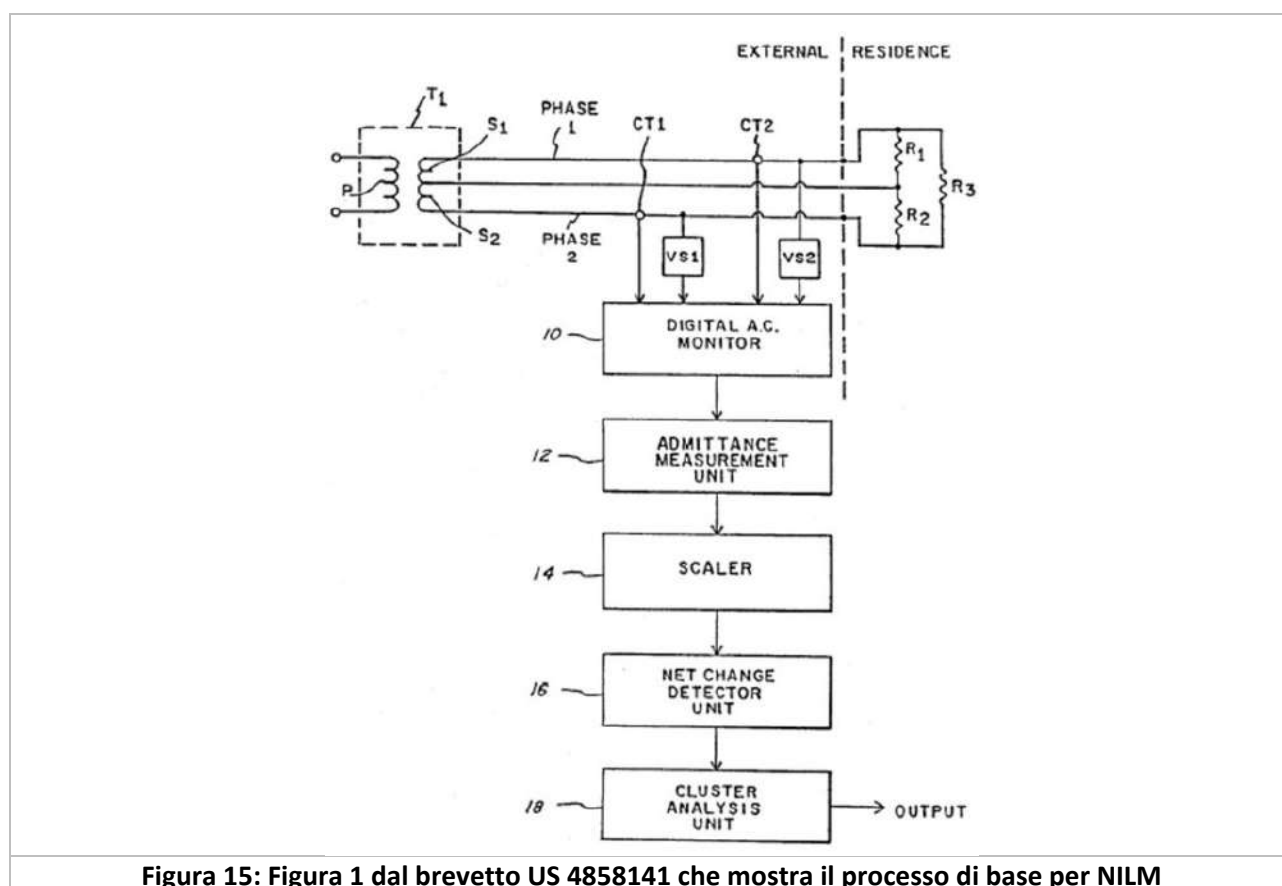
2.1.1 Approccio NILM per la disaggregazione delle misure elettriche

Il monitoraggio del carico non intrusivo (NILM), o monitoraggio del carico degli apparecchi non intrusivi (NIALM), è un processo per analizzare i cambiamenti nella tensione e nella corrente che entrano in una casa e dedurre quali apparecchi sono utilizzati in casa e il loro consumo energetico individuale. I contatori elettrici con tecnologia NILM vengono utilizzati dalle società di servizi

pubblici per rilevare gli usi specifici dell'energia elettrica nelle diverse abitazioni. NILM è considerato un'alternativa a basso costo al collegamento di singoli monitor su ogni appliance. Tuttavia, presenta problemi di privacy.

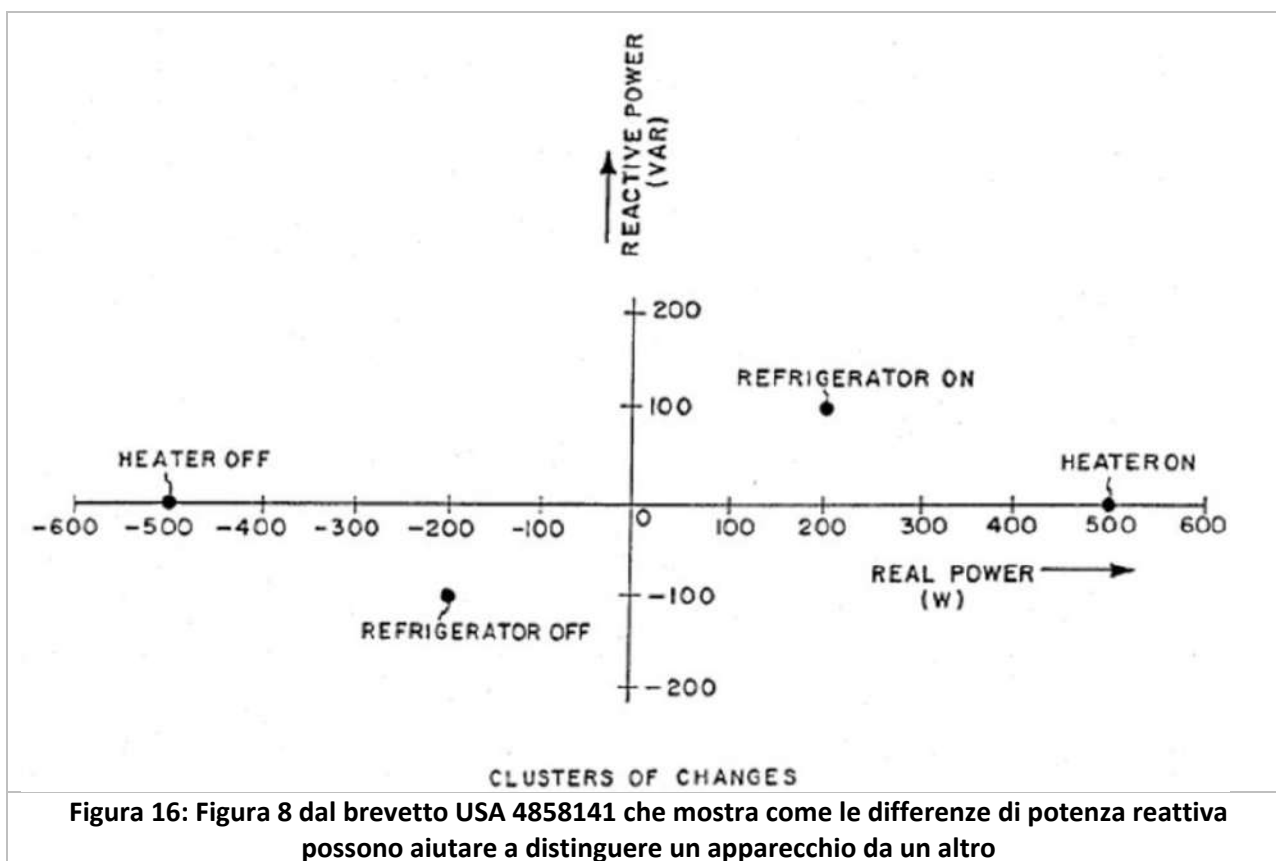
Il monitoraggio del carico non intrusivo è stato inventato da George W. Hart, Ed Kern e Fred Schwebpe del MIT all'inizio degli anni '80 con il finanziamento dell'Electric Power Research Institute.

Il processo di base è descritto nel brevetto statunitense 4,858,141. Come mostrato nella figura seguente, un monitor CA digitale è collegato all'alimentazione monofase che entra in una residenza. Le variazioni della tensione e della corrente vengono misurate (ad es. unità di misura dell'ammettenza), normalizzate (scaler) e registrate (unità di rilevamento della variazione netta). Viene quindi eseguita un'analisi del cluster per identificare quando le diverse appliance vengono accese e spente.



Ad esempio, se si accende una lampadina da 60 watt, seguita dall'accensione di una lampadina da 100 watt, seguita dallo spegnimento della lampadina da 60 watt e dallo spegnimento della lampadina da 100 watt, l'unità NIALM corrisponderà i segnali di accensione e spegnimento della lampadina da 60 watt e i segnali di accensione e spegnimento della lampadina da 100 watt per determinare la quantità di energia utilizzata da ciascuna lampadina e quando. Il sistema è sufficientemente sensibile da poter discriminare le singole lampadine da 60 watt a causa delle normali variazioni nell'assorbimento di potenza effettivo delle lampadine con la stessa potenza nominale (ad esempio, una lampadina potrebbe assorbire 61 watt, un'altra 62 watt).

Il sistema può misurare sia la potenza reattiva che la potenza reale. Quindi due apparecchi con lo stesso assorbimento di potenza totale possono essere distinti dalle differenze nella loro impedenza complessa. Come mostrato nella figura riportata di seguito, ad esempio, un motore elettrico del frigorifero e un riscaldatore puramente resistivo possono essere distinti in parte perché il motore elettrico presenta variazioni significative di potenza reattiva quando si accende e si spegne, mentre il riscaldatore non ne ha quasi nessuna.



I sistemi NILM possono anche identificare apparecchi con una serie di modifiche individuali nell'assorbimento di potenza. Queste appliance sono modellate come macchine a stati finiti. Una lavastoviglie, ad esempio, ha riscaldatori e motori che si accendono e si spengono durante un tipico ciclo di lavaggio delle stoviglie. Questi verranno identificati come cluster e verrà registrato il consumo di energia per l'intero cluster. È quindi possibile identificare l'assorbimento di potenza della "lavastoviglie" in contrapposizione a "gruppo di riscaldamento a resistenza" e "motore elettrico".

Il monitoraggio del carico non intrusivo ha ampie prospettive di applicazione a causa del suo basso costo di implementazione e della scarsa interferenza per gli utenti di energia, cosa molto attesa nel settore industriale di recente a causa dello sviluppo di algoritmi di apprendimento. Mirando all'indagine sul monitoraggio del carico pratico e affidabile nelle implementazioni sul campo, in [11] viene proposto un approccio di disaggregazione del carico non intrusivo basato su un **algoritmo di apprendimento della rete neurale potenziato**. L'approccio di monitoraggio dell'appliance

presentato stabilisce inizialmente il modello di rete neurale seguendo la strategia di apprendimento supervisionato e quindi utilizza l'ottimizzazione basata sull'apprendimento non supervisionato per migliorare la flessibilità e l'adattabilità per diversi scenari, portando al miglioramento delle prestazioni di disaggregazione.

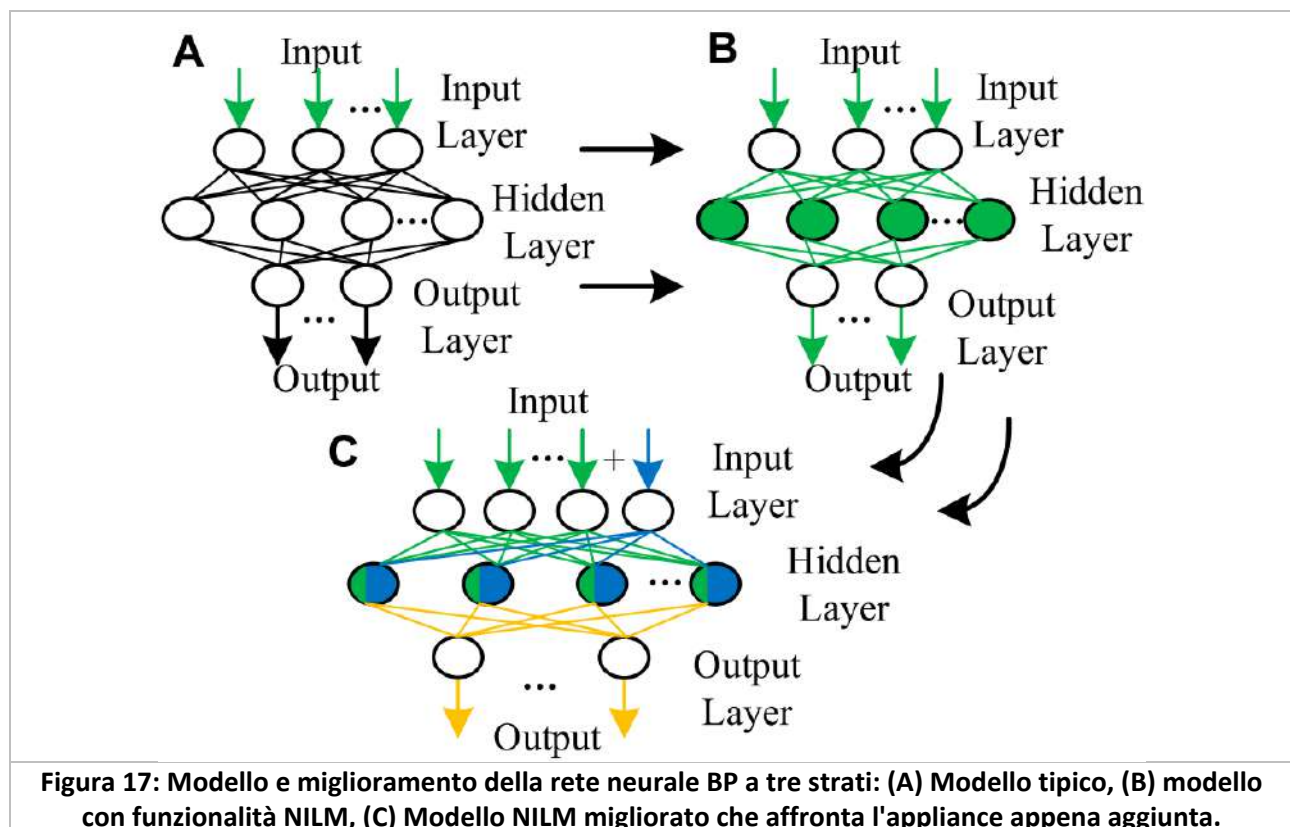


Figura 17: Modello e miglioramento della rete neurale BP a tre strati: (A) Modello tipico, (B) modello con funzionalità NILM, (C) Modello NILM migliorato che affronta l'appliance appena aggiunta.

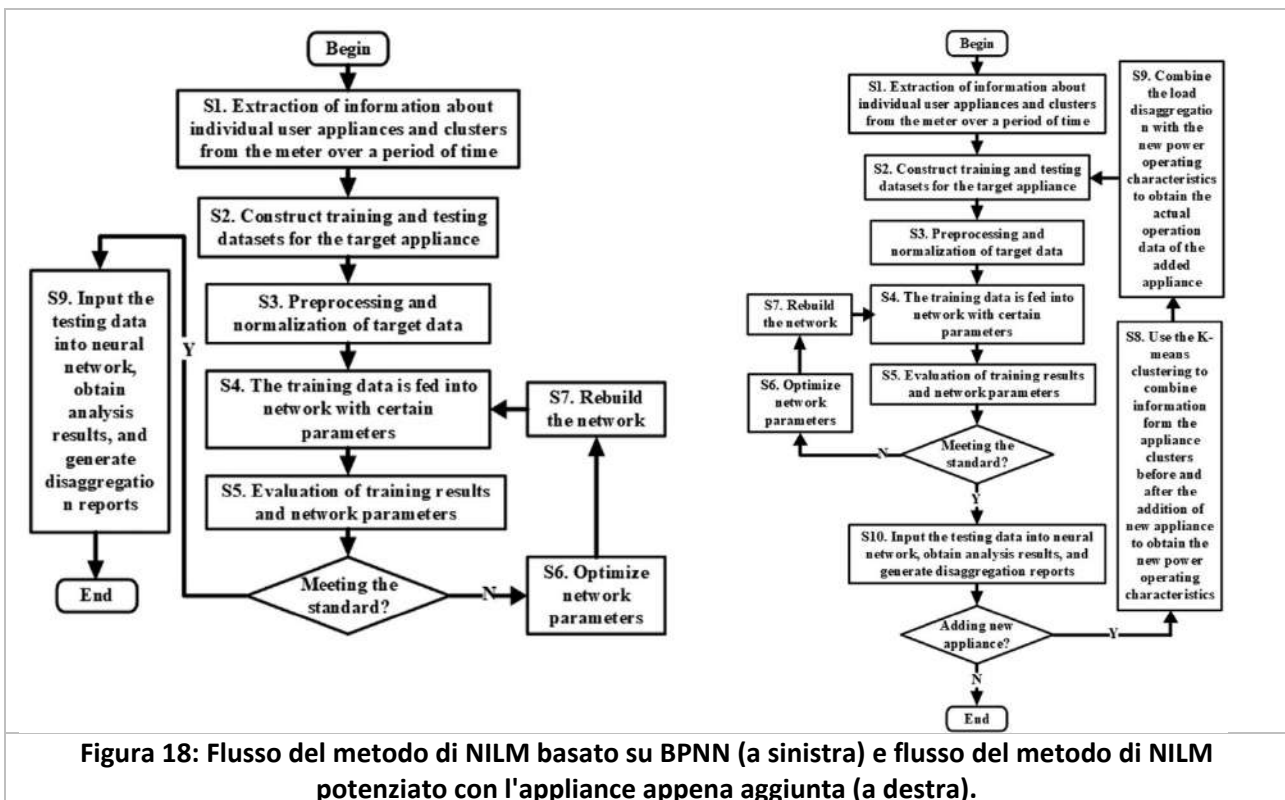


Figura 18: Flusso del metodo di NILM basato su BPNN (a sinistra) e flusso del metodo di NILM potenziato con l'appliance appena aggiunta (a destra).

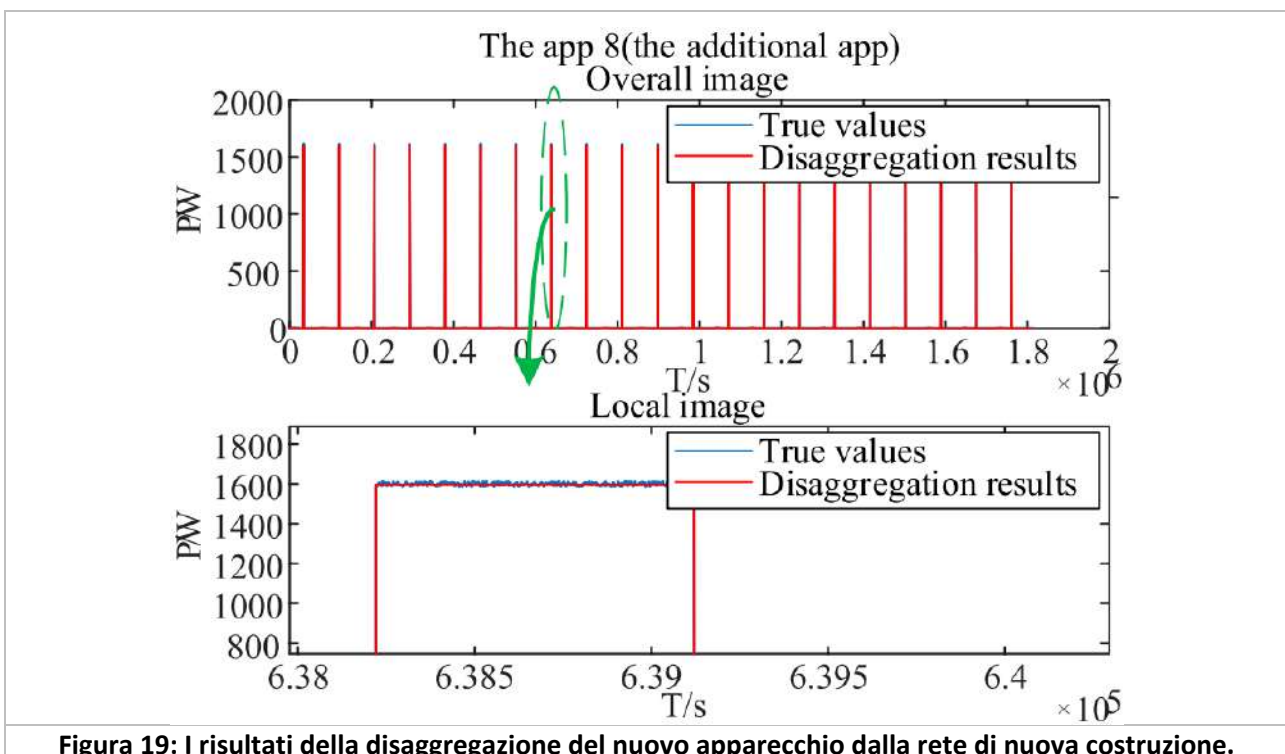


Figura 19: I risultati della disaggregazione del nuovo apparecchio dalla rete di nuova costruzione.

La soluzione basata su BPNN viene studiata a fondo e migliorata in questo articolo, verso la pratica di disaggregazione del carico non intrusiva nelle misurazioni sul campo, in particolare per affrontare l'appliance appena aggiunta con campioni carenti nel campo NILM. Nello specifico, viene proposto un approccio NILM basato sul **Deep Learning**, in cui il modello viene stabilito sulla base della rete neurale BP supervisionata e potenziato dall'ottimizzazione non supervisionata, e la soluzione è risultata dotata di scalabilità adattiva.

Grazie al progetto di ottimizzazione congiunta basato su grandi campioni di apparecchi originali e piccoli campioni di apparecchi appena aggiunti, le nuovissime funzionalità di apparecchi sconosciuti vengono gestite efficacemente dal modello BPNN potenziato, portando alla capacità di riconoscere l'apparecchio fuori portata.

Dalle verifiche sul set di dati pubblico REDD, l'approccio proposto ha dimostrato di avere buone prestazioni nel monitoraggio del carico non intrusivo. Oltre al miglioramento della precisione, l'approccio proposto presenta anche una buona scalabilità, che è efficiente nel riconoscere l'appliance appena aggiunta.

La disaggregazione energetica o il monitoraggio del carico non intrusivo (NILM), un problema di discriminazione della sorgente cieca a ingresso singolo, mira a interpretare il consumo di elettricità dell'utente di rete nella misurazione del livello dell'apparecchio. In [12] viene presentato un nuovo approccio per la disaggregazione di potenza utilizzando **una Long Short-Term Memory (LSTM) combinata con Convolutional Neural Network (CNN)**. Le reti neurali profonde hanno dimostrato di essere un'applicazione significativa per questi tipi di problemi a causa della loro complessità e dell'enorme numero di parametri addestrabili. Il metodo ibrido proposto nell'articolo potrebbe aumentare significativamente l'accuratezza complessiva di NILM perché beneficia di entrambi i vantaggi della rete. Il metodo proposto utilizzava l'apprendimento da sequenza a sequenza, in cui l'input è una finestra della rete e l'output è una finestra dell'appliance di destinazione.

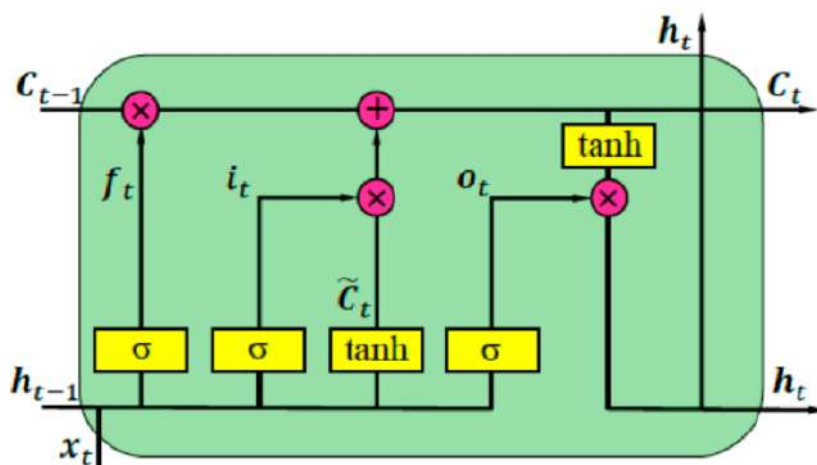


Figura 20: Struttura interna della cella LSTM

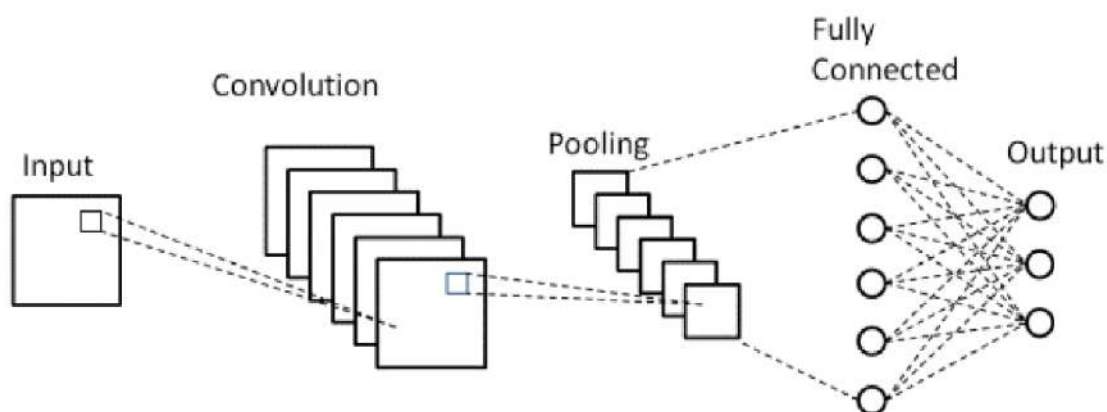


Figura 21: Architettura CNN

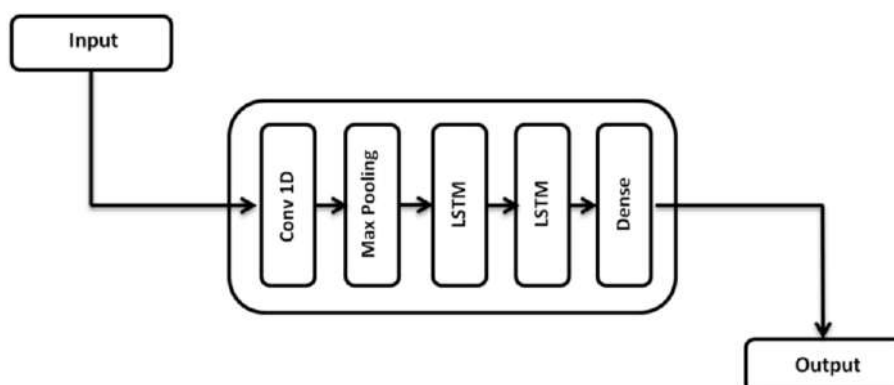


Figura 22: Architettura della rete neurale profonda proposta

L'approccio di rete neurale profonda proposto è stato applicato al set di dati sull'energia domestica reale "REFIT". Il set di dati sulle misurazioni del carico elettrico REFIT descritto in questo articolo include i carichi aggregati di tutta la casa e nove misurazioni di singoli elettrodomestici a intervalli di 8 secondi per casa, raccolti continuamente per un periodo di due anni da 20 case nel Regno Unito.

Il metodo proposto ha raggiunto prestazioni significative, migliorando l'accuratezza e le misure del punteggio F1 rispettivamente del 95,93% e dell'80,93%, il che dimostra l'efficacia e la superiorità dell'approccio proposto per il monitoraggio dell'energia domestica.

I principali vantaggi del metodo proposto sono stati pertanto la sua accuratezza e la capacità di disaggregare le caratteristiche di ciascun apparecchio.

Infine, il confronto tra il metodo proposto e altri metodi pubblicati di recente è stato presentato e discusso in base all'accuratezza, al numero di dispositivi considerati e alle dimensioni dei parametri addestrabili della rete neurale profonda. Il metodo proposto mostra prestazioni notevoli rispetto ad altri metodi precedenti.

Il monitoraggio del carico non intrusivo fornisce agli utenti informazioni dettagliate sul consumo di elettricità dei loro apparecchi e offre ai fornitori di energia una migliore visione dell'utilizzo dei loro clienti. Può anche essere utilizzato per migliorare l'assistenza agli anziani, i servizi legali e l'ottimizzazione del consumo energetico. In [13] vengono esaminati i compromessi di progettazione nel processo di sviluppo di un modello di classificazione basato sull'**apprendimento automatico (approcci Support Vector Machine (SVM), Multilayer Perceptron (MLP) e Random Forest (RF))** dal punto di vista dell'ingegneria delle funzionalità, della selezione del modello e dell'ottimizzazione. Nell'articolo viene mostrato che funzionalità ben progettate hanno un impatto maggiore sulle prestazioni del modello rispetto alla selezione della tecnica di apprendimento automatico.

Feature set	Precision	Recall	f1
raw	0.486	0.460	0.421
min, max, median	0.746	0.664	0.661
mean, std	0.714	0.646	0.642
all	0.851	0.837	0.835
best	0.851	0.835	0.834

Figura 23: Confronto dei set di funzionalità utilizzando SVM

class	precision	recall	f1
computer monitor	0,882	0,720	0,780
laptop computer	0,897	0,800	0,838
television	0,958	0,927	0,941
washer dryer	0,952	0,700	0,804
microwave	0,679	0,710	0,687
boiler	0,945	0,937	0,940
toaster	0,705	0,940	0,806
kettle	0,684	0,806	0,739
fridge	0,961	0,980	0,970

Figura 24: Prestazioni per classe, SVM e miglior set di funzionalità

Algorithm	Precision	Recall	f1
SVM default	0.704	0.672	0.647
SVM optimised	0.851	0.835	0.834
KNN default	0.784	0.776	0.772
KNN optimised	0.787	0.776	0.776
MLP default	0.786	0.769	0.766
MLP optimised	0.849	0.830	0.828
LR default	0.742	0.691	0.674
LR optimised	0.788	0.770	0.768
RF default	0.822	0.815	0.813
RF optimised	0.822	0.815	0.813

Figura 25: Impatto dell'ottimizzazione dell'algoritmo

Nell'articolo sono stati mostrati i compromessi nello sviluppo di un accurato sistema di classificazione dell'appliance e viene proposto un nuovo set di funzionalità per migliorare le prestazioni del classificatore. In primo luogo, viene dimostrato che utilizzando quantità statistiche e personalizzate in grado di catturare la forma delle serie temporali grezze è possibile migliorare le prestazioni del punteggio F1 fino a 42 punti percentuali rispetto alla linea di base delle serie temporali grezze. Gli esperimenti condotti hanno mostrato che le caratteristiche con le migliori prestazioni erano media, asimmetria, curtosi, picco-picco, varianza, il numero di volte in cui la derivata ha cambiato segno sull'intervallo osservato, somma dei valori assoluti della derivata su ogni intervallo all'interno della serie temporale e deviazione standard. In secondo luogo, viene dimostrato che anche la scelta della tecnica di apprendimento automatico e dei parametri ottimali sono importanti. Il modello con la migliore preforming che utilizza SVM ottimizzato ha ottenuto un punteggio F1 migliore di 0,187 rispetto all'SVM non ottimizzato con le peggiori prestazioni. Tuttavia, tutti i modelli SVM, MLP e random forest ottimizzati hanno funzionato bene con punteggi F1 compresi tra 0,81 e 0,84 per il problema considerato.

Riepilogando, in base ai risultati il miglioramento del punteggio F1 tra le caratteristiche non ingegnerizzate e quelle proposte è arrivato al 42%, mentre il miglioramento tra il peggior modello non ottimizzato e quello meglio ottimizzato è stato del 19%.

I metodi riportati in [14] hanno incluso il rilevamento degli eventi, l'estrazione delle caratteristiche per ottenere la corrente differenziale basata sugli eventi, l'aumento dei dati e la classificazione delle applicazioni.

Viene introdotto innanzitutto il background, il suo significato e le recenti ricerche del monitoraggio del carico non invasivo e viene descritto brevemente il set di dati BLUED e la **Back Propagation Neural Network (BP-NN)**. Nello specifico, viene utilizzata la rete BP con un singolo livello nascosto per raggiungere il compito di NILM e ottenere ottimi risultati.

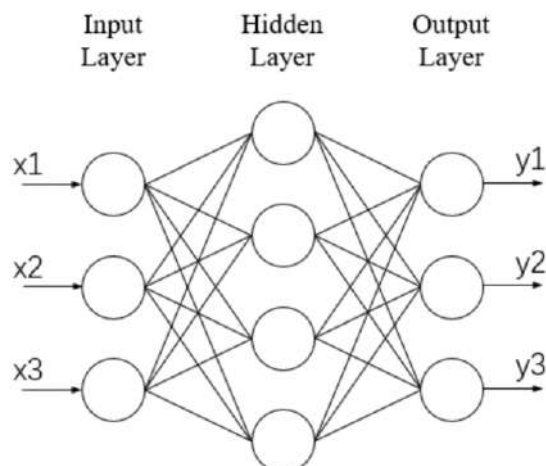


Figura 26: Modello semplificato della rete neurale BP.

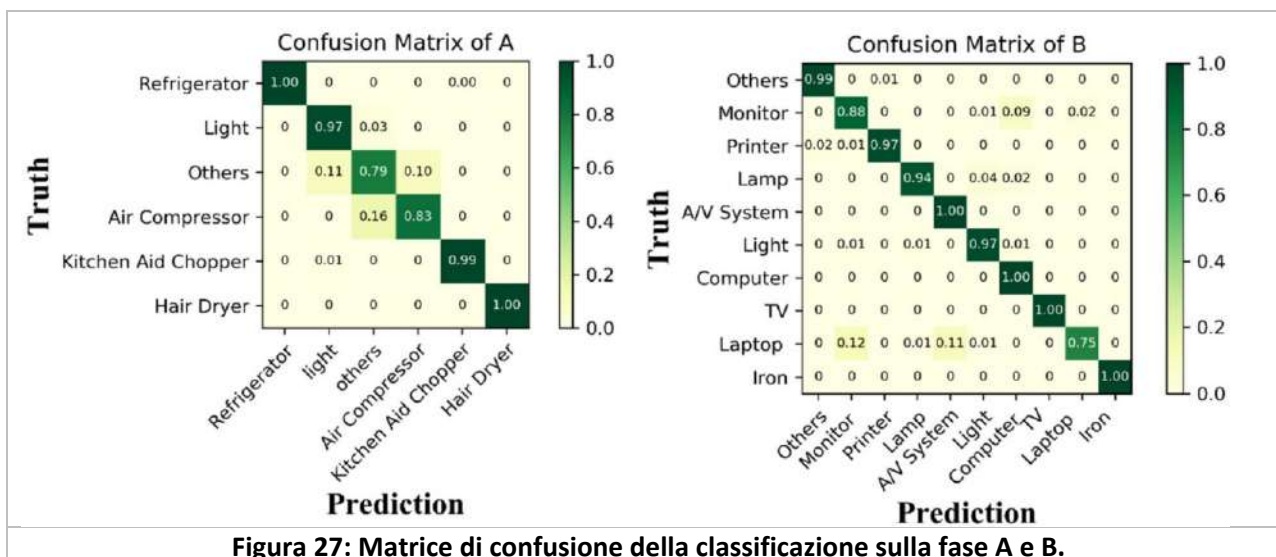


Figura 27: Matrice di confusione della classificazione sulla fase A e B.

In secondo luogo, l'attività di rilevamento degli eventi viene completata utilizzando la rete neurale BP con potenza attiva campionata a 60 Hz, ed è stato ottenuto il 99% su accuratezza, precisione, richiamo e punteggio F1 su tutte e quattro le metriche. Per risolvere il problema della classificazione, vengono utilizzate le caratteristiche della corrente differenziale basata sugli eventi ed è stato utilizzato l'allineamento di fase per ridurre l'influenza causata dai disturbi della rete elettrica.

La quantità di dati in alcuni dispositivi è insufficiente, quindi si utilizza il metodo di aumento dei dati per superare questo problema.

Per questo esperimento viene utilizzato il set di dati BLUED e il metodo proposto ha raggiunto una buona prestazione nel set di dati, con una precisione superiore al 95% sulla fase A e B per la classificazione dell'applicazione.

All'interno dell'articolo in [15] viene proposto un nuovo approccio seq2seq che utilizza **Graph Neural Networks (GNN)** come encoder e un decoder basato su Transformer. Nello specifico, l'approccio proposto utilizza **Graph Convolutional Networks (GCN)** per codificare il segnale aggregato in nodi grafici distinti, basati sulla metrica dell'entropia del segnale dell'appliance corrispondente che rappresenta i suoi diversi stati operativi. Utilizzando questo approccio vengono create codifiche che incorporano dipendenze invarianti nel tempo all'interno del segnale aggregato, formando una rappresentazione più olistica e significativa dei dati.

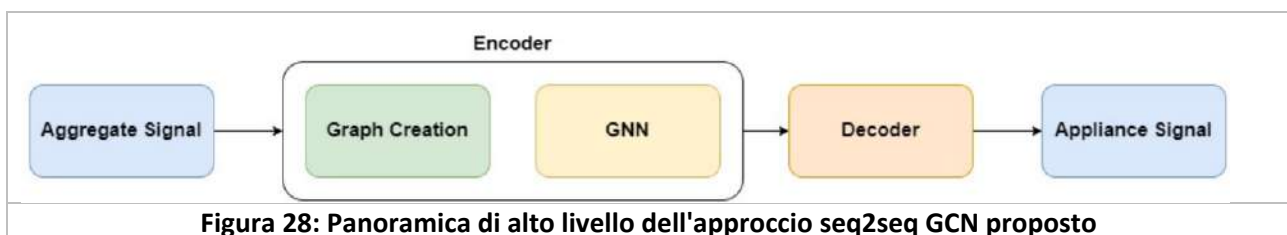
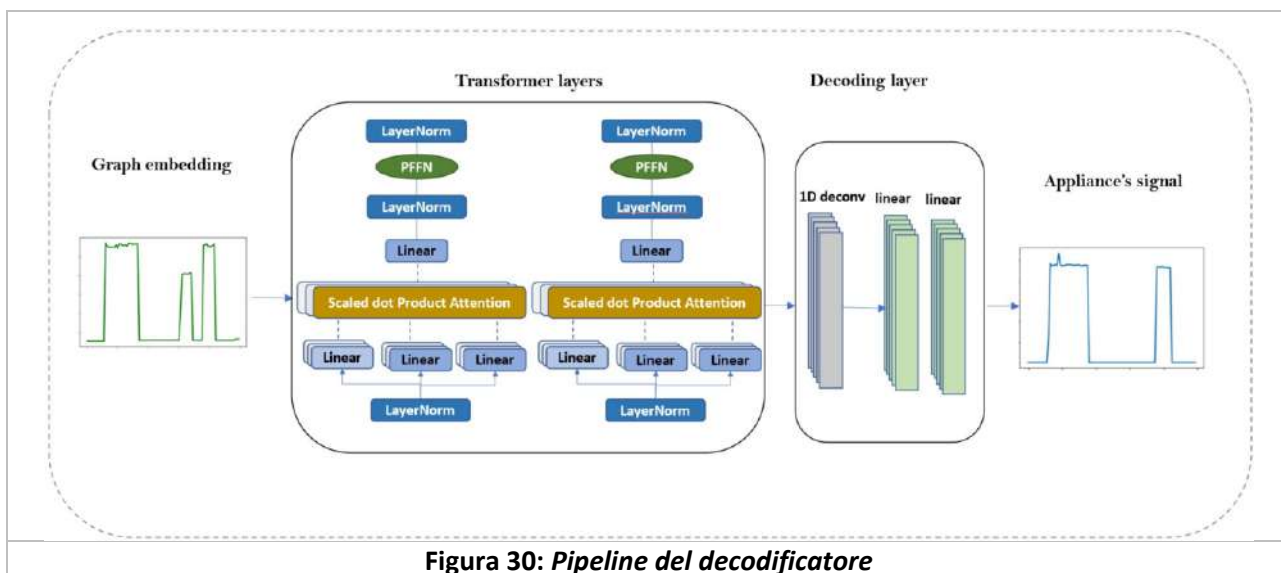
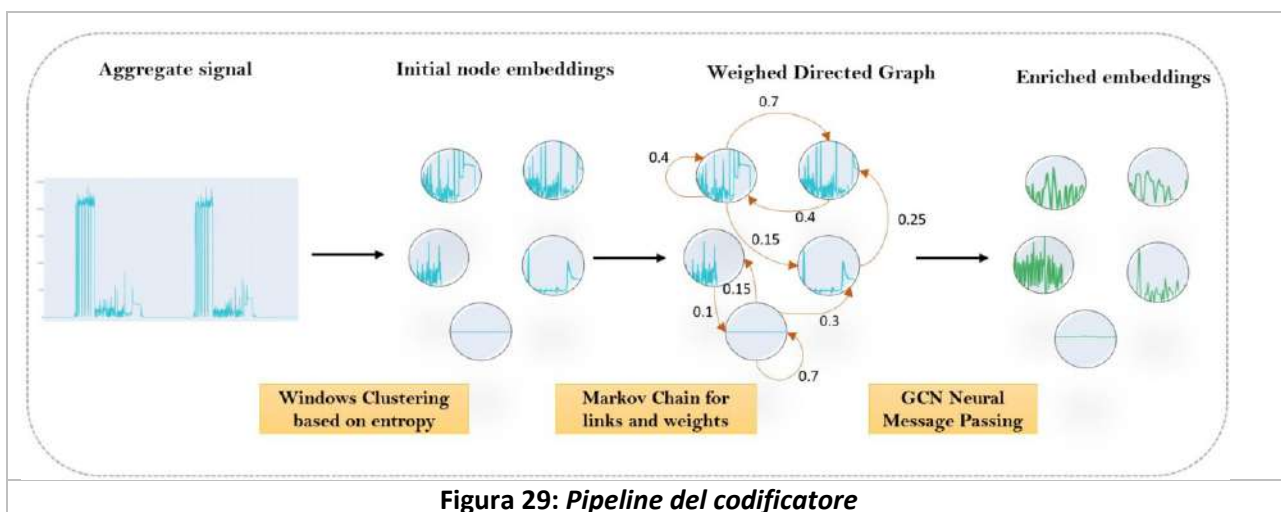


Figura 28: Panoramica di alto livello dell'approccio seq2seq GCN proposto



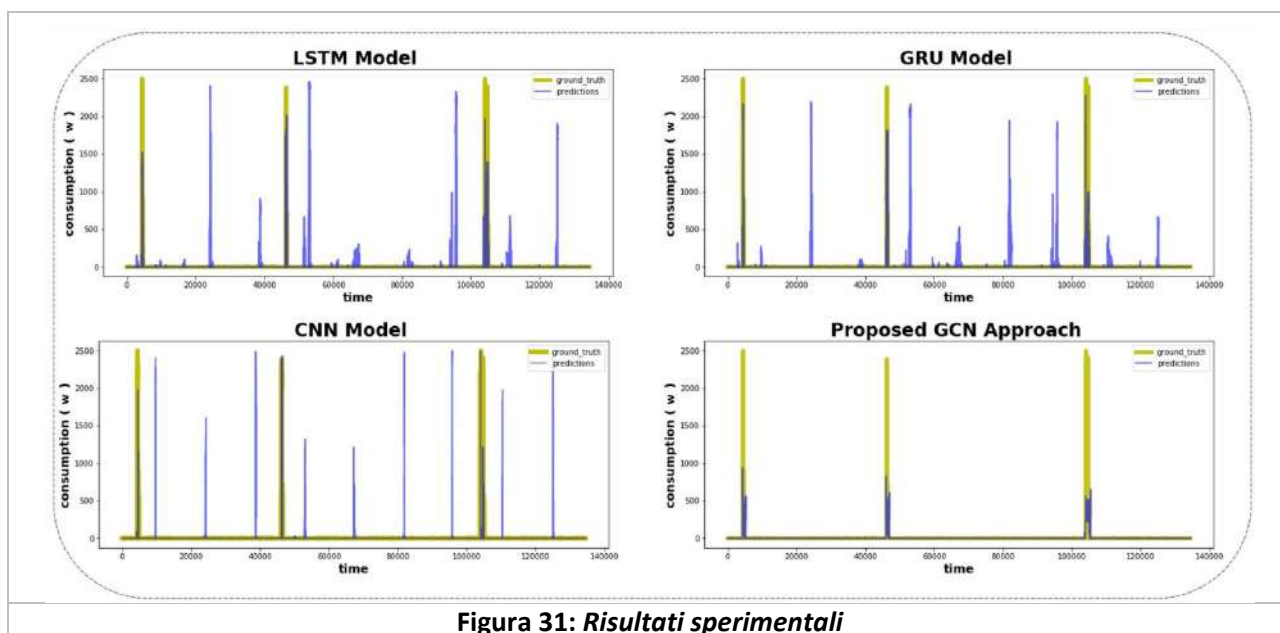


Figura 31: Risultati sperimentali

I risultati sperimentali ottenuti indicano che l'approccio proposto ha prodotto risultati soddisfacenti o addirittura migliori su apparecchi multi-stato come la lavatrice, rispetto ai metodi all'avanguardia. Sebbene l'approccio funzioni bene solo su dispositivi multi-operativi, dà una dimensione interessante al problema e propone, secondo gli autori, il primo approccio basato su grafi verso NILM, aprendo la strada all'applicazione di più approcci GNN su questo dominio.

Riepilogando, la metodologia GCN proposta rappresenta il primo approccio che utilizza le reti neurali a grafo nella disaggregazione energetica di dispositivi multi-operativi laddove, secondo gli autori, sembra superare i modelli sequenziali allo stato dell'arte senza richiedere alcun bilanciamento dei dati.

In [16] viene proposta un'architettura **ENSM (Machine Learning Network) supervisionata non parametrica** efficiente con dimensioni ridotte e un tempo di inferenza rapido senza sacrificare le prestazioni. L'architettura proposta può massimizzare le prestazioni di disaggregazione energetica e prevedere nuove osservazioni basate su quelle passate.

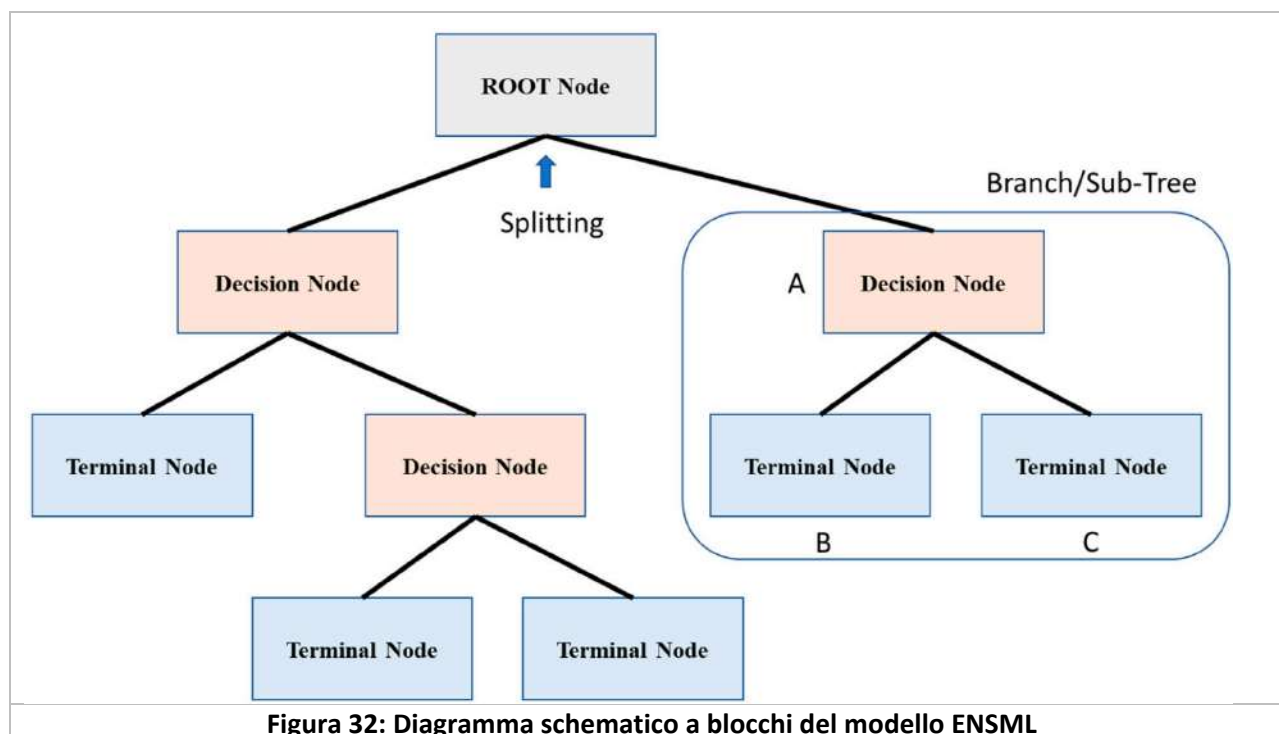
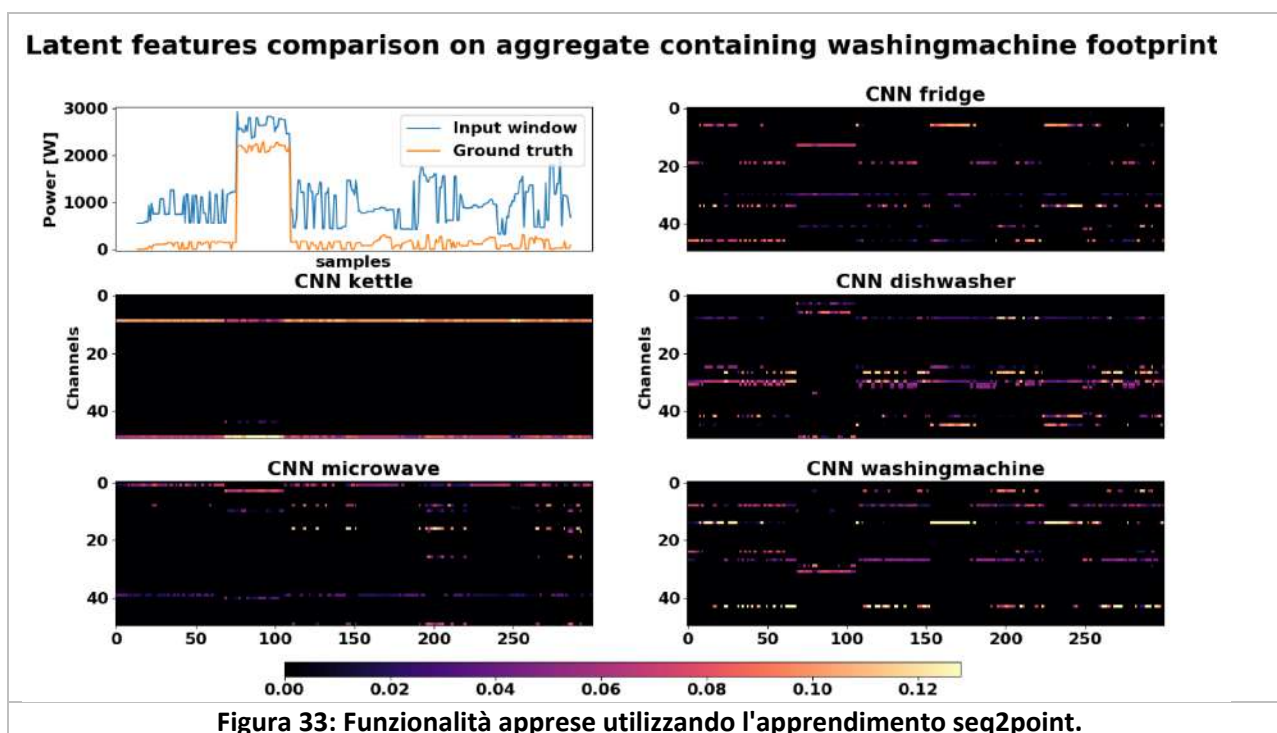


Figura 32: Diagramma schematico a blocchi del modello ENSML

Riepilogando, utilizzando i dati dei contatori intelligenti domestici, questa ricerca ha proposto un sistema di riconoscimento e previsione dell'apparecchio che utilizza la tecnica NILM basata su una rete ENSML semplice e di bassa complessità, in grado di identificare le apparecchiature elettriche domestiche comuni da una tipica lettura del contatore intelligente domestico. La ricerca si è concentrata sulla previsione di diversi apparecchi che hanno quantità e modelli diversi di consumo di energia. I risultati hanno mostrato che l'utilizzo del modello ENSML ha notevolmente aumentato l'accuratezza della previsione dell'energia nel 99% dei casi.

Per mitigare il problema della non identificabilità, sono stati proposti vari metodi che incorporano la conoscenza del dominio nel NILM che si sono dimostrati efficaci sperimentalmente. Di recente, tra questi metodi, le **deep neural networks** si sono dimostrate le migliori. Probabilmente, l'apprendimento da sequenza a punto (seq2point) recentemente proposto è promettente per NILM. Tuttavia, i risultati sono stati eseguiti solo sullo stesso dominio di dati. Non è chiaro se il metodo possa essere generalizzato o trasferito a domini diversi, ad esempio, i dati di test sono stati ricavati da un paese diverso rispetto ai dati di formazione. In [17] viene affrontato questo problema e

vengono proposti due schemi di **transfer learning**, vale a dire **l'appliance transfer learning (ATL)** e **il cross-domain transfer learning (CTL)**. Per ATL, i risultati ottenuti nell'articolo mostrano che le caratteristiche latenti apprese da un apparecchio "complesso", ad esempio una lavatrice, possono essere trasferite a un apparecchio "semplice", ad esempio un bollitore. Per CTL, la conclusione proposta è che l'apprendimento seq2point è trasferibile. Precisamente, quando i dati di addestramento e di test si trovano in un dominio simile, l'apprendimento seq2point può essere applicato direttamente ai dati di test senza messa a punto; quando i dati di addestramento e di test si trovano in domini diversi, l'apprendimento seq2point necessita di una messa a punto prima di essere applicato ai dati di test. Inoltre, solo i livelli completamente connessi necessitano di una messa a punto per l'apprendimento del trasferimento.



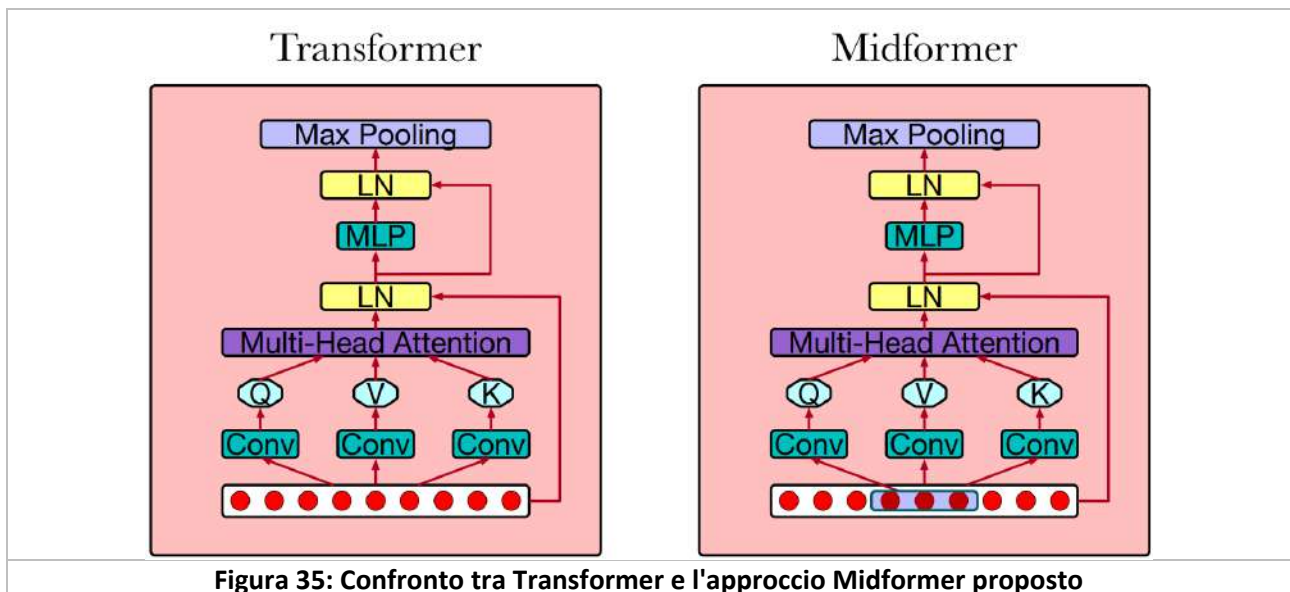
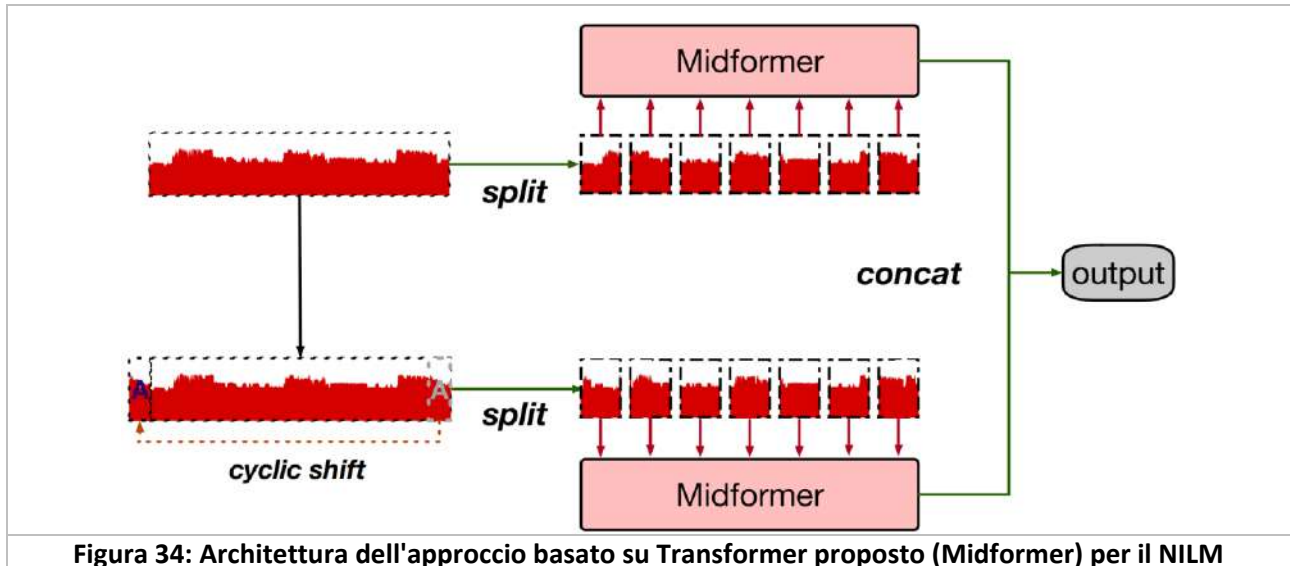
Negli studi precedenti, l'apprendimento e il test seq2point erano considerati solo nello stesso dominio. Al contrario, nell'articolo viene considerato di far apprendere il modello in un dominio e testarlo in un dominio diverso. L'ipotesi degli autori è stata che le funzionalità dell'appliance estratte utilizzando le reti neurali convoluzionali fossero invarianti tra le appliance (e anche tra i domini di

dati) e gli esperimenti condotti sul trasferimento dell'apprendimento hanno infine supportato l'ipotesi. Due diversi approcci di transfer learning sono stati studiati sperimentalmente, che sono **l'appliance transfer learning** e il **cross-domain transfer learning**. Nell'ottica degli autori, mediante il trasferimento del deep learning, potrebbe essere possibile trovare un modello unico per l'appliance residenziale in tutto il mondo che possa essere poi adattato utilizzando dati di addestramento ridotti, ottenendo così l'oracle NILM, notevoli risparmi computazionali e una riduzione dipendente da specifici dati etichettati dall'appliance. I modelli Seq2point sono stati addestrati su REFIT e poi testati su REDD e UK-DALE.

Le conclusioni sono state le seguenti: in primo luogo, per quanto riguarda l'ATL, gli strati CNN addestrati sulle lavatrici possono essere applicati ad altri apparecchi. Fondamentalmente, questo indica che tutte le appliance utilizzano firme simili per scopi NILM. Pertanto, si possono utilizzare dati di grandi dimensioni per addestrare i livelli CNN e quindi applicare i livelli CNN addestrati a qualsiasi altro dato invisibile per ottenere le firme. I livelli completamente connessi possono quindi essere addestrati esclusivamente sui dati invisibili. In secondo luogo, per quanto riguarda CTL, è stato scoperto che se i domini dei dati di addestramento e test sono simili, il modello addestrato potrebbe non richiedere una messa a punto; se i domini sono diversi, la messa a punto aiuta a migliorare le prestazioni del modello addestrato. L'intuizione è che se i domini sono simili, la regolazione sottile può portare a un adattamento eccessivo poiché solo un piccolo sottoinsieme è stato utilizzato per questa. D'altra parte, se i domini sono diversi, la messa a punto aiuta il modello ad adattarsi al dominio invisibile.

In [18] viene proposto un modello di **Middle Window Transformer**, chiamato Midformer, per NILM. I modelli esistenti sono limitati dall'elevata complessità computazionale, dalla dipendenza dai dati e dalla scarsa trasferibilità. In Midformer, per prima cosa viene sfruttato l'incorporamento patch per accorciare la lunghezza dell'input, quindi viene ridotta la dimensione delle query nell'attention level utilizzando solo la global attention su alcune posizioni di input selezionate al centro della finestra, al fine di acquisire il contesto globale. La tecnica della finestra spostata ciclicamente è stata utilizzata per preservare la connessione tra le patch. Viene seguito anche il paradigma di pre-addestramento

e messa a punto per alleviare la dipendenza dai dati, ridurre il calcolo nell'addestramento alla modellazione e migliorare la trasferibilità del modello a compiti e domini sconosciuti.



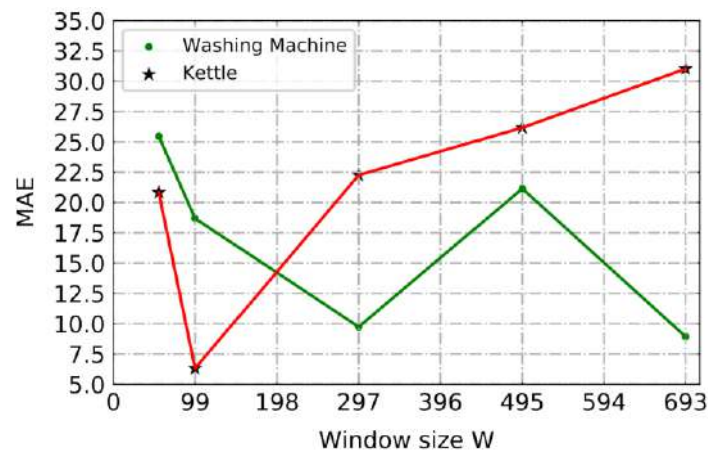


Figura 36: MAE ottenuti da modelli Midformer con differenti dimensioni delle finestre testando il bollitore nella casa 9 e la lavatrice nella casa 15. Il numero di patch è impostato su 11

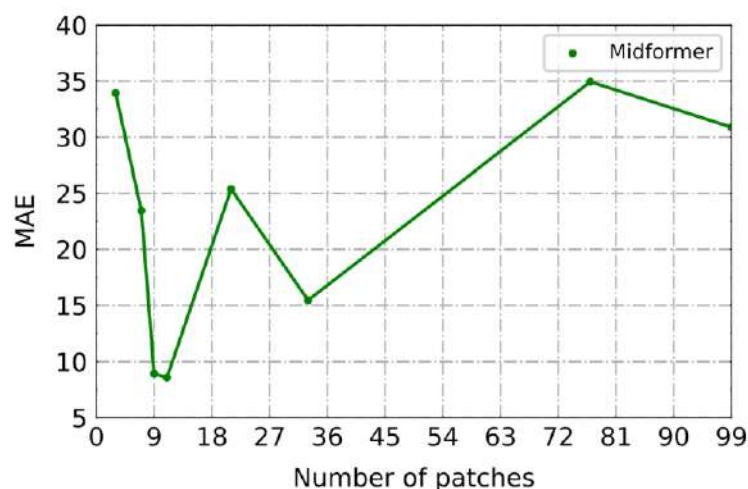


Figura 37: MAE raggiunti dai modelli Midformer con differenti numeri di patch per la lavatrice, quando la dimensione della finestra è impostata su 693.

Riepilogando, nell'articolo è stato proposto il modello Midformer per affrontare il problema NILM. È stata utilizzata l'attenzione sul patch-wise e ridotta la dimensione della query per ridurre la complessità temporale quadratica nei modelli Transformer tradizionali alla complessità lineare. Inoltre, si è posto l'accento sulle prestazioni di trasferibilità dei modelli, che hanno contribuito a ridurre i costi di formazione del modello e facilitato l'implementazione del modello in vari ambienti. Lo studio sperimentale proposto utilizza due set di dati del mondo reale e ha mostrato le prestazioni superiori e la maggiore trasferibilità del modello Midformer proposto su tre modelli di base all'avanguardia per affrontare il problema NILM.

I modelli di deep learning per il monitoraggio del carico non intrusivo (NILM) tendono a richiedere una grande quantità di dati etichettati per l'addestramento. Tuttavia, è difficile generalizzare i modelli addestrati a siti invisibili a causa delle diverse caratteristiche di carico e modelli operativi delle apparecchiature tra i set di dati. Per affrontare tali problemi, in [19] viene proposto il metodo **Self-Supervised Learning (SSL)**, in cui non sono richiesti dati a livello di apparecchio etichettati dal set di dati di destinazione o dalla casa. Inizialmente, solo le letture della potenza aggregata dal set di dati target sono necessarie per pre-addestrare una rete generale tramite un'attività auto-supervisionata per mappare le sequenze di potenza aggregata ai rappresentanti derivati. Quindi, vengono eseguite attività a valle supervisionate per ogni categoria di appliance per mettere a punto la rete pre-addestrata, dove vengono trasferite le funzionalità apprese nell'attività. L'utilizzo di set di dati di origine etichettati consente alle attività a valle di apprendere come ogni carico viene disaggregato, mappando l'aggregato alle etichette. Infine, la rete ottimizzata viene applicata alla disaggregazione del carico per i siti di destinazione.

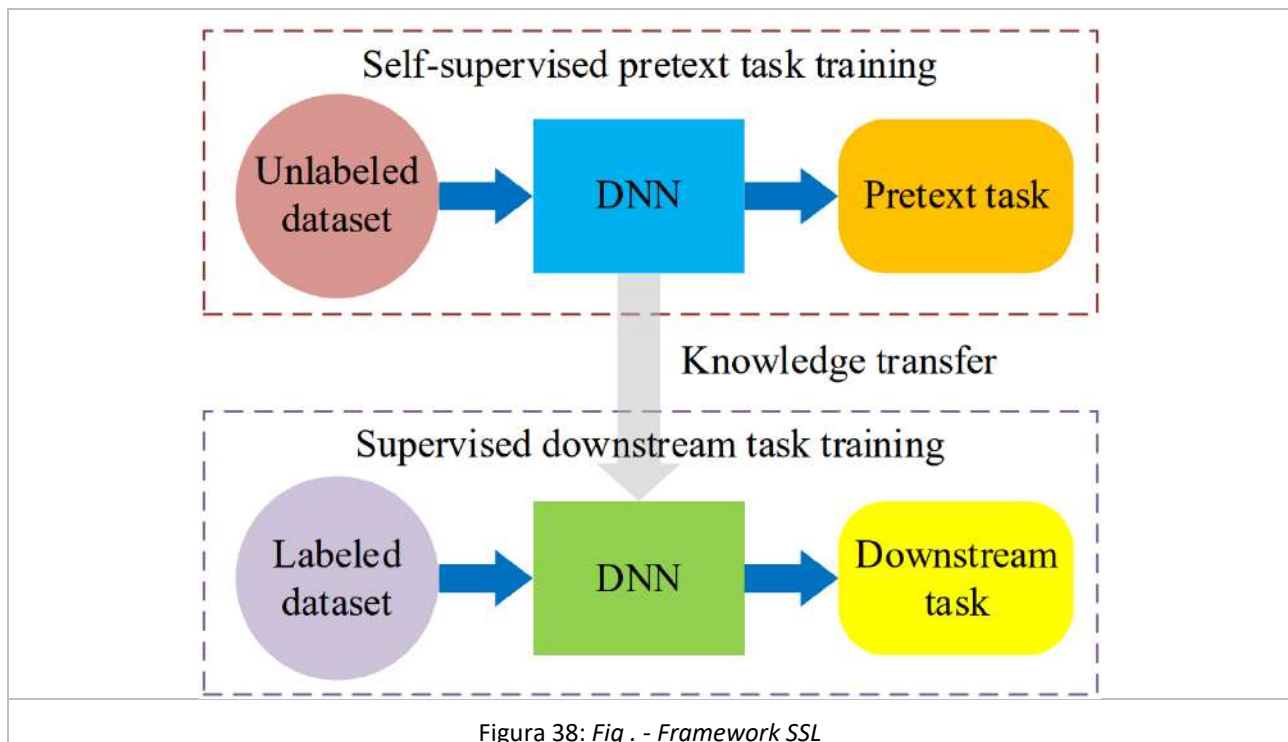


Figura 38: Fig . - Framework SSL

L'SSL viene applicato alle soluzioni NILM basate su S2p e Bi-GRU all'avanguardia. Viene proposta un'attività pretesto per pre-addestrare una rete generale per mappare sequenze di potenza aggregate a rappresentanti derivati, dove non sono richiesti dati etichettati per il set di dati di destinazione. Eseguendo un'attività a valle supervisionata basata sui dati etichettati dai set di dati di origine, la rete pre-addestrata viene messa a punto con il trasferimento delle funzionalità e quindi applicata per disaggregare i carichi per i siti di destinazione.

Per la convalida vengono utilizzati set di dati REDD, UK-DALE e REFIT accessibili al pubblico, con più casi sperimentali progettati per valutare le prestazioni di SSL utilizzando dati all'interno dello stesso set e tra set di dati. Inoltre, negli esperimenti vengono utilizzate reti neurali all'avanguardia per eseguire attività NILM.

Per una dimostrazione completa delle prestazioni NILM, vengono utilizzate varie metriche che mostrano che SSL (in particolare il **Federated Semi-Supervised Learning o FSSL**) generalmente supera lo schema ZSL.

Case No.	App.	S2p-ZSL			S2p-FSSL (S2p-PSSL)			Bi-GRU-ZSL			Bi-GRU-FSSL (Bi-GRU-PSSL)		
		MAE (W)	SAE	EpD (kWh)	MAE (W)	SAE	EpD (kWh)	MAE (W)	SAE	EpD (kWh)	MAE (W)	SAE	EpD (kWh)
1	F	28.10	0.06	0.12	26.69 (32.55)	0.12 (0.39)	0.16 (0.33)	32.10	0.04	0.24	31.54 (32.29)	0.01 (0.09)	0.23 (0.28)
	M	11.94	1.66	0.16	11.27 (8.79)	1.51 (0.66)	0.15 (0.08)	8.49	0.61	0.09	11.02 (8.88)	1.30 (0.70)	0.13 (0.09)
	DW	54.22	0.39	0.63	39.57 (43.64)	0.13 (0.03)	0.38 (0.34)	87.83	0.59	1.22	95.15 (105.56)	0.67 (0.70)	1.19 (1.19)
	WM	35.06	0.13	0.40	45.64 (39.55)	0.83 (0.41)	0.52 (0.43)	47.48	0.80	0.63	35.86 (41.69)	0.16 (0.45)	0.41 (0.46)
	K	25.87	0.59	0.30	18.55 (18.00)	0.26 (0.16)	0.16 (0.15)	18.30	0.23	0.13	18.36 (24.04)	0.20 (0.07)	0.13 (0.12)
	mean	31.04	0.57	0.32	28.34 (28.51)	0.57 (0.33)	0.27 (0.27)	38.84	0.45	0.46	38.39 (42.49)	0.47 (0.40)	0.42 (0.43)
2	F	28.10	0.06	0.12	27.62 (28.73)	0.01 (0.13)	0.14 (0.19)	32.10	0.04	0.24	31.19 (32.33)	0.01 (0.14)	0.24 (0.27)
	M	11.94	1.66	0.16	7.25 (9.46)	0.25 (0.84)	0.07 (0.09)	8.49	0.61	0.09	5.57 (10.91)	0.24 (1.24)	0.05 (0.13)
	DW	54.22	0.39	0.63	45.48 (49.45)	0.14 (0.26)	0.39 (0.46)	87.83	0.59	1.22	89.70 (104.85)	0.60 (0.72)	1.21 (1.21)
	WM	35.06	0.13	0.40	34.00 (33.92)	0.18 (0.09)	0.42 (0.38)	47.48	0.80	0.63	34.89 (43.23)	0.08 (0.54)	0.43 (0.49)
	K	25.87	0.59	0.30	19.08 (20.23)	0.22 (0.19)	0.14 (0.15)	18.30	0.23	0.13	19.87 (28.75)	0.20 (0.03)	0.13 (0.14)
	mean	31.04	0.57	0.32	26.69 (28.36)	0.16 (0.30)	0.23 (0.25)	38.84	0.45	0.46	36.24 (44.01)	0.23 (0.53)	0.41 (0.45)
3	F	41.74	0.15	0.21	43.62 (36.70)	0.23 (0.09)	0.29 (0.16)	46.30	0.48	0.57	45.59 (50.51)	0.48 (0.54)	0.57 (0.64)
	M	18.65	0.26	0.13	18.78 (17.79)	0.26 (0.55)	0.12 (0.24)	16.88	0.58	0.24	17.32 (18.86)	0.65 (0.63)	0.26 (0.26)
	DW	57.08	0.18	0.20	51.13 (61.05)	0.02 (0.25)	0.18 (0.25)	40.04	0.51	0.24	42.62 (66.92)	0.41 (0.54)	0.27 (0.48)
	WM	30.49	0.76	0.65	30.11 (37.00)	0.63 (0.55)	0.61 (0.73)	38.19	0.45	0.68	32.58 (40.25)	0.60 (0.39)	0.62 (0.70)
	K	36.99	0.34	0.30	35.91 (38.14)	0.29 (0.36)	0.30 (0.35)	35.35	0.51	0.43	34.53 (44.14)	0.54 (0.53)	0.43 (0.52)
	mean	36.99	0.34	0.30	35.91 (38.14)	0.29 (0.36)	0.30 (0.35)	35.35	0.51	0.43	34.53 (44.14)	0.54 (0.53)	0.43 (0.52)
4	F	41.74	0.15	0.21	38.74 (35.46)	0.32 (0.13)	0.39 (0.21)	46.30	0.48	0.57	47.32 (50.91)	0.54 (0.57)	0.64 (0.68)
	M	18.65	0.26	0.13	18.51 (18.10)	0.42 (0.47)	0.18 (0.20)	16.88	0.58	0.24	17.16 (19.27)	0.45 (0.63)	0.19 (0.26)
	DW	57.08	0.18	0.20	40.82 (56.97)	0.43 (0.07)	0.17 (0.14)	40.04	0.51	0.24	40.11 (64.71)	0.52 (0.43)	0.26 (0.42)
	WM	30.49	0.76	0.65	28.14 (36.18)	0.71 (0.62)	0.60 (0.75)	38.19	0.45	0.68	33.16 (37.20)	0.46 (0.34)	0.61 (0.65)
	K	36.99	0.34	0.30	31.55 (36.68)	0.47 (0.32)	0.34 (0.33)	35.35	0.51	0.43	34.44 (43.02)	0.49 (0.49)	0.43 (0.50)
	mean	36.99	0.34	0.30	31.55 (36.68)	0.47 (0.32)	0.34 (0.33)	35.35	0.51	0.43	34.44 (43.02)	0.49 (0.49)	0.43 (0.50)
5	F	19.96	0.23	0.25	21.32 (18.40)	0.26 (0.15)	0.29 (0.18)	25.45	0.24	0.27	25.97 (28.21)	0.27 (0.25)	0.29 (0.28)
	M	12.78	0.68	0.11	11.34 (11.43)	0.54 (0.14)	0.09 (0.06)	8.86	0.17	0.06	7.71 (8.88)	0.40 (0.38)	0.08 (0.08)
	DW	36.11	0.34	0.35	18.54 (26.66)	0.21 (0.29)	0.25 (0.33)	14.65	0.15	0.19	18.48 (28.38)	0.02 (0.09)	0.14 (0.23)
	WM	17.21	0.15	0.22	15.84 (17.03)	0.11 (0.20)	0.24 (0.26)	17.60	0.26	0.29	15.75 (15.68)	0.04 (0.05)	0.25 (0.27)
	K	17.57	0.08	0.10	13.37 (13.84)	0.13 (0.25)	0.10 (0.18)	17.67	0.50	0.36	18.42 (20.52)	0.53 (0.54)	0.37 (0.38)
	mean	20.73	0.30	0.21	16.08 (17.47)	0.25 (0.21)	0.19 (0.20)	16.85	0.26	0.23	17.27 (20.33)	0.25 (0.26)	0.23 (0.25)
6	F	19.96	0.23	0.25	21.19 (18.33)	0.25 (0.11)	0.27 (0.15)	25.45	0.24	0.27	24.31 (27.29)	0.18 (0.22)	0.20 (0.24)
	M	12.78	0.68	0.11	6.93 (11.15)	0.28 (0.10)	0.06 (0.06)	8.86	0.17	0.06	7.16 (8.69)	0.50 (0.41)	0.09 (0.08)
	DW	36.11	0.34	0.35	25.11 (36.63)	0.04 (0.02)	0.17 (0.24)	14.65	0.15	0.19	21.02 (24.19)	0.01 (0.17)	0.15 (0.25)
	WM	17.21	0.15	0.22	15.49 (32.73)	0.16 (1.59)	0.26 (0.49)	17.60	0.26	0.29	18.17 (15.90)	0.30 (0.08)	0.29 (0.27)
	K	17.57	0.08	0.10	13.33 (17.86)	0.12 (0.19)	0.10 (0.15)	17.67	0.50	0.36	18.14 (25.92)	0.47 (0.43)	0.33 (0.31)
	mean	20.73	0.30	0.21	16.41 (23.34)	0.17 (0.40)	0.17 (0.22)	16.85	0.26	0.23	17.76 (20.40)	0.29 (0.26)	0.21 (0.23)

Figura 39: risultati sperimentali

In particolare, le funzionalità apprese tramite SSL aiutano a migliorare la sottovalutazione, ridurre l'identificazione errata e mitigare l'influenza di etichette errate. Tali conclusioni sono supportate anche dai segnali di potenza disaggregati per ciascun apparecchio e dai risultati della stima del consumo totale di energia. Pertanto, SSL contribuisce al perfezionamento della rete e al miglioramento delle prestazioni di disaggregazione. Inoltre, FSSL offre prestazioni migliori rispetto a PSSL in termini di precisione di disaggregazione, sebbene richieda tempi di esecuzione più lunghi.

2.2 Modellazione di generatori

In questa sezione si riportano i più recenti esempi di applicazione di algoritmi di intelligenza artificiale per la modellazione di generatori.

In questa sezione si riportano i più recenti esempi di applicazione di algoritmi di intelligenza artificiale per la modellazione digital twin di generatori di energia, sia basati sulla trasformazione di

fonti rinnovabili, che su fonti tradizionali. In particolare, si considereranno le applicazioni per l'ottimizzazione delle prestazioni o il calcolo delle attività manutentiva con approccio predittivo (predictive maintenance).

2.2.1 Fonti rinnovabili

2.2.1.1 Impianti eolici

In [20], utilizzando un DT costruito per la valutazione della fatica di flange ad anello imbullonate su strutture di supporto OWT, vengono esplorate le sfide e le opportunità nella definizione delle incertezze di interesse. Vengono propagate queste incertezze attraverso il DT in modo coerente utilizzando un approccio di modellazione surrogata del **Gaussian Process (GP)**, emulando in modo efficiente un simulatore numerico computazionalmente costoso. Il GP è un interessante metodo di modello surrogato data questa efficienza computazionale, oltre a fornire una stima dell'incertezza di previsione in punti non osservati nello spazio di output. L'uso del modello GP surrogato è incluso in una definizione del Surrogate DT, un quadro che include processi di gemellaggio "veloci" e "lenti". Vengono, inoltre, definiti sei requisiti per applicare i DT alle strutture OWT che forniscono linee guida pratiche per la modellazione di questo complesso asset in condizioni di incertezza.

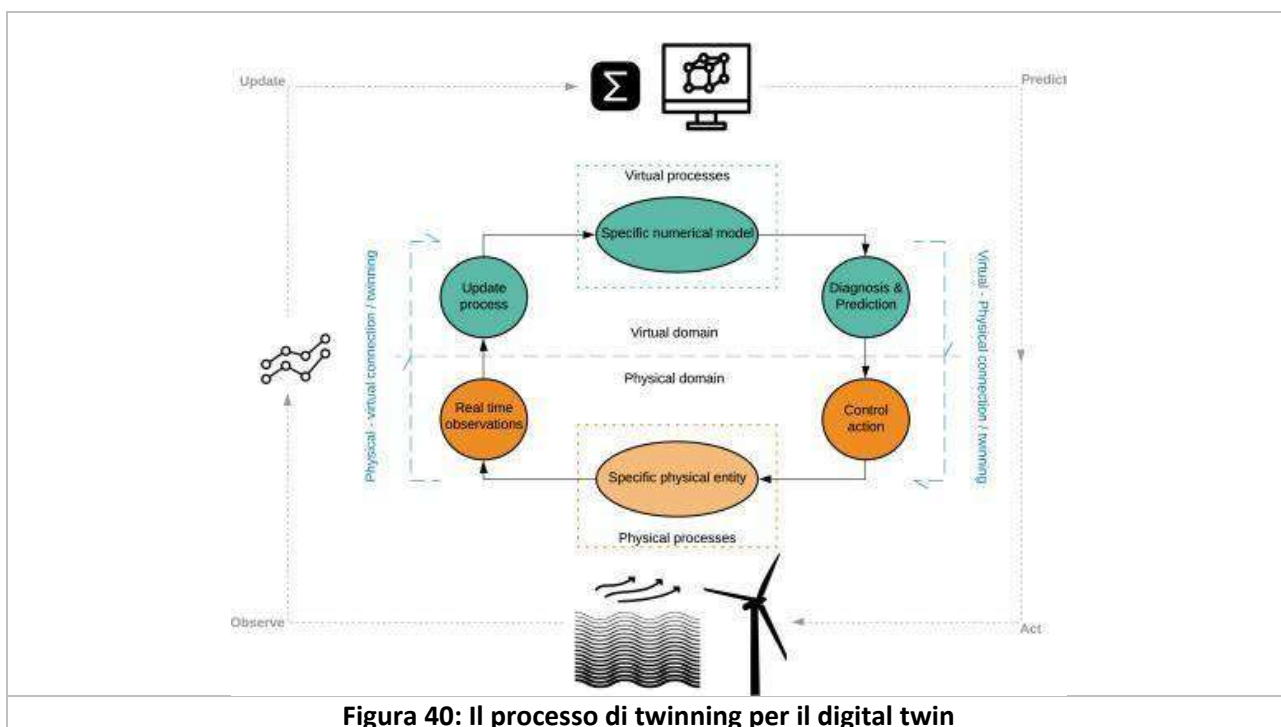


Figura 40: Il processo di twinning per il digital twin

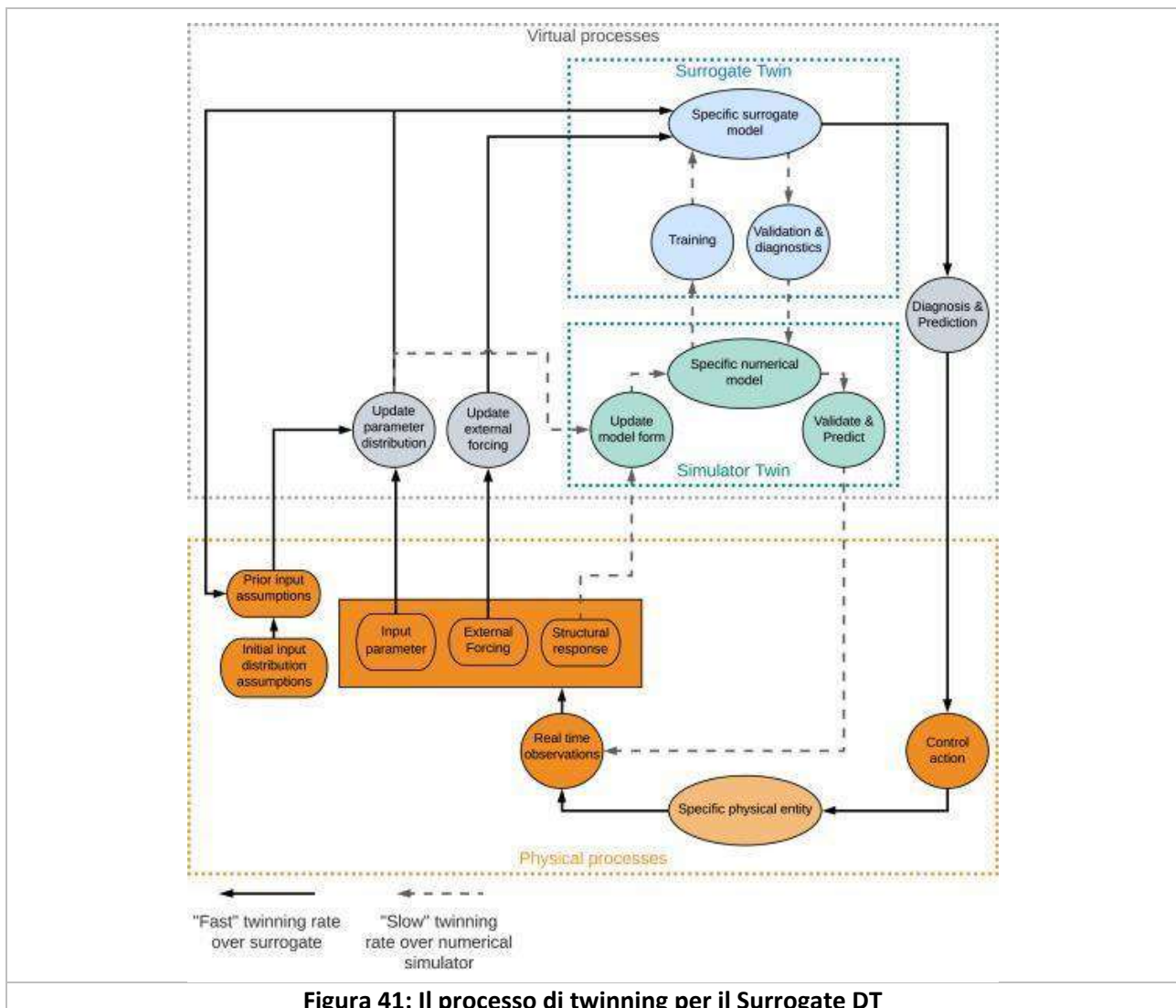
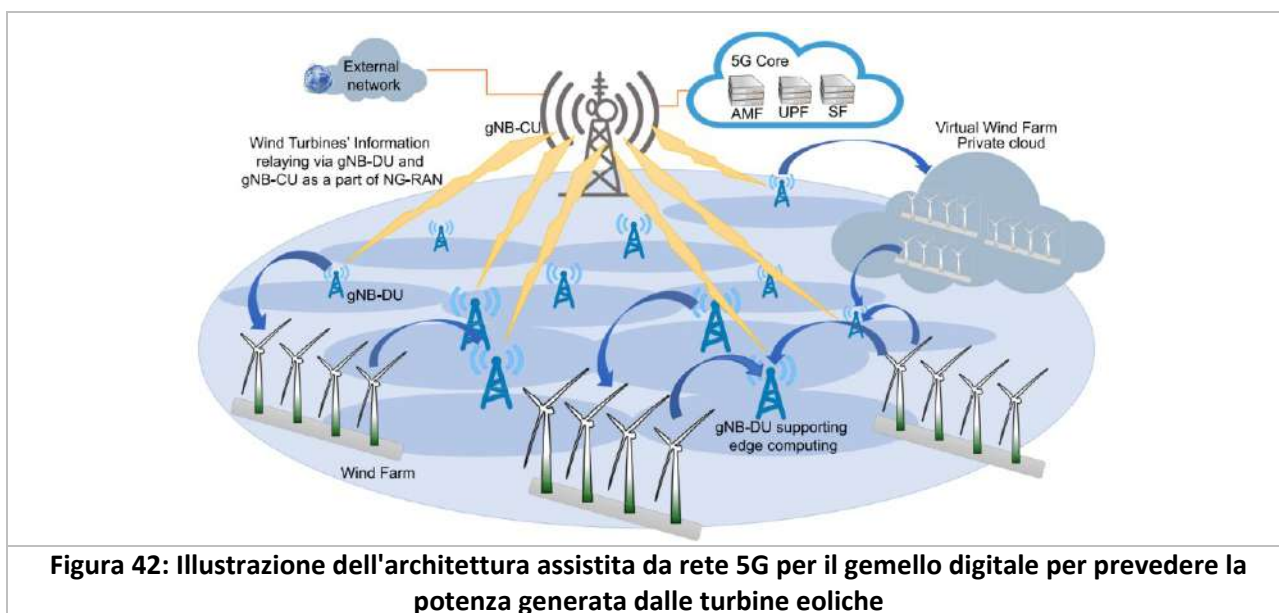


Figura 41: Il processo di twinning per il Surrogate DT

Nell'articolo sono state esplorate le opportunità disponibili quando si assimilano i dati di osservazione con l'incertezza associata alla stima dei parametri inerente alla simulazione. L'introduzione di un esempio motivante, l'analisi della fatica di una flangia ad anello imbullonata su una struttura OWT, ha fornito un banco di prova per il quadro proposto. Vengono discusse le fonti di incertezza e i metodi per quantificare e propagare tale incertezza in modo probabilistico. In particolare, è stata esplorata l'applicazione di modelli surrogati computazionalmente efficienti, ad esempio il modello surrogato GP. L'uso del modello surrogato GP in questo lavoro ha motivato la definizione del Surrogate DT, in cui vengono utilizzati processi di doppio gemellaggio ("veloce" e "lento") per aggiornare rispettivamente il simulatore numerico sottostante e il modello surrogato.

Il processo di "twinning veloce" viene applicato su parametri di input e forzature esterne, per informare le previsioni del modello surrogato GP e, a sua volta, guidare le azioni di controllo sull'entità fisica. Il processo di "slow twinning" viene applicato al simulatore numerico su una scala temporale più lunga, aggiornando il simulatore rispetto alle osservazioni e (in sequenza) riaddestrando e riconvalidando il modello surrogato GP rispetto al simulatore aggiornato.

Il monitoraggio dei parchi eolici e la previsione della generazione di energia è un problema complesso a causa dell'imprevedibilità della velocità del vento. Di conseguenza, limita il potere decisionale del team di gestione per pianificare il consumo energetico in modo efficace. Il modello proposto in [21] risolve questa sfida utilizzando un framework di Digital Twin basato su cloud assistito da 5G-Next Generation-Radio Access Network (5G-NG-RAN) per monitorare virtualmente le turbine eoliche e formare un modello predittivo per prevedere la velocità del vento e la generazione energia. Il modello sviluppato si basa sull'infrastruttura di Digital Twin di Microsoft Azure come piattaforma di Digital Twin a 5 dimensioni. La modellazione predittiva si basa su un **approccio di Deep Dearning, la Convolutional Temporal Network (CNN) (TCN) seguita da una regressione K-nearest Neighbors (KNN) non parametrica**. La modellazione predittiva ha due componenti. In primo luogo, elabora i dati delle serie temporali univariate del vento per prevederne la velocità. In secondo luogo, stima che la produzione di energia per ogni trimestre dell'anno vada da una settimana a un mese intero (ovvero, previsione a medio termine).

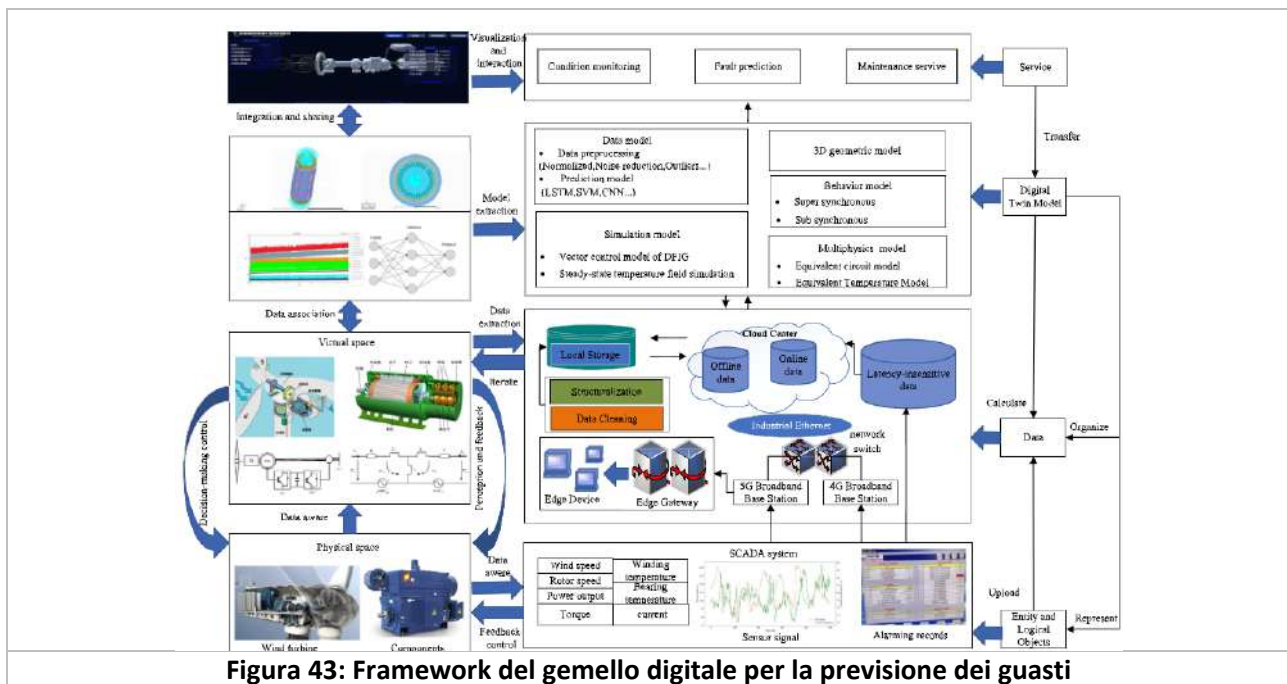


I risultati ottenuti confermano l'applicabilità del quadro proposto. Inoltre, l'analisi comparativa con i modelli di previsione classici esistenti mostra che l'approccio progettato ha ottenuto risultati migliori. Il modello può aiutare il team di gestione a monitorare i parchi eolici da remoto e a stimare in anticipo la produzione di energia.

Ricapitolando, nell'articolo è descritto lo sviluppo di un framework di Digital Twin basato su cloud assistito 5G-NG-RAN per monitorare le unità SCADA delle turbine eoliche. I Digital Twin hanno consentito il monitoraggio in tempo reale dei parchi eolici senza visitarli fisicamente. Il cloud assistito 5G-NG-RAN ha consentito ai servizi a bassa latenza di supportare il gemello digitale in tempo reale per la modellazione predittiva nelle turbine eoliche. Inoltre, una pipeline di apprendimento automatico è stata progettata sulla rete neurale convoluzionale temporale e sulla regressione kNN per prevedere la generazione di energia eolica. La valutazione empirica del set di dati pubblicamente disponibile conferma l'applicabilità in scenari di parchi eolici reali.

La previsione accurata dello stato è fondamentale per rilevare potenziali guasti del generatore di induzione a doppia alimentazione (DFIG). In [22] viene presentato un paradigma unico di previsione dei guasti Digital Twin (DT), che utilizza un paradigma di edge intelligence per garantire una previsione dei guasti ad alte prestazioni. Digital Twin basato su Edge offre capacità di elaborazione e archiviazione su dispositivi edge per consentire un'elaborazione dei dati efficace. Pertanto, il framework Digital Twin ha adottato un approccio di calcolo collaborativo edge cloud per costruire attività decisionali e di elaborazione dei dati di deep learning supervisionate distribuite. Questo approccio ha fornito un modo fattibile per implementare la tecnologia del gemello digitale nelle reti cloud edge industriali dinamiche per garantire la previsione dei guasti ad alte prestazioni.

Lo spettro di corrente dello statore e lo spettro di potenza istantanea del DFIG vengono simulati e valutati, seguiti dallo sviluppo di un prototipo di verifica sul DFIG e dall'osservazione qualitativa delle prestazioni interne del DFIG. Il codificatore è stato costruito utilizzando la **Convolutional Neural Network Squeeze-and-Excitation (SE) (SE-CNN)** con il meccanismo di attenzione del canale, e il decodificatore è stato costruito sulla base di una rete di memoria a lungo termine (LSTM) per formare il modello di autocodifica denominato SE-CNN-LSTM. Per prevedere i problemi e fornire una base per la manutenzione predittiva, l'indicatore di stato è stato implementato sulla base dei residui di previsione delle variabili di stato.



Sulla base dei risultati della co-simulazione edge-cloud e dei dati raccolti dal sistema di controllo di supervisione e acquisizione dati presso il parco eolico di Da Bancheng nello Xinjiang, in Cina, è stato dimostrato che il framework e la tecnologia sono utilizzabili e in grado di prevedere i guasti del generatore. Infatti, l'efficacia del framework proposto e dello schema di analisi dei dati è stata dimostrata prevedendo lo stato operativo del DFIG.

In [23] si presenta un monitoraggio delle condizioni basato sull'intelligenza artificiale e un quadro di manutenzione predittiva per le turbine eoliche. Lo scopo di questo framework è stato il rilevamento precoce automatizzato di guasti operativi nei sistemi e sottosistemi delle turbine eoliche. I dati estesi raccolti da WinJi attraverso i suoi servizi digitali sono stati sfruttati in questo lavoro per sviluppare e testare un accurato schema di rilevamento dei guasti e dimostrarne le prestazioni nelle turbine eoliche funzionanti commercialmente. Più algoritmi di apprendimento automatico sono stati confrontati per i modelli del normale funzionamento delle turbine, inclusi il **Random Forest (RF)**, il **Support Vector Machine (SVM)** e il **Gradient Boosting Machine (GBM)**.

Confrontando gli algoritmi applicati, gli autori hanno scoperto che i modelli addestrati in base al GBM si sono comportati in modo più affidabile nel rilevamento automatico senza supervisione dei guasti in base al normale comportamento operativo del componente. Inoltre, mentre hanno raggiunto una precisione simile nel descrivere il normale comportamento operativo, i modelli addestrati con il GBM hanno superato gli altri algoritmi di riferimento in termini di tempi di addestramento del modello. Quest'ultima è una proprietà dell'algoritmo rilevante e vantaggiosa per ridimensionare l'approccio a un gran numero di componenti monitorati e parchi eolici come nel portafoglio WinJi. L'RMSE dei modelli confrontati e i tempi in secondi necessari per l'addestramento su un processore Apple M1 sono mostrati nella seguente tabella.

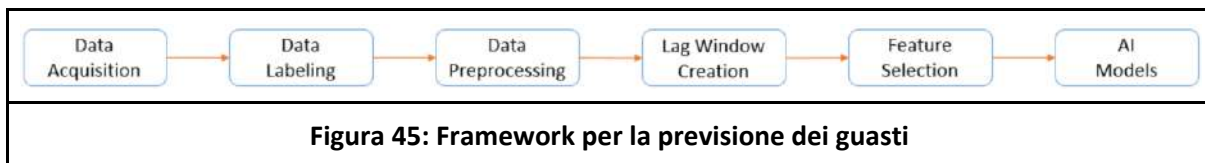
	RMSE on test dataset for variable generator temp 1 [° C]	Model training time [sec] for variable generator temp 1
Support vector machine	0.062	1320
Random Forest	0.064	2040
Gradient boosting	0.060	130

Figura 44: Panoramica degli algoritmi di apprendimento automatico per l'addestramento di modelli del normale comportamento operativo delle turbine

Nell'articolo si presentano i risultati della convalida di due parchi eolici onshore e si dimostra un'accuratezza del 97% per un orizzonte di rilevamento di 2 mesi dello sviluppo di eventi di guasto che richiedono attenzione da parte del personale di manutenzione. Pertanto, il rilevamento precoce delle anomalie ha consentito una manutenzione basata sulle condizioni e una migliore pianificazione delle riparazioni. Ciò può prevenire danni consequenziali, ridurre i tempi di inattività delle turbine e prolungare la durata delle turbine monitorate. Inoltre, il framework presentato è stato in grado di rilevare tempestivamente i guasti indipendentemente dal produttore e ha consentito la notifica ai gestori e ai proprietari del parco eolico.

In [24] viene sviluppato per un cliente reale un nuovo framework basato su tecniche di apprendimento automatico per la previsione dei guasti dei generatori di parchi eolici. La

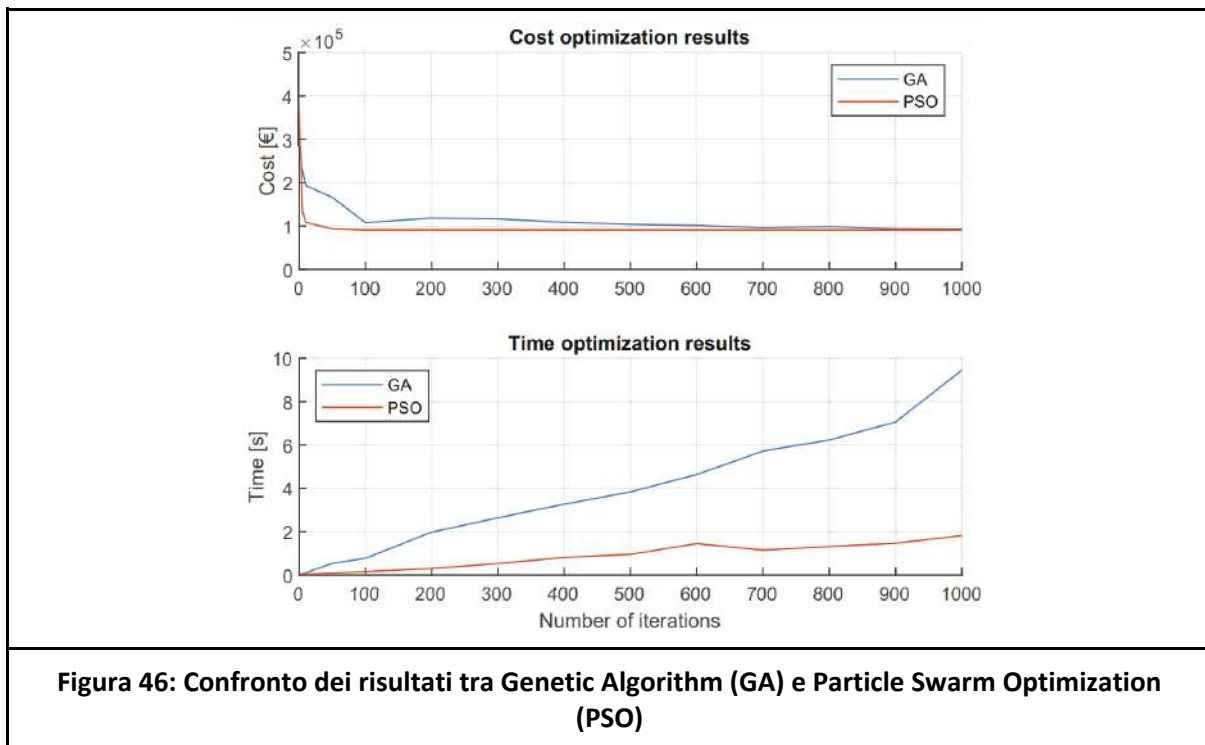
metodologia di prognosi dei guasti proposta affronta la limitazione dei dati come lo squilibrio di classe e i dati mancanti, esegue test statistici sulle serie temporali per verificarne la stazionarietà, seleziona le caratteristiche con il potere predittivo più elevato e applica modelli di apprendimento automatico per prevedere un guasto con un orizzonte di 1 ora. Le tecniche proposte vengono testate e convalidate utilizzando dati storici per un parco eolico a Summerside, Isola del Principe Edoardo (PEI), Canada, e i modelli vengono valutati in base a metriche di valutazione appropriate.



La ricerca condotta si è concentrata sullo sviluppo della metodologia per prevedere i guasti in un orizzonte di 1 ora. Questo obiettivo è stato raggiunto mediante diversi modelli di apprendimento automatico, ovvero **Artificial Neural Network (ANN)**, **Long Short-Term Memory (LSTM)**, **Support Vector Machine (SVM)**, i **modelli di insieme Adaptive Boosting (AdaBoost)** e classificatori con voto di maggioranza. È stato condotto un caso di studio per testare e dimostrare l'efficacia degli algoritmi proposti. I dati di un anno, registrati con una granularità di 10 minuti forniti dalla città di Summerside, PEI, sono stati utilizzati per sviluppare, testare, dimostrare e confrontare vari modelli di machine learning ed è stato selezionato un modello di previsione finale. Le principali conclusioni e risultati della tesi presentata sono le seguenti:

- I risultati ottenuti con il framework proposto dimostrano il potenziale della sua effettiva implementazione per ridurre al minimo i tempi di inattività del parco eolico studiato.
- I risultati presentati dimostrano che disponendo di dati storici sufficienti, i modelli di apprendimento automatico possono acquisire correttamente i modelli di errore sottostanti per la previsione dei guasti.
- I modelli sviluppati possono essere addestrati off-line e utilizzati on-line per prevedere in modo rapido e adeguato in tempo reale la probabilità che si verifichi un guasto nell'orizzonte di un'ora successiva.

In [25] viene definito e risolto un nuovo problema di ottimizzazione multi-obiettivo per casi di studio reali utilizzando **Genetic Algorithms (GA)** e **Particle Swarm Optimization (PSO)** per ridurre al minimo i costi operativi e massimizzare le prestazioni delle turbine eoliche. Vengono confrontati i risultati di entrambi gli algoritmi considerando diversi scenari in un caso di studio reale.



Il **Genetic Algorithm (GA)** è risultato generalmente meno costoso dal punto di vista computazionale, ma ha fornito risultati meno ottimali rispetto a **Particle Swarm Optimization (PSO)**. L'uso dell'una o dell'altra tecnica dovrebbe essere discusso a seconda dell'applicazione, della complessità e dell'obiettivo del problema. Il Particle Swarm Optimization (PSO) è stato selezionato come il miglior algoritmo in quanto, in confronto, ha raggiunto la soluzione ottimale più velocemente e con un minor costo computazionale. Lo studio di sensibilità ha mostrato come i tempi di manutenzione preventiva possono essere aumentati con un aumento minimo dei costi rispetto alla manutenzione correttiva. Anche la produzione di energia ha provocato variazioni di costo, ma è naturalmente mitigata dai proventi dell'energia generata stessa.

In [26] vengono valutati e confrontati diversi modelli di **Artificial Neural Network (ANN)** per la diagnosi e il monitoraggio delle turbine eoliche (WT) che prevedono il guasto del sistema del macchinario in base ai segnali dei sensori dei componenti interni e alla temperatura di generazione. Sono stati sviluppati diversi modelli ML con diversi dati di sensori basati sui dati SCADA raccolti da nove WT in dieci anni ricevuti dal Westmill Windfarm situato a Swindon, nel Regno Unito, per prevedere in anticipo il guasto del generatore nei WT utilizzando il riconoscimento del modello basato su dati storici. I risultati dell'accuratezza di ciascun modello in termini di offset minimo, massimo e deviazione standard tra i valori di temperatura del generatore previsti e quelli effettivi sono stati confrontati e contrastati ed è stato esplorato l'effetto dei dati del sensore di input.

Model	Minimum error offset value	Maximum error offset value	Standard deviation error offset value
1	0.0	73.51	4.10
2	0.0	55.19	3.80
3	0.0	66.67	4.32
4	0.0	67.18	3.99
5	0.0	73.89	4.40

Figura 47: Precisione dei modelli per gli offset di deviazione minima, massima e standard tra i valori registrati previsti e quelli effettivi.

Viene utilizzato un approccio al modello di machine learning con due livelli nascosti utilizzando una **Multilayer Perceptron (MLP)** per raggiungere l'obiettivo.

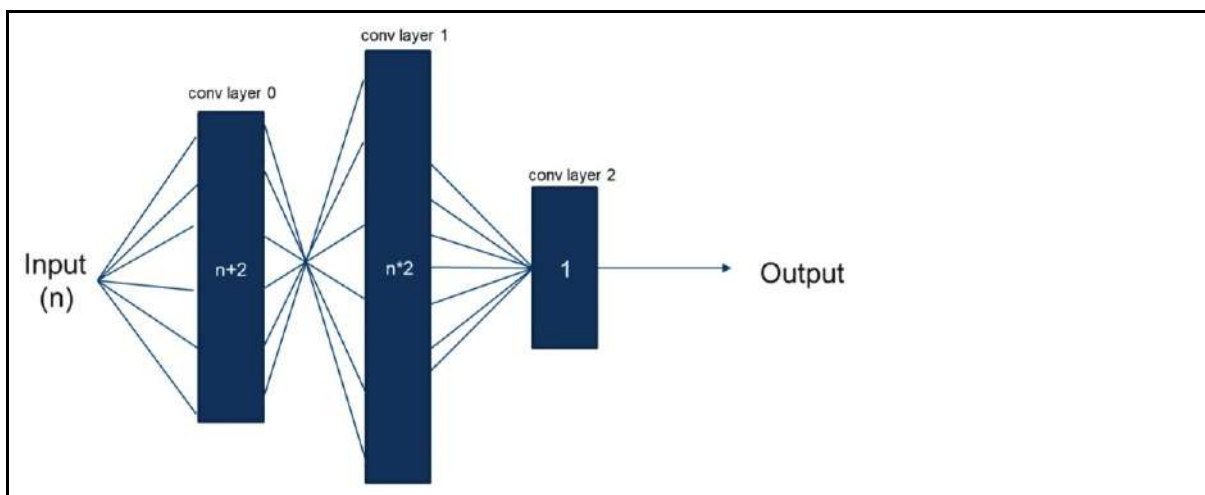


Figura 48: Il layout schematico dell'algoritmo di regressione.

Il sistema sviluppato ha previsto l'uscita della temperatura del generatore del WT con una precisione del 99,8% con 2 mesi di previsione della misurazione in anticipo. Nel complesso, nell'articolo viene mostrata la possibilità di utilizzare algoritmi di regressione guidati da ML per prevedere il guasto del generatore di WT causato dal calore, riducendo i costi di manutenzione relativi ai tempi di inattività e al personale e, allo stesso tempo, migliorando la disponibilità operativa degli apparati.

Utilizzando dati di serie temporali e valutazioni di ispezione delle pale delle turbine, in [27] viene presentato un modello di classificazione per prevedere la durata residua delle pale delle turbine eoliche. L'obiettivo del modello di apprendimento automatico è prevedere il punteggio di gravità delle turbine utilizzando le caratteristiche aggregate create per catturare l'affaticamento delle pale. Il modello utilizzato è il **Random Forest (RF)** potenziato con convalida incrociata dieci volte. La seguente figura mostra i risultati delle singole turbine colorati per sito.

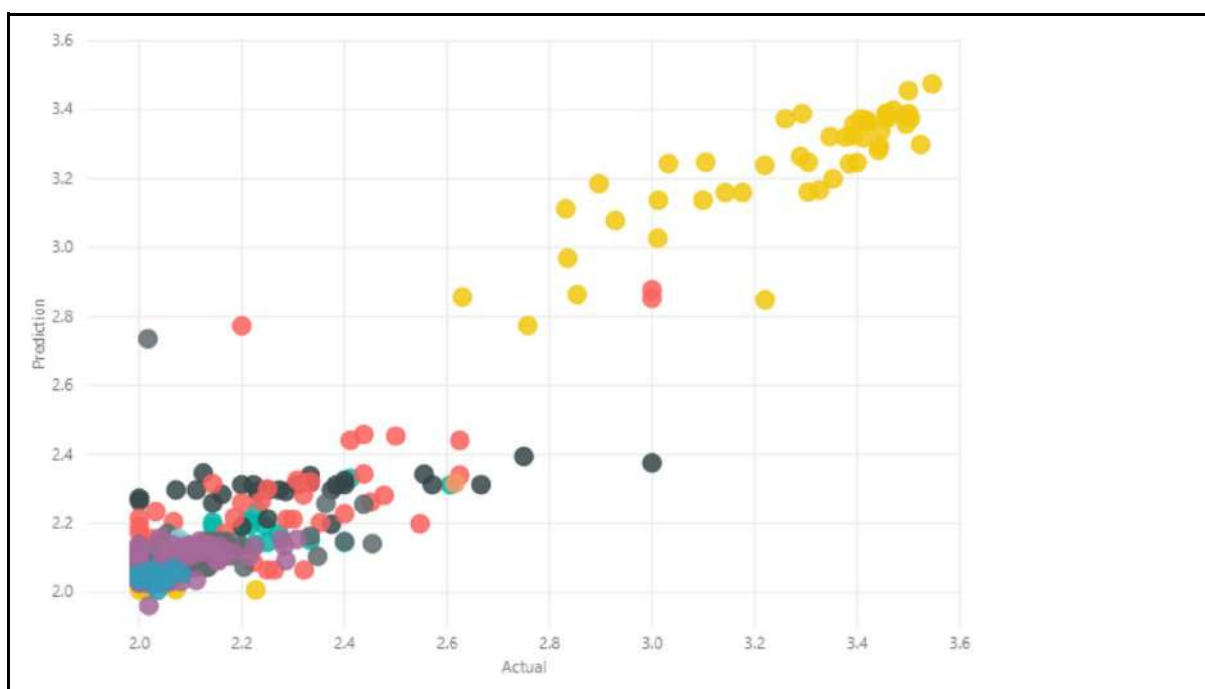
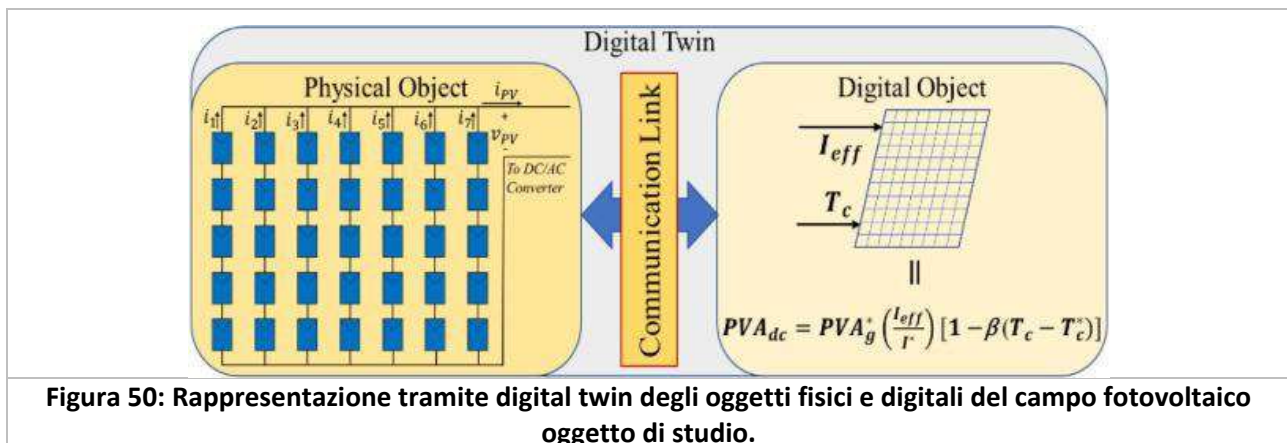


Figura 49: Effettivo vs. previsto per sito.

Nell'articolo viene dimostrato che, campionando le condizioni effettive delle pale con la fotografia dei droni, è possibile prevedere le condizioni del resto delle pale delle turbine eoliche nel sito con un errore R-quadro dell'86%. Inoltre, viene confermato che l'effetto della pioggia (ad alta velocità della punta della lama) risulta essere la forza motrice significativa per l'erosione della lama. Questo risultato non solo ha consentito di risparmiare denaro sulle ispezioni con i droni, ma ha consentito anche una stima più precisa delle entrate perse e una comprensione più chiara degli investimenti di riparazione.

2.2.1.2 Impianti fotovoltaici

In [28] è stato sviluppato un metodo di diagnosi dei guasti FV che prevede le seguenti tre fasi: (1) Rilevamento dei guasti in un campo fotovoltaico mediante DT, (2) classificazione dei guasti rilevati mediante ConvMixer e (3) notifica dei guasti rilevati e classificati utilizzando un sistema LoRa (a lungo raggio). Viene implementata la generazione di energia CC basata su modello di array fotovoltaici utilizzando DT. Viene presentato un nuovo **metodo di Deep Learning (DL), chiamato mixer convoluzionale (ConvMixer)** basato sull'incorporamento di patch e sulla combinazione di convoluzioni in senso profondo e puntuale per classificare i guasti fotovoltaici. Gli input di ConvMixer sono immagini 2D generate dai dati sull'alimentazione dell'array CC FV utilizzando una trasformazione del campo di transizione di Markov (MTF).



La trasformazione Markov Transition Field (MTF) viene applicata alla potenza dell'array DC PV per generare un'immagine 2D come input per ConvMixer nella fase di classificazione. Il noioso processo di utilizzo di un'altra trasformazione, che richiede tentativi ed errori, viene così evitato. La replica dell'oggetto fisico come controparte di un oggetto digitale utilizzando il gemello digitale si è dimostrata efficace nel rilevamento dei guasti. I risultati della classificazione ottenuti dal ConvMixer proposto, che ha raggiunto un'accuratezza complessiva del 97,00% nei test, sono più accurati di quelli ottenuti da **altri metodi ML classici**, di cui la **Random Forest (RF)** ha avuto l'accuratezza complessiva più alta del 92,27%. Il metodo proposto supera anche altri ben noti metodi basati sulla CNN, di cui ResNet aveva la massima precisione complessiva del 94,00%. Con il minor numero di parametri addestrabili (603.143), il ConvMixer proposto richiede il tempo di addestramento più breve (30,057 min); tutti gli altri metodi basati sulla CNN richiedono almeno il doppio del tempo per l'addestramento. I meccanismi di incorporamento delle patch in ConvMixer aiutano notevolmente a generalizzare le funzionalità estratte mediante convoluzioni sia in profondità che in senso puntuale. È incluso un sistema di notifica che utilizza LoRa per trasmettere informazioni sui tipi e le caratteristiche dei guasti. Inoltre, viene utilizzato un simulatore digitale in tempo reale per confermare l'applicabilità pratica del sistema. Il sistema di notifica LoRa è molto adatto per l'uso in reti WAN a bassa potenza (LPWAN), come quelle necessarie nei grandi parchi fotovoltaici.

I risultati della simulazione dimostrano che il ConvMixer proposto supera altri metodi classici di machine learning (ML), come decisional tree, k-nearest neighbor, random forest e support vector machine, nonché altri metodi classici basati su CNN, come AlexNet, ResNet50, VGG16 e VGG19. Un

simulatore digitale in tempo reale (Opal-RT eMegsim) viene utilizzato per verificare l'applicabilità in tempo reale dell'integrazione di DT, ConvMixer e il sistema di notifica LoRa.

A causa della natura limitante della corrente e delle caratteristiche di uscita non lineari degli array fotovoltaici, il rilevamento dei guasti non è così semplice e si propone l'applicazione dell'intelligenza artificiale per il rilevamento dei guasti nei sistemi fotovoltaici. L'idea alla base dell'approccio in [29] è quella di confrontare il modulo fotovoltaico difettoso con il suo modello accurato (impronta digitale di fabbrica) controllando le curve I-V e P-V di ogni campo fotovoltaico utilizzando l'algoritmo della **Artificial Neural Network (ANN)** come sottosezione dell'intelligenza artificiale (Tecniche AI).

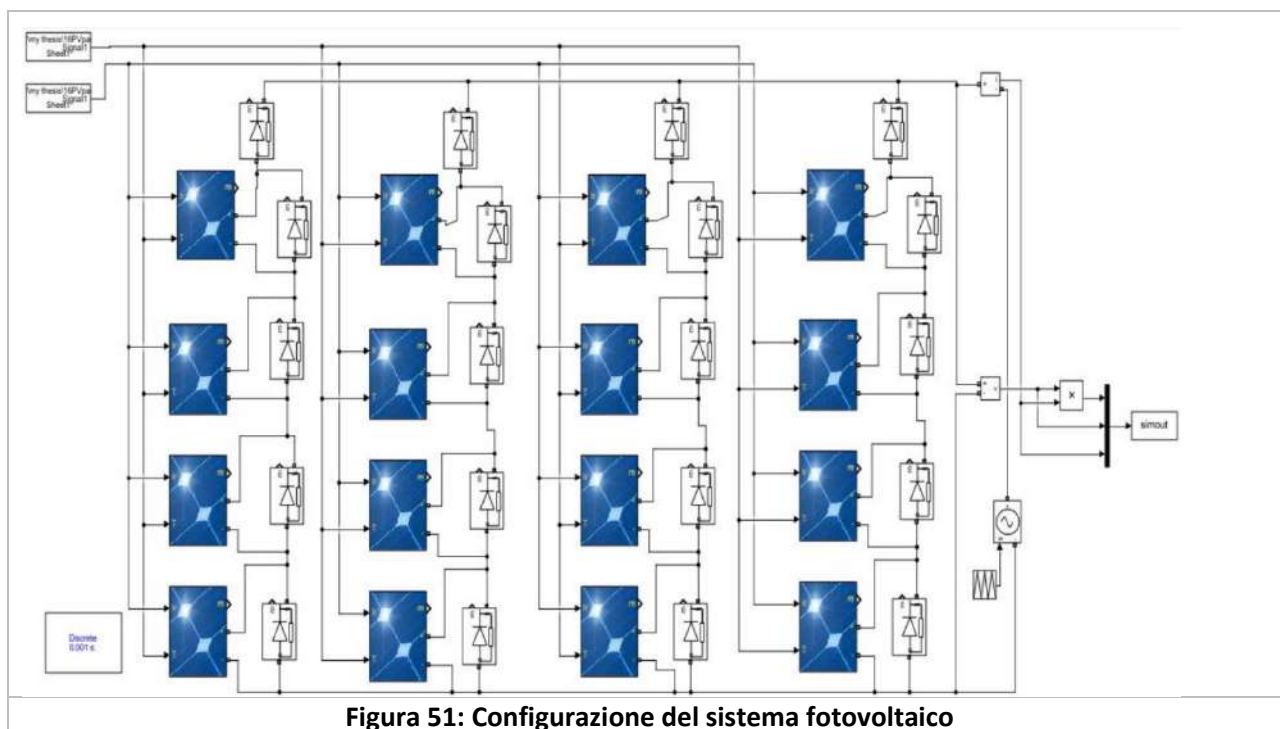


Figura 51: Configurazione del sistema fotovoltaico

La difficoltà di trovare il guasto e di classificarne la tipologia in un impianto fotovoltaico di grandi dimensioni può essere risolta utilizzando la tecnica proposta. Si è riscontrato che questa soluzione potrebbe facilitare il compito di manutenzione degli impianti fotovoltaici, in particolare quelli di grandi dimensioni. Inoltre, la tecnica può risolvere il problema dell'improvvisa riduzione di potenza a causa di guasti imprevisti. L'algoritmo ANN, come sottosezione delle tecniche AI, viene applicato

per classificare i guasti. Sono stati creati tre scenari basati su una serie di parametri utilizzati come input per l'algoritmo. Questi parametri sono stati estratti dalle curve di output. Matlab e Simulink sono stati i software utilizzati per simulare il sistema fotovoltaico modellato. I codici sono stati creati da Matlab per raggiungere lo scopo del documento. Per confrontare i tre scenari sono stati utilizzati diversi strumenti di valutazione (matrici di confusione, curve ROC, istogramma degli errori e migliori curve delle prestazioni di convalida). Il terzo scenario ha avuto le migliori performance di rilevamento e classificazione con una percentuale del 99,9%. Il secondo scenario e il primo scenario hanno avuto percentuali di performance rispettivamente del 94,7% e dell'83,2%. Inoltre, l'errore dell'istogramma, le curve ROC e le migliori curve di convalida hanno enfatizzato lo stesso risultato. I tre scenari hanno raggiunto un'elevata precisione in base ai dati di input reali. L'algoritmo ANN è pertanto in grado di rilevare i guasti rapidamente e con elevate prestazioni e affidabilità.

Riepilogando, l'approccio proposto ha raggiunto un'elevata prestazione di rilevamento dei guasti offrendo il vantaggio di determinare quale tipo di guasto si è verificato. I risultati hanno confermato che le prestazioni proposte migliorano con l'aumentare del numero di punti distintivi, fornendo un grande valore all'industria solare fotovoltaica.

La potenza erogata dal modulo fotovoltaico cambia continuamente nel tempo a seconda delle condizioni climatiche. Questi cambiamenti sono più importanti nelle aree tropicali come il Senegal a causa della variazione delle stagioni (secca e piovosa). Inoltre, in letteratura sono presentati diversi tipi di algoritmi di inseguimento del punto di massima potenza (MPPT) per ottenere il massimo rendimento dal sistema fotovoltaico. Possono essere riassunti in due categorie: metodi classici e metodi intelligenti. I metodi classici in condizioni meteorologiche non uniformi non sono efficienti e si evidenzia un'importante perdita di energia. Tuttavia, di fronte a queste problematiche come la perdita di energia e l'assenza di condizioni meteorologiche uniformi, viene utilizzato un metodo adattivo per ottimizzare l'energia fotovoltaica. In [30] vengono proposti due controller intelligenti basati su **Artificial Neural Network (ANN)** e **Adaptive Neuro Fuzzy Inference System (ANFIS)** per ottimizzare la produzione di PV in condizioni meteorologiche non uniformi e confrontati per mostrare il modello più potente.

Per l'ANN, la sfida principale è trovare l'optimal neural nello strato nascosto e nell'articolo esso è stato ottenuto utilizzando un fattore di valutazione come l'errore quadratico medio (MSE). Queste tecniche che utilizzano algoritmi di intelligenza artificiale (AI) vengono utilizzate per l'ottimizzazione della potenza di un sistema fotovoltaico e vengono addestrate e convalidate con dati reali da una micro-centrale fotovoltaica nella stagione secca e piovosa installata presso il liceo politecnico di Dakar. Le prestazioni dei regolatori per ottimizzare la potenza FV vengono valutate durante le stagioni secche e piovose.

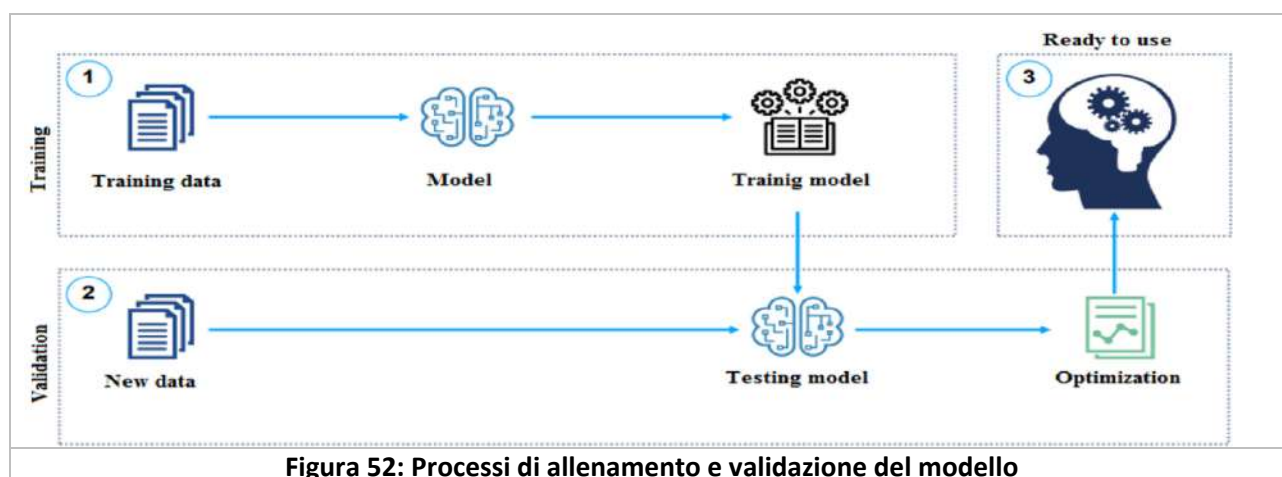


Figura 52: Processi di allenamento e validazione del modello

Diversi fattori influenzano le prestazioni dei pannelli fotovoltaici, come condizioni di ombreggiamento parziale e condizioni meteorologiche non uniformi. Nell'area tropicale, le condizioni climatiche variano in ogni stagione e sono più significative nella stagione delle piogge a causa delle nuvole. Pertanto vengono analizzate le prestazioni di un regolatore ANN e ANFIS MPPT di un PV in condizioni climatiche non uniformi in due mesi (marzo e agosto). Per ottenere il controllore efficiente più significativo, sono stati adottati diversi coefficienti matematici per determinare la validità di risultati accurati come MAPE, MAE e RMSE. I risultati mostrano che ANFIS è più potente di ANN grazie alla sua capacità di adattarsi a condizioni meteorologiche non uniformi e alla sua capacità di generalizzare.

Le caratteristiche corrente-tensione (curve I-V) dei moduli fotovoltaici (PV) contengono molte informazioni sulla loro salute. In letteratura, per la diagnosi vengono utilizzate solo informazioni parziali dalle curve IV. In [31] viene sviluppata una metodologia per sfruttare appieno le curve I-V per la diagnosi dei guasti FV. Nella fase di pre-elaborazione, la curva I-V viene prima corretta e ricampionata. Quindi le caratteristiche di guasto vengono estratte in base all'uso diretto del vettore ricampionato della corrente o alla trasformazione per campo di differenza angolare Gramiano o grafico di ricorrenza. **Sei tecniche di apprendimento automatico, ovvero Artificial Neural Network (ANN), Support Vector Machine (SVM), Decision Trees (DT), Random Forest (RF), K-nearest Neighbors algorithm (KNN) e classificatore bayesiano naive (ingenuo)** vengono valutate per la classificazione delle otto condizioni (sane e sette difettose) dell'array fotovoltaico. Viene prestato uno sforzo particolare per scoprire le migliori prestazioni (precisione e tempo di elaborazione) quando si utilizzano diverse funzionalità di input combinate con ciascuno dei classificatori. Inoltre, viene affrontata anche la robustezza rispetto al rumore ambientale e agli errori di misurazione.

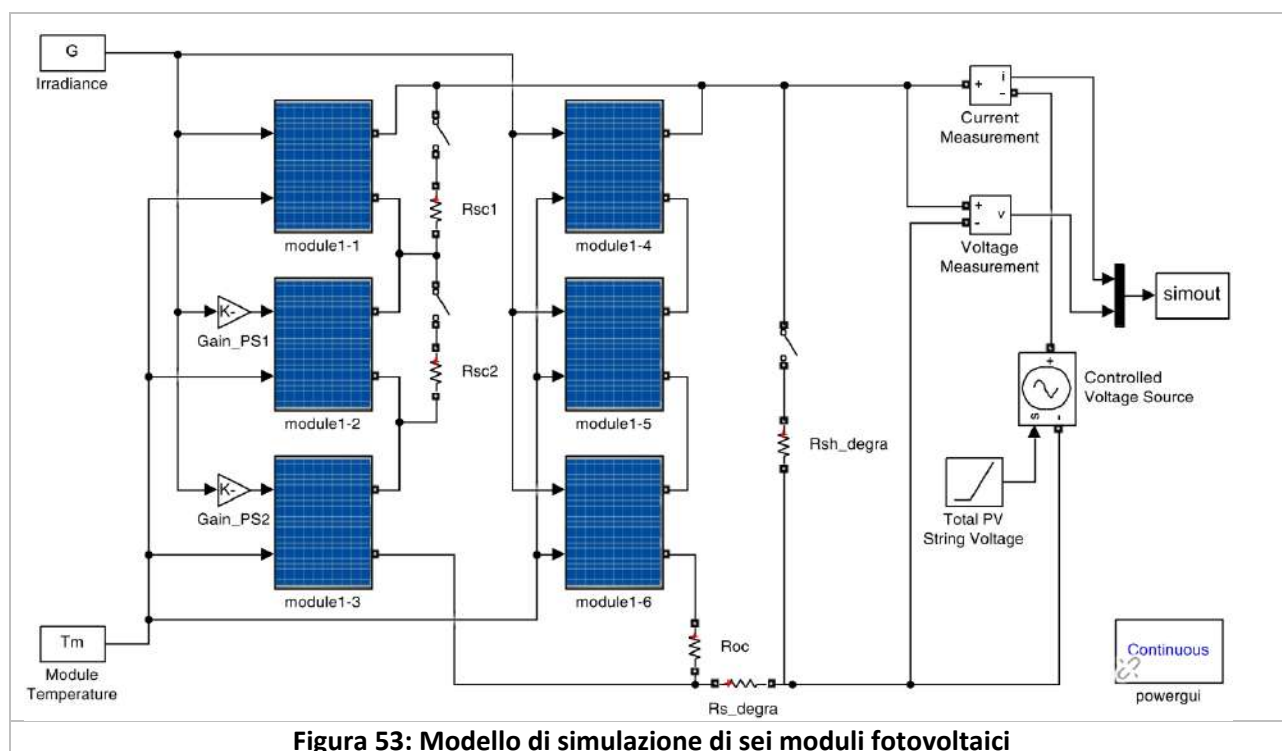


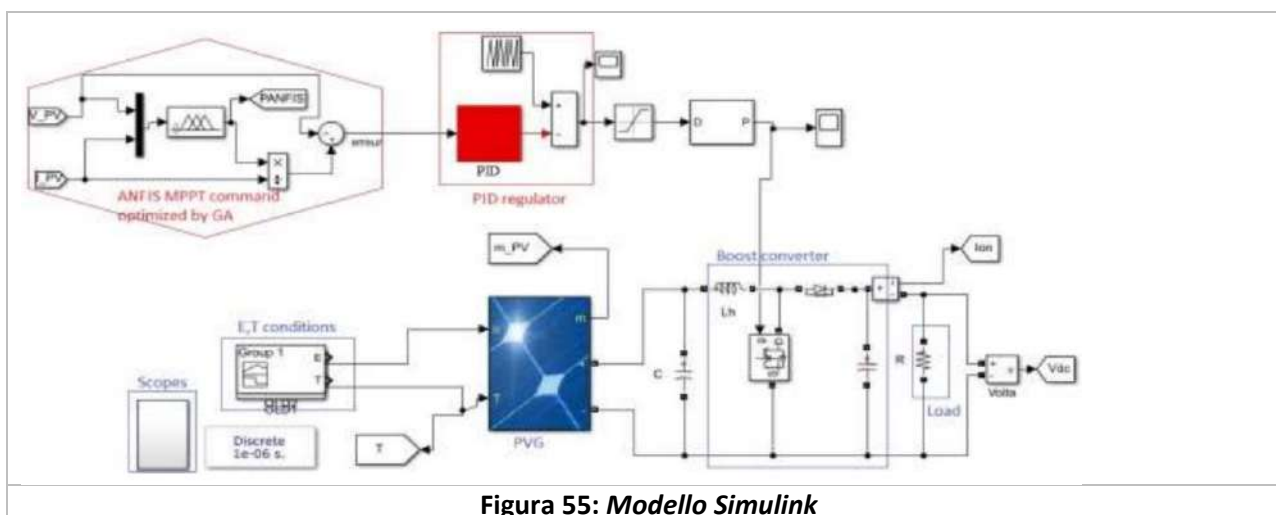
Figura 53: Modello di simulazione di sei moduli fotovoltaici

Nell'articolo viene pertanto introdotta una metodologia che utilizza curve I-V complete e tecniche di apprendimento automatico per la diagnosi dei guasti del campo fotovoltaico in otto condizioni.

Viene dimostrato che i tre metodi di estrazione delle caratteristiche che utilizzano informazioni complete sulla curva I-V hanno prestazioni migliori rispetto a quelli che utilizzano informazioni parziali di una curva I-V. Tra i tre metodi completi basati sull'informazione, attraverso la trasformazione a 1 o 2 dimensioni, il campo di differenza angolare Gramiano ha mostrato la più alta capacità di discriminare le otto condizioni e ha mostrato la massima robustezza agli errori di misurazione e ai disturbi ambientali. I classificatori sintonizzati con campioni simulati sono stati convalidati con curve I-V misurate sul campo.

Si è scoperto che il miglior classificatore ha raggiunto un'accuratezza di classificazione del 100% sia con la simulazione che con i dati sul campo. Sono stati inoltre analizzati la riduzione dimensionale delle caratteristiche, la robustezza dei classificatori al disturbo e l'impatto della trasformazione.

In [32] è presentato un nuovo approccio ibrido basato sulla combinazione di **Adaptive Neuro Fuzzy Inference System (ANFIS)** e **Genetic Algorithm (GA)** per la continuazione del punto di massima potenza (MPP) di un generatore fotovoltaico (PVG). I PVG, considerati troppo dipendenti dalle condizioni climatiche, necessitano di strumenti che gli permettano di erogare in ogni momento il massimo della potenza disponibile. A tale scopo, viene eseguita una combinazione tra GA e ANFIS per una migliore ottimizzazione della potenza del PVG. Il GA è utilizzato in questo contesto per ottimizzare la scelta delle funzioni di appartenenza (MsF) di ANFIS per ottimizzare la potenza fotovoltaica.



I risultati ottenuti hanno dimostrato l'efficacia di GA nella ricerca e ottimizzazione di un controller ibrido neuro-fuzzy di tipo ANFIS. Hanno mostrato che l'ANFIS ottimizzato da GA è migliore dell'ANFIS non ottimizzato. Il metodo proposto, infatti, ha superato i metodi ANFIS non ottimizzati trovati in letteratura con un tempo di risposta più breve (8,24 ms contro 83 ms per un controller ANFIS non ottimizzato).

Questo è anche il caso della validazione sperimentale dei due modelli MPPT. Le prestazioni e la robustezza del controller dipendono fortemente dalla scelta dei parametri GA (dimensione della

popolazione, tasso di selezione, tasso di ricombinazione e numero di generazioni), ma anche dall'apprendimento ANFIS.

I metodi di apprendimento automatico sono stati utilizzati per risolvere complicati problemi pratici in diverse aree e stanno diventando sempre più popolari oggi. Lo scopo dell'articolo [33] è stato quello di valutare la previsione della produzione energetica di tre diversi sistemi fotovoltaici e la supervisione dei sensori di misura, tramite Machine learning e data mining in risposta al comportamento delle variabili climatiche del luogo oggetto di studio. D'altra parte, l'articolo ha incluso anche l'implementazione dei modelli risultanti nel sistema SCADA attraverso indicatori, che consentiranno all'operatore di gestire attivamente la rete elettrica. Viene offerta inoltre una strategia di simulazione e previsione in tempo reale di sistemi fotovoltaici e sensori di misurazione nel concetto di reti intelligenti.

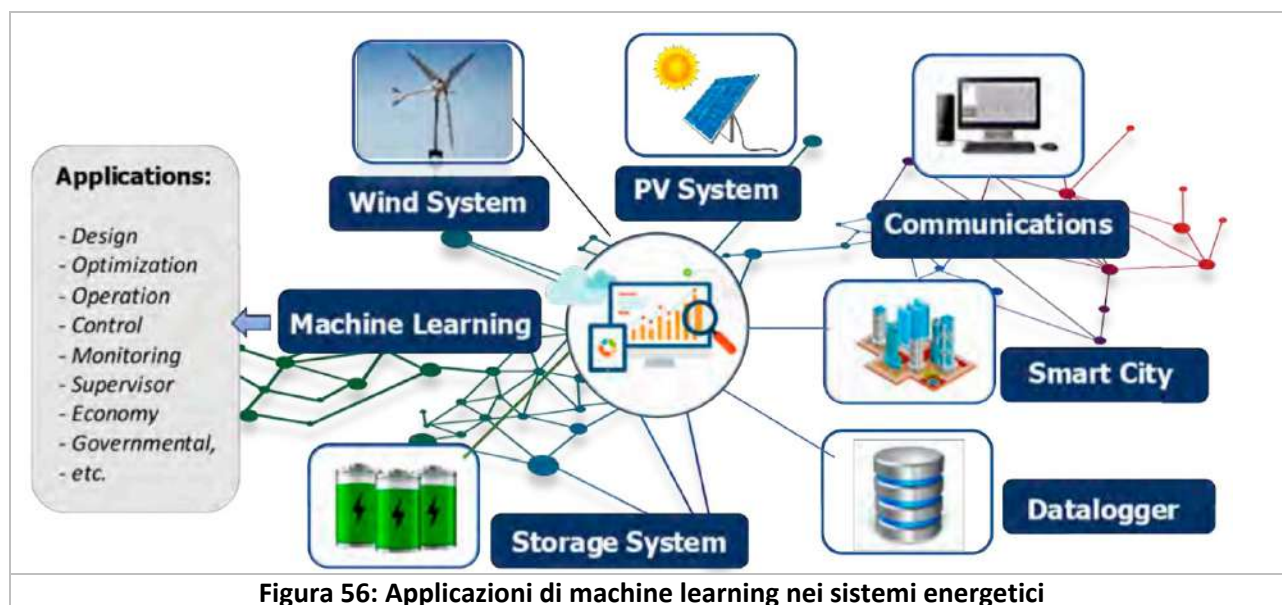


Figura 56: Applicazioni di machine learning nei sistemi energetici

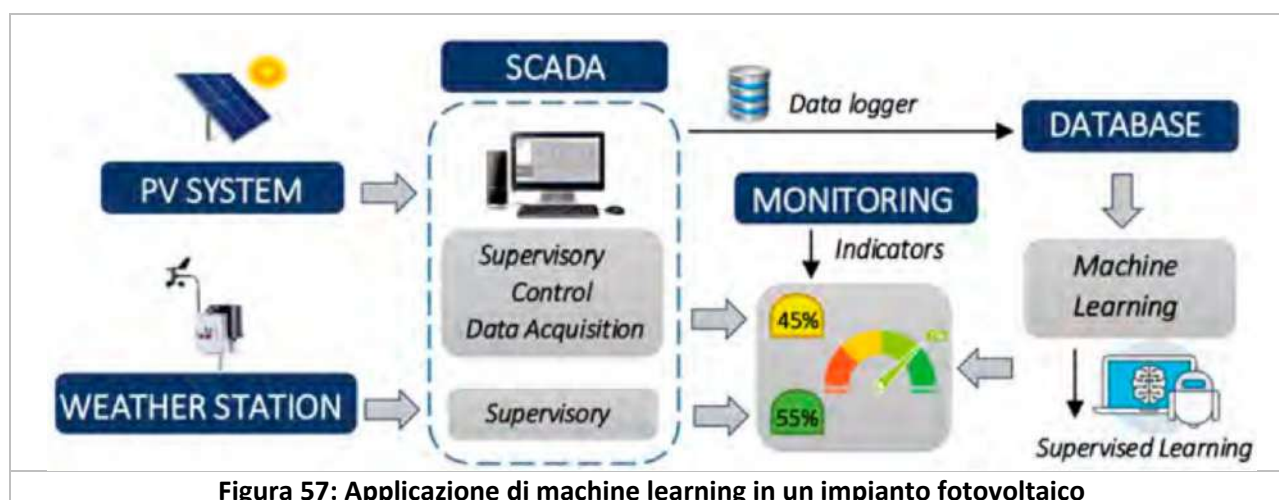


Figura 57: Applicazione di machine learning in un impianto fotovoltaico

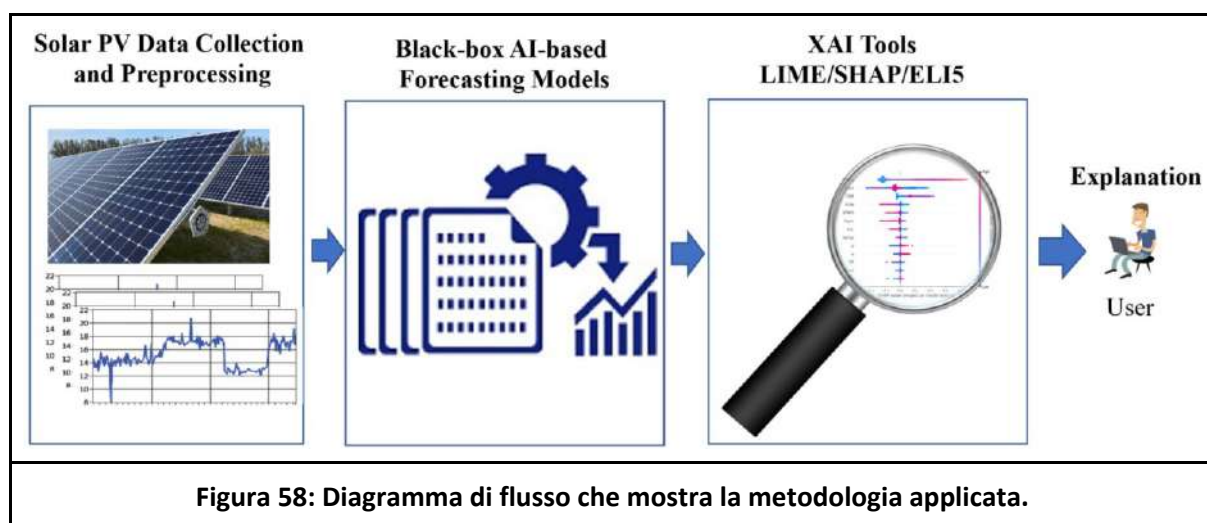
Nell'articolo viene dimostrato che è possibile prevedere la potenza fotovoltaica dei tre sistemi studiati mediante modelli di regressione stabiliti con ottima approssimazione ed in particolare **Linear Support Vector Machines (LSVM)**, **Logistic Regression (LR)**, **Locally Weighted Linear Regression (LWLR)**, **Gaussian Discriminant Analysis (GDA)**, **Back-propagation Neural Network (BPNN)**, **Expectation Maximization (EM)**, **Naive Bayes (NB)** and **Value-Added Tax (VAT)**. L'importanza del monitoraggio delle variabili tramite i sensori di misura ha permesso di ottenere un adeguato controllo dell'impianto fotovoltaico. L'applicazione di una strategia di rilevamento dei guasti è stata dimostrata attraverso tecniche di modelli predittivi negli impianti fotovoltaici, in modo tale da consentire di monitorare gli impianti fotovoltaici confrontando la potenza fotovoltaica generata in tempo reale con i valori calcolati in base alle variabili meteorologiche dell'irraggiamento e della temperatura. Sono stati ottenuti coefficienti di correlazione dell'83,27%, 82,36% e 85,76% nei risultati del modello rispettivamente per i sistemi PVS1, PVS2 e PVS3.

Viene stabilito un range del 20% che permette il confronto dei valori di calcolo con i valori misurati in tempo reale. Le equazioni del modello di misura della temperatura e del limite di misura della radiazione hanno consentito inoltre di monitorare il comportamento dei sensori ed escludere possibili errori di misura in funzione del tempo. In questo modo si è ottenuto un ulteriore mezzo di monitoraggio e controllo delle equazioni utilizzando questi parametri. L'implementazione dei modelli di previsione degli impianti fotovoltaici nello SCADA ha consentito il monitoraggio del gestore dell'impianto elettrico in modo ottimale. Infine, è stata dimostrata l'importanza

dell'applicazione delle tecniche di apprendimento automatico e la sua ampia varietà di sviluppo nel campo della gestione dell'energia e la sua importanza nelle reti intelligenti.

In [34] vengono presentati diversi casi d'uso della previsione dell'energia solare fotovoltaica utilizzando strumenti **XAI (Explainable Artificial Intelligence)**, come **LIME (Local Interpretable Model-agnostic Explanations)**, **SHAP (SHapley Additive exPlanations)** ed **ELI5 (Explain Like I'm 5)**, che possono contribuire all'adozione di strumenti XAI per applicazioni di reti intelligenti.

Viene presentata l'applicazione di un modello di previsione **Random Forest (RF)** per prevedere la generazione di energia solare fotovoltaica. Inoltre, gli strumenti XAI, come LIME, SHAP ed ELI5, vengono applicati al modello AI **Random Forest** per comprendere e spiegare le ragioni di una particolare previsione e per contribuire all'adozione degli strumenti XAI nelle applicazioni smart grid. I dati utilizzati per questo articolo sono stati un set di dati pubblico di GEFCOM 2014.



In base ai risultati, gli strumenti XAI possono fornire informazioni dettagliate per interpretare le caratteristiche e i risultati del modello, nonché migliorare i risultati del modello attraverso la spiegabilità e la trasparenza. Nell'articolo vengono fornite alcune indicazioni su come scegliere uno strumento XAI, come LIME, SHAP ed ELI5. Ogni strumento/pacchetto XAI ha i suoi punti di forza e

limiti, in termini di costo di elaborazione, spiegazione a livello locale/globale, peso delle funzionalità, ecc. Le utility sono state disposte per creare centri di controllo di nuova generazione con tecnologie di visualizzazione e pedagogi di business analytics, che supportino tecnologie emergenti, come AI e realtà mista, ma allo stesso tempo per semplificare l'usabilità per i dipendenti che hanno meno esperienza con queste tecnologie. I sistemi di previsione solare fotovoltaica basati su XAI e strumenti simili hanno fornito un terreno di gioco molto produttivo per le utility.

Per implementare l'energia solare come fonte di energia rinnovabile affidabile su vasta scala, l'impianto di operazioni e manutenzione efficienti (O&M) funge da gioco di ruolo cruciale. In [35] viene dimostrato l'utilizzo della tecnologia di intelligenza artificiale (AI) per ottimizzare le attività O&M per oltre 150 siti di progetto situati dal centro al sud di Taiwan fino a 54 MW con 11 marchi di inverter e 9 fornitori di moduli contemporaneamente. Senza modulo specifico, inverter e parametri di posizione come input, il modello di previsione della potenza per ogni inverter-MPPT viene addestrato sulla base dei propri dati storici di produzione e stabilito come propria impronta digitale. Ogni comportamento dell'inverter-MPPT di tutti i progetti viene monitorato e analizzato dal machine learning ogni cinque minuti. In totale, oltre 7.583 elaborazioni di dati MPPT sono state completate ogni 5 minuti e vi sono stati almeno 6.369.720 input di dati batch per l'addestramento del modello per tutti i siti del progetto. Per il calcolo di dati così massicci, sono stati testati vari algoritmi di apprendimento automatico dell'intelligenza artificiale **come Convolutional Neural Network (CNN), classificatore K-Nearest Neighbor (KNN), Non Linear Regression (NLR) e Long Short-Term Memory (LSTM)**. Sono state modificate e combinate queste metodologie per il motore del modello di previsione della potenza.

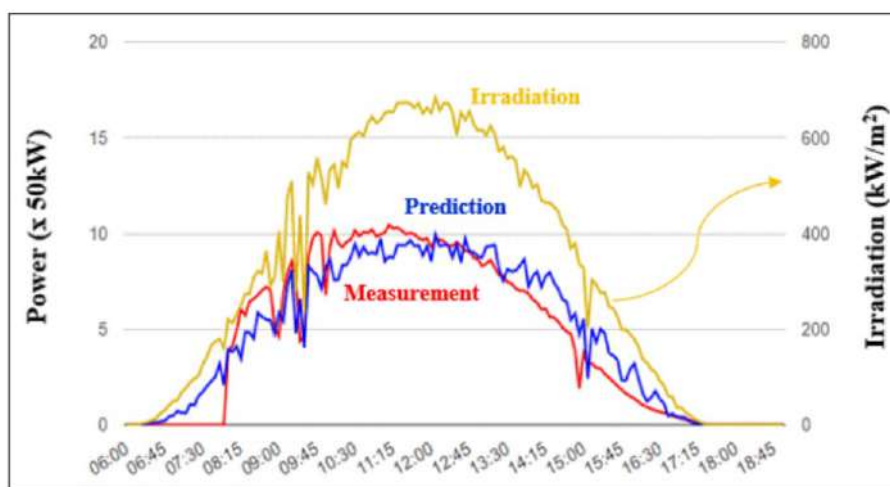
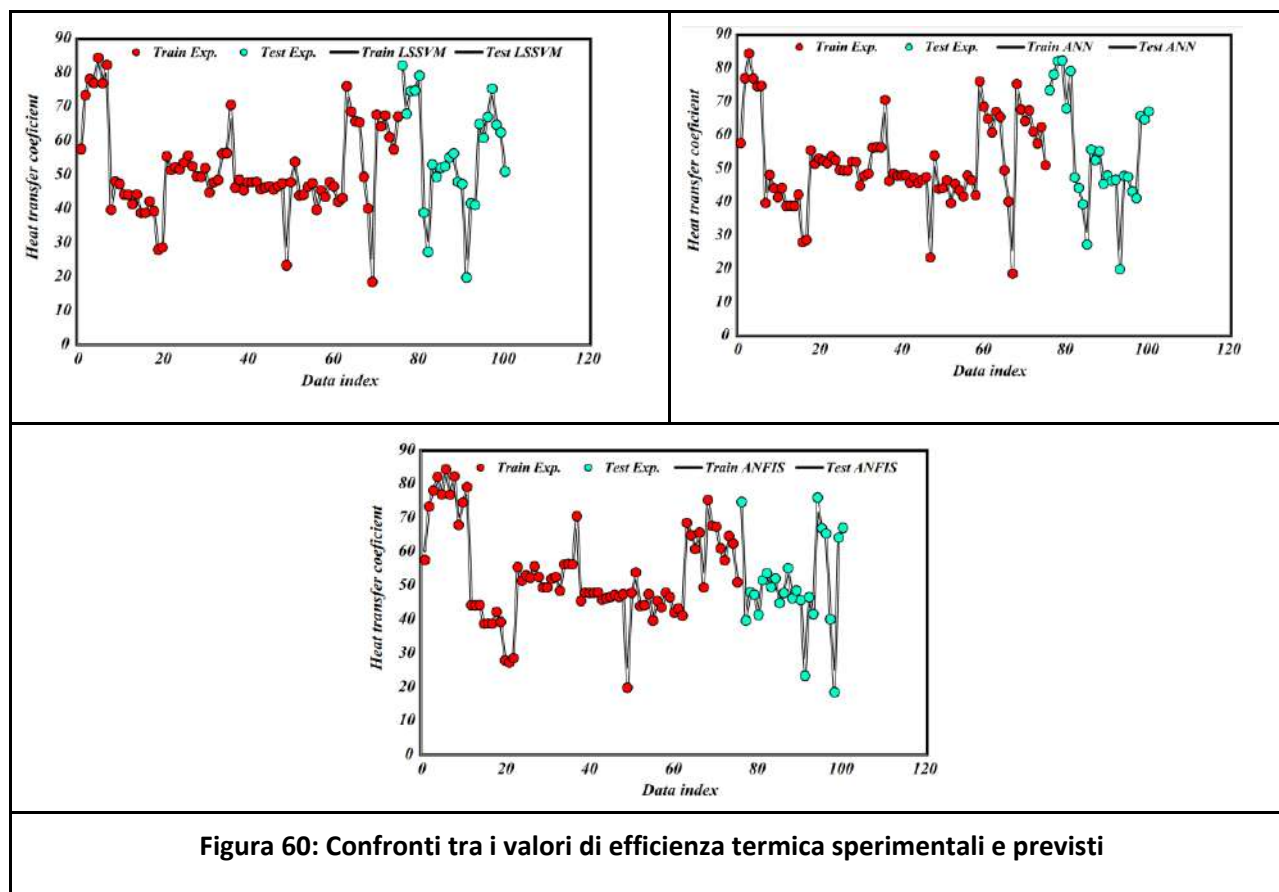


Figura 59: Modellazione della previsione della potenza come impronta digitale per ogni sito del progetto fotovoltaico.

L'avviso di rilevamento dei guasti con modalità di errore viene giudicato automaticamente e la notifica tempestiva viene inviata all'utente tramite dispositivo mobile o e-mail. La precisione per il rilevamento dei guasti è stata del 99,2% e ha raggiunto anche il 92,3% per la diagnosi della modalità di guasto. I falsi allarmi di solito sono svaniti rapidamente dopo aver impiantato la soluzione AI, dimostrando che il sistema apprende davvero rapidamente. Inoltre, la riduzione dei tempi di inattività per apparecchiature malfunzionanti ha garantito una pratica O&M più efficiente. Dal risultato dell'implementazione di un anno, il rendimento energetico medio di 0,16 kWh/kWp/giorno (4,7%) è migliorato per i siti del progetto con questo sistema di intelligenza artificiale rispetto a quelli senza.

Le energie rinnovabili, in particolare l'energia solare, sono state impiegate in numerose applicazioni pur essendo energia senza emissioni di CO₂ rispetto alle risorse di combustibili fossili. Lo scopo principale dell'articolo [36] è stato quello di prevedere l'efficienza termica delle configurazioni fotovoltaiche-termiche (PV/T) in relazione alla temperatura di ingresso, alla portata di ricircolo e all'irradiazione solare, modificando gli algoritmi di **Multilayer Perceptron Artificial Neural Network (MLP-ANN)**, **Adaptive Neuro-Fuzzy Inference System (ANFIS)** e **Least-Squares Support Vector**

Machines (LSSVM). Per questo obiettivo, sono state eseguite più di 100 misurazioni empiriche su una configurazione fabbricata PV/T raffreddata ad acqua. Sono stati impiegati metodi grafici e statistici per determinare la credibilità dei modelli proposti nella previsione accurata dell'efficienza termica, ed è stata confermata l'esistenza di una grande concordanza tra modelli predittivi e dati effettivi.



Il modello **Artificial Neural Network (ANN)** proposto ha fornito le migliori prestazioni rispetto ai modelli **ANFIS** e **LSSVM** a causa dei valori dell'errore quadratico medio (MSE) e del coefficiente di determinazione (R^2) rispettivamente pari a 0,009 e 1,00. Inoltre, sono stati eseguiti confronti numerici con altri modelli sviluppati di recente.

È stata inoltre effettuata l'implementazione dell'analisi di rilevamento dei valori anomali allo scopo di determinare possibili punti di dati periferici, mostrando la fattibilità dei modelli proposti per la maggior parte dei punti dati. Quasi tutti i dati si trovano nella regione ammissibile. È stato inoltre

realizzato uno studio di sensibilità per specificare gli effetti di diversi parametri di ingresso sull'efficienza termica, in cui la temperatura di ingresso ha rivelato il massimo impatto sull'efficienza termica del collettore solare per quanto riguarda il fattore di rilevanza più elevato pari a 0,6. Di conseguenza, questi modelli sono di facile utilizzo e possono essere un valore affidabile per studiare i parametri efficaci nei collettori solari.

In [37] viene presentato uno studio sull'impianto di desalinizzazione dell'acqua di mare (SWDP) situato in Egitto alimentato dalla rete elettrica. La sfida principale in un tale sistema non lineare con un alto livello di variabilità è stata il dimensionamento ottimale del SWDP con l'intero solare proposto, pur mantenendo buone prestazioni dinamiche. Nell'articolo viene presentata una solida analisi del carico elettrico per valutare la progettazione ottimale di SWDP alimentata dal sistema fotovoltaico. Inoltre, sono stati sviluppati regolatori di inseguimento del punto di massima potenza (MPPTC) per migliorare le prestazioni dinamiche dell'SWDP alimentato dal sistema fotovoltaico. Per realizzare questo studio, è stato implementato un vero e proprio impianto di desalinizzazione dell'acqua di mare connesso alla rete situato in Egitto. L'SWDP selezionato produce 700 m³/giorno. I dati sperimentali reali dell'impianto sono stati estratti attraverso letture giornaliere dei contatori di consumo di energia elettrica e di produzione di acqua. Questi dati sperimentali sono stati quindi introdotti nel programma HOMER per suggerire i componenti ottimali del sistema fotovoltaico in base al costo attuale netto minimo. Inoltre, per affrontare la sfida della bassa efficienza di conversione del sistema fotovoltaico, sono stati studiati tre MPPTC per migliorare le prestazioni dinamiche dell'SWDP proposto alimentato dal sistema fotovoltaico. La conduttanza incrementale in combinazione con tre tecniche di ottimizzazione dell'intelligenza artificiale (AI) è stata sviluppata per la valutazione delle prestazioni dinamiche dell'SWDP presentato alimentato da PV. Le tecniche utilizzate riguardano gli algoritmi di **Particle Swarm Optimization (PSO)**, **Grey Wolf Optimization (GWO)** e **Harris Hawks Optimization (HHO)**. Il sistema è stato costruito, modellato e simulato tramite MATLAB/SIMULINK.

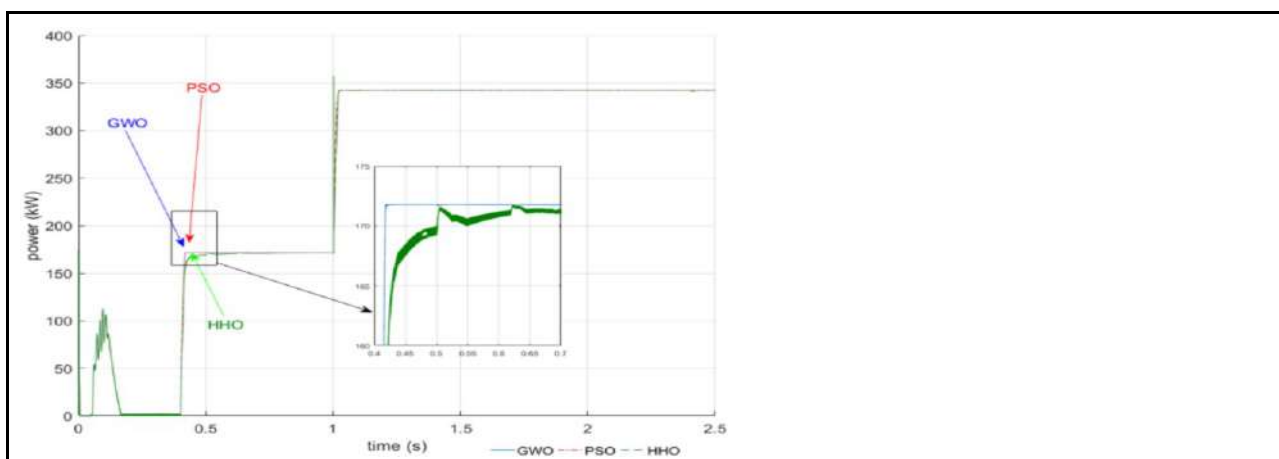


Figura 61: Profilo di potenza PV dopo aver utilizzato l'ottimizzazione (AI).

I risultati raggiunti dai tre metodi sono stati promettenti nell'estrarre la massima potenza con il minimo errore dal sistema fotovoltaico, migliorando al contempo le prestazioni di SWDP. Le tecniche sono state confrontate e, dopo l'utilizzo delle stesse, il risultato è risultato instabile nel caso dell'utilizzo della tecnica **Hawks Optimization (HHO)**. Quando si è utilizzato il regolatore basato su **Particle Swarm Optimization (PSO)** si è ottenuta una risposta dinamica moderata. Le migliori prestazioni si sono ottenute applicando la tecnica **Grey Wolf Optimization (GWO)**. La tecnica GWO ha raggiunto le migliori prestazioni in termini di potenza estratta. In conclusione, la simulazione ottenuta (così come i risultati sperimentali) hanno dimostrato l'efficacia della strategia di progettazione ottimale suggerita. Il metodo proposto può essere applicato su larga scala rispetto al sistema in studio aumentando la taglia del campo fotovoltaico adatto ai carichi.

2.2.2 Fonti non rinnovabili

Gli studiosi cercano di riciclare l'energia sprecata per produrre elettricità integrando generatori termoelettrici (TEG) con motori a combustione interna (ICE), che si basano sulla conducibilità elettrica, β , delle strisce conduttrici termiche.

L'obiettivo principale dell'articolo in [38] è stato quello di tracciare il trasferimento di calore, che è dovuto alla differenza di temperatura tra le guarnizioni calde e fredde attraverso i conduttori termici che rappresentano le gambe del TEG, che devono essere installate in un percorso sicuro sulla superficie della guarnizione senza ostacoli che impediscono il trasferimento di calore (ad es. crepe).

Le gambe TEG sono in lega di ferro, alluminio e rame a forma di nastro con caratteristiche specifiche che garantiscono la massima trasformazione termoelettrica, che ha oscillato tra una distribuzione uniforme, gaussiana ed esponenziale a seconda della struttura della lega. I cancelli di scarico e aspirazione ICE sono stati scelti come lati TEG. Il modello gemello del simulatore digitale controlla l'efficienza dell'integrazione attraverso due fasi sequenziali, iniziando con la registrazione delle cause del guasto della conduttività termica tramite riprese ed estraendo i loro dati mediante **procedure di neural network** nell'alimentazione del secondo stadio, che rivelano che le crepe sono un ostacolo importante nella riduzione della potenza generata dal TEG. Pertanto, l'interesse della seconda fase è prevedere le posizioni delle crepe, $P_{i,j}$, e la loro intensità, Q_P , sulla base dell'ant colony algorithm che recluta i dati di imaging (STTF-NN-ACO) per installare i conduttori termici lontano dalle posizioni delle fessure.

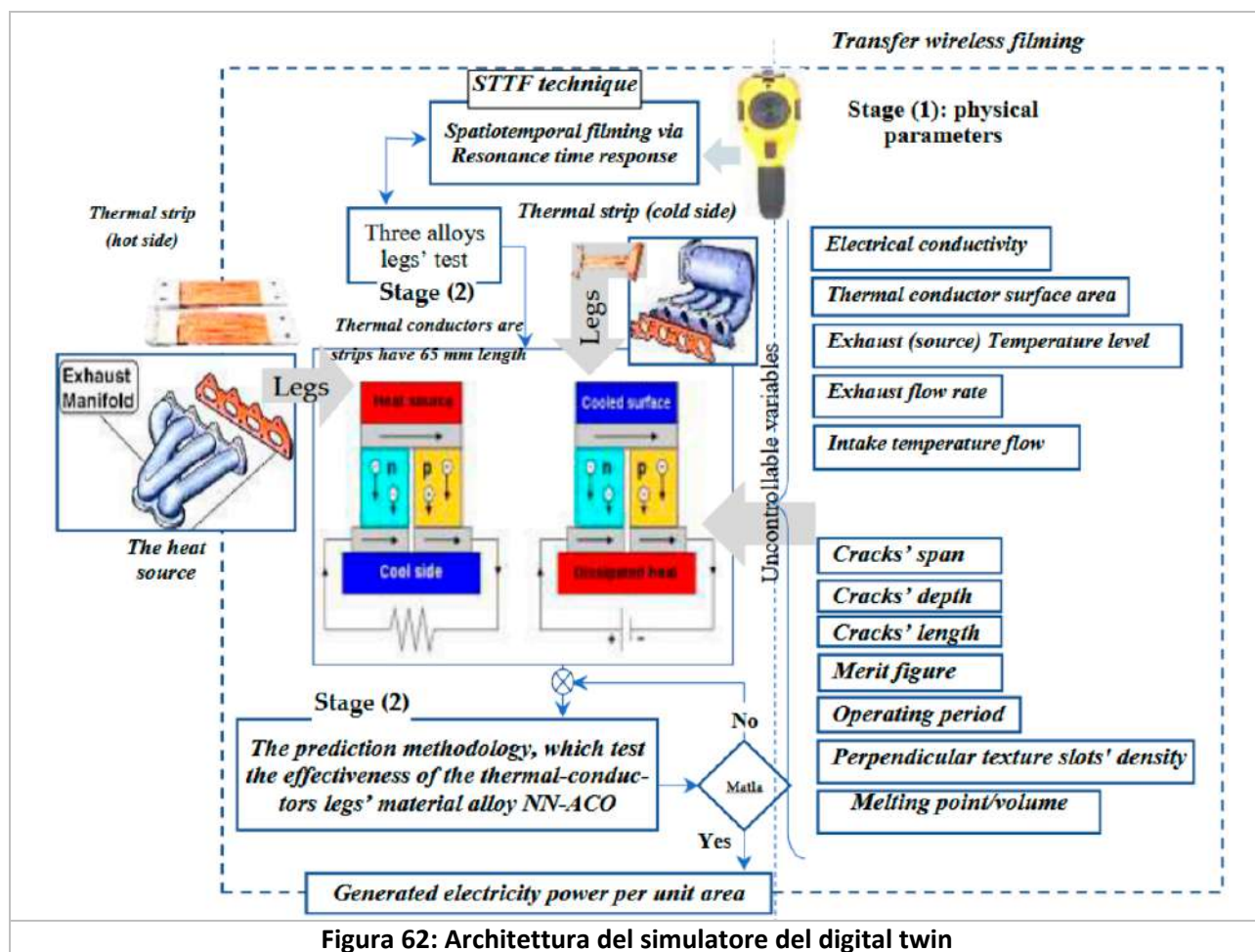


Figura 62: Architettura del simulatore del digital twin

La verifica metaeuristica proposta (STTF-NN-ACO) ha mostrato una superiorità nella previsione su [Mat-ACO] dell'8,2% e ha aumentato l'efficienza dei TEG del 32,21%. Inoltre, l'aumento della potenza totale generata del 12,15% e delle ore di lavoro del TEG del 20,39%, riflette una riduzione del consumo di carburante fino al 19,63%.

Tutte le equazioni sono state formulate per servire gli obiettivi attraverso il simulatore digital twin che verifica l'efficienza dell'ibridazione di TEG e un altro generatore ICE. L'efficienza della previsione per le posizioni delle crepe e la loro intensità si è riversata positivamente sul costo del consumo di carburante a causa della dipendenza del veicolo dall'elettricità generata. L'interferenza dell'estrazione dei dati utilizzando la rete neurale ha consentito all'addestramento di previsione di determinare con precisione il percorso migliore per l'installazione delle strisce di conducibilità termica.

Nello studio presentato nell'articolo [39] è stata integrata una scala appropriata di cella a combustibile a ossido solido con una GT e sono state condotte analisi comparative complete basate su criteri energetici, exergetici, exergoeconomici ed ambientali (4E) tra un sistema GT da 10 MW ed un sistema ibrido SOFC-GT. Sono stati considerati tre diversi scenari per l'analisi comparativa, che comprendevano l'operazione di base, il comportamento parametrico e le condizioni ottimali. È stata utilizzata un'adeguata **Artificial Neural Network (ANN)** per modellare i sistemi e poi le ANNs ottenute sono state introdotte nell'**ottimizzatore Grey wolf**. Utilizzando l'ottimizzatore Grey wolf, è stato risolto un problema di ottimizzazione multi-obiettivo considerando l'efficienza exergetica, le emissioni di CO₂ e il costo unitario del prodotto come obiettivi.

	Parameter	Value
GT and SOFC-GT	Net power	10 MW
	Pressure ratio of compressor	8
	Effectiveness of A.H.E	0.8
	Combustor efficiency	0.98
	Gas turbine isentropic efficiency	0.8
	power turbine isentropic efficiency	0.8
	Turbine inlet temperature (of G.T)	974.5 °C
	Generator efficiency	0.98
	Effectiveness of F.H.E	0.95
	Effectiveness of W.H.E	0.95
	Isentropic efficiency of water pump	0.89
	Isentropic efficiency of fuel compressor	0.89
	Lower heating value of methane	50 050 kJ/kg
Pressure drops	A.H.E	2%
	F.H.E	2%
	W.H.E	2%
	A.B or C.C	3%
	SOFC	2%
Ambition condition	Temperature	25 °C
	Pressure	101.3 kPa
SOFC	Density of output current	6000 A/m ²
	Anode exchange current density	6500 A/m ²
	Cathode exchange current density	2500 A/m ²
	Fuel utilization factor	0.85
	Steam to carbon ratio	2.5
	Cell area	10 ⁻² m ²
	Thickness of anode	5 × 10 ⁻⁴ m
	Thickness of cathode	5 × 10 ⁻⁵ m
	Thickness of electrolyte	1 × 10 ⁻⁵ m
	Thickness of interconnect	3 × 10 ⁻³ m
Economic data	Effective gaseous diffusivity inside anode	2 × 10 ⁻⁵ m ² /s
	Effective gaseous diffusivity inside cathode	5 × 10 ⁻⁶ m ² /s
	Interest rate	12%
	Life time of system	20 years
	Annual operation hours of system	8000 h
	Cost of fuel per energy unit	14 \$/GJ

Figura 63: Dati di input di base per la simulazione di SOFC-GT e GT

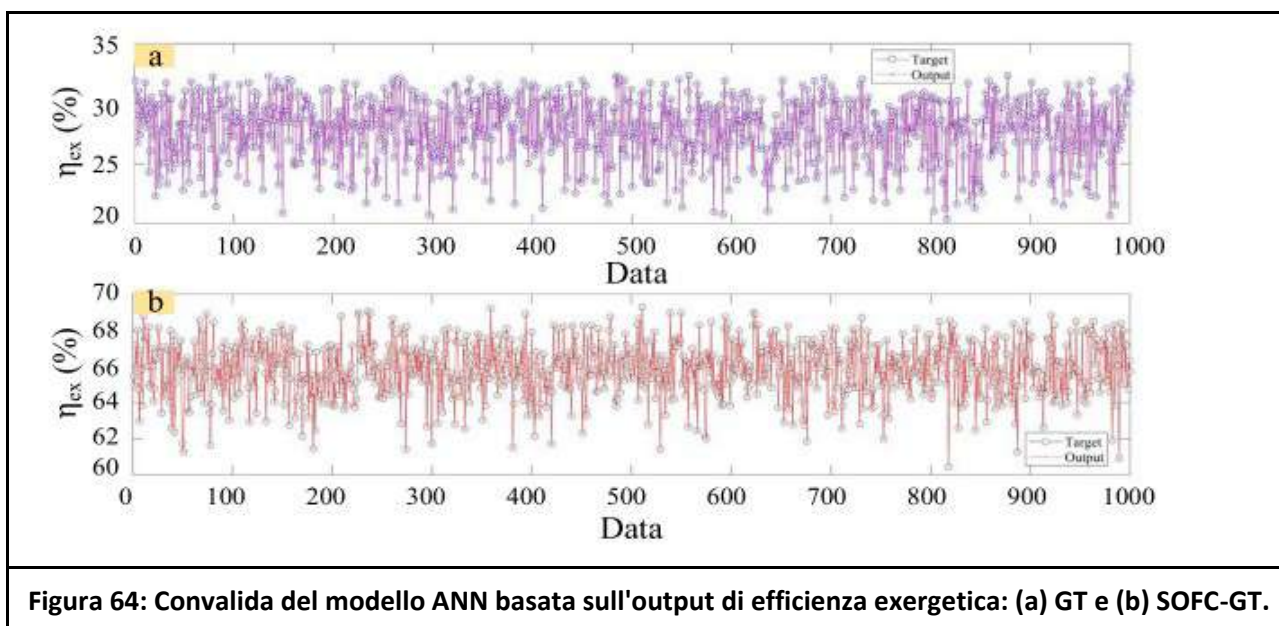


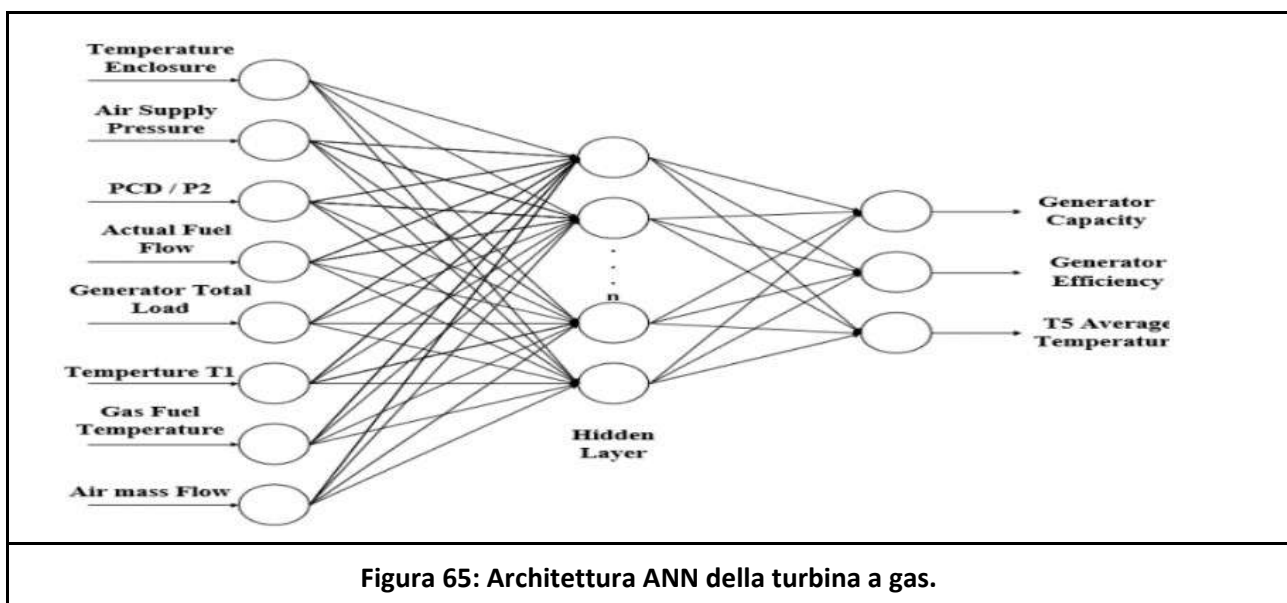
Figura 64: Convalida del modello ANN basata sull'output di efficienza exergetica: (a) GT e (b) SOFC-GT.

Dopo la modellazione tramite Artificial Neural Network (ANN), il sistema ibrido SOFC-GT ha avuto una migliore performance rispetto al sistema GT in termini di performance ambientali ed exergoeconomiche come hanno mostrato i risultati di tutti e tre gli scenari. Tuttavia, questa integrazione ha leggermente aumentato il costo di investimento del sistema nello scenario di base. Dopo l'ottimizzazione, al miglior compromesso tra gli obiettivi, l'efficienza exergetica SOFC-GT, le emissioni di CO₂ ed il costo del prodotto unitario sono state ottenute rispettivamente al 68,5%, 277,7 kg/MWh e 22,7 \$/GJ, mentre per la GT erano rispettivamente del 32,38%, 588,4 kg/MWh e 26,9 \$/GJ.

Si cerca di riciclare l'energia sprecata per produrre elettricità integrando generatori termoelettrici (TEG) con motori a combustione interna (ICE), che si basano sulla conducibilità elettrica, β , delle strisce conduttrici termiche.

Nell'articolo [40] la funzione obiettivo è la massimizzazione dell'efficienza della turbina a gas, con alcune limitazioni operative come vincoli, manipolando il rapporto aria/combustibile. Il modello è stato sviluppato utilizzando un **Artificial Neural Network (ANN)**, e un **Genetic Algorithm (GA)** è stato selezionato come tecnica di ottimizzazione stocastica per risolvere il problema. Il modello di

rete neurale è stato creato direttamente utilizzando i dati operativi di un generatore di turbina a gas parametro reale. I dati necessari per la modellizzazione basata su ANN sono circa 8150 set di dati che verranno utilizzati per addestrare e convalidare il modello ANN. I set di dati variabili sono stati divisi in due parti; per scopi di addestramento è del 87,5% e per la convalida è del 12,5%. La gestione dei pesi per le reti neurali è stata effettuata utilizzando l'algoritmo di Levenberg-Marquardt che potrebbe dare buoni risultati con $RMSE = 7,3 \times 10^{-3}$.



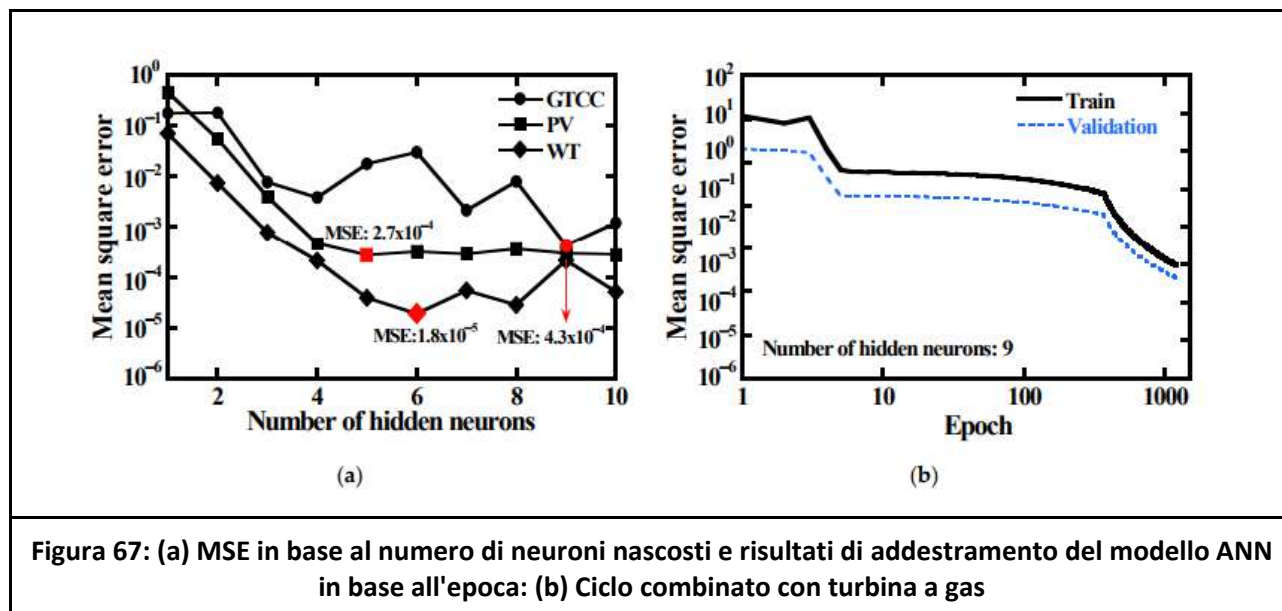
Before Optimization	Air Supply Pressure	PCD (P2)	Actual Fuel Flow	Generator Total kW	T1 Temperature	Gas Fuel Temp	Turbine Air Inlet DP	Air mass flow	Generator Capacity	Generator Efficiency	T5 Average Temperature	AFR
MIN	6.67	6.68	0.0775	283	14.9	35.2	11.3	0.670934	2856	0.8925	332	6.51041
MAX	7.39	7.53	0.242222	1673	33.2	62.1	12.9	2.447537	3127	0.9771875	488	19.08412
After Optimization	Air Supply Pressure	PCD (P2)	Actual Fuel Flow	Generator Total kW	T1 Temperature	Gas Fuel Temp	Turbine Air Inlet DP	Air mass flow	Generator Capacity	Generator Efficiency	T5 Average Temperature	AFR
A wide range of operating data	6.8	6	0.4	1341	16.6	30.2	10.5	1.9	3.20E+03	1	440.4008	4.6885
Constraint Variable Output T5 range 680 - 740	7.1	7.4	0.1	2377.2	31	27.1	12.5	2.7	3.20E+03	1	695.015	24.7279
Constraint Variable Output Eff: 0.9 - 1, Generator Total KW: 2000 - 2800 KW and AFR: 14-20	6.9	7	0.2	2567.3	20	54.9	11.3	2.2	3193.3	0.9916	409.9	19.8

Figura 66: Valori ottimali dopo l'ottimizzazione

Sulla base dei risultati e dell'analisi dei dati dello studio che è stato condotto, le conclusioni ottenute sono le seguenti: (a) i risultati della simulazione del modello del Generatore di Turbine a Gas rappresentano già le condizioni reali in campo; (b) l'applicazione dei controlli basati sull'Intelligenza Artificiale su un sistema di generatore di turbine a gas è in grado di produrre un'immagine dell'input-output di un sistema GTG; (c) l'applicazione del GAO al sistema di generatore di turbine a gas, quando implementata, sarà in grado di risparmiare costi operativi fino a 280,8 kg/h.

Il documento [41] ha adottato un metodo di ottimizzazione che utilizza uno schema di intelligenza artificiale che combina un **Artificial Neural Network (ANN)** e un **Genetic Algorithm (GA)**. In particolare, l'ANN è stato costruito sulla base dei risultati di un modello basato sulla fisica per ottenere una grande quantità di dati di simulazione precisi in un breve periodo di tempo. Un sistema di generazione distribuita (DG) basato su una turbina a gas (GT) è stato simulato per dimostrare l'utilità del metodo di ottimizzazione. Tenendo conto della capacità e delle prestazioni a carico parziale della GT, l'ottimizzazione della dimensione e della strategia di funzionamento del sistema DG è stata eseguita per tre scenari di progettazione del sistema. I criteri di ottimizzazione erano l'efficienza economica e l'ecosostenibilità. Il metodo ha ridotto il tempo di calcolo di oltre tre ordini

di grandezza mantenendo la stessa precisione del modello basato sulla fisica. L'approccio e la metodologia, si prevede siano applicabili all'ottimizzazione accurata e veloce di vari sistemi energetici sofisticati.



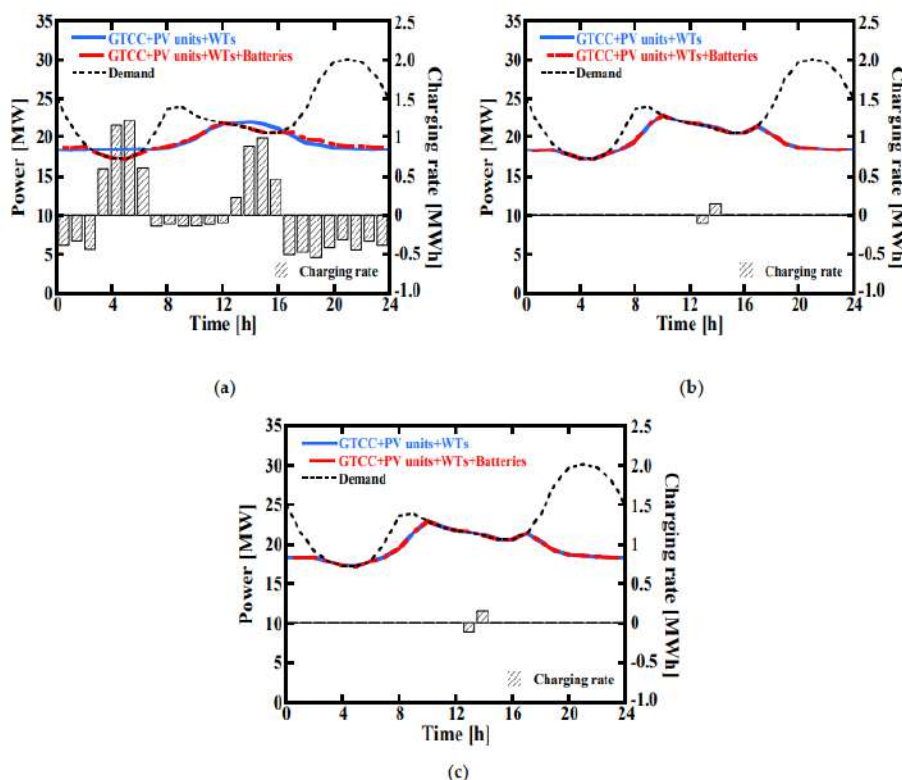


Figura 68: Risultati della strategia operativa secondo le modalità operative della turbina a gas: (a) Caso 1; (b) Caso 2; (c) Caso 3.

Per semplificare il processo di calcolo del complesso sistema DG, è stato costruito un modello ANN. Il modello ANN attuale, basato sul modello fisico, a differenza del modello ANN convenzionale basato sui dati misurati, non richiedeva un dataset di test per l'overfitting. Di conseguenza, una grande proporzione del dataset è stata utilizzata per addestrare il modello ANN, quindi il modello ANN ha imitato molto bene il modello fisico: l'MSE aveva un massimo di $2,7 \times 10^{-4}$ e un minimo di $1,8 \times 10^{-5}$. Di conseguenza, il modello ANN ha mostrato un miglioramento di almeno 5200 volte e al massimo 22.300 volte per il tempo di calcolo e ha ridotto la memoria richiesta per il calcolo fino al 62,3%. Pertanto, il modello ANN è adatto per l'utilizzo nel calcolo di ottimizzazione del sistema DG.

Nell'articolo [42] è stato sviluppato un modello basato sull'apprendimento automatico per prevedere l'efficienza netta ottimizzata. È stato sviluppato un modello matematico di impianto di potenza a microturbina a gas (MGTPPMM) che poteva generare una potenza elettrica di 6 kW, il quale è stato convalidato utilizzando le condizioni operative di Jet Cat PHT3. I dati risultanti sono stati utilizzati per ottimizzare e prevedere l'efficienza netta delle turbine a gas e anche le caratteristiche del compressore e della turbina utilizzando un modello di algoritmo di apprendimento automatico (MLAM). Il MGTPPMM ha un'accuratezza del 92% nella simulazione di varie condizioni operative, mentre sono stati sviluppati due MLAM Predictor 1 **Linear Regression (LR)** e Predictor 2 **Random Forest (RF) Decision Tree (DT)**. Il Predictor 1 ha un'accuratezza del 76% e il Predictor 2 ha un'accuratezza inferiore all'1% di errore alle condizioni operative del livello del mare. Tuttavia, a quote più elevate di 17 km, l'accuratezza del MGTPPMM è dell'88%, mentre quella del MLAM - Predictor 2 ha un errore inferiore al 6% a causa della mancanza di dati validati sufficienti per i set di addestramento. I modelli proposti potrebbero aiutare a progettare microturbine a gas più efficienti a livello di componente e potrebbero anche prevedere le caratteristiche delle prestazioni per monitorare la salute del sistema virtualmente mediante l'implementazione di un sistema di monitoraggio del degrado.

Parameters	Values
The coefficient of recovery of the total pressure at the inlet	0.99
Compressor type	Centrifugal
Relative current density in the cross-section at the inlet to the compressor	0.75
The base value of the polytropic efficiency of the compressor	0.86
Overall pressure ratio of the compressor	3
Stagnation temperature at turbine inlet	1500 K
Fuel injection temperature	293.15 K
Composition and properties of fuel	Kerosene T I
Combustion efficiency ratio	0.98
Turbine Type	Axial
The basic value of the isentropic efficiency of the compressor loaded turbine	0.91
The basic value of the isentropic efficiency of the compressor free turbine.	0.91
Excess power capacity on the rotor shaft for generator output.	8kW
The design value of the available pressure reduction at the output	1.05
Rate factor (diffuser)	0.5

Figura 69: Parametri di input di progetto in MGTPPMM

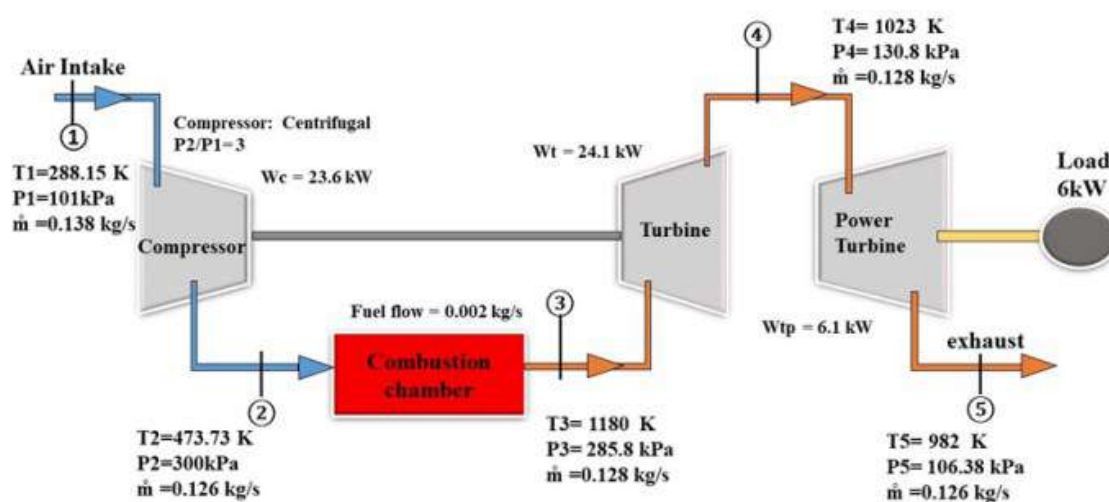
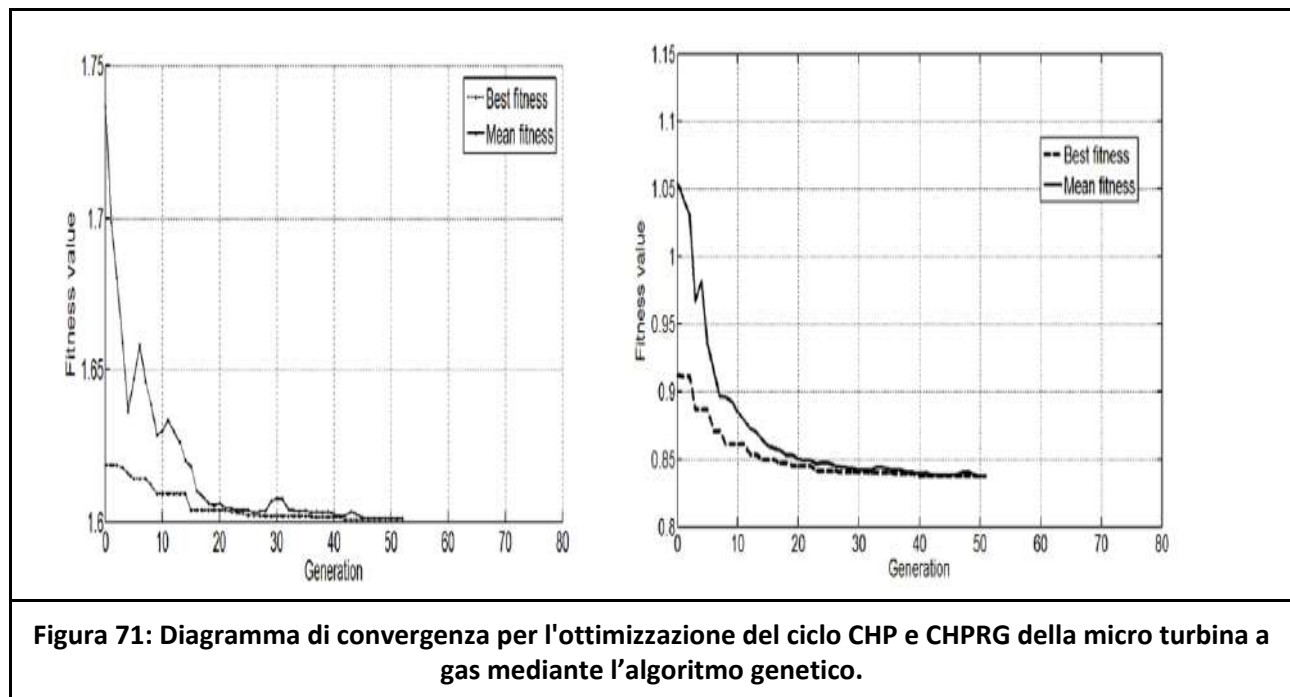


Figura 70: Risultato finale del ciclo ottimizzato per un sistema da 6 kW al livello del mare

Il paper scientifico [43] riguarda l'ottimizzazione di una microturbina a gas alimentata a gas naturale in quattro configurazioni: semplice, con rigeneratore, produzione combinata di calore ed energia (CHP) e CHP con rigeneratore. La funzione obiettivo considera l'efficienza energetica ed exergetica e il costo dell'elettricità, tenendo conto degli effetti dell'inquinamento ambientale. Sono stati utilizzati tre algoritmi di ottimizzazione: **Genetic Algorithm, Bee Colony Algorithm e Searching Algorithm.**



Cycle mode	r_a	r_c	$\eta_i(\%)$	$\eta_c(\%)$	$\eta_{th}(\%)$	$\eta_{II}(\%)$	$s_{gen}(kW/K)$	CE(US\$/kWh)
Simple mode	1.40	6.4	87	80	29.01	61.38	0.0583	0.1024
RG mode	1.21	6.5	80	85	34.39	74.86	0.0369	0.0864
CHP mode	1.42	6.5	81	84	39.42	85.42	0.0529	0.0754
CHPRG mode	1.82	6.5	85	82	38.22	57.25	0.0652	0.0286

Figura 72: Risultati della modalità ottimale mediante algoritmo genetico.

I risultati mostrano che il rapporto aria/combustibile ottimale calcolato dal metodo di ricerca per la micro turbina a gas con i cicli sopra menzionati è stato rispettivamente di 1,7, 1,3, 1,6 e 2,3.

Applicando il **Genetic Algorithm**, i rapporti aria/combustibile ottimali sono stati rispettivamente di 1,40, 1,21, 1,42 e 1,82. In questi punti, l'efficienza energetica ottenuta è stata del 29, 34,4, 39,4 e 38,2%, l'efficienza di seconda legge ottenuta è stata del 61,4, 74,9, 85,4 e 57,2%, e il costo dell'elettricità ottenuto è stato rispettivamente di 0,102, 0,086, 0,075 e 0,029 US\$/kWh. Utilizzando il **Bee Colony Algorithm**, i rapporti aria/combustibile ottimali sono stati rispettivamente di 1,36, 1,13, 1,32 e 1,61. In questi punti, l'efficienza energetica è stata del 28,2, 34,3, 40,5 e 35,6%, l'efficienza di seconda legge è stata del 60,9, 75,67, 81,80 e 56,65%, e il costo dell'elettricità è stato rispettivamente di 0,105, 0,087, 0,073 e 0,03 US\$/kWh. Tra questi metodi, il **Genetic Algorithm (GA)** è stato selezionato come il miglior metodo in termini di miglior risultato di ottimizzazione.

L'articolo [44] propone un nuovo approccio di diagnosi del gas-path di una turbina a gas, basato su intelligenza artificiale basata sui dati di conoscenza. L'obiettivo è garantire che le turbine a gas funzionino in modo sicuro, affidabile, ecologico ed efficiente. Il metodo proposto è un'ibridazione di deep learning e analisi del gas-path.

In primo luogo, viene costruito un modello termodinamico della turbina a gas da diagnosticare attraverso una strategia di modellizzazione adattativa. In secondo luogo, viene simulato un grande numero di dati di conoscenza corrispondenti ai parametri di salute dei componenti e ai parametri di condizione di confine della turbina a gas e ai parametri misurabili del gas-path. Infine, viene progettato un modello di deep learning per la modellizzazione della regressione di questo database di conoscenza, utilizzando il vettore di input composto dai parametri di condizione di confine e dai parametri misurabili del gas-path e il vettore di output dei parametri di salute dei componenti. Durante il funzionamento della turbina a gas, il modello addestrato restituisce i parametri di salute dei componenti in tempo reale. Per la modellazione, vengono utilizzati degli algoritmi di intelligenza artificiale, tra cui il **Deep Neural Network (DNN)**.

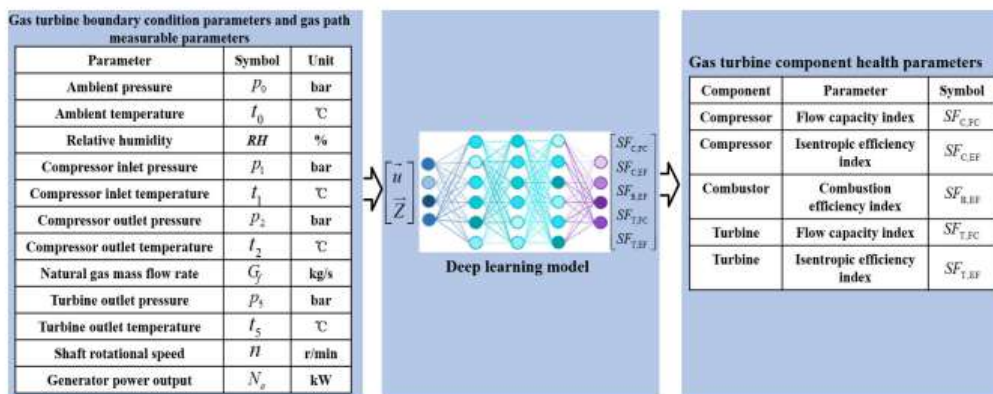


Figura 73: Progetto di un modello di deep learning per la modellazione della regressione di questo knowledge database

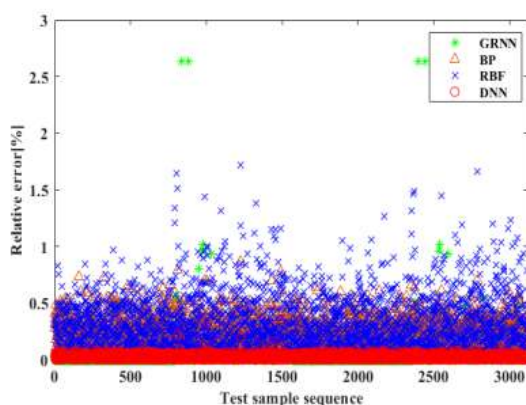
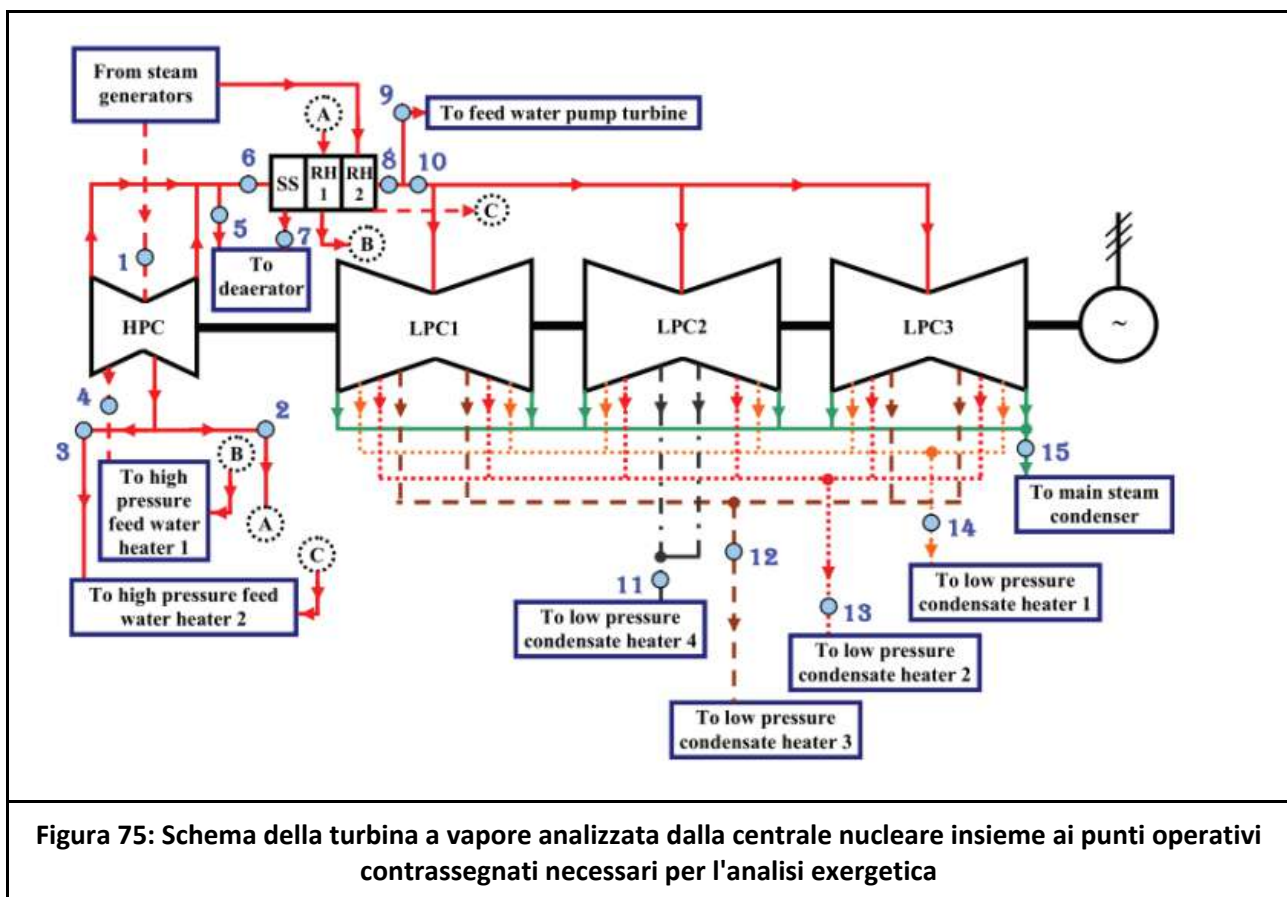


Figura 74: Gli errori relativi dell'indice di capacità di flusso della turbina

I risultati della simulazione sperimentale mostrano che il metodo proposto può fornire parametri di salute accurati e quantificati per ciascun componente del gas-path, con un errore quadratico medio complessivo inferiore al 0,033% e un errore relativo massimo inferiore allo 0,36%, dimostrando un grande potenziale di applicazione.

Nell'articolo [45] si presenta un'analisi exergetica dell'intera turbina, dei cilindri della turbina e delle parti dei cilindri in quattro diversi regimi di funzionamento. La turbina analizzata opera in una

centrale nucleare mentre tre dei quattro regimi di funzionamento sono ottenuti utilizzando algoritmi di ottimizzazione: **Simplex Algorithm (SA)**, **Genetic Algorithm (GA)** e **Improved Genetic-Simplex Algorithm (IGSA)**. Il regime di funzionamento IGSA fornisce la più alta potenza meccanica sviluppata dall'intera turbina pari a 1022,48 MW, seguito da GA (1020,06 MW) e SA (1017,16 MW), mentre nel regime di funzionamento originale l'intera turbina sviluppa una potenza meccanica pari a 996,29 MW. Inoltre, IGSA provoca l'aumento più elevato della potenza meccanica sviluppata in quasi tutti i cilindri e le parti dei cilindri rispetto al regime di funzionamento originale.



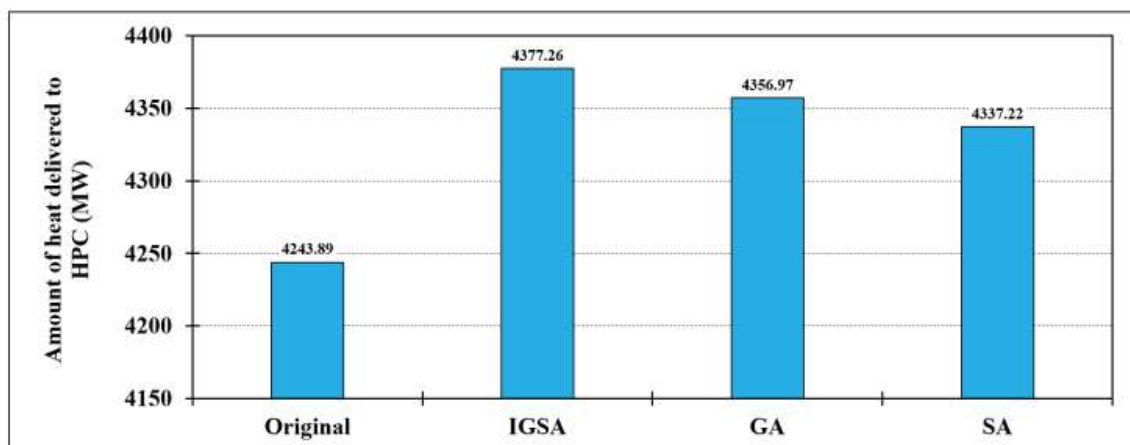


Figura 76: Quantità di calore fornita all'HPC – Regime operativo originario e regimi operativi ottenuti da tre algoritmi di ottimizzazione (IGSA, GA e SA)

Tutti gli algoritmi di ottimizzazione osservati aumentano la distruzione exergetica dell'intera turbina rispetto al regime di funzionamento originale: il minor aumento è causato da IGSA, seguito da GA e infine da SA. L'efficienza exergetica più elevata dell'intera turbina, pari all'85,92%, è ottenuta da IGSA, seguita da GA (85,89%) e SA (85,82%), mentre la più bassa efficienza exergetica è ottenuta nel regime di funzionamento originale (85,70%). La turbina analizzata, che utilizza vapore umido, è poco influenzata dalla variazione della temperatura ambiente. IGSA, che mostra una prestazione dominante nei parametri di analisi exergetica della turbina analizzata, in alcune situazioni viene superato da GA. Pertanto, nell'ottimizzazione delle prestazioni della turbina a vapore, IGSA e GA possono essere consigliati.

Nel documento [46] si propone un framework che fa uso della tecnologia IIOT per raccogliere i dati in modo affidabile e costruire modelli basati su tecniche matematiche e di apprendimento automatico per prevedere il tempo e la quantità ottimale di carburante. Il framework proposto è stato implementato e testato con dataset reali. Nel paper vengono utilizzati diversi algoritmi di machine learning, tra cui il **Random Forest (RF) Regressor**, l'**eXtreme Gradient Boosting (XGBoost)** e il **Long Short-Term Memory Neural Network (LSTM)**.

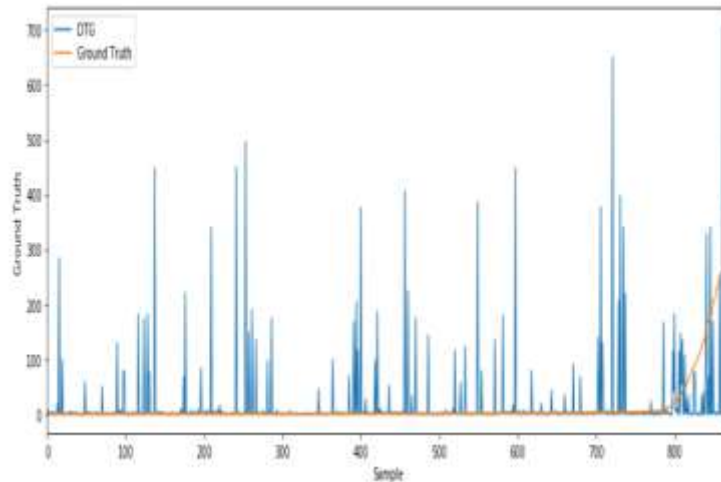


Figura 77: La misura delle prestazioni del Decision Tree Regressor

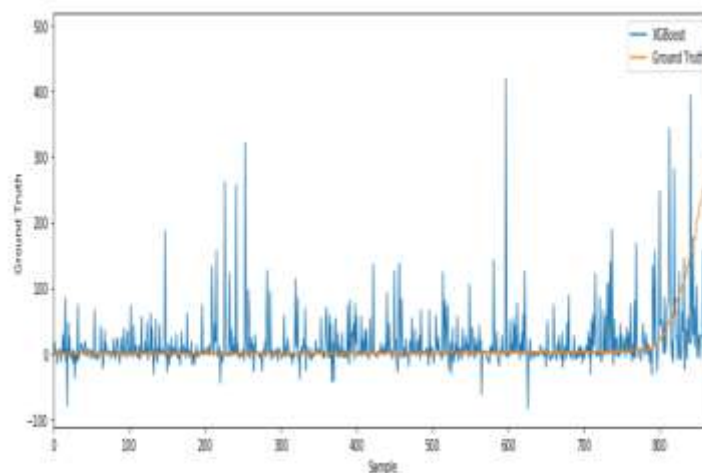


Figura 78: La misura delle prestazioni del Random Forest

Si è osservato che le prestazioni di Random Forest (RF) Regressor e dell'eXtreme Gradient Boosting (XGBoost) non sono abbastanza buone, a causa del fatto che il loro MAE e MSE non sono abbastanza vicini allo zero. Per ottenere decisioni di qualità, è importante costruire un modello con prestazioni previste più elevate. Per questa ragione, i dati vengono trattati come dati di serie temporali e utilizzano una rete neurale **Long Short Term Memory (LSTM)** per prevedere il consumo di carburante.

2.2.3 Accumulatori

In [47] viene proposta la struttura di un digital twin di una batteria per un veicolo elettronico, che ha il potenziale per migliorare notevolmente la consapevolezza situazionale del BMS e consentire il funzionamento ottimale delle unità di accumulo della batteria. Viene proposto un Digital Twin (DT) come soluzione alla difficoltà del calcolo integrato per lo stato di salute (SOH) e lo stato di carica (SOC) incrementali utilizzando il modello **Extreme Gradient Boost (XGBoost)** e il **filtro di Kalman esteso (EKF)** per prevedere la stima dello stato della batteria EV. La condizione della batteria è stata determinata utilizzando l'EKF, che può fornire informazioni vitali per la manutenzione. Innanzitutto, la durata utile della batteria può essere estesa con una stima accurata del SOC; quindi viene suggerito un approccio di previsione basato sull'apprendimento per misurare lo stato di salute della batteria al fine di aumentare la durata della batteria.

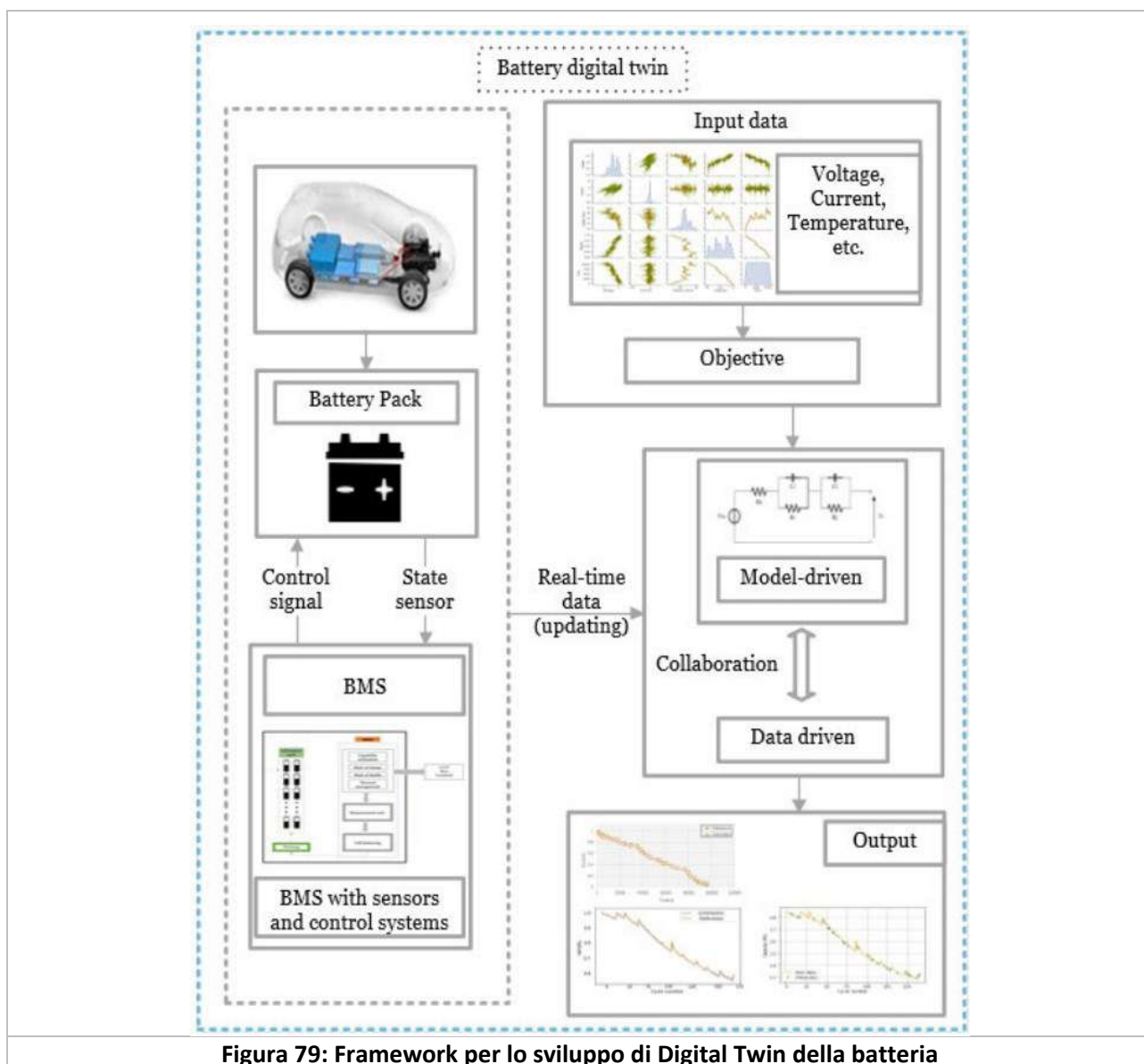


Figura 79: Framework per lo sviluppo di Digital Twin della batteria

Combinando le capacità di filtraggio dell'algoritmo EKF con l'eccezionale capacità di previsione della regressione appropriata non lineare dell'algoritmo XGBoost, l'algoritmo suggerito può avere la stabilità degli errori di stima nell'intera gamma di errori di stima, migliorando al contempo l'accuratezza del filtraggio. I risultati della simulazione hanno dimostrato che l'algoritmo XGBoost basato su EKF offre prestazioni migliori nella previsione. Oltre a fornire l'affidabilità e l'accuratezza della stima dello stato, l'algoritmo proposto presenta chiari vantaggi in termini di generalità e robustezza.

La previsione accurata della durata e la valutazione dell'affidabilità delle batterie agli ioni di litio sono di grande importanza per la manutenzione predittiva. Nell'intero ciclo di vita di una batteria, la descrizione accurata delle caratteristiche dinamiche e stocastiche della vita è sempre stata un problema fondamentale. In [48] viene proposto un gemello digitale per l'affidabilità basato sulla previsione del ciclo di vita utile rimanente per le batterie agli ioni di litio. Il modello di degradazione della capacità, il modello di degradazione stocastico, la previsione della durata e il modello di valutazione dell'affidabilità sono stabiliti per descrivere la casualità del degrado della batteria e la dispersione della durata di più celle. Sulla base **dell'algoritmo bayesiano**, viene proposto un metodo di evoluzione adattivo per il modello del digital twin per migliorare l'accuratezza della previsione, seguito dalla verifica sperimentale. Infine, vengono implementate la previsione della vita, la valutazione dell'affidabilità e la manutenzione predittiva della batteria basata sul digital twin.

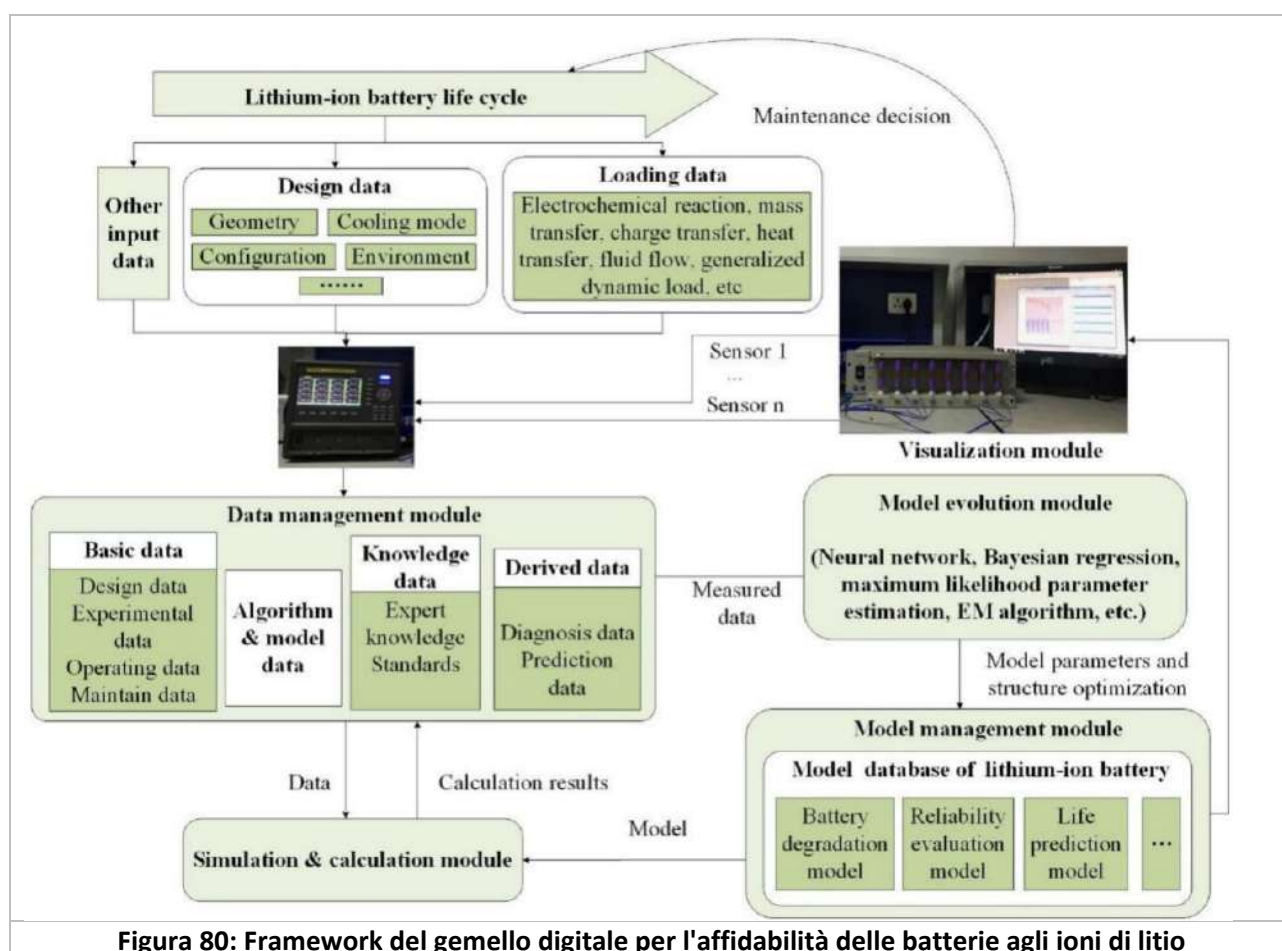


Figura 80: Framework del gemello digitale per l'affidabilità delle batterie agli ioni di litio

Gli esperimenti di degradazione di diverse celle sono stati effettuati per verificare l'accuratezza dei modelli. Infine, è stata condotta la previsione di vita e la valutazione dell'affidabilità basata sui Digital Twin, seguita dall'analisi del ciclo di evoluzione e dalla decisione di manutenzione predittiva. I risultati hanno mostrato che il gemello digitale per l'affidabilità ha una buona precisione nell'intero ciclo di vita. L'errore può essere controllato a circa il 5% con l'algoritmo di evoluzione adattiva. Per le batterie L1 e L6 in questo caso, i costi di manutenzione predittiva dovrebbero diminuire rispettivamente del 62,0% e del 52,5%.

Nell'articolo [49] viene proposta una nuova classificazione degli indicatori di salute (Health Index HI) in cui questi sono divisi in variabili misurate e variabili calcolate. Per illustrare l'importanza della preparazione dei dati, vengono valutati quattro metodi di riduzione del rumore nel processo di estrazione degli HI; vengono applicati diversi metodi di selezione delle caratteristiche, tra cui il metodo basato su filtro, il metodo basato su wrapper e il metodo basato sulla fusione, per selezionare i sottoinsiemi degli HI. I metodi utilizzati sono il **Support Vector Machine (SVM)**, **l'Artificial Neural Network (ANN)**, **il Gaussian Process Regression (GPR)** e **il Random Forest (RF)**. Al fine di valutare le prestazioni di stima in potenziali utilizzi reali nell'era dei big data futuri, i tre metodi di selezione degli HI e i quattro metodi di apprendimento automatico vengono valutati utilizzando tre set di dati pubblici e due strategie di stima.

Battery	Index	ANN	SVM	RVM	GPR
B0005	RMSE	1.12	0.60	1.64	0.55
	MAE	3.15	1.23	3.01	2.07
B0006	RMSE	2.71	1.98	7.36	1.27
	MAE	6.57	6.04	13.13	3.55
B0007	RMSE	1.76	1.05	3.16	0.39
	MAE	3.65	2.25	4.69	1.03
B0018	RMSE	2.12	1.40	8.41	1.63
	MAE	6.24	2.91	15.32	5.14
Time	(ms)	256.31	13.43	12.21	10.48

Figura 81: Errori stimati (%) e tempo di calcolo (ms) di diversi algoritmi.

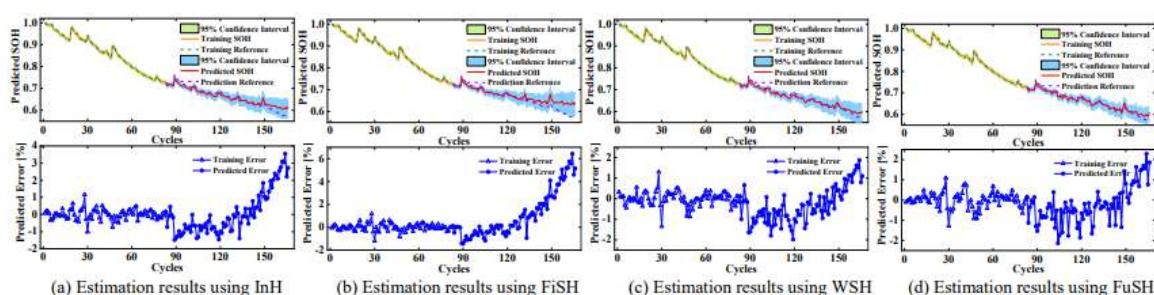


Figura 82: Risultati della stima SOH di B0006 utilizzando il 50% di dati per l'addestramento del modello

I risultati mostrano che la combinazione del metodo di selezione basato sulla fusione e il **Gaussian Process Regression (GPR)** ha prestazioni di stima complessivamente superiori in termini di accuratezza ed efficienza computazionale.

Nell'articolo [50] viene presentato un metodo dettagliato per effettuare previsioni a lungo termine per la RUL (Remaining-Useful-Life) per le batterie al litio basato su una **combinazione di Gated Recurrent Unit Neural Network (GRU NN) e un metodo di soft-sensing**. In primo luogo, è stato estratto un indicatore di salute (Health index HI) indiretto dai processi di carica utilizzando un metodo di soft-sensing che può descrivere con precisione la degradazione della potenza invece della capacità. Successivamente, è stata applicata una **GRU NN** con una finestra scorrevole per

apprendere lo sviluppo delle prestazioni a lungo termine. Il metodo utilizza anche un drop-out e un metodo di early stopping per prevenire l'overfitting. Per costruire i modelli e convalidare l'efficacia del metodo proposto, è stato utilizzato un set di dati sulle batterie della NASA del mondo reale con diverse misurazioni della batteria.

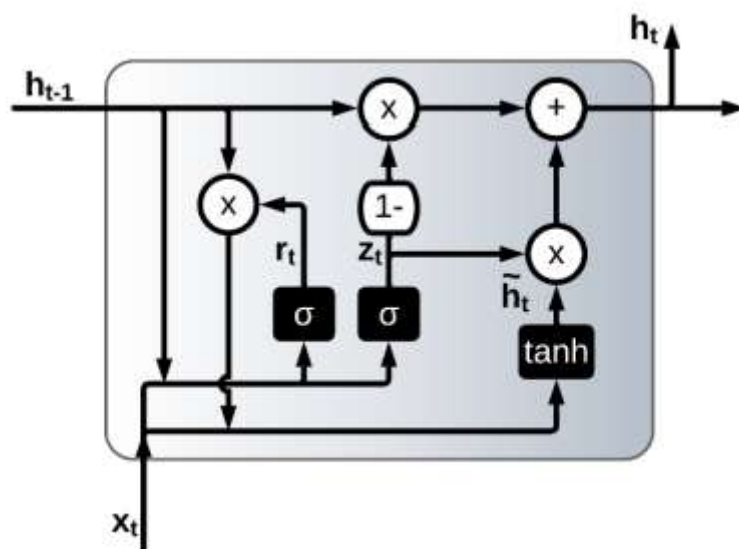


Figura 83: L'architettura della proposta GRU NN.

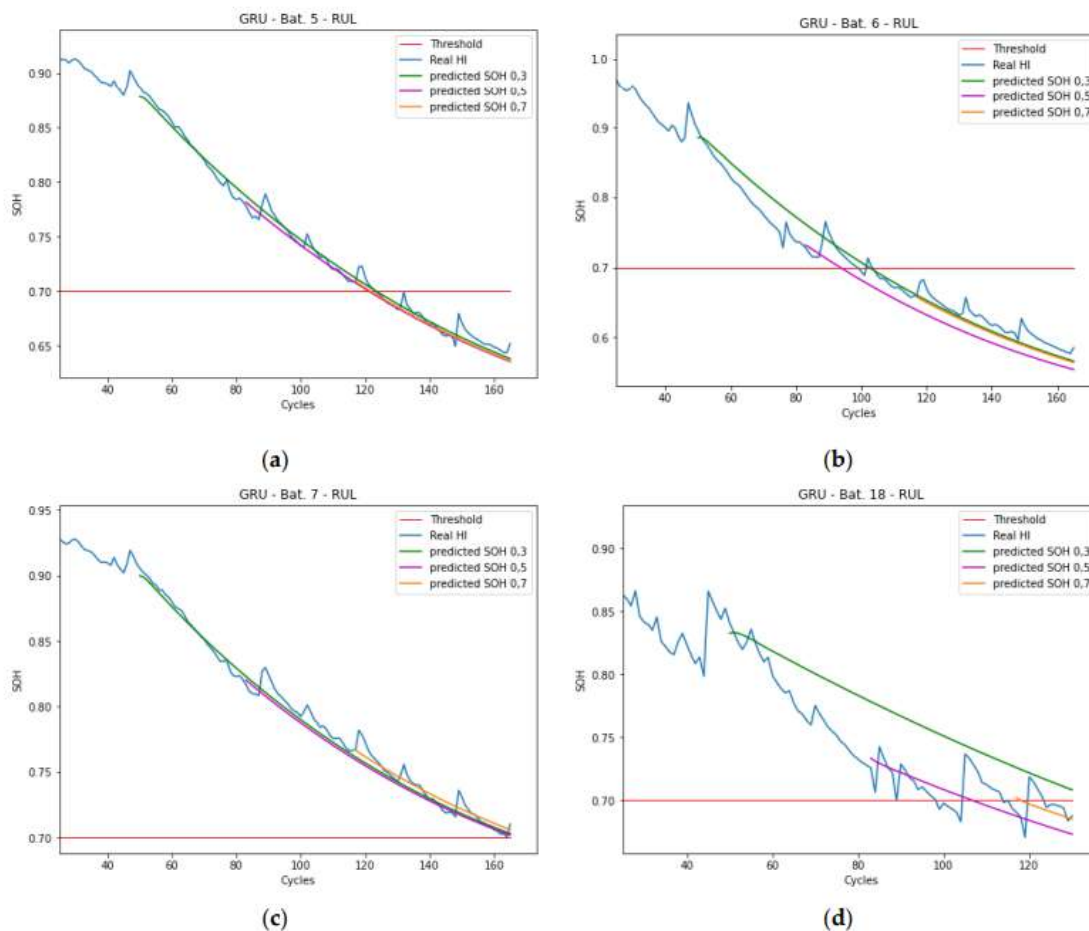


Figura 84: I risultati di previsione della Remaining-Useful-Life (RUL) per diverse posizioni di partenza delle batterie

Ciò porta a risultati precisi per la stima della SOH e risultati accurati per le tendenze di previsione a lungo termine della RUL. Nella realtà, i carichi variabili di carica effettuati dagli utenti sulle batterie sono una delle principali sfide. Le previsioni mostrano l'andamento discendente delle singole batterie. Le singole previsioni in ciascun set di dati della batteria iniziano nella posizione designata e ciascuno mostra successivamente una curva simile. Da ciò si può dedurre che l'attenzione del modello è sull'evoluzione a lungo termine del degrado delle prestazioni.

Il documento [51] presenta lo sviluppo preliminare di uno strumento di prognostica basato sui dati, per prevedere il SoH e il RUL della batteria al litio. Questo lavoro mira a completare due compiti. In primo luogo, viene effettuato un benchmark completo del modello data-driven utilizzando un algoritmo di apprendimento automatico sui dati di prognosi della batteria. In secondo luogo, viene sviluppato un modello preliminare data-driven utilizzando un **algoritmo di deep learning** per i dati di prognosi. L'efficacia dell'approccio proposto è stata implementata in un caso di studio con un set di dati di batteria ottenuto dal database del Centro di Eccellenza per la Prognostica del NASA Ames.

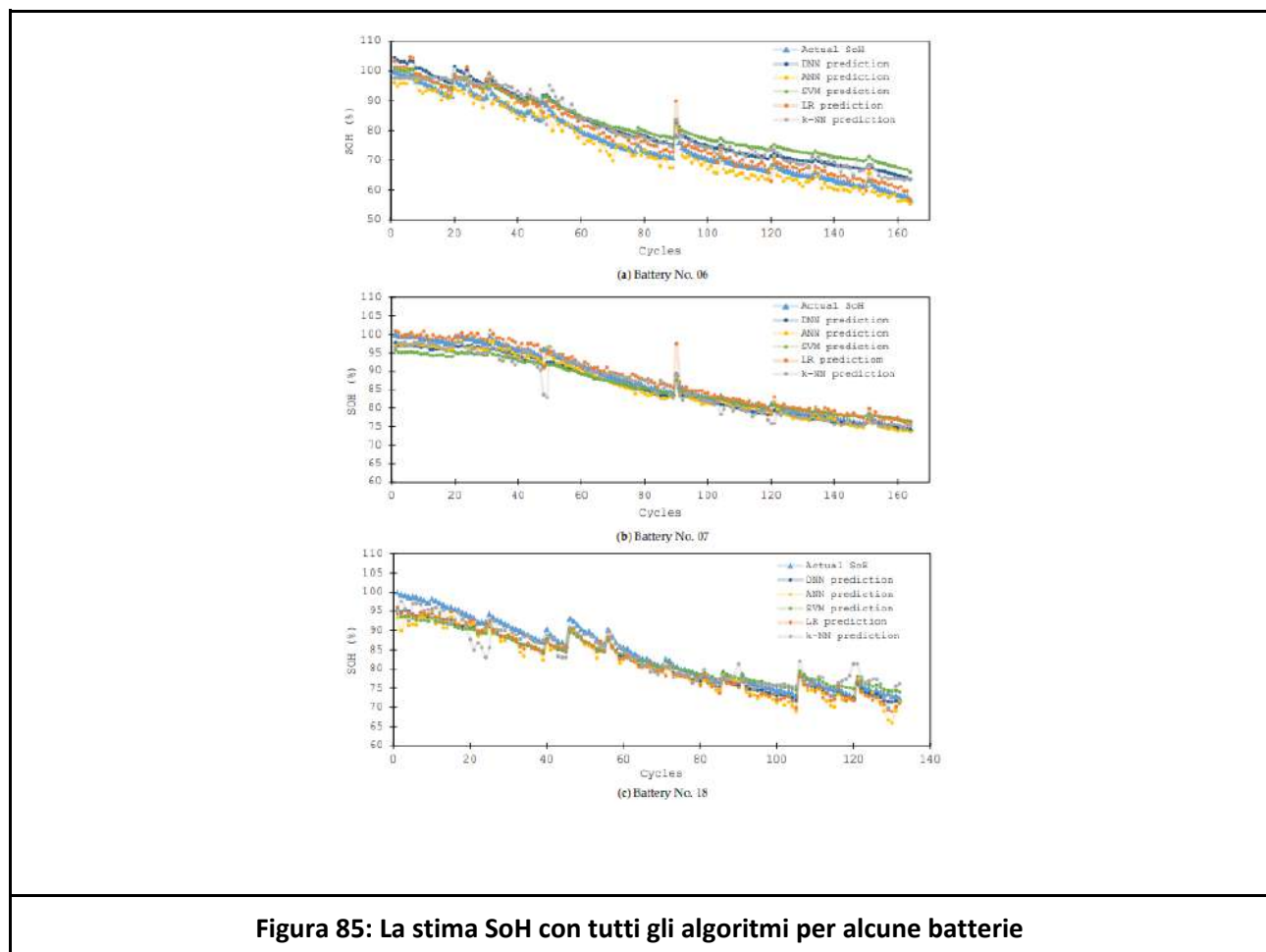
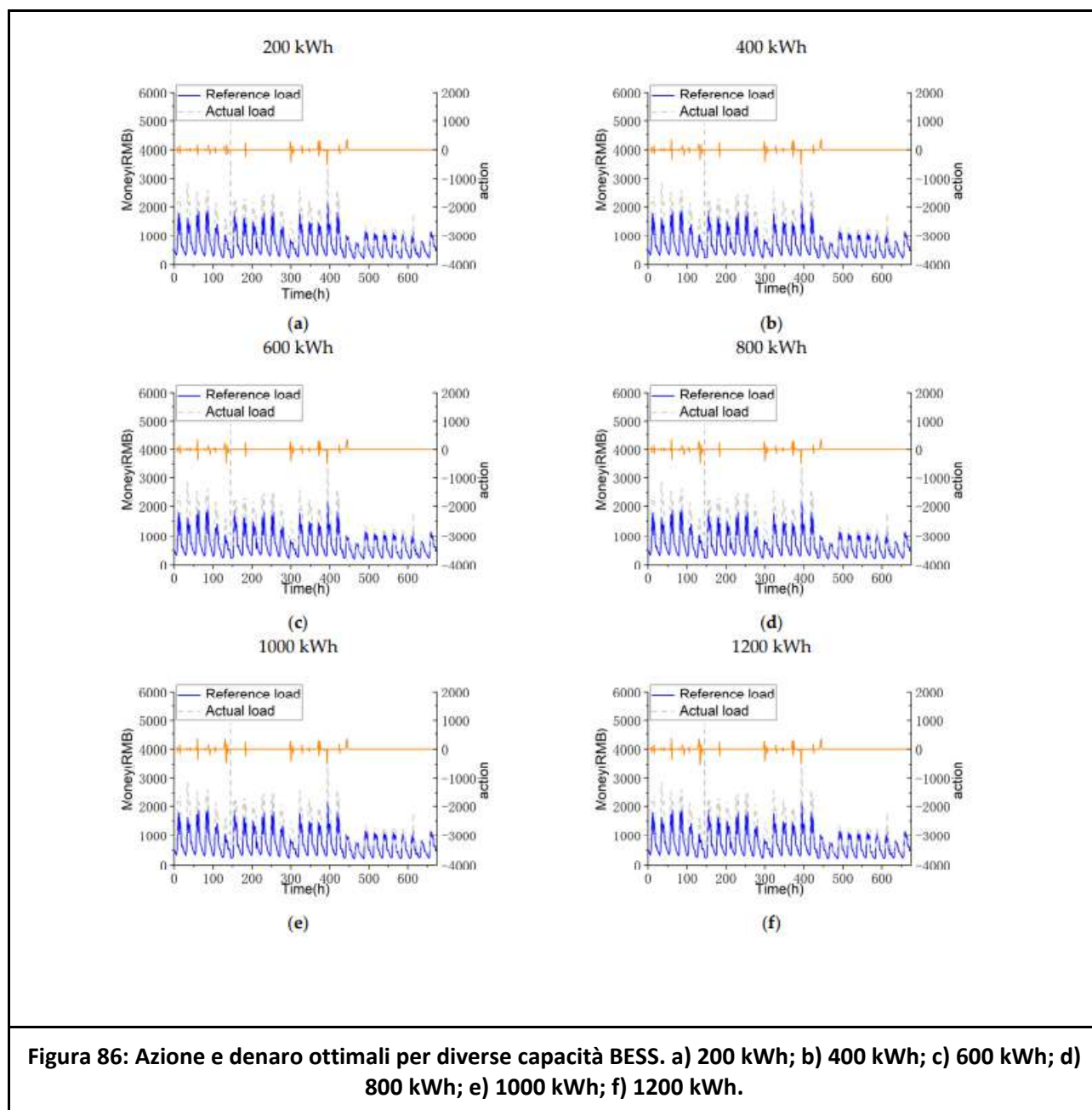


Figura 85: La stima SoH con tutti gli algoritmi per alcune batterie

L'algoritmo **Deep Neural Networks (DNN)** proposto è stato confrontato con altri algoritmi di machine learning, ovvero **Support Vector Machine (SVM)**, **k-Nearest Neighbors (k-NN)**, **Artificial Neural Networks (ANN)** e **Linear Regression (LR)**. I risultati sperimentali rivelano che le prestazioni

dell'algoritmo DNN potrebbero essere pari o superiori ad altri algoritmi di machine learning. Questo articolo ha raggiunto il suo obiettivo di fungere da benchmark per il modello data-driven di prognosi dei dati di batteria utilizzando algoritmi di apprendimento automatico e, sulla base dei risultati dei casi di studio, dimostra che l'algoritmo di deep learning fornisce un risultato promettente per la previsione e la modellizzazione dei dati di prognosi, soprattutto nelle applicazioni di prognosi e gestione dello stato di salute delle batterie. Sulla base dell'accuratezza raggiunta, si ritiene anche che il modello basato sui dati possa sostituire in futuro il modello tradizionale basato sulla fisica in vari campi e applicazioni.

Il paper scientifico [52] propone un modello che considera il ciclo di vita di una batteria al litio e i parametri di installazione della batteria, e i dati di consumo di energia e di generazione di energia fotovoltaica di un parco industriale sono stati utilizzati per stabilire un modello di gestione dell'energia. Il sistema di gestione dell'energia mirava a ridurre i costi operativi e ottenere la capacità di stoccaggio energetico ottimale, che è vincolata dalle prestazioni della batteria al litio e dalla domanda della rete. Questo articolo propone infatti una strategia di gestione delle batterie per ottimizzare il consumo energetico per gli utenti che possiedono dispositivi di generazione di energia elettrica e desiderano ridurre i costi di consumo energetico. In particolare, l'utente firma un accordo con il mercato spot in base ai risultati previsionali e poi l'energia viene distribuita in base alla situazione effettiva del giorno.



Con i costi operativi e la capacità della batteria ragionevole come obiettivi di ottimizzazione, sono stati adottati il metodo **Deep Deterministic Policy Gradient (DDPG)**, il **Greedy Dynamic Programming Algorithm** e il **Genetic Algorithm (GA)**, dove le prestazioni della batteria al litio e il requisito della rete elettrica erano i vincoli. Rispetto al greedy algorithm e al genetic algorithm, DDPG ha uno spazio di funzionamento continuo e non è facile cadere nella soluzione ottimale locale durante l'esplorazione della soluzione ottimale globale. I risultati mostrano che senza considerare il costo della batteria, il metodo di gestione energetica greedy, il metodo di gestione energetica GA e

il metodo di gestione energetica DDPG possono ridurre la capacità e i costi operativi del sistema di stoccaggio di energia del 18,9%, del 36,1% e del 35,9%, rispettivamente, rispetto ai costi attuali dell'elettricità.

L'articolo [53] presenta approcci euristici per valutare lo stato di carica (SOC) della batteria al piombo-acido attraverso la generazione di potenza ottimale in un sistema di erogazione di energia rinnovabile ibrido eolico-solare autonomo. L'analisi matematica è stata effettuata applicando tecniche euristiche modificate, ovvero il **Tuned Genetic Algorithm (TGA)** e l'**Improved Colliding Body Optimization (ICBO)**, per una specifica condizione di carico in India. Più energia ottimizzata viene generata dalle fonti di energia distribuita rinnovabile quando valutata da ICBO rispetto a quella ottenuta da TGA, sia per le richieste di carico invernale che estiva.

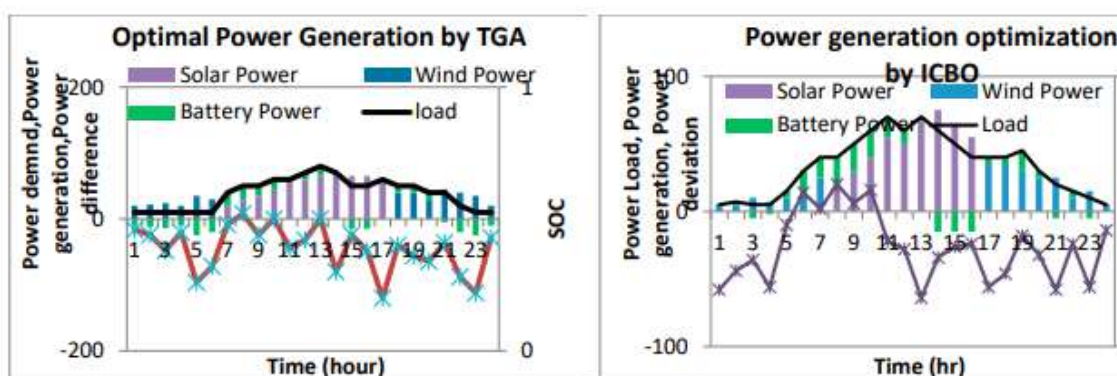


Figura 87: Generazione di energia rinnovabile e SOC, come calcolato da TGA per carico estivo e generazione ottimale di energia e SOC valutati mediante ICBO per la richiesta di carico invernale data

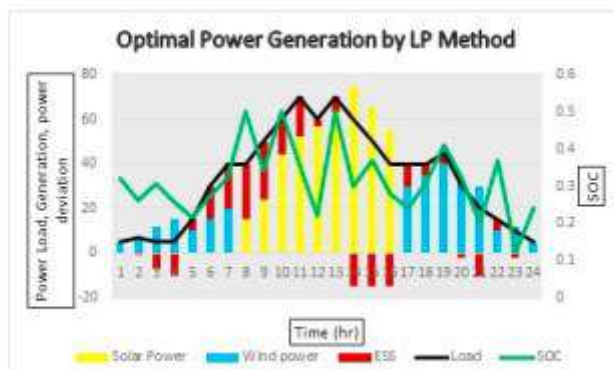


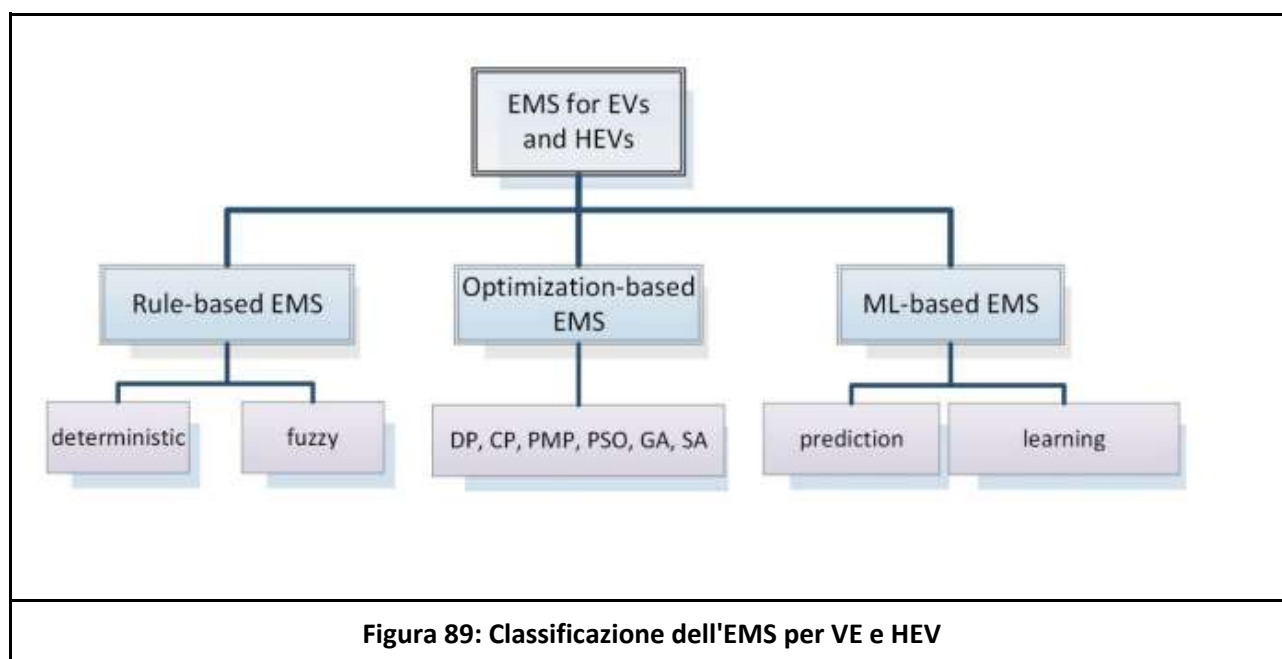
Fig. 6. Renewable Power generation and SOC, as computed by LP for winter load

Figura 88: Produzione di energia rinnovabile e SOC, come calcolato da LP per il carico invernale

È stata effettuata un'operazione ottimale di un sistema di fornitura di energia rinnovabile solare-eolica, per una specifica condizione di carico indiana. A questo riguardo, è stata calcolata la caratteristica comportamentale della batteria al piombo, ovvero lo stato di carica (SOC), attraverso l'applicazione di tecniche di intelligenza artificiale. Gli **algoritmi TGA e ICBO** sono stati modificati e implementati per ottenere la deviazione di potenza minima, correlata alla minima perdita di durata della batteria. Si è riscontrato che la generazione di potenza più ottimizzata è stata ottenuta mediante l'applicazione di **ICBO** piuttosto che **TGA**, sia per i carichi invernali che estivi. Inoltre, il valore SOC raggiunto da ICBO era inferiore del 4,3% rispetto a quello ottenuto da TGA. Pertanto, il **metodo ICBO** con una generazione di potenza più ottimizzata ha portato a meno cicli di carica-scarica, che infine hanno aiutato il sistema di accumulo di energia con una migliore durata operativa. Tuttavia, con un valore SOC più elevato, la performance elettrica della batteria potrebbe essere negativamente influenzata se mantenuta per un periodo più lungo. Al fine di ottenere risultati migliori, è stato effettuato uno studio parametrico degli algoritmi modificati, ovvero **TGA e ICBO**, e sono stati ottenuti un insieme di parametri operativi critici attraverso lo studio parametrico di **TGO e ICBO** per un profilo di carico in tempo reale specifico.

Il documento [54] esamina in modo completo le tecniche e gli algoritmi più avanzati che mirano a gestire in modo efficiente l'energia di carica/scarica degli Electric Vehicles (EV) basati su tecniche di

apprendimento automatico (ML), e si identificano le somiglianze e le principali modifiche degli approcci proposti. Ciò consente di derivare un insieme di lezioni utili e di identificare una serie di sfide, vantaggi e limitazioni che derivano dall'implementazione di questi metodi di gestione dell'energia degli EV basati su ML, che dovrebbero essere affrontati al fine di promuovere una rete elettrica più efficiente ed economica.



In questo paper è stata presentata una revisione estesa delle Energy Management Strategies (EMS) per veicoli elettrici e ibridi, e la si è classificata in tre categorie: EMS basate su regole, EMS basate sull'ottimizzazione e EMS basate sul machine learning. Si sono esaminate una vasta gamma di strategie EMS per ogni categoria e si sono evidenziati i pro e i contro di ciascun approccio, nonché la promettente robustezza degli algoritmi e dei metodi basati sul machine learning. Sulla base di questa revisione, ci si può concentrare sulle sfide, i vantaggi e anche le limitazioni che sorgono dall'implementazione di questi metodi EMS, che dovrebbero essere affrontati al fine di rendere gli EMS più efficienti con una risposta in tempo reale. Tra i tanti algoritmi di machine learning utilizzati, ci sono: **Neural Network, Decision Tree, Support Vector Machines (SVMs), Linear Regression (LR), Logistic Regression, Clustering e Genetic Algorithms (GAs)**. Gli approcci basati sulla previsione

aumentano significativamente le prestazioni di un EMS fornendo previsioni che possono essere aggiornate in tempo reale, migliorando così l'efficienza complessiva del sistema grazie alla capacità di adattarsi alle incertezze del traffico. La principale sfida di questi approcci è l'aumento dell'accuratezza della previsione, che può essere realizzato ottimizzando la selezione dell'orizzonte di previsione appropriato e aumentando l'utilizzo dei dati forniti da fonti esterne (ad esempio, dati meteorologici o dati di traffico da ITS). D'altra parte, un'altra sfida riguarda la complessità computazionale di questi approcci, specialmente nei metodi basati su MPC, che è cruciale per la fornitura di implementazioni in tempo reale dell'EMS dei veicoli elettrici.

Nell'articolo [55] l'abilità di prevedere in anticipo la vita residua (RL) della batteria è considerata critica per garantire un'offerta affidabile di energia e l'uso più efficiente di tale energia. Quando si tratta di prevedere con precisione il livello di carica delle batterie (SoC), i sistemi di gestione delle batterie devono essere durevoli e affidabili. A causa della natura non lineare della degradazione delle batterie, è estremamente difficile prevedere la stima del SoC con una degradazione considerevolmente inferiore rispetto a quanto attualmente possibile.

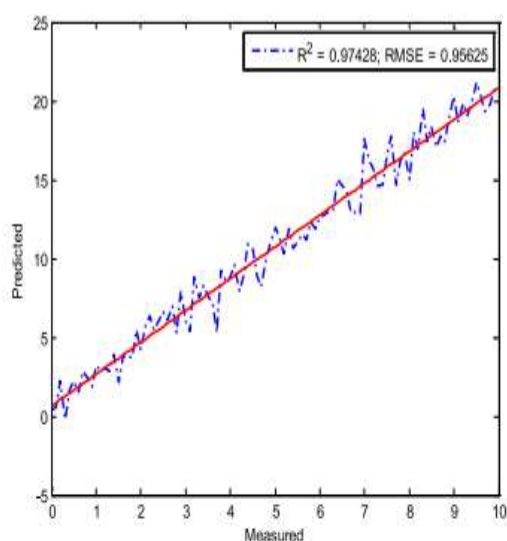


Fig. 4. Distribution of proposed random forest for SoC.

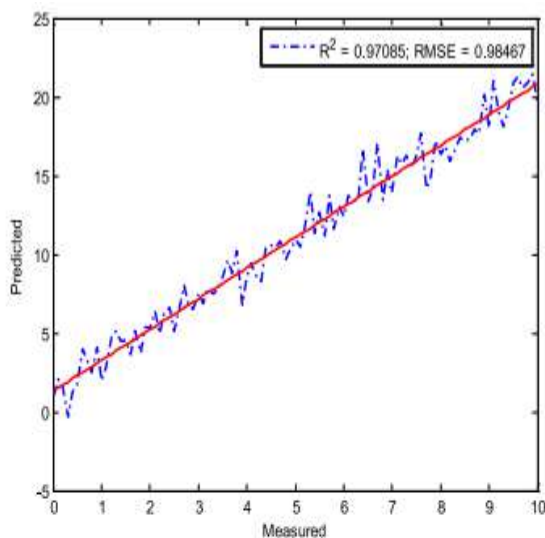


Fig. 6. Distribution of proposed random forest for capacity losses.

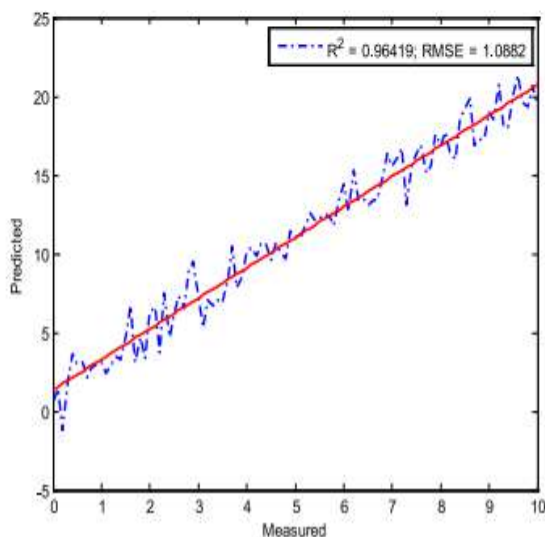
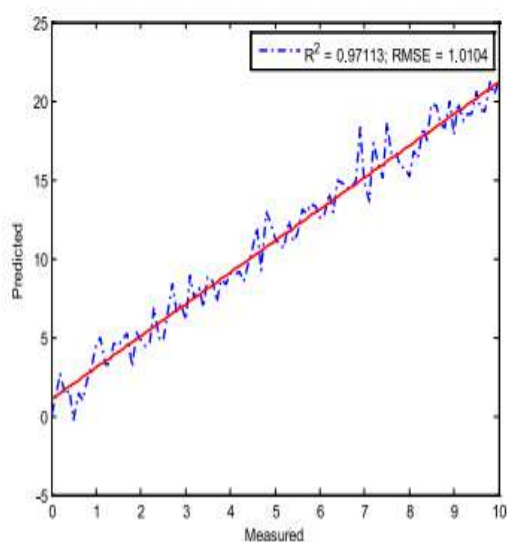


Figura 90: Distribuzione delle proposte con RD e ERD per SoC e distribuzione delle proposte con RD e ERD per perdite di capacità

Nel presente lavoro, si utilizza un **modello di Random Forest di tipo Ensemble** per prevedere la durata residua della batteria (RL) al fine di ridurre al minimo la degradazione dei dati. Il modello raccoglie, pre elabora e classifica i dati utilizzando **tecniche di Random Forest (RF) e Ensemble Random Forest** per la previsione della RL. La simulazione fa uso sia dell' R^2 che dell'errore quadratico medio (RMSE). L'uso di un modello di **Ensemble Random Forest** migliora l'accuratezza delle

previsioni. Utilizzando la tecnica proposta, le stime del SoC in tempo reale possono essere effettuate una volta che gli iperparametri sono stati regolati. La tecnica di ensemble RF proposta supera la tecnica di RF con un MAE dell'85%. Ciò potrebbe essere dovuto alla capacità del kernel GPR di recuperare i dati sequenziali dalle strutture temporali complesse, il che potrebbe spiegare la sua efficacia. Secondo i ricercatori, l'uso della strategia di ensemble RF invece dell'approccio RF individuale, consente di ottenere una riduzione del MAE del 10%. Pertanto, si può dedurre che l'approccio proposto basato su **Ensemble Random Forest**, anziché migliorare l'accuratezza della stima puntuale, consente di migliorare le stime del SoC stimando una distribuzione di probabilità.

3 Previsione di domanda, produzione e costo dell'energia

L'energia è essenziale per lo sviluppo sostenibile in qualsiasi paese, sia esso sociale, economico o ambientale. Le emissioni ambientali possono essere ridotte attraverso una minore dipendenza dai combustibili fossili, il posizionamento di fonti di generazione di energia rinnovabile e la gestione del consumo di elettricità. La dipendenza dai combustibili fossili, l'impatto negativo sull'ambiente e la volatilità dei prezzi hanno incoraggiato e incrementato le attività di efficienza energetica e cambiamenti nei sistemi di alimentazione elettrica. Uno di questi cambiamenti è stato lo sviluppo della generazione decentralizzata, locale e di microgenerazione da fonti energetiche rinnovabili (FER). Sebbene le reti intelligenti (SG) risolvono molti dei problemi contemporanei, danno origine a un nuovo controllo e risolvono i problemi soprattutto con il ruolo crescente delle FER. L'altro cambiamento è stato che la liberalizzazione del mercato dell'energia ha portato maggiore competitività al settore e il crescente utilizzo delle FER ha anche portato maggiore complessità al processo di transizione energetica. Sia la liberalizzazione che la complessità della transizione energetica possono essere affrontate attraverso la digitalizzazione e l'integrazione del sistema energetico per creare un Internet of Energy (IoE). Per sfruttare appieno le opportunità offerte dalla digitalizzazione, sarà necessario garantire sia una cooperazione coordinata tra tutti i principali stakeholder, sia rigorosi sistemi di monitoraggio. La digitalizzazione e le potenzialità fornite dalle reti intelligenti, come SG, Internet-of-Thing o Internet-of-People, offrono l'opportunità di affrontare queste sfide e rappresenteranno, quindi, un importante punto di svolta per i futuri sistemi

energetici. Durante l'ultimo decennio, diverse nuove tecniche sono state utilizzate per la gestione della domanda di energia per prevedere con precisione il fabbisogno energetico futuro. I sistemi di modellazione vengono sempre più utilizzati per prevedere energia e prezzi. Per un futuro prospero, un ambiente economico sostenibile e sicuro è fondamentale una gestione energetica che richieda una corretta allocazione delle risorse disponibili.

3.1 Previsione della domanda di energia

Il settore energetico mondiale sta affrontando sfide crescenti, come l'aumento della domanda e dell'efficienza, il cambiamento del modello di domanda e offerta e l'assenza di una migliore analisi di gestione. Nei paesi in via di sviluppo, questa sfida è ancora più intensa. Il trasferimento dei dati del settore energetico all'intelligenza artificiale può gradualmente risolvere questo problema. Gli algoritmi di intelligenza artificiale possono analizzare i dati delle apparecchiature, risolvere problemi e quindi risparmiare denaro, tempo e vita. Le applicazioni di intelligenza artificiale sono utilizzate anche per il consumo energetico intelligente. Le case moderne con algoritmi di apprendimento automatico possono rispondere automaticamente alle fluttuazioni dei prezzi dell'elettricità e controllare il consumo di energia. I sistemi basati sull'apprendimento automatico possono aiutare i fornitori di energia a prepararsi a tenere il passo con le fluttuazioni delle forniture di energia rinnovabile. Gli algoritmi di intelligenza artificiale possono migliorare continuamente il monitoraggio, analizzare il consumo di energia, scoprire nuovi problemi ed eseguire analisi per migliorare le prestazioni. Pertanto, l'intelligenza artificiale e l'apprendimento automatico possono svolgere un ruolo cruciale nella gestione pratica delle sfide del settore energetico.

I dati sul consumo di energia possono rivelare tendenze e prevedere modelli futuri al fine di ridurre i rischi, diminuire le emissioni di carbonio e aumentare la redditività.

Gli approcci sono volti a determinare i consumi principalmente sulla scala temporale a breve, medio e lungo termine:

- Short-term load forecast (STLF): il periodo di tempo di STLF dura da pochi minuti o ore a un giorno o una settimana in avanti. La previsione STLF mira al dispatch economico e all'impegno ottimale dell'unità di generazione, affrontando al contempo il controllo in tempo reale e la valutazione della sicurezza.
- Medium-term load forecast (MTLF): il periodo di tempo di MTLF va da una settimana a un anno (possibilmente due anni). La previsione MTLF mira alla programmazione della manutenzione, al coordinamento della spedizione del carico e alla definizione dei prezzi in modo che la domanda e la generazione siano bilanciate.
- Long-term load forecast (LTLF): il periodo di tempo di LTLF varia da pochi anni fino a 10–20 anni in avanti. La previsione LTLF mira alla pianificazione dell'espansione del sistema, i.e. generazione, trasmissione e distribuzione. In alcuni casi influisce anche sull'investimento in nuove unità di generazione.

3.1.1 Short-term load forecast (STLF)

Nell'articolo [56] viene proposto un approccio di **transfer learning** basato sulla correlazione incrociata per ottenere dati da diverse parti del mondo per ottenere risultati predittivi più efficaci con dati limitati, utile nell'ambito dell'edilizia, dove le nuove costruzioni e i nuovi contatori, offrono dati storici relativamente piccoli. Risultati di applicazione sul sistema energetico reale convalidano i vantaggi del modello proposto.

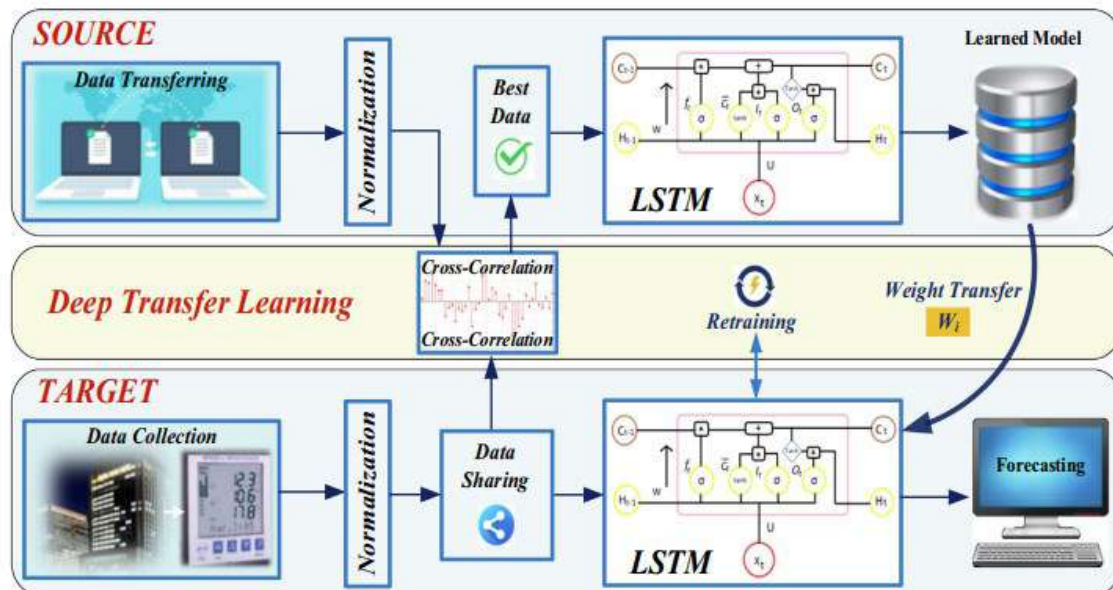


Figura 91: Schema di approccio all'apprendimento del trasferimento basato su XCORR

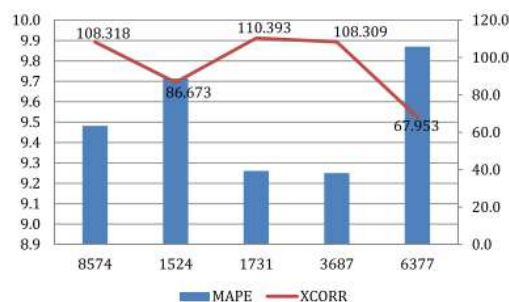
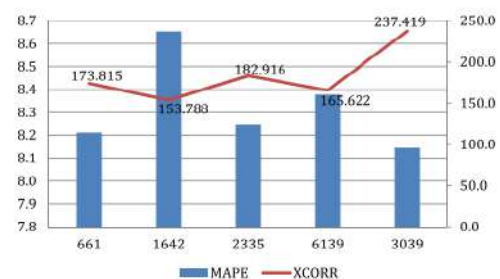
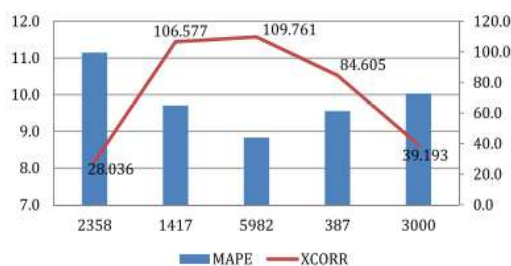


Figura 92: Grafici XCORR e MAPE per New York, Austin e California

Lo scopo principale di questo studio è stimare il consumo negli edifici con la minore quantità di dati energetici. Il problema critico è stato determinare i dati di trasferimento dell'apprendimento più adatti. In questo contesto, è stato proposto un approccio basato su XCORR per trovare la somiglianza

tra i dati di trasferimento dell'apprendimento e i dati originali. Nei test eseguiti senza *transfer learning* con dati originali, nel caso peggiore sono stati ottenuti valori di 1021.671, 711.703 e 22.916; nel caso migliore, invece, sono stati ottenuti i valori 852.912, 494.496 e 13.913 rispettivamente per RMSE, MAE e MAPE.

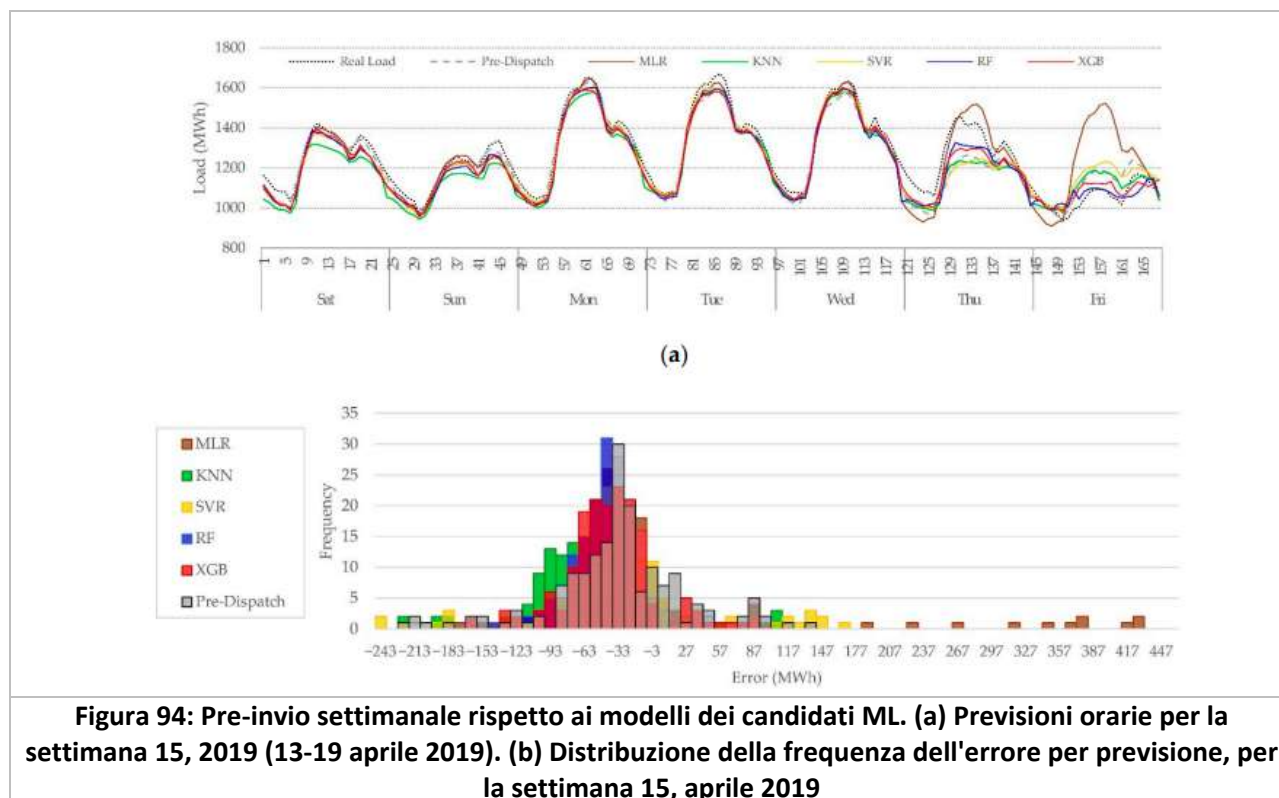
Come risultato del processo di transfer learning con gli edifici adatti selezionati, usando XCORR sono stati ottenuti valori di 736.706, 352.176 e 8.145 valori, rispettivamente per RMSE, MAE e MAPE. Pertanto, rispetto all'assenza di transfer learning, è stato ottenuto un miglioramento del 5,768% nei valori MAPE, corrispondente al tasso di riduzione dell'errore relativo del 41,45%.

Nell'articolo [57] si propongono dei set di modelli di machine learning (ML) per migliorare l'accuratezza delle previsioni a 168 ore. I modelli sviluppati utilizzano funzionalità da più fonti, come il carico storico, il tempo e le festività. Dei cinque modelli di ML sviluppati e testati in vari contesti di profilo di carico, l'algoritmo **Extreme Gradient Boosting Regressor (XGBoost)** ha mostrato i migliori risultati, superando lo storico precedente di previsioni settimanali basate su reti neurali. Inoltre, poiché i modelli XGBoost si basano su un insieme di alberi decisionali, questo ha facilitato l'interpretazione del modello, che ha fornito un rilevante risultato aggiuntivo, l'importanza caratteristica della previsione.

Model	Hyperparameter	Randomized Search	RS ¹ Selection	Grid Search	GS ² Selection
KNN	n_neighbors	(3, 6, 9, 12, 15, 18, 21, 24, 27)	29	(23, 26, 29, 32, 35)	29
	weights	(uniform, distance)	distance	(distance)	distance
	metric	(minkowski, euclidean, manhattan)	manhattan	(manhattan)	manhattan
	leaf_size	(5, 15, 25, 35)	35	(29, 32, 35, 38, 41)	29
SVR	kernel	(linear, rbf)	rbf	(rbf)	rbf
	epsilon	(0.0001, 0.3), step 0.01	0.0701	(0.03505, 0.09113), step 0.01402	0.03505
	C	(0.01, 0.1, 1, 10, 100)	100	(50, 70, 90, 110, 130)	50
	tol	(0.001)	0.001	(0.001)	0.001
RF	criterion	(mse)	mse	(mse)	mse
	n_estimators	(50, 100, 150, 200)	200	(200)	200
	max_samples	(0.5, 0.6, 0.7, 0.8, 0.9)	0.5	(0.5)	0.5
	max_depth	(10, 20, 30, 40, 50)	40	(20, 35, 50, 65)	35
	ccp_alpha	(5×10^{-2} , 0.01), step 0.0005	0.0095005	(0.0094, 0.0099, 0.0104, 0.0109, 0.0114)	0.0094
	random_state	(123)	123	(123)	123
XGB	eval_metric	(rmse)	rmse	(rmse)	rmse
	n_estimators	(100, 150, 200, 250, 300, 350, 400)	300	(300, 310, 320)	300
	max_depth	(3, 5)	5	(5, 7, 9)	7
	subsample	(0.6, 0.7, 0.8, 0.9)	0.8	(0.7)	0.7
	colsample_bytree	(0.7, 0.8, 0.9, 1.0)	0.9	(0.7, 0.75)	0.7
	colsample_bylevel	(0.7, 0.8, 0.9, 1.0)	1.0	(0.7)	0.7
	colsample_bynode	(0.7, 0.8, 0.9, 1.0)	0.8	(0.7, 0.75)	0.7
	learning_rate	(0.0001, 0.2501), step 0.05	0.0501	(0.05, 0.045)	0.045
	min_child_weight	(1, 3, 5, 7)	7	(7)	7
	gamma	(0, 5, 10, 15)	15	(15)	15.00
	random_state	123	123	(123)	123

¹ RS stands for randomized search. ² GS stands for grid search.

Figura 93: Ricerca e risultati nello spazio iperparametrico



Feature	MLR	KNN	SVR	RF	XGB
L_{h-168}	23.70%	3.00%	16.00%	73.60%	24.70%
L_{h-336}	12.90%	2.40%	6.90%	1.70%	12.80%
L_{h-504}	14.50%	4.70%	15.30%	10.70%	12.10%
L_{h-672}	14.10%	2.40%	12.80%	3.60%	16.40%
LMA_{h-168}	3.00%	0.80%	0.30%	1.00%	0.50%
LMA_{h-336}	4.20%	1.40%	0.30%	1.00%	0.50%
month of the year $_h$	0.50%	0.00%	0.00%	0.70%	0.70%
day of the week $_h$	0.80%	7.30%	4.10%	0.50%	2.20%
weekend indicator $_h$	3.30%	22.10%	11.00%	0.10%	2.70%
holiday indicator $_h$	11.30%	16.40%	11.20%	3.40%	9.40%
hour of the day $_h$	2.20%	27.40%	2.70%	0.70%	14.40%
temperature in Panama City $_h$	9.20%	9.10%	16.50%	2.10%	3.10%
relative humidity in Panama City $_h$	0.30%	3.20%	2.80%	1.00%	0.40%

Figura 95: Importanza media delle funzionalità per modello ML

Questa ricerca ha presentato un nuovo modello che migliora la previsione STLF settimanale con un gap di 72h. Il modello è stato sviluppato e confrontato con dati e previsioni precedenti del sistema energetico di Panama. I risultati di 14 settimane di test hanno confermato l'idoneità dell'algoritmo XGB. In primo luogo, per l'efficienza del tempo nell'addestramento e nella previsione, in secondo luogo, per la flessibilità a causa dello spazio degli iperparametri.

Per le ragioni sopra esposte, questa ricerca apporta diversi importanti contributi al campo di studio. Innanzitutto, mostra che i modelli costruiti con XGB hanno prestazioni superiori a modelli costruiti con altri algoritmi. Inoltre, conferma che la temperatura gioca un ruolo critico in STLF. Poi, dimostra l'eccellente performance dei modelli ML mediante previsioni per un orizzonte più lungo rispetto alla ricerca tipica e anche con un intervallo di dati di tre giorni prima della previsione. Infine, questa ricerca identifica i dati pubblici che possono essere utilizzati da altri ricercatori per migliorare le previsioni STLF o anche presi come parametri di riferimento.

In [58] si propone un modello per STLF che utilizza la **Deep Neural Network (DNN)** con selezione delle funzionalità. Un'ampia gamma di modelli di previsione intelligenti è stata progettata e testata sulla base di molteplici funzioni di attivazione, come la tangente iperbolica (tanh), diverse varianti di rectifier linear unit (ReLU) e sigmoid. Tra gli altri, DNN con leaky ReLU ha prodotto la miglior accuratezza della previsione. Per quanto riguarda la precisione dei metodi utilizzati in questo lavoro di ricerca, vengono utilizzati l'errore percentuale (MAPE), l'errore quadratico medio (MSE) e l'errore quadratico medio (RMSE). C'era anche una dipendenza da più dettagli parametrici e variabili per determinare la capacità delle tecniche di previsione del carico.

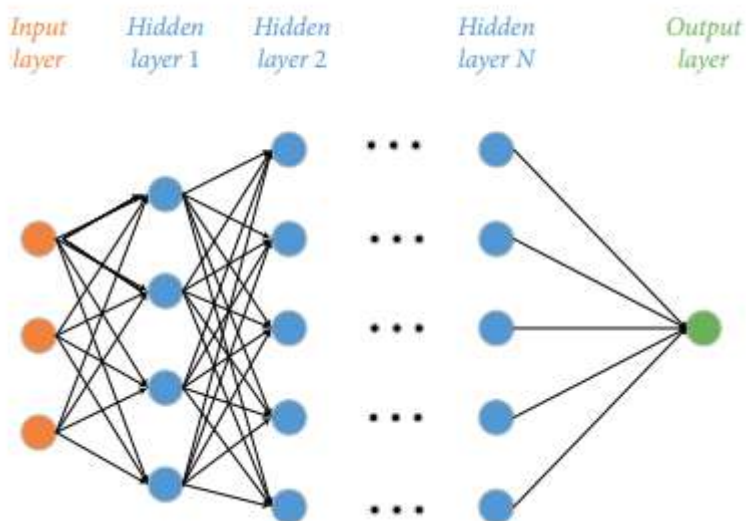


Figura 96: Una tipica architettura di rete neurale profonda

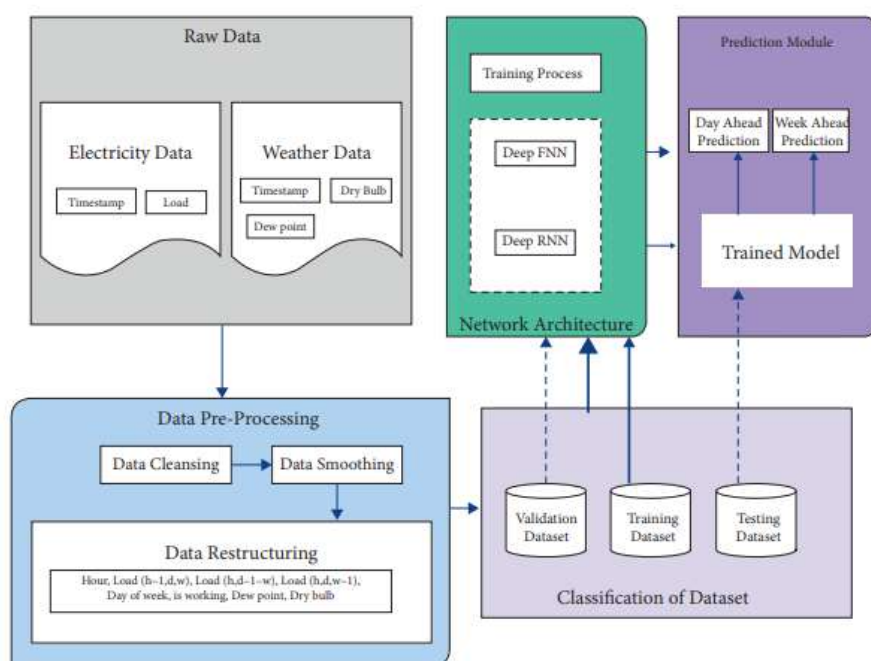


Figura 97: Modellazione per la previsione del carico

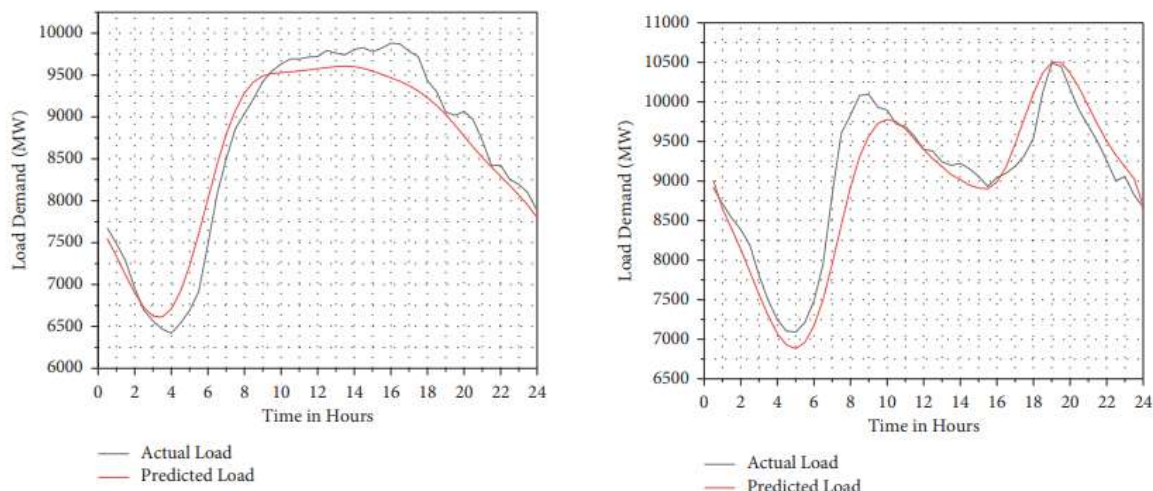


Figura 98: Risultati della previsione di carico con un giorno di anticipo del modello LSTM per la stagione autunnale e estiva

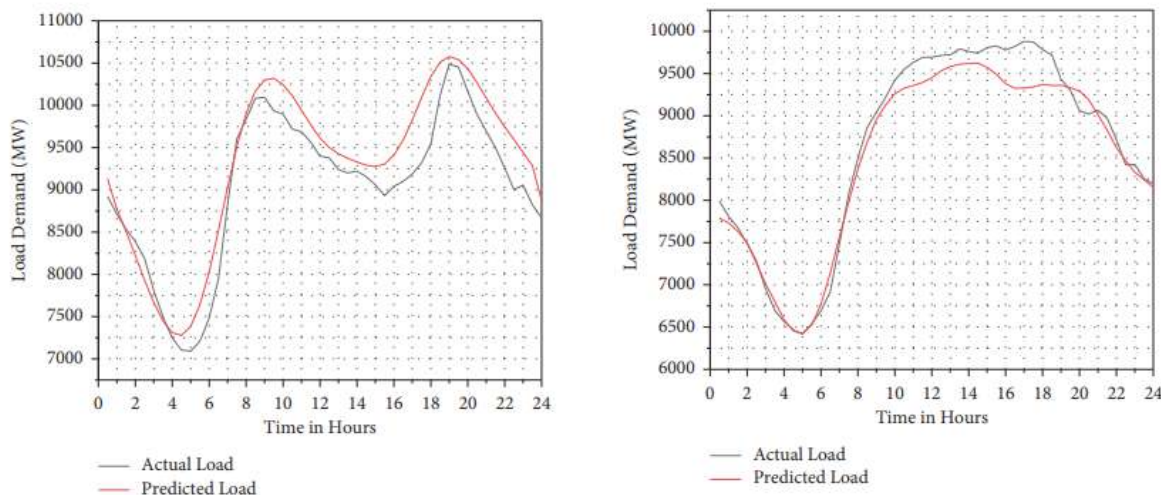


Figura 99: Risultati della previsione del carico con un giorno di anticipo del modello RNN per la stagione estiva e autunnale

Sono stati utilizzati tre tipi di modellazione come la **Long Short-Term Memory (LSTM)**, la **Recurrent Neural Network (RNN)** l'**RNN** e la **Neural Network (NN)** per la previsione di un giorno nelle stagioni. C'è stata più accuratezza usando la funzione di attivazione leaky ReLU con RNN. Buoni risultati per dati di previsione annuali sono stati ottenuti utilizzando anche le tecniche sopracitate. I dati di addestramento effettuano una fase di pre-elaborazione per prevedere le nuove funzionalità che

saranno importanti per l'uso di elettricità. I modelli di previsione ibridi proposti hanno mostrato elevata precisione di previsione e generalizzazione che porterebbe a minori costi operativi e funzionamento sicuro di aziende elettriche.

Nell'articolo [59] vengono confrontati due benchmark (**Autoregressive Integrated Moving Average (ARIMA)** e una tecnica manuale esistente utilizzato presso il sito del caso) rispetto a tre modelli di deep learning (**Recurrent Neural Network (RNN)**, **Long Short-Term Memory (LSTM)** e **Gated Recurrent Unit (GRU)**) e due modelli di machine learning (**Support Vector Regression (SVR)** e **Random Forest (RF)**) per la previsione del carico a breve termine (STLF) utilizzando i dati di un impianto di produzione di resina termoplastica brasiliano. È utilizzato il metodo di ricerca della griglia per identificare le migliori configurazioni per ciascun modello e quindi utilizzare il test Diebold-Mariano per confermare i risultati. I risultati suggeriscono che l'approccio legacy utilizzato nel sito del caso è quello con le prestazioni peggiori e che il modello GRU ha superato tutti gli altri modelli che sono stati testati.

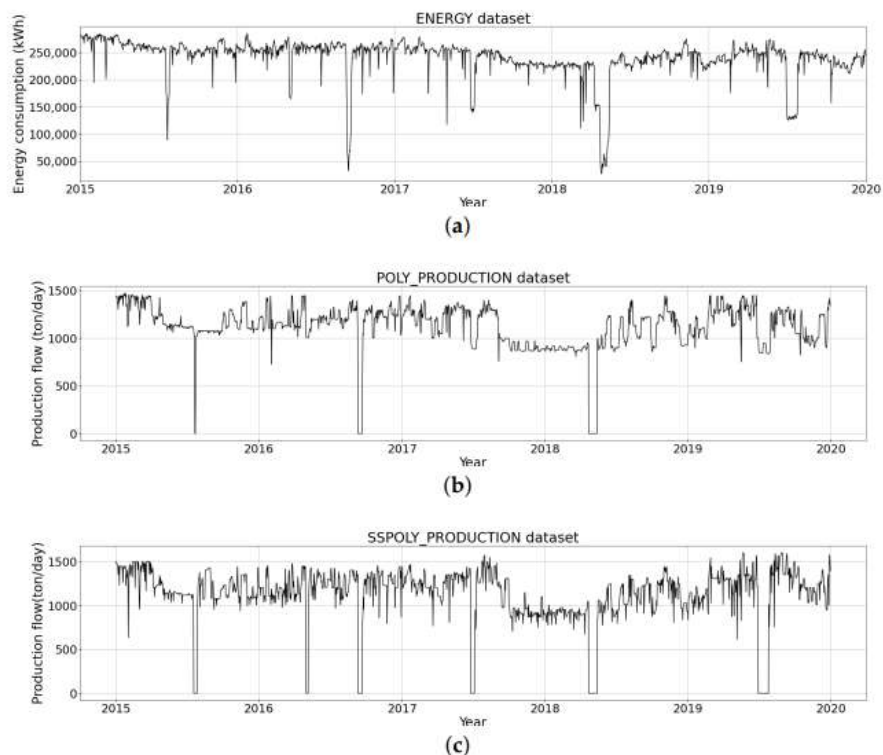


Figura 100: Set di dati giornalieri del (a) consumo di energia, (b) flusso di produzione per la fase di polimerizzazione e del (c) flusso di produzione per la fase di polimerizzazione allo stato solido.

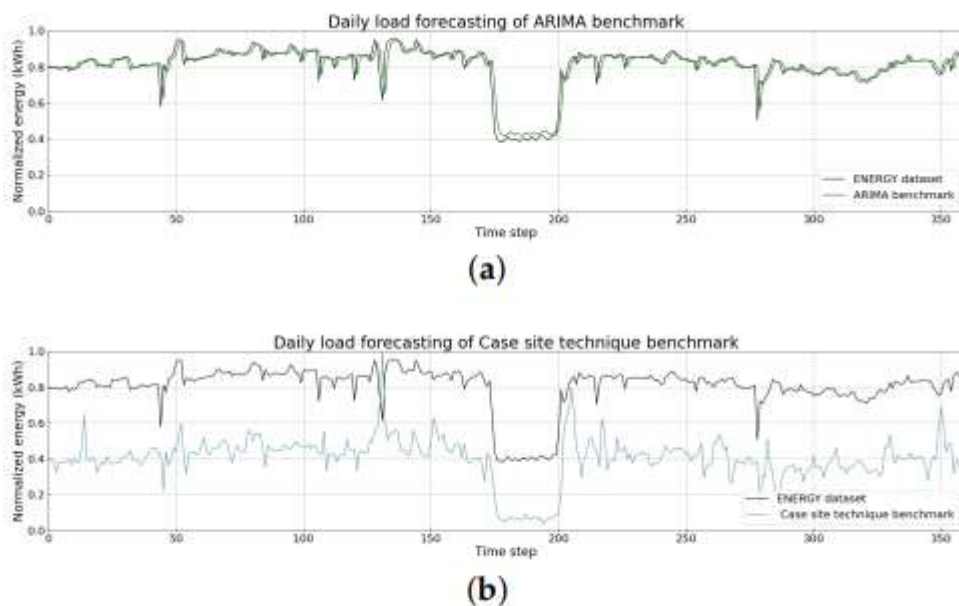


Figura 101: Previsione del carico giornaliero utilizzando (a) la tecnica del caso in loco e (b) il modello ARIMA

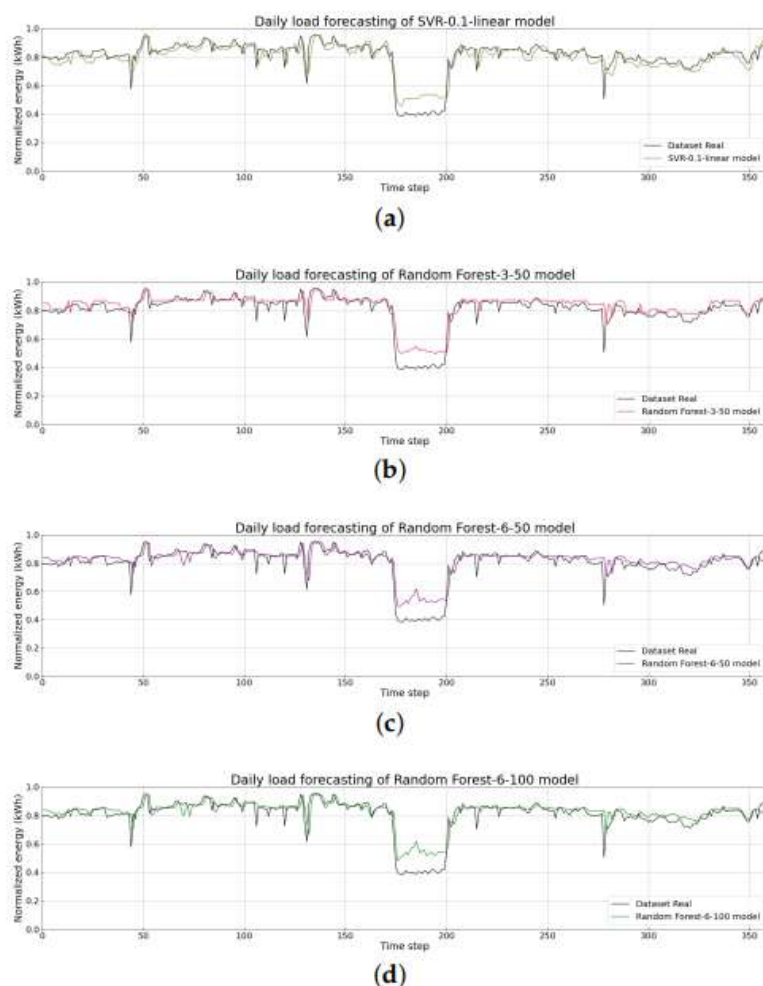


Figura 102: Previsione del carico giornaliero utilizzando (a) il modello lineare SVR-0.1, (b) il modello Random Forest-3-50, (c) il modello Random Forest-6-50 e (d) il modello Random Forest-6-100

Nell'articolo si confronta l'efficacia del deep-learning e dei modelli di machine learning per la previsione del carico a breve termine per impianti di produzione ad alta intensità energetica. Inoltre, si confrontano questi modelli con la tecnica di previsione manuale incombente e una tecnica di previsione classica delle serie temporali, la **Autoregressive Integrated Moving Average (ARIMA)**. A differenza degli studi esistenti, si considerano i dati reali di un terreno pluriennale, compresi i dati sul consumo totale di energia dell'impianto e i dati provenienti da due fasi di un processo produttivo intensivo di energia complessa. L'utilizzo dei dati di produzione ha contribuito in modo significativo al miglioramento della precisione STLF riducendo l'RMSE.

Come soluzione al fatto che i CNN non sono in grado di estrarre in modo efficiente le caratteristiche essenziali del carico elettrico, il documento [60] propone un nuovo modello di ensemble, che è una struttura in quattro fasi composta da un modulo di estrazione di caratteristiche, un blocco residuo densamente connesso (DCRB), uno strato di **Bidirectional Long Short-Term Memory (Bi-LSTM)** e un **Ensemble Thinking**. Il modello prima estrae le funzionalità di base e derivate da righe di dati utilizzando il modulo di estrazione delle caratteristiche. Le caratteristiche derivate comprendono la temperatura media oraria e caratteristiche del carico di elettricità, che possono catturare caratteristiche di enorme casualità e tendenza nel carico di elettricità. Il DCRB può estrarre efficacemente le caratteristiche essenziali dall'input multiscala rispetto ai modelli basati sulla CNN. I risultati dell'esperimento mostrano che il metodo proposto può fornire prestazioni di previsione più elevate rispetto ai modelli esistenti, di quasi lo 0,9–3,5%.

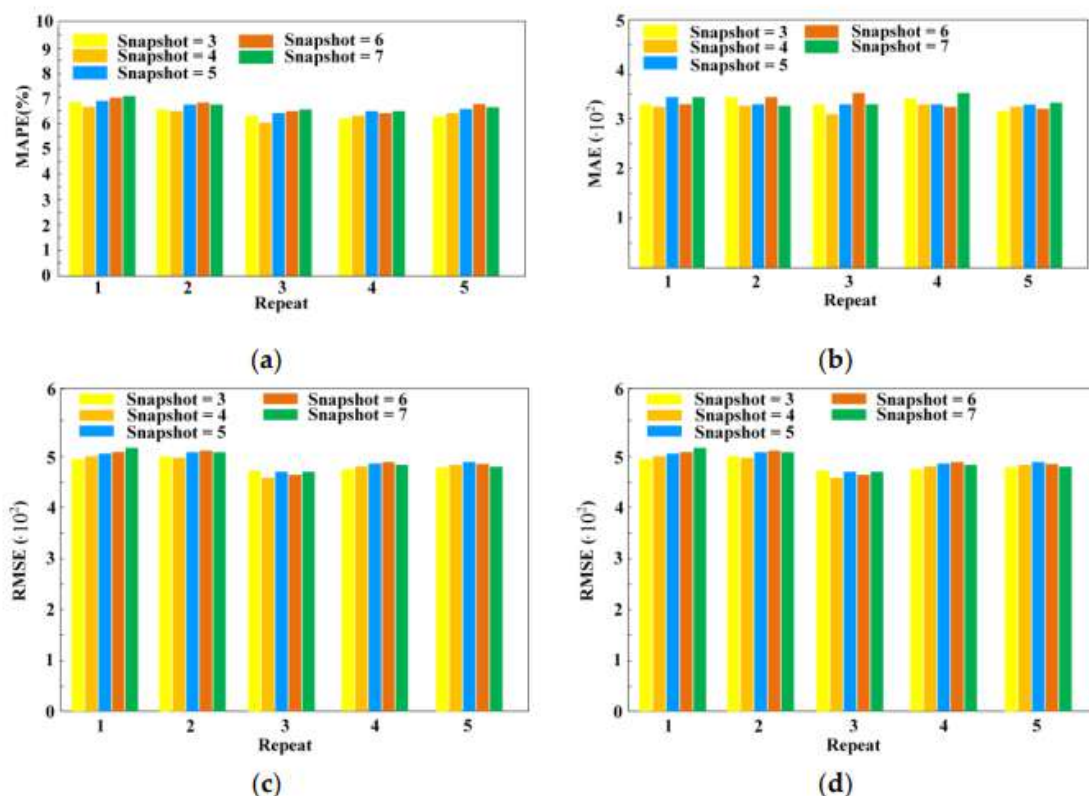


Figura 103: Il confronto delle prestazioni con differenti set di dati iperparametrici del Nord America : (a) errore percentuale medio assoluto (MAPE), (b) errore medio assoluto (MAE), (c) root-mean squared error (RMSE), (d) mean-squared error (MSE)

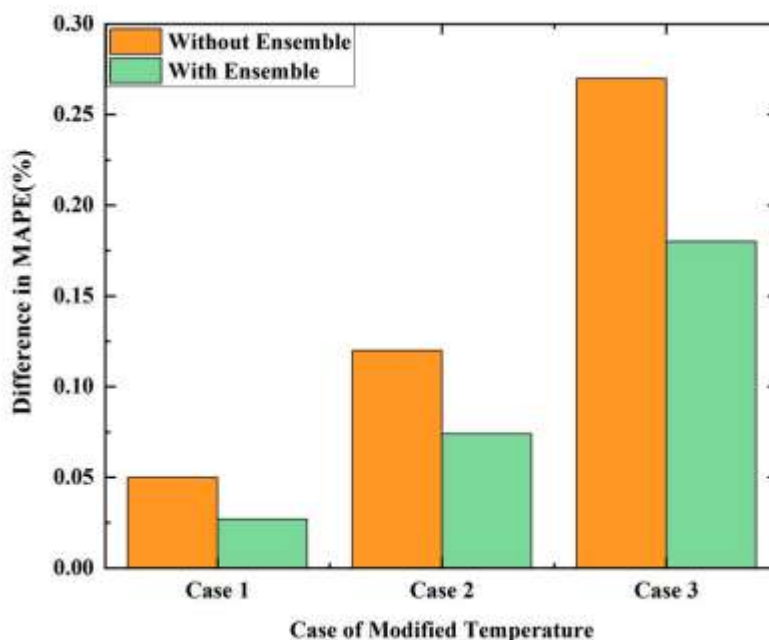


Figura 104: Confronto del modello proposto con ensemble thinking e il modello proposto senza ensemble thinking in diversi casi di modifica della temperatura. Caso 1: media 0 °F e deviazione standard 1 °F; caso 2: media 0 °F e deviazione standard 2 °F; caso 3, media 0 °F e deviazione standard 3°F; MAPE.

Case	With Ensemble	Without Ensemble
Case 1	18.78 s	12.93 s
Case 2	18.80 s	12.91 s
Case 3	18.75 s	12.98 s

Figura 105: Confronto del tempo di calcolo relativo al modello proposto con l'ensemble thinking e il modello proposto senza ensemble thinking.

La previsione del carico svolge un ruolo essenziale nella gestione dei sistemi di alimentazione, e questo può alleviare le limitazioni causate dalla mancanza di energia elettrica. In questo documento, si è progettato un nuovo modello di ensemble con istantanee multiple per prevedere il carico elettrico a breve termine. Ogni snapshot comprendeva più livelli separati completamente connessi, un DCRB e uno strato Bi-LSTM.

Inoltre, i risultati sperimentali indicano che il metodo proposto era di superiore precisione di previsione rispetto ai modelli esistenti e il miglioramento massimo è stato superiore al 3,5% in MAPE. In un lavoro futuro, verranno studiati diversi metodi di deep learning per sfruttare i loro

vantaggi di previsione. Inoltre, simulazioni sull'ottimizzazione dei parametri della rete neurale Bi-LSTM verranno eseguite per costruire un modello più accurato.

L'articolo [61] esplora le prestazioni di tre modelli di machine learning (**Support Vector Regression (SVR)**, **Random Forest** e **Extreme Gradient Boosting (XGBoost)**), tre modelli di deep learning (**Recurrent reti neurali (RNN)**, **Long Short-Term Memory (LSTM)** e **Gated Recurrent Unit (GRU)**) e un classico modello di serie temporali, **Autoregressive Integrated Moving Average (ARIMA)** per la previsione del consumo energetico giornaliero. Il set di dati comprende 8040 record generati in un periodo di 11 mesi da gennaio a novembre 2020 da una struttura logistica non refrigerata situata in Irlanda. È stato utilizzato il metodo di ricerca a griglia per identificare le migliori configurazioni per ciascun modello. L'XGBoost ha sovraperformato altri modelli sia per la previsione del carico a brevissimo termine (VSTLF) che per la previsione a breve termine (STLF); il modello ARIMA ha ottenuto i risultati peggiori.

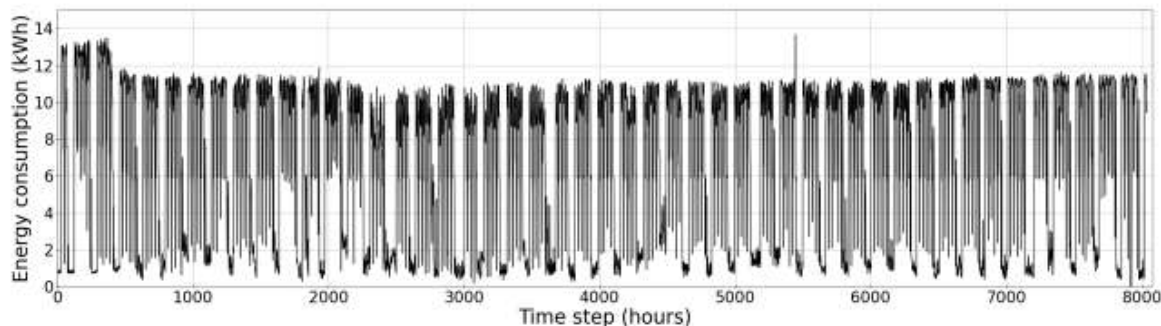


Figura 106: Panoramica delle serie temporali del consumo energetico del sito (gennaio-novembre 2020).

Models	RMSE	MAPE (%)	MAE
ARIMA	0.1114	28.8016	0.0615
SVR-10-rbf	0.1001	42.8923	0.0705
SVR-10-linear	0.1028	42.9745	0.0610
Random Forest-9-200	0.0863	22.771	0.0469
XGBoost-5-100	0.0844	21.659	0.0461
XGBoost-7-100	0.0847	21.573	0.0453
RNN-3-400	0.0947	28.325	0.0592
RNN-4-200	0.0909	29.122	0.0552
LSTM-3-200	0.0912	28.805	0.0553
LSTM-3-300	0.0911	29.03	0.0560
LSTM-4-400	0.0919	28.117	0.0564
GRU-3-100	0.0918	28.94	0.0558
GRU-4-300	0.0933	28.297	0.0585

Figura 107: Risultati RMSE, MAPE e MAE per i vari modelli selezionati.

Models	Prediction (h)	RMSE	MAPE (%)	MAE
XGBoost-5-100	1	0.0844	21.659	0.0461
XGBoost-7-100	1	0.0847	21.573	0.0453
XGBoost-5-100	12	0.1835	21.580	0.1009
XGBoost-7-100	12	0.1717	21.749	0.0926
XGBoost-5-100	24	0.2342	21.603	0.1370
XGBoost-7-100	24	0.2149	21.639	0.1232

Figura 108: Risultati RMSE, MAPE e MAE per i modelli VSTLF e STLF.

Confrontando i risultati ottenuti con i modelli STLF e VSTLF, si è osservato un aumento massimo rispettivamente del 177% e del 202% in RMSE e MAE. Questo aumento può essere spiegato con l'implementazione del modello STLF poiché esiste una dipendenza tra le uscite.

Questo lavoro ha esplorato la previsione del carico elettrico a breve e brevissimo termine per un tipo di edificio poco studiato: i magazzini. Lo studio ha confrontato le prestazioni di modelli di machine learning e deep learning per prevedere il consumo energetico della prossima ora sulla base dei dati delle 48 ore precedenti e ha confrontato questa prestazione rispetto a una classica tecnica di previsione delle serie temporali, ARIMA. A differenza degli studi esistenti, non sono stati considerati i dati solo dal punto di vista di un proprietario e di un operatore di magazzino, ma di una

ESCO che opera nell'ambito di un contratto di rendimento energetico e di dati a disposizione da parte della ditta. I risultati suggeriscono che i modelli XGBoost hanno superato tutti gli altri modelli di apprendimento e apprendimento profondo, nonché ARIMA, per la previsione del carico a brevissimo termine.

Tutti i modelli di machine learning e deep learning hanno superato ARIMA quando si utilizza RMSE come misura; i modelli SVR hanno presentato i peggiori risultati MAPE e MAE. Non sorprende, inoltre, che XGBoost è meno preciso per orizzonti temporali più lunghi, ovvero 12 ore e 24 ore.

Nel documento [62] l'obiettivo è stato quello di implementare un modello predittivo di carico mirato sui singoli consumatori tenendo conto della sua casualità. Per questo scopo, una metodologia è quella di determinare e selezionare le caratteristiche predittive per i singoli STLF. La previsione del carico di un singolo consumatore viene simulata sulla base delle quattro principali tecniche di apprendimento automatico utilizzate nella letteratura. Una riduzione del 2,73% dell'errore di previsione si ottiene dopo la corretta selezione delle caratteristiche predittive. Rispetto al modello di base (metodo di previsione persistente), l'errore è ridotto fino al 19,8%. Tra le tecniche analizzate, la **Support Vector Regression (SVR)** ha mostrato gli errori più piccoli (8,88% e 9,31%).

Feature
Energy Consumption (t-25 to t-1)
Day (t)
Hour (t)
Rush hour (t)
Maximum
Minimum
Average
Median
Quantile (15%)
Range
Variance
Standard deviation
Obliquity
Kurtosis

Figura 109: Insieme di caratteristiche valutate

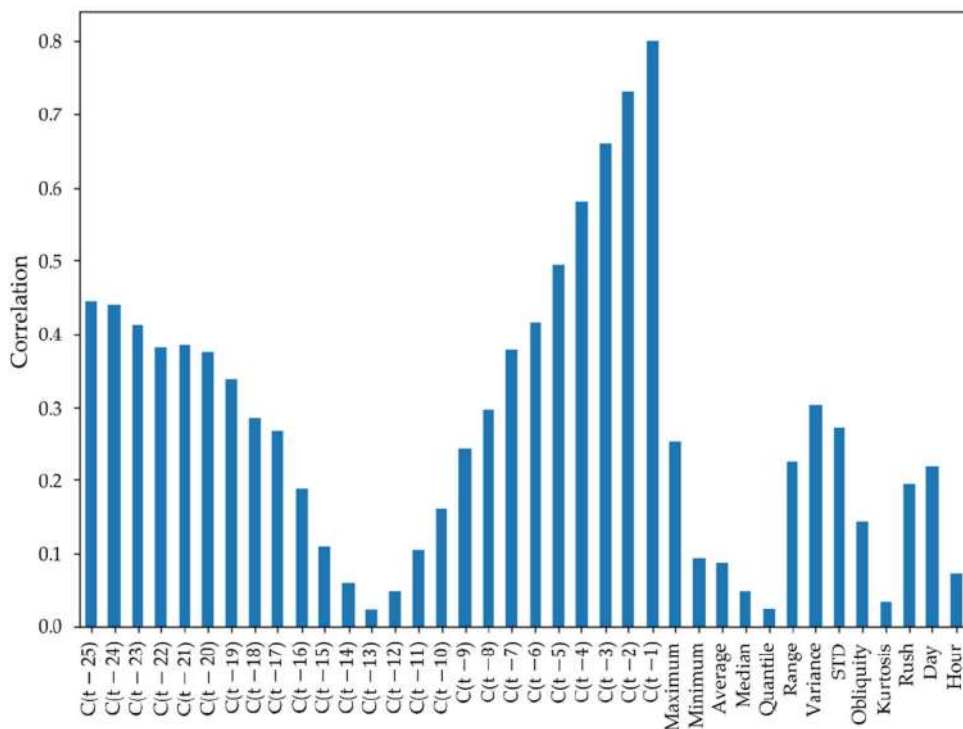


Figura 110: Grafico a barre della correzione di ogni caratteristica con la previsione del consumo energetico

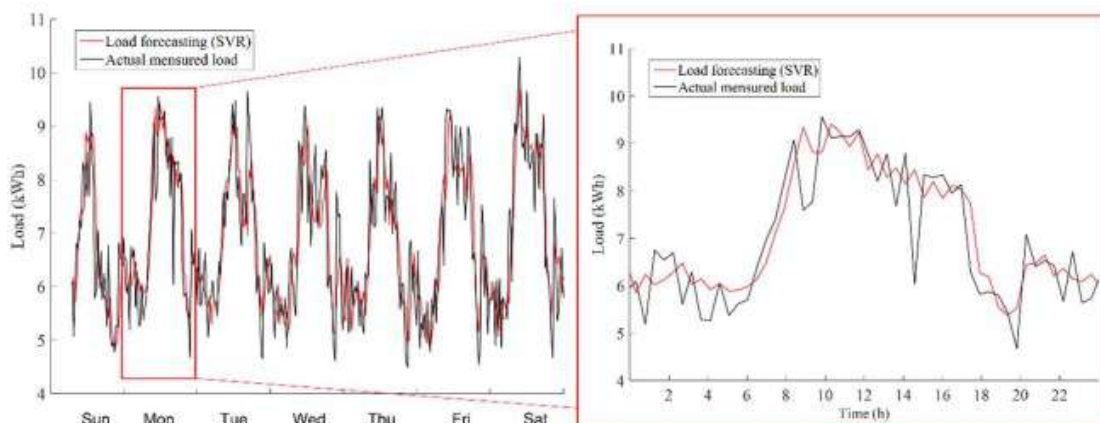


Figura 111: Previsione del carico in base al modello SVR con orizzonte temporale di 30 minuti

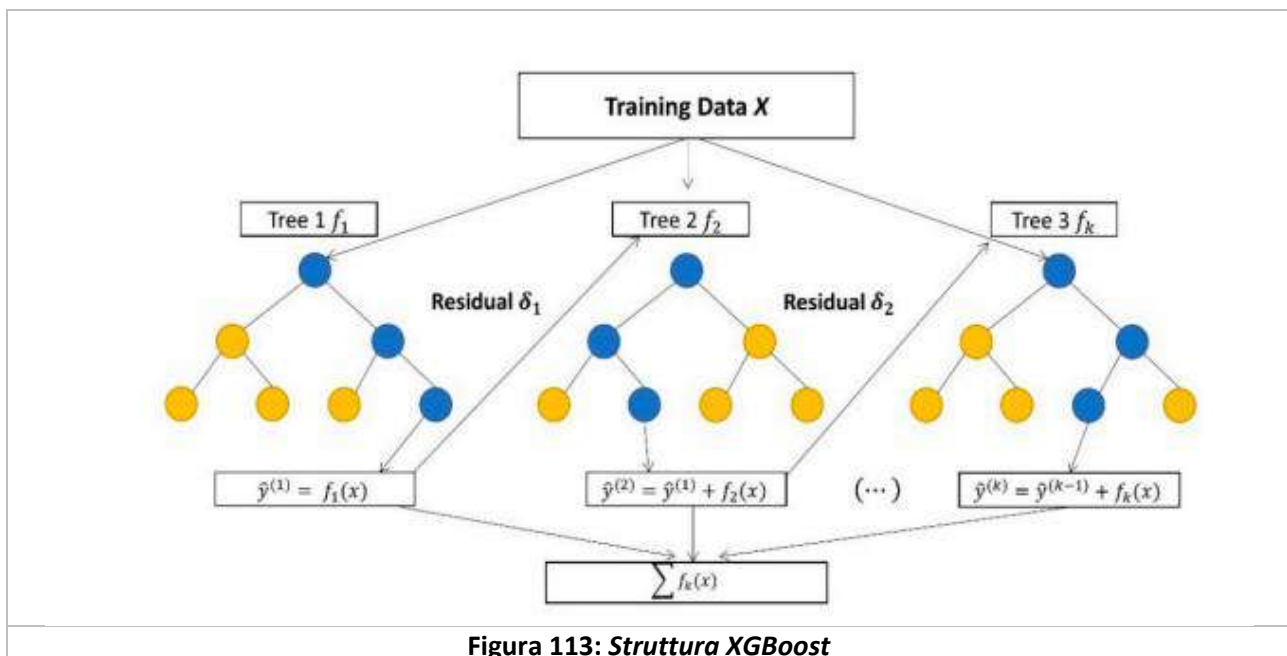
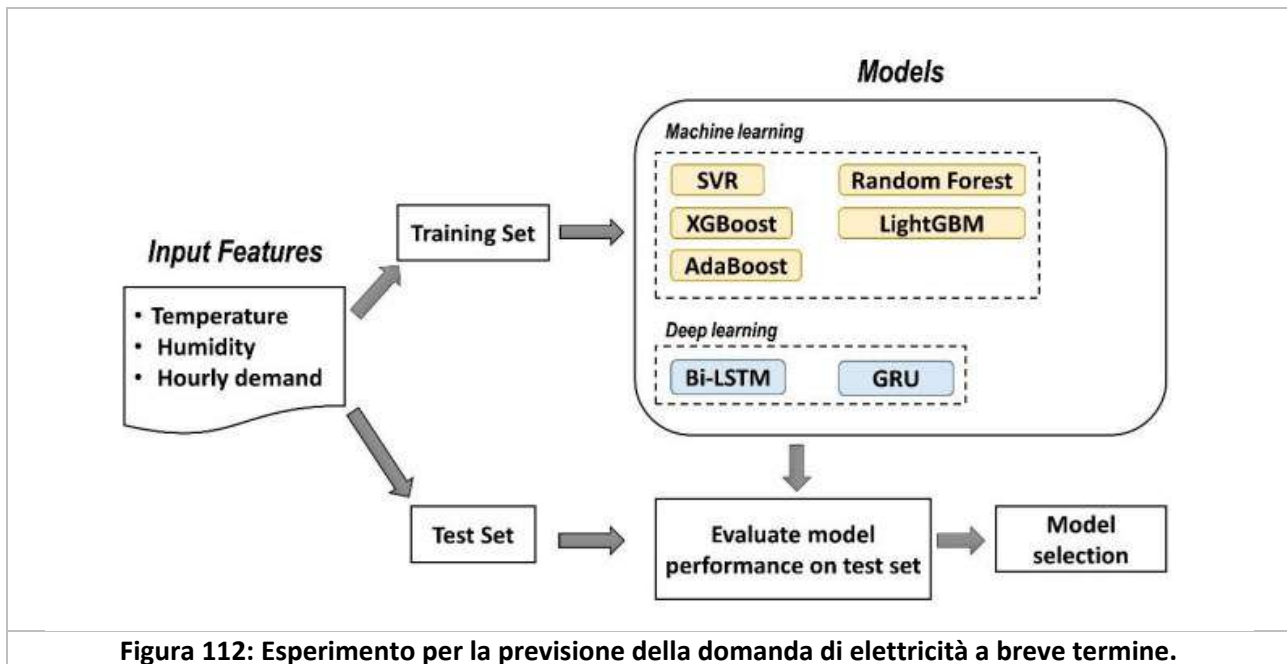
Nel documento, è stata sviluppata una metodologia per estrarre e selezionare il miglior set di attributi per la previsione del carico a breve termine. L'estrazione di attributi da un database è stata effettuata da un'analisi statistica che tiene conto della posizione, della dispersione e della distribuzione della serie storica. Il processo di selezione è stato effettuato per analisi di correlazione ed MI, indicando un insieme di caratteristiche esplicative e indipendenti. La metodologia è stata applicata insieme a quattro delle ben note tecniche di previsione delle serie temporali in cui prevedere il consumo di un singolo consumatore.

I risultati ottenuti indicano un miglioramento nella previsione del carico a breve termine. L'errore di previsione (MAPE) è diminuito del 19,8% e l'adattamento della curva (R2) è aumentato del 36,31% nel caso migliore (SVR) rispetto all'approccio di riferimento.

Considerando solo l'estrazione e la selezione degli attributi e mantenendo costanti le tecniche di previsione, è stato ottenuto un miglioramento dell'1,78% nell'orizzonte di 30 minuti e del 2,37% nell'orizzonte di 6 ore della previsione del carico. In questo modo, si può dire che il dataset selezionato ha contribuito al miglioramento dei modelli finali, rendendo evidente il contributo e l'importanza del trattamento e della selezione delle caratteristiche dal database per una corretta previsione del consumo di energia.

L'articolo [63] si occupa della previsione del carico a breve termine (STLF), diventata un'area di ricerca attiva negli ultimi anni. STLF si occupa di prevedere la domanda da un'ora a 24 ore in avanti. Sono state sperimentate diverse metodologie dell'apprendimento automatico in un caso di studio complesso a Panama. Il **Deep Learning (DL)** è un paradigma di apprendimento più avanzato nel campo dell'apprendimento automatico che continua a dare progressi significativi in aree di dominio come la previsione dell'elettricità, il rilevamento di oggetti, il riconoscimento vocale e tanto altro. Si sono identificati i principali predittori della domanda di elettricità a breve termine: il carico della settimana precedente, il carico del giorno precedente e la temperatura. Si è scoperto che il modello di regressione del deep learning ha ottenuto le migliori prestazioni, che ha prodotto un R2 di 0,93 e

un errore percentuale assoluto medio (MAPE) di 2,9%, mentre il modello **Adaptive Boosting (AdaBoost)** ha ottenuto la peggiore performance con un R2 di 0,75 e MAPE del 5,70%.



Method	R ²	RMSE	MSE	MAPE (%)	Training Time (s)
Bi-LSTM	0.40	70.54	4976.10	4.25	3600
GRU	0.31	73.26	5366.68	4.43	7560
Deep learning regression	0.93	50.34	2534.69	2.90	1200
SVR	0.81	82.38	6786.33	4.90	22
XGBoost	0.91	58.01	3365.01	3.45	3.3
AdaBoost	0.75	94.57	8943.16	5.70	4.8
Random forest	0.92	55.05	3030.90	3.22	7.1
LightGBM	0.92	51.69	2672.08	3.07	0.8

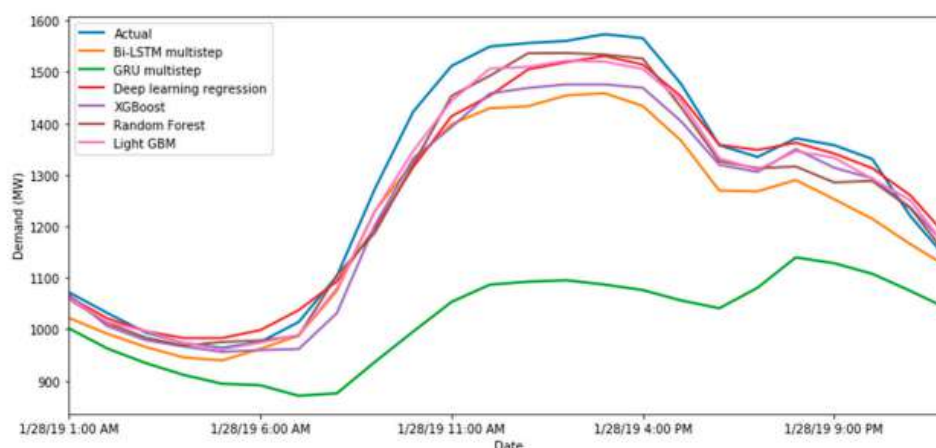


Figura 114: Risultati dei modelli per la previsione elettrica a breve termine (Tabella) e prestazioni sui modelli sul set di prova (Grafico)

Per convalidare la metodologia, questa ricerca ha introdotto un caso di studio su un sistema di alimentazione a Panama. Questo caso di studio è stato significativo per capire dove si trovano i sistemi di alimentazione attualmente, le loro sfide e come stanno iniziando a prepararsi per il futuro. Il caso di studio ha rivelato che i gestori di energia stanno diventando sempre più preoccupati per l'affidabilità della rete.

La metodologia ha innanzitutto affrontato due domande di ricerca: (1) Quali caratteristiche sono le più significative per prevedere la domanda di energia elettrica a breve termine? Inoltre, (2) quali sono i metodi più efficaci per acquisire la domanda oraria? Pertanto, si sono valutate nove funzioni di input per la previsione della domanda oraria. Queste erano il mese, il giorno della settimana, l'ora del giorno, il carico medio delle 24 ore precedenti, l'indicatore giorno lavorativo/fine settimana, la temperatura, l'umidità relativa, il carico alla stessa ora del giorno precedente e il carico alla stessa ora della stessa settimana della settimana precedente. L'analisi dell'importanza delle caratteristiche

basata sul **Random Forest Regression** ha rivelato che le caratteristiche più significative erano il carico allo stesso giorno della settimana precedente alla stessa ora (0,72), il carico alla stessa ora del giorno precedente (0,14) e la temperatura (0,04).

Diversi modelli sono stati proposti per la complessa mappatura non lineare tra le nove caratteristiche di input. Il modello di regressione **Deep Learning (DL)** ha performato al meglio per la predizione della domanda un'ora avanti, con un valore R^2 di 0,93 e un MAPE del 2,90%. Inoltre, il deep learning tende a funzionare meglio se addestrato con set di dati di grandi dimensioni. Pertanto, per migliorare le prestazioni predittive, è stato utilizzato un modello di deep learning consistente di tre strati densi con 95 unità nascoste ciascuno. Sfortunatamente, il modello richiedeva un tempo medio di addestramento di 20 min a causa del numero di livelli nascosti utilizzati. Tra i modelli di deep learning, il **modello multi-step Gated Recurrent Unit (GRU)** ha ottenuto i risultati peggiori perché utilizza una lunga sequenza di input di 72 h per prevedere la domanda di elettricità 24 h in avanti, che può portare al problema del gradiente evanescente. Così, è diventato evidente che la previsione in più fasi rimane un'area di ricerca impegnativa. Lo studio ha anche scoperto che **Adaptive Boosting (AdaBoost)** aveva la peggiore performance tra i modelli di machine learning, con un valore R^2 di 0,75 e MAPE del 5,70%.

Il modello proposto nell'articolo [64] utilizza la **Voting Regression (VR)**, producendo previsioni con medie ponderate delle previsioni da cinque modelli individuali: tre regressori lineari multipli parametrici e due modelli non parametrici di apprendimento automatico. I regressori sono modelli di regressione lineare con gradiente discendente (LR), stimatori dei minimi quadrati ordinari (OLS) e modelli di auto-regressione generalizzata dei minimi quadrati (GLSAR). Al contrario, i modelli di apprendimento automatico sono i **Decision Trees (DT)** e **Random Forest (RF)**. Per selezionare le variabili e gli iperparametri del modello migliori, si sono utilizzate le prestazioni di convalida incrociata (CV), invece delle prestazioni dei dati di test, che hanno prodotto prestazioni di test eccessivamente buone. Si sono confrontati vari schemi di convalida e si è scoperto che lo schema Blocked-CV fornisce l'errore di convalida più vicino all'errore di prova. Utilizzando Blocked-CV, i risultati del test mostrano che il modello VR supera tutti i predittori individuali.

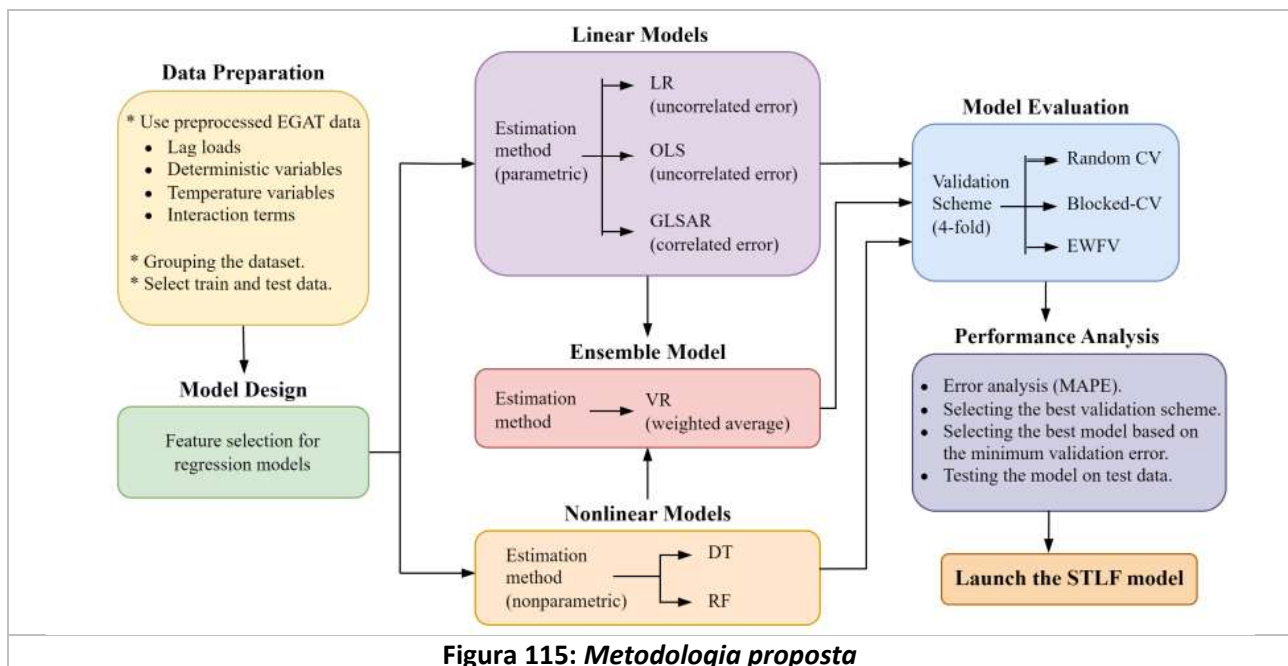


Figura 115: Metodologia proposta

Group #	Validation Scheme	LR	OLS	GLSAR (rho = 1)	DT	RF	VR
1	Random CV	2.0111×10^6	0.7687	0.7093	0.6645	0.4204	0.7440
	Blocked-CV	0.7058	0.7057	0.6481	0.6643	0.4258	0.6819
	EWFV	3.4955×10^6	1.2838	1.223	0.6788	0.6917	6.9911×10^5
2	Random CV	0.7443	0.7437	0.7613	0.6533	0.3775	0.7518
	Blocked-CV	0.7345	0.7344	0.7334	0.6059	0.3119	0.7207
	EWFV	0.9988	0.9971	0.9873	0.5640	0.3023	0.8853
3	Random CV	3.7437	3.744	3.9505	1.6675	0.8918	3.6341
	Blocked-CV	2.7438	2.7446	2.7855	1.7104	0.6485	2.7701
	EWFV	5.3703	5.3706	5.4704	1.0700	0.6997	5.1507
4	Random CV	1.8061×10^6	0.7761	0.6569	0.5045	0.1410	8.1892×10^3
	Blocked-CV	0.7303	0.7253	0.6062	0.5528	0.1951	0.6361
	EWFV	3.7727×10^6	1.2577	1.1346	0.4116	0.3069	7.5455×10^5

Figura 116: Differenze tra validazione MAPE e test MAPE per tutti i gruppi

Il set di dati fornito da EGAT è stato diviso in quattro gruppi e in mezz'ora per comodità di modellazione. Sono stati implementati tre diversi schemi di convalida del modello per valutare i modelli e il Blocked-CV è stato quindi selezionato come lo schema migliore in base ai risultati. Il modello di Ensemble Learning proposto ha prodotto i valori minimi di Blocked-CV con MAPE di 2,5636%, 2,665%, 7,9559% e 3,0910% rispettivamente per i gruppi 1, 2, 3 e 4; e i valori minimi del test MAPE di 1,8817%, 1,9458%, 5,1858% e 2,4549% per i gruppi 1, 2, 3 e 4. I valori MAPE per la

previsione dei carichi festivi (Gruppo 3) sono piuttosto elevati a causa della quantità limitata di dati disponibili. Tuttavia, sono ancora competitivi rispetto ai modelli della ricerca precedente. Il raggruppamento del set di dati di giorni feriali (Gruppo 1) e di fine settimana (Gruppo 2) si sono dimostrati validi, dal momento che ne è risultata una diminuzione di valori sia nel MAPE di convalida che nel MAPE di test rispetto al set di dati completo (Gruppo 4). Inoltre, per i dati, quando si considera l'accuratezza delle previsioni dei predittori individuali, tutti i modelli regressivi parametrici basati su MLR hanno superato quelli non parametrici.

Nell'articolo [65] viene eseguita la previsione a breve termine del carico dell'impianto di fertilizzanti Yara India Private Limited installato a Babrala, Uttar Pradesh, utilizzando una **multi-layer Feed Forward Neural Network (FFNN)** implementata in MATLAB. L'algoritmo utilizzato è l'algoritmo di Levenberg Marquardt. Tuttavia, la formazione e i risultati ottenuti dall'ANN sono molto veloci e precisi. Gli input forniti alla rete neurale sono il tempo, la temperatura dell'aria, l'ambiente del compressore, la temperatura dell'aria fresca al compressore e l'apertura dell'IGV. La necessità, i benefici e la crescita delle centrali elettriche captive in India e l'uso di **Artificial Neural Network (ANN)** per la previsione a breve termine del carico delle centrali elettriche captive sono spiegati dettagliatamente nel lavoro.

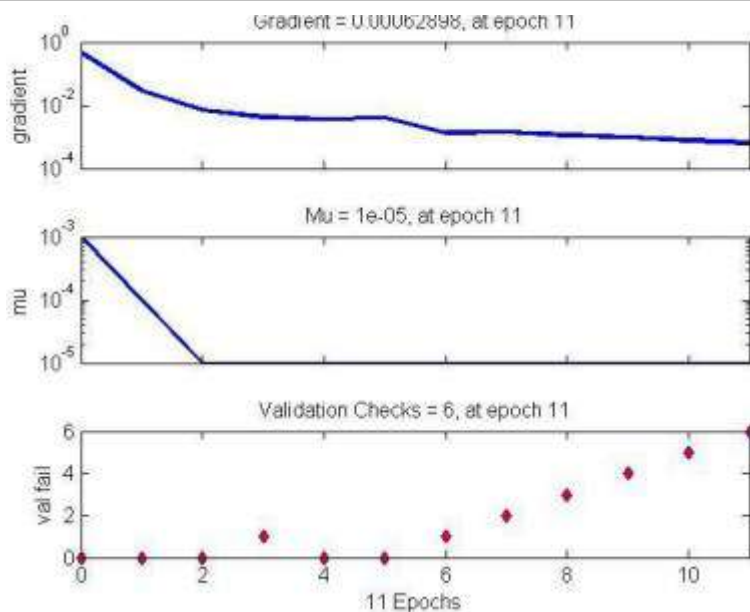


Figura 117: Grafici dello stato di addestramento.

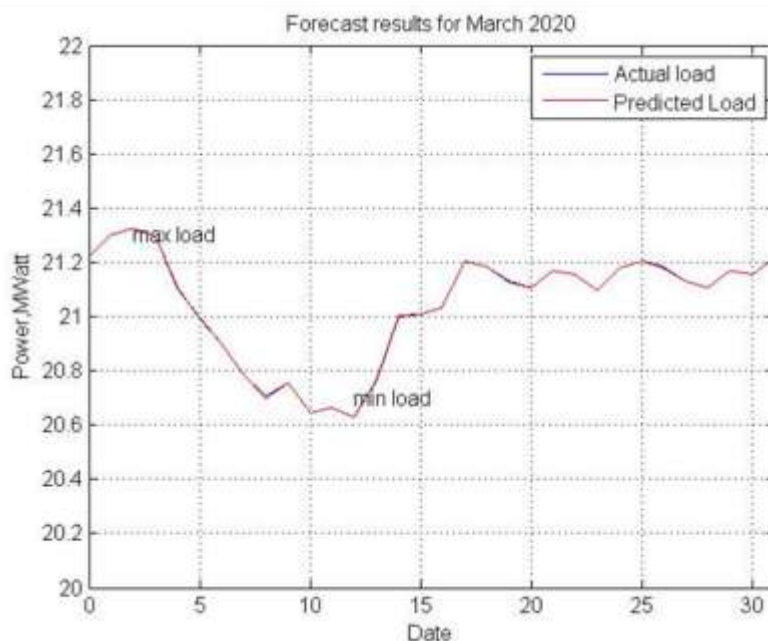


Figura 118: Risultati della predizione per il mese di marzo.

I dati sono stati raccolti per tutte le 24 ore. I modelli includono caratteristiche come carico (MW), data, ora, temperatura dell'aria ambiente dal compressore, temperatura dell'aria fresca al compressore e apertura dell'IGV. La rete progettata è addestrata con un algoritmo supervisionato

(algoritmo di retropropagazione). L'implementazione di questa rete sta mostrando risultati ragionevolmente buoni con un errore percentuale medio assoluto (MAPE) dell' 1,210%, che risulta molto basso. L'errore sarà inferiore utilizzando un grande set di dati con una variazione molto bassa. Il risultato del modello ANN riflette una buona performance di previsione con un errore minimo. Pertanto, la previsione del carico a breve termine utilizzando l'ANN è un metodo efficace per la previsione del carico.

3.1.2 Medium-term load forecast (MTLF)

Il lavoro proposto nell'articolo [66] considera la questione della costruzione di un modello di previsione a medio termine dei grafici di carico per EPS con specifiche proprietà, basate sull'uso di metodi di apprendimento automatico dell'insieme. Questo documento implementa l'approccio di identificazione delle funzionalità più significative per applicare modelli di machine learning per la previsione del carico a medio termine in un sistema di alimentazione isolato.

È stato condotto uno studio comparativo dei seguenti modelli: **Linear Regression (LR)**, **Support Vector Regression (SVR)**, **Decision Trees (DT)**, **Random Forest (RF)**, **Extreme Gradient Boosting (XGBoost)**, **Adaptive Boosting (AdaBoost)**, **AdaBoost su regressione lineare**. L'isolamento delle funzionalità da una serie temporale consente l'implementazione di modelli più semplici e resistenti all'overfitting. Tutto ciò consente di aumentare l'affidabilità delle previsioni e di espandere l'uso delle tecnologie dell'informazione nella pianificazione, gestione e funzionamento di EPS isolati. Calcoli dell'errore totale di previsione hanno dimostrato che le caratteristiche dei modelli proposti sono di alta qualità e accurate, e quindi possono essere utilizzate per prevedere il carico reale di un sistema di alimentazione.

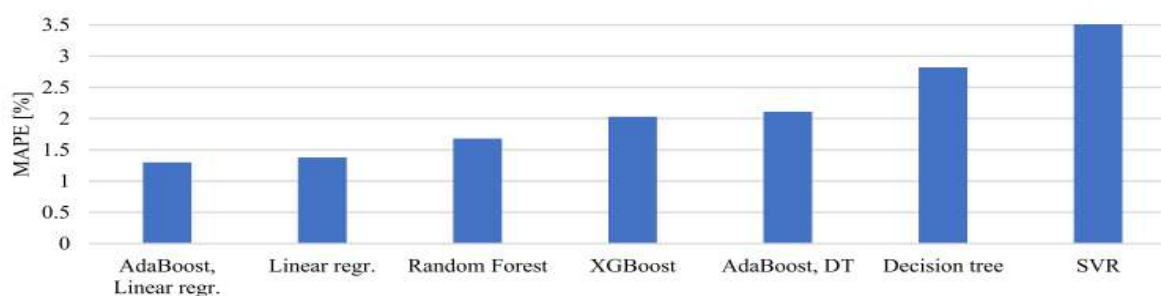


Figura 119: Risultati per il dataset di test

Model	MAPE on the training set, %	MAPE on validation set, %	MAPE on the test set, %
Linear regression	1.29	1.40	1.38
SVR	2.89	3.28	3.52
Decision tree	0.45	1.96	2.82
Random Forest	0.73	1.68	1.68
XGBoost	0.71	1.66	2.03
AdaBoost over decision trees	0.56	1.64	2.11
AdaBoost over linear regression	1.32	1.40	1.30

Figura 120: Risultati per modello

Questo studio presenta un metodo di previsione che si basa sulla selezione delle caratteristiche più significative delle serie temporali e sui modelli di apprendimento automatico dell'insieme, che consente di prevedere il carico nell'EPS isolato di GBAO per una settimana in avanti. Sono stati analizzati sette diversi modelli per prevedere il carico dell'area oggetto di studio sulla base dei dati di consumo del carico precedente. Il miglior risultato è stato ottenuto con un modello ensemble (AdaBoost) che combinato quattro regressioni lineari in un modello.

I risultati ottenuti nel corso di questo studio possono essere utilizzati per migliorare la qualità della previsione del carico del GBAO EPS quando si prendono decisioni informate sulla struttura del carico nella regione. Inoltre, i metodi proposti possono essere consigliati per altre società di alimentazione che gestiscono soccorritori isolati. L'uso dell'ensemble forecasting model è finalizzato ad aumentare l'affidabilità delle previsioni al fine di garantire alta qualità e potenza affidabile ai consumatori in EPS isolati simili.

Il documento [67] affronta il problema della previsione del carico elettrico a medio termine. Risolvere questo problema è necessario per il funzionamento e la pianificazione del sistema di alimentazione, nonché per la negoziazione di contratti a termine nei mercati dell'energia liberalizzati. Si dimostra che l'approccio di modellazione della rete neurale profonda proposto basato sull'architettura neurale profonda è efficace nel risolvere il problema di previsione del carico elettrico a medio termine.

La **Neural Network** proposta ha un elevato potere espressivo per risolvere problemi di previsione stocastica non lineare con serie temporali che includono tendenze, stagionalità e significative fluttuazioni casuali. Allo stesso tempo, è semplice da implementare e addestrare, non richiede la pre-elaborazione del segnale ed è dotato di un meccanismo di riduzione del bias di previsione. Si confronta l'approccio rispetto a dieci linee di base di metodi statistici classici, apprendimento automatico e approcci ibridi, su 35 serie temporali mensili della domanda di elettricità per i paesi europei. Lo studio empirico mostra che la rete neurale proposta supera chiaramente tutte le prestazioni concorrenti in termini sia di accuratezza che di bias di previsione.

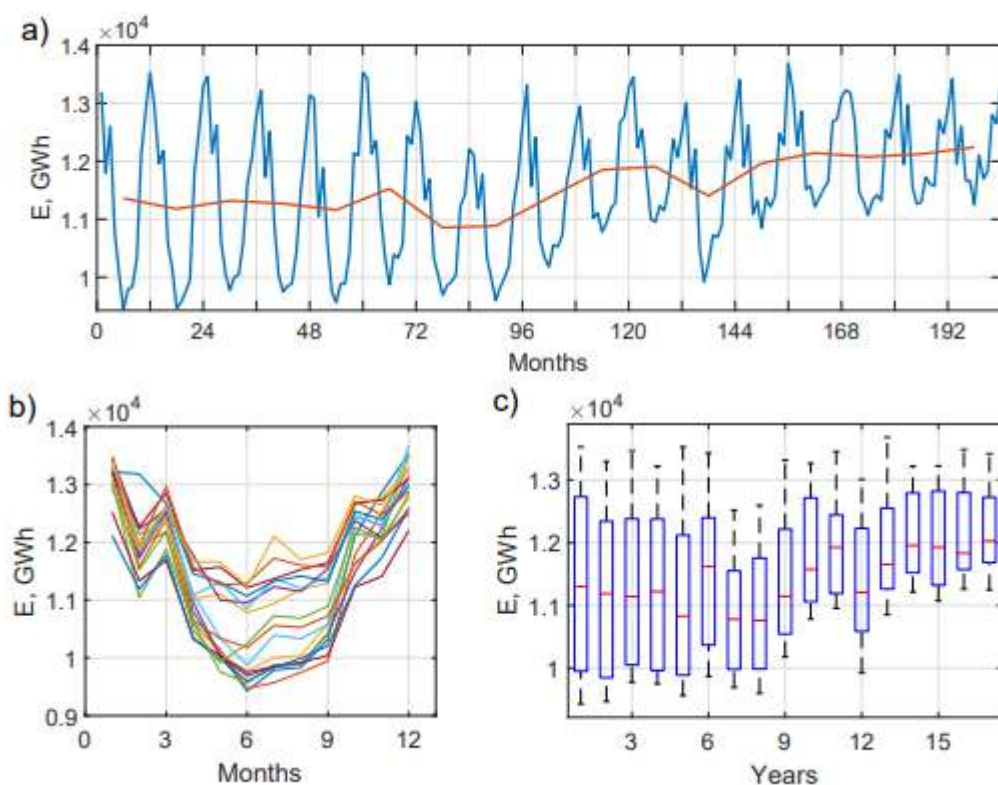


Figura 121: Serie storiche della domanda mensile di elettricità per la Polonia (a), suoi andamenti annuali (b) e la dispersione negli anni successivi (c)

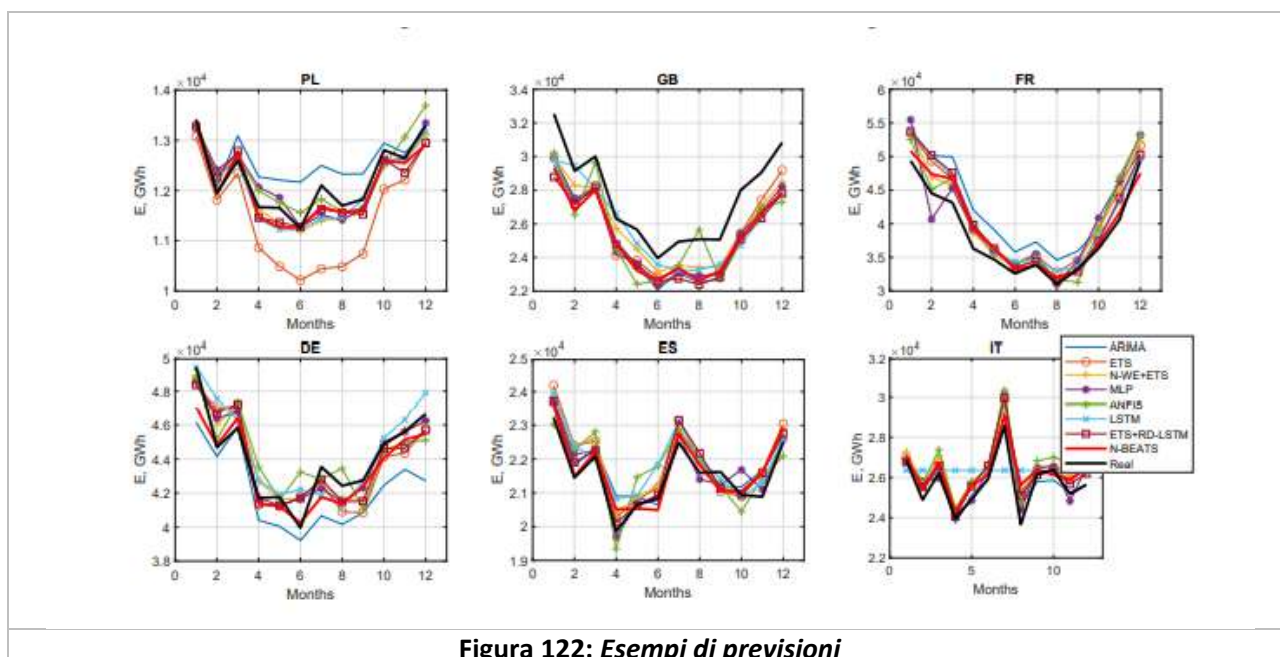
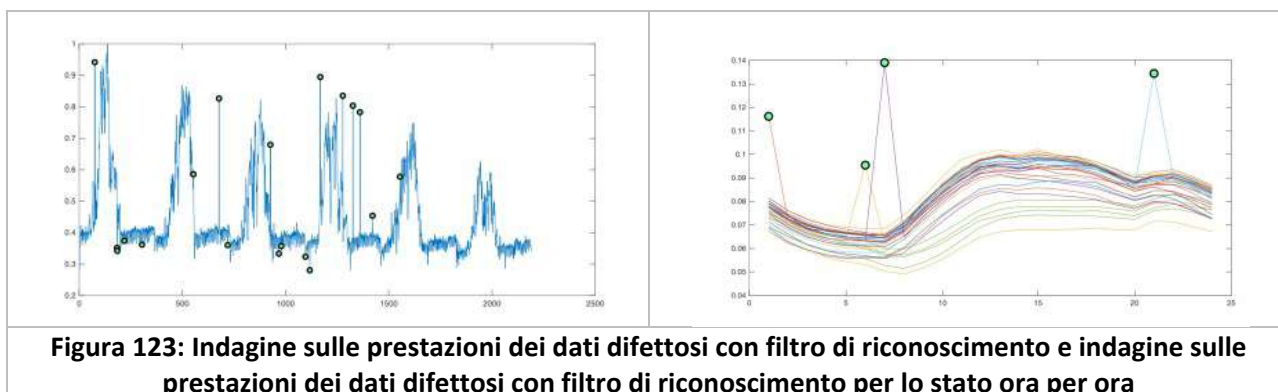


Figura 122: Esempi di previsioni

La previsione del carico a medio termine considerata in questo lavoro è una sfida al problema che richiede che il modello previsionale sia altamente flessibile e gestibile con la non stazionarietà e stagionalità. In questo studio, si è proposto e convalidato empiricamente una nuova architettura per il problema di previsione del carico elettrico a medio termine che risponde a queste aspettative: la **Neural Network N-BEATS**. Lo studio empirico della previsione del carico elettrico a medio termine per 35 paesi europei ha mostrato che N-BEATS ha ottenuto le migliori prestazioni rispetto ai metodi statistici, di machine learning e ibridi. N-BEATS ha chiaramente sovraperformato la concorrenza in termini sia di accuratezza che di distorsione delle previsioni. Per quanto riguarda la percentuale assoluta di errore medio, il nostro metodo ha fornito un guadagno relativo del 25% rispetto ai metodi statistici, 28% rispetto ai metodi di apprendimento automatico e 14% rispetto ai metodi ibridi avanzati e all'avanguardia. Il suo successo è dovuto a un'unica architettura che combina una pila profonda di strati completamente connessi, backward e forward link residuali, aggregazione delle previsioni parziali in un flash ion gerarchico ed ensembling. L'apprendimento incrociato, ovvero l'apprendimento su più serie temporali, consente a N-BEATS di catturare le caratteristiche e le componenti condivise delle serie temporali individuali.

Nello studio proposto nell'articolo [68] si sviluppa un approccio basato sulla **Neural Network (NN)** è progettato per la previsione del carico a medio termine (MTLF). La struttura e gli iperparametri sono sintonizzati per ottenere la migliore accuratezza di previsione su un anno in avanti. L'approccio suggerito è praticamente applicato a una regione dell'Iran mediante l'uso di set di dati del mondo reale di 10 anni. I fattori influenti come i fattori economici, meteorologici e sociali vengono studiati e il loro impatto sulla precisione viene analizzato numericamente. I dati errati vengono rilevati da un metodo efficace suggerito. Oltre al picco di carico, viene previsto anche il modello di carico di 24 ore, che aiuta ad una migliore pianificazione a medio termine. Le simulazioni mostrano che l'approccio suggerito è pratico e l'accuratezza è addirittura superiore al 95% quando ci sono drastici cambiamenti climatici.



Questo studio presenta una nuova direzione pratica per MTLF. In primo luogo, le variabili e i fattori influenti sulla precisione sono analizzati mediante studi numerici e calcolando le correlazioni. Per gli studi numerici, in 16 scenari, gli input di NN sono considerati condizioni economiche, condizioni meteorologiche, e diversi fattori sociali. Gli ingressi più affidabili vengono selezionati analizzando i vari fattori. Poi per migliorare la precisione, anche la struttura NN è sintonizzata per i pesi di NN. Inoltre, si suggerisce un approccio semplice per rilevare dati insufficienti e ridurre l'impatto sull'apprendimento. Le simulazioni mediante l'uso di set di dati reali verificano l'applicabilità dello schema progettato.

L'obiettivo principale dello studio proposto nell'articolo [69] è sviluppare e confrontare precisi modelli a livello distrettuale per la previsione della domanda di carico elettrico basati sulle tecniche di apprendimento automatico, tra cui **Support Vector Machine (SVM)** e **Random Forest (RF)** e **metodi di deep learning come la Nonlinear AutoRegressive network with eXogenous inputs (NARX)** e **le reti neurali Long Short-Term Memory (LSTM)**. Un set di dati che include nove anni di carico storico per Bruce County, Ontario, Canada, fuso con le informazioni climatiche (temperatura e velocità del vento), viene utilizzato per addestrare i modelli dopo aver completato le fasi di preelaborazione e pulizia.

I risultati mostrano che utilizzando il deep learning, il modello potrebbe prevedere maggiormente la domanda di carico in modo più accurato rispetto a SVM e RF, con un R-quadrato di circa 0,93–0,96 e un MAPE di circa il 4-10%. Il modello può essere utilizzato non solo dai comuni, ma anche da società di servizi pubblici e distributori di energia nella gestione e nell'ampliamento delle reti elettriche.

Model	RMSE		R ²		Computation Time (min)
	Mean $\pm$ SD	Range	Mean $\pm$ SD	Range	
LSTM	9.569 $\pm$ 3.050	[7.16,15.00]	0.921 $\pm$ 0.0341	[0.85,0.95]	93
NARX	6.987 $\pm$ 1.433	[5.52,9.75]	0.953 $\pm$ 0.0142	[0.93,0.98]	16
SVM	11.878 $\pm$ 2.447	[7.89,15.22]	0.886 $\pm$ 0.0345	[0.84,0.94]	244
RF	13.028 $\pm$ 3.483	[10.05,21.49]	0.853 $\pm$ 0.0489	[0.78,0.91]	50

Figura 124: Statistiche delle metriche di valutazione (accuratezza ed errore) per ciascun modello basate su una convalida incrociata di 10 volte

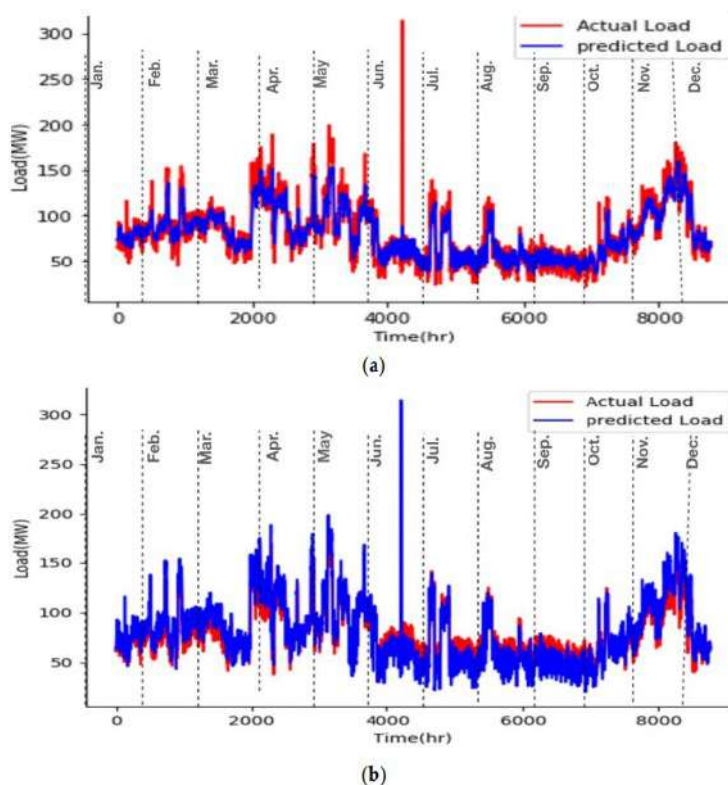


Figura 125: Modelli di previsione, carico previsto vs. dati osservati (per il 2019). (a) Modello di regressione RF; (b) modello SVR.

In questo articolo ci si è concentrati su varie tecniche di machine learning, ensemble learning e deep learning per prevedere la richiesta di energia elettrica con risoluzione oraria a scala regionale. RF, SVM, LSTM e la rete neurale NARX sono considerate quattro tecniche che vanno dal semplicistico al complesso in ottica computazionale. Come dimostrato dall'esperimento, tutti e quattro i modelli hanno mostrato fattibilità nella previsione e possono prevedere con precisione le tendenze a vari gradi. Il rumore disponibile nei dati storici (come i valori negativi per il consumo di elettricità) è stato gestito nella fase di pre-elaborazione. Alcune fluttuazioni della domanda, che statisticamente appaiono come rumore, non sono state rimosse dai dati di addestramento, così com'erano ripetute frequentemente nel corso di vari anni; quindi, si riteneva che riflettessero la natura del consumo elettrico della regione. Il modello LSTM è stato ottimizzato parametricamente per inibire i cambiamenti improvvisi una tantum (a causa di estremi ambientali imprevedibili o improvvise interruzioni di corrente) nel processo di apprendimento, mantenendo gli effetti della stagionalità. È stato dimostrato da questo studio che la complessità dei moduli LSTM potrebbe avere successo per

distinguere i due diversi comportamenti. Pertanto, mentre gli estremi stagionali potevano essere previsti con successo dal modello di deep learning, il modello non è stato distratto dalle fluttuazioni ad hoc.

Nell'articolo [70] si prende un "distribution transformer" come oggetto di ricerca e si propone un metodo di previsione di carico a medio e lungo termine per oggetti di gruppo basato su **Image Representation Learning (IRL)**.

In primo luogo, i dati del trasformatore di distribuzione vengono pre-elaborati per ripristinare la variazione di carico allo stato naturale. E poi, il processo di previsione del carico viene disaccoppiato in due parti: la previsione dell'andamento del carico del prossimo anno e la previsione numerica del tasso di scambio del carico. In secondo luogo, cambiano le immagini di caricamento relative ai dati annuali e interannuali su cui le informazioni sono costruite. Nel frattempo, una previsione dell'apprendimento della rappresentazione del modello di immagine basato sulla **Convolutional Neural Network (CNN)** viene utilizzato per prevedere il trend di sviluppo del carico, ottenendo le immagini di caricamento per l'addestramento. Basato sui dati di previsione e il risultato della previsione dell'andamento del carico, il modello di previsione del gruppo corrispondente ai dati di previsione può essere selezionato per realizzare la previsione numerica del tasso di variazione del carico. A causa dell'elevato numero di oggetti predittivi, questo documento introduce l'indice di valutazione della previsione di gruppo per misurare l'effetto previsionale di metodi diversi. Infine, i risultati sperimentali mostrano che, rispetto agli esistenti metodi di previsione del trasformatore di distribuzione, il metodo proposto in questo documento ha un migliore effetto di previsione generale e fornisce una nuova idea e soluzione per la previsione intelligente del carico a medio e lungo termine della rete di distribuzione.

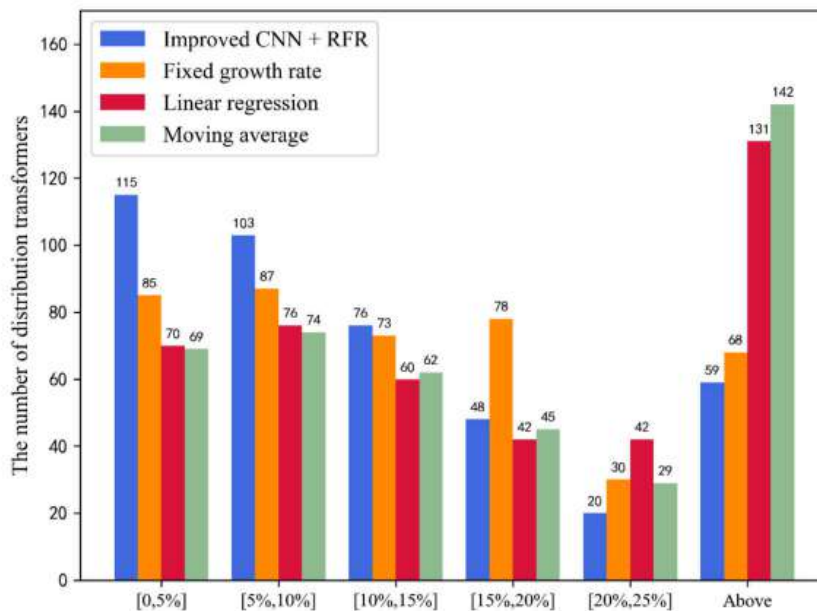


Figura 126: Il numero di trasformatori di distribuzione in diversi intervalli di errore

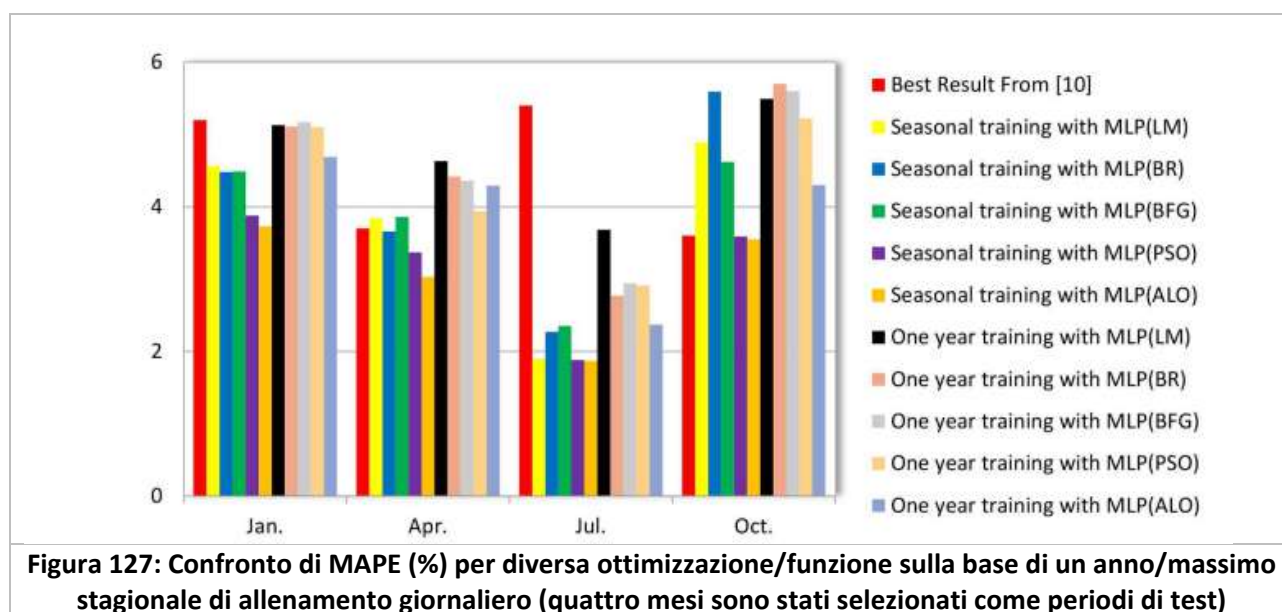
Il metodo in questo documento ha i seguenti vantaggi:

- 1) Viene proposta una nuova idea di previsione del disaccoppiamento per realizzare la suddivisione del compito di previsione. Si possono migliorare i metodi di previsione utilizzati in diversi collegamenti in base a diversi obiettivi di previsione, o sostituire i metodi esistenti con altri metodi che possono raggiungere gli stessi obiettivi di previsione con condizioni di dati diversi, in modo da realizzare la combinazione flessibile di metodi di previsione e migliorare, appunto, l'effetto di previsione.
- 2) Può integrare e utilizzare completamente i dati di carico del quotidiano e dimensioni annuali, e utilizzare il modello di apprendimento di rappresentazione dell'immagine per migliorare la capacità di previsione dell'andamento del carico del trasformatore di distribuzione.
- 3) Sulla base della stazionarietà del cambiamento complessivo degli oggetti di gruppo, i dati di addestramento vengono espansi attraverso la classificazione della forma dei dati. Il metodo di apprendimento automatico del **Random Forest (RF) regression** è introdotto per stabilire un modello di previsione di popolazione di oggetti multipli, che risolve il problema dei dati che limita l'applicazione del metodo di apprendimento automatico alla previsione a medio e lungo termine

delle reti di distribuzione. Allo stesso tempo, può realizzare la rapida previsione di un grande numero di oggetti di trasformazione della distribuzione e riduce la soggettività della previsione.

4) Rispetto ad altre previsioni a medio e lungo termine, il metodo proposto utilizza meno funzionalità di dimensioni di dati e richiede meno tempo di registrazione. Impostando il processo di previsione ragionevolmente, l'effetto di previsione è migliore rispetto agli attuali metodi di applicazione ingegneristica.

Nell'articolo [71] è proposto un nuovo metodo composito basato su una **rete neurale Multilayer Perceptron (MLP)** e tecniche di ottimizzazione per risolvere il problema MTLF. Il metodo proposto ha un algoritmo di addestramento ottimale composto da due algoritmi di ricerca e una rete neurale perceptron multistrato. L'accuratezza del metodo di previsione è ampiamente valutata sulla base di diversi set di dati di riferimento.

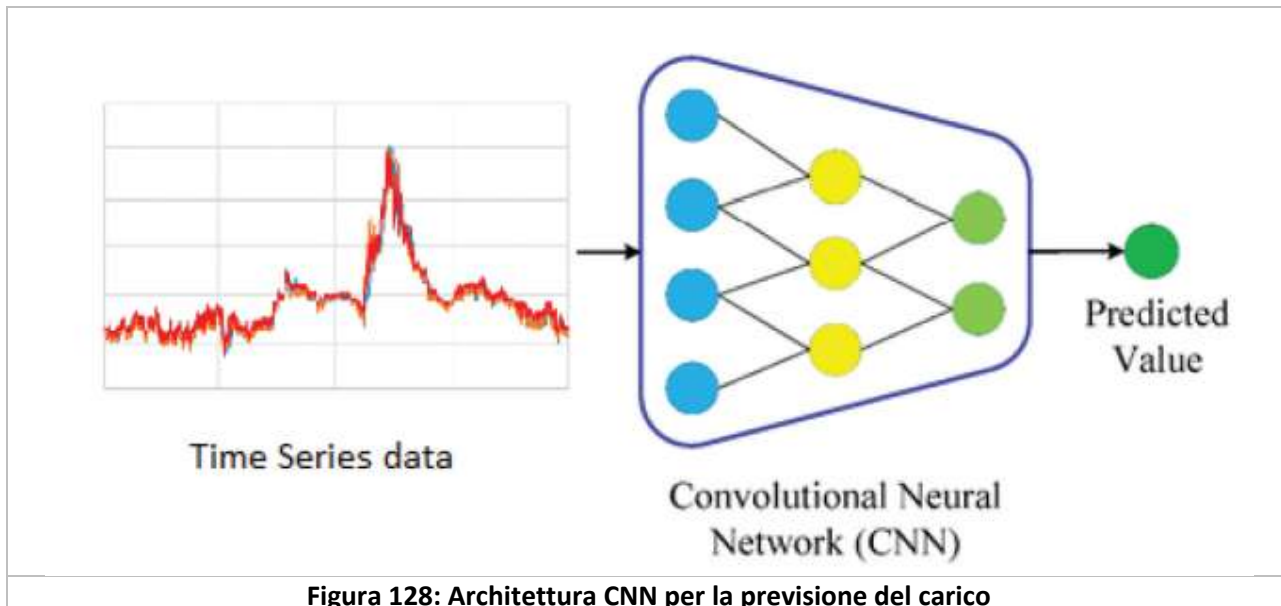


Uno degli obiettivi importanti di MTLF è la pianificazione nei sistemi di alimentazione. D'altra parte, la programmazione delle ore di punta è molto importante per il funzionamento e l'affidabilità del sistema. Pertanto, per il mensile sono stati utilizzati l'orizzonte di previsione e la fase di previsione giornaliera (picco di carico giornaliero) per MTLF. In questo studio, l'ALO (Ant Lion

Optimisation)/PSO (Particle Swarm Optimisation) è stato proposto come trainer stocastico al MLP. Il problema della formazione del MLP è stato formulato per l'algoritmo ALO/PSO per trovare i valori ottimali per i pesi. Altri obiettivi in questo studio sono stati il determinare i dati ottimali per applicare **la Neural Network** e per creare la struttura ottimale nella rete neurale in modo che l'errore di convalida sia minimizzato. Ottimizzando i dati di input aumentando al contempo la precisione delle previsioni, si riduce anche il tempo di calcolo del problema. Creare la struttura ottimale della rete neurale è in realtà la determinazione del numero di neuroni nello strato di mezzo che è considerato come una variabile decisionale nell'ALO e negli algoritmi di ottimizzazione PSO. Secondo lo studio, gli algoritmi ALO e PSO producono risultati promettenti grazie all'elevata esplorazione e sfruttamento dell'ALO/PSO ma i risultati di ALO sono migliori di PSO e tra i metodi studiati di addestramento stocastico come formazione stagionale/annuale con gli algoritmi ALO o PSO, il modello di addestramento stagionale ha ottenuto il più basso MSE e MAPE in due casi di studio esaminati. Infine, i risultati di questo studio dal primo set di dati sono stati confrontati con i migliori risultati in un altro articolo e hanno dimostrato che il metodo proposto è stato utilizzato in questo studio per dare risultati migliori e più accettabili. Inoltre, il carico di picco è stato utilizzato per testare maggiormente e dimostrare l'efficacia del metodo combinato.

Nell'articolo [72], invece, si mira a fornire una previsione di carico efficiente di un sistema progettato per diversi alimentatori locali per la previsione del carico a medio e lungo termine. Questo modello migliora i requisiti futuri per espansioni, attrezzature, vendita al dettaglio o assunzione di personale presso aziende elettriche. Si è puntato a migliorare le previsioni anticipate utilizzando l'approccio ibrido e ottimizzando i parametri dei modelli. Si è utilizzato **Long Short-Term Memory (LSTM), Convolutional Neural Network (CNN), Multilayer Perceptron (MLP) e metodi ibridi**. Si è utilizzato l'errore quadratico medio (RMSE), il MAPE, l'errore assoluto medio (MAE) e l'errore quadratico per il confronto. Nella previsione del carico a 3 mesi, l'errore di previsione più basso è stato acquisito utilizzando LSTM MAPE (2,70). Per 6 mesi in avanti sulla previsione della previsione, MLP dà il massimo con MAPE (2.36). Inoltre, per prevedere la previsione del carico con 9 mesi di anticipo, la previsione più alta è stata ottenuta utilizzando LSTM in termini di MAPE (2.37). Allo stesso modo, il MAPE avanti di 1 anno (2.25) è stato prodotto utilizzando LSTM e il MAPE (2.49) per 6 anni avanti è

stato fornito utilizzando MLP. I metodi proposti raggiungono prestazioni stabili e migliori per la previsione della previsione del carico.



Method	R^2	MAPE	MAE	MSE	RMSE
MLP	0.9901	2.30	313.63	179418.40	423.58
LSTM	0.9905	2.25	304.34	172385.71	415.19
CNN	0.9732	3.81	512.08	485420.82	696.72
CNN+LSTM	0.9850	2.75	372.47	271644.37	521.19

Figura 129: Previsione del prossimo anno di previsione del carico

Method	R^2	MAPE	MAE	MSE	RMSE
MLP	0.9815	2.49	311.64	222351.94	471.54
LSTM	0.9619	3.03	416.75	458797.66	677.34
CNN	0.9283	5.42	673.84	862976.07	928.96
CNN+LSTM	0.9533	3.51	471.05	561931.59	749.62

Figura 130: Previsione della previsione del carico di sei anni

Dal momento che viene utilizzato nel processo decisionale e nelle operazioni della rete elettrica, la previsione del carico elettrico è importante. La previsione accurata del carico elettrico farà assistere gli operatori nello sviluppo di una strategia aziendale efficace per massimizzare i vantaggi economici della gestione dell'energia. Pertanto, è importante sviluppare tali modelli di previsione del carico che potrebbero fornire stabilità, robustezza e precisione della previsione del carico. Si è utilizzato un dataset di consumo di energia elettrica raccolto su una base oraria per questo studio. Vengono utilizzati modelli di deep learning come **Long Short-Term Memory (LSTM)**, **Multilayer Perceptron (MLP)**, **Convolutional Neural Network (CNN) + LSTM** e **1D-CNN** per prevedere il set di dati per 3, 6, 9 mesi, 1 e 6 anni in avanti. Le prestazioni di previsione sono state valutate in termini di diverse metriche di errore standard. Complessivamente, LSTM ha generato il miglior carico anticipato, seguiti da MLP, CNN + LSMT e CNN. Le scoperte mostrano che la soluzione attuale potrebbe essere più appropriata per i potenziali problemi di espansione. Di conseguenza, il modello proposto risulta migliore dei modelli di riferimento in termini di risultati previsionali.

Lo studio dell'articolo [73] è stato condotto per determinare il miglior modello nella previsione della domanda di carico di elettricità per i successivi 1, 2 e 3 mesi in un'università in Malesia, confrontando 4 misure di errore, come *Mean Squared Error*, errore quadratico medio, MAPE e *Geometric Root Mean Squared Error*. In questo sono stati confrontati due metodi di previsione per ottenere i migliori risultati previsionali. I metodi sono il **metodo Exponential Smoothing (HWES) di Holt-Winter** e la **media mobile integrata autoregressiva stagionale (SARIMA)**. Questo studio ha utilizzato 68 dati secondari da Gennaio 2010 fino ad agosto 2015. Microsoft Excel e JMP sono stati utilizzati per determinare il tipo di componenti delle serie temporali, stima del modello e valutazione del modello. I risultati hanno mostrato che esistono componenti stagionali e di tendenza nel set di dati. Il modello migliore era SARIMA poiché denotava la misurazione dell'errore più basso rispetto all'HWES. Quindi, questo modello verrà utilizzato per la previsione dei prossimi 1, 2 e 3 mesi nella domanda di carico di elettricità.

SARIMA Model	AIC	SBC	Significant	Comment
(0,0,1) (1,0,0)12	1104.2488	1110.9073	All significant.	Yes
(0,0,1) (2,0,0)12	1104.0053	1112.8834	AR2,24 not significant.	No
(0,1,1) (0,0,2)12	1104.3797	1113.1985	All significant.	No
(0,2,2) (0,2,1)12	712.6560	719.6067	MA2,12 not significant.	No
(0,2,2) (1,2,1)12	713.3501	722.0384	AR2,12 and MA2,12 not significant.	No
(0,2,2) (2,2,0)12	713.6221	722.3104	AR2,24 not significant.	No
(1,2,2) (0,2,1)12	711.1422	719.8305	AR1,1 and MA2,12 not significant.	No
(1,2,2) (0,2,2)12	711.9432	722.3692	AR1,1,MA2,12 and MA2,24 not significant.	No
(1,2,2) (2,2,0)12	713.5767	725.7404	AR1,1, AR2,12, AR2,24 and MA2,12 not significant.	No
(2,0,1) (2,0,2)12	1104.0559	1121.8119	AR2,12, AR2,24 and MA2,24 not significant.	No

Figura 131: Parti dell'output di SARIMA

Questo studio dimostra l'applicazione di successo di livellamento esponenziale e il metodo Box-Jenkin in previsione della ELD a medio termine. In conclusione, dopo tutti i metodi scelti per il test, questo studio ha scoperto che SARIMA è il miglior modello per prevedere ELD in UUM. In questo studio, ELD in UUM ha sia componenti di tendenza che stagionali. Il modello SARIMA (0,0,1) (1,0,0)12 è il miglior modello per essere utilizzato per prevedere i prossimi 1, 2 e 3 mesi. È mostrato che il modello SARIMA di Box-Jenkins ha la metodologia adeguata alle previsioni ELD a medio termine a causa dei valori MAPE in m passi in avanti, che sono inferiori al 5%. Questo studio aiuterà la gestione universitaria nella previsione dell'ELD del campus per il prossimo mese. Può anche aiutare nella pianificazione della futura ELD per evitare sprechi. Lo studio utilizza solo dati mensili, altrimenti noti come dati a medio termine, per prevedere l'ELD.

Il documento [74] presenta una previsione del carico elettrico a medio termine per una rete di distribuzione dell'Abuja Municipal Area Council (AMAC) basata su **Artificial Neural Network (ANN)**. I risultati tecnici vengono confrontati con quelli di un metodo convenzionale (metodo di regressione lineare multipla), per gli stessi dati. Il metodo proposto da ANN tiene conto dell'effetto della

temperatura, del tempo e del tasso di crescita della popolazione e delle attività di diverse regioni delle aree urbane per quanto riguarda lo stile di vita e le tipologie di consumatori. I dati dei valori di picco da mensili a annuali sono raccolti per il periodo dal 2012 al primo trimestre del 2018. Quindi, il metodo della rete neurale artificiale ha presentato un risultato con MAPE medio di 0,00197 mentre la regressione lineare multipla ha una MAPE media di 0,004545. Il valore R della deviazione è stato rispettivamente dell'8,06% e del 34,42% per i metodi ANN e MLR.

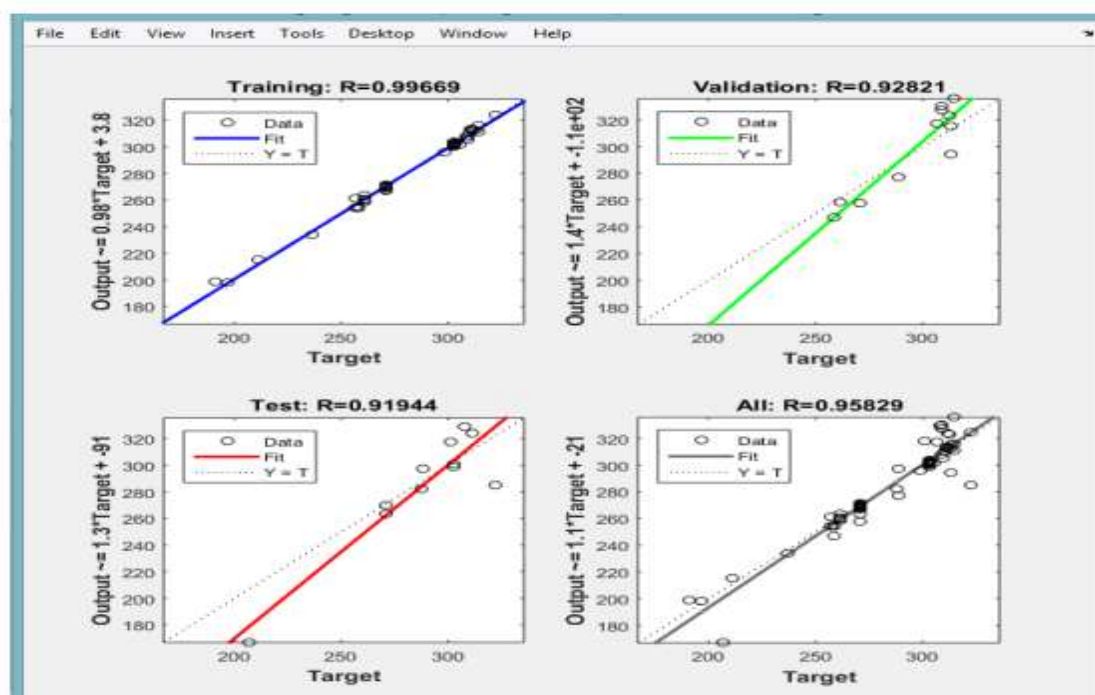


Figura 132: Il grafico di regressione della previsione del carico

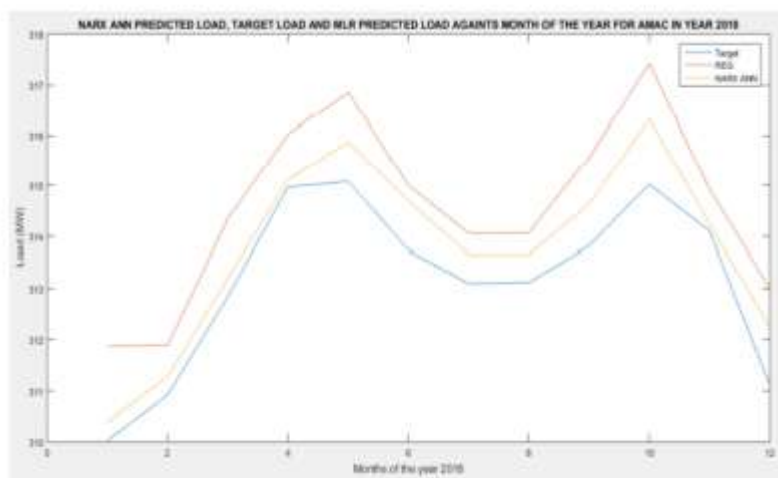


Figura 133: Carico previsto di AMAC utilizzando ANN e utilizzando regressioni lineari multiple per l'anno 2018

Questo lavoro ha presentato un modello, in grado di prevedere la richiesta di carico del consiglio municipale dell'area di Abuja (AMAC) a medio termine. Questo **Nonlinear AutoRegressive network with exogenous inputs (NARX) -ANN** basato sulla tecnica di previsione del carico a medio termine, utilizza una **Feed Forward Neural Network (FFNN) a tre strati e l'algoritmo di propagazione ridimensionato per l'addestramento**, che è in grado di prevedere la domanda di carico di AMAC per i prossimi 48 mesi. I dati storici del meteo, la popolazione e la domanda di carico precedente del consiglio comunale di zona di Abuja (AMAC) è stato considerato. Per verificare l'efficacia e l'esattezza del modello NARX-ANN sviluppato, è stata applicata anche la tecnica di regressione lineare multipla. Il modello NARX-ANN è stato in grado di eseguire meglio del modello di regressione lineare multiplo, perché doveva imparare da dati di addestramento la relazione tra variabili e il loro effetto sulla domanda di carico, rendendo così la previsione relativamente precisa. Maggiore è il numero di cicli di formazione, maggiore è l'errore di previsione migliorato. Tuttavia, il modello conduce al sovrallenamento di risultati sfavorevoli. Si può vedere anche dal modello NARX-ANN che, se il numero di variabili di input sono aumentate, aumenta anche l'accuratezza della previsione. Si può concludere quindi che il modello di rete neurale artificiale è migliore per la previsione del carico a medio termine.

3.1.3 Long-term load forecast (LTLF)

L'articolo [75] presenta un modello di previsione del carico elettrico a lungo termine per il sistema energetico del Camerun. Il modello utilizza tecniche di machine learning, in particolare **una Artificial Neural Network (ANN)**, per prevedere la domanda di energia elettrica fino al 2035. I dati utilizzati per addestrare il modello sono stati raccolti dal sistema energetico del Camerun e comprendono informazioni sul consumo di energia elettrica, il numero di abitazioni, la popolazione e il reddito pro capite. Inoltre, il modello è stato validato utilizzando i dati reali del consumo di energia elettrica del Camerun nel periodo tra il 2010 e il 2017. I risultati mostrano che il modello di previsione basato sul machine learning ha una buona capacità predittiva, con un errore medio del 5% rispetto ai dati reali del consumo di energia elettrica. Questo tipo di modello può essere utilizzato per supportare la pianificazione del sistema energetico del Camerun, aiutando a prevedere la domanda futura di energia elettrica e a pianificare adeguatamente la produzione di energia.

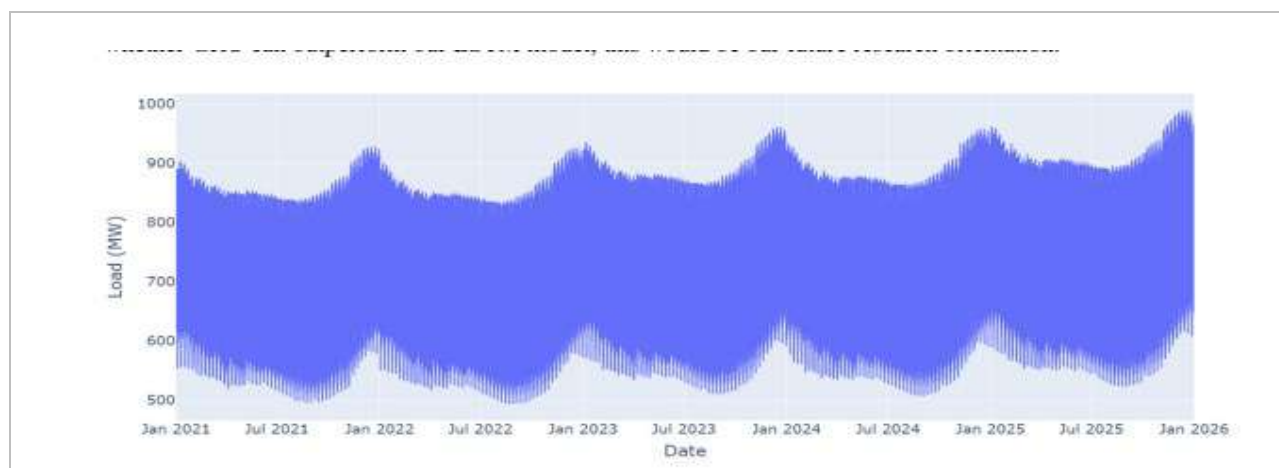


Figura 134: Proiezioni di previsione del carico dal 2021 al 2025, basate sul modello LSTM-RNN proposto

Il documento propone un modello **LSTM-RNN (Long Short-Term Memory Recurrent Neural Network)** con un MAPE soddisfacente del 5,4297%, che supera il metodo corrente (modello di regressione lineare) utilizzato per la previsione del carico a SONATREL per il sistema di potenza del Camerun. Il modello risulta estremamente accurato e può essere utilizzato in modo affidabile per prevedere gli anni futuri di carico. Si considera ora un'altra variante di LSTM, il **gated recurrent unit (GRU)** che si concentra solo sulle funzioni di gate piuttosto che sulla memoria a lungo termine. GRU può ridurre il numero di parametri rispetto a LSTM poiché ci sono due porte in GRU: reset gate e

update gate. L'obiettivo degli autori è mantenere la capacità di riserva delle informazioni a lungo termine di LSTM mentre l'intera rete è resa il più semplice possibile.

L'articolo [76] presenta un nuovo modello per la previsione a lungo termine della domanda di energia elettrica nelle abitazioni, insieme a un metodo innovativo per rilevare il furto di energia elettrica. Il modello di previsione utilizza una combinazione di tecniche di deep learning, in particolare una **Recurrent Neural Network (RNN)** e un **Gradient Boosted Regression Trees (GBRT)**, per prevedere la domanda di energia elettrica a livello di singola abitazione per un periodo di tempo fino a un anno. Il metodo di rilevamento del furto di energia elettrica si basa su un algoritmo di clustering, che utilizza i dati di consumo di energia elettrica per identificare i pattern di consumo anomali e sospetti. Questo metodo può aiutare le compagnie energetiche a individuare e prevenire il furto di energia elettrica, migliorando la loro efficienza operativa e riducendo le perdite di reddito.

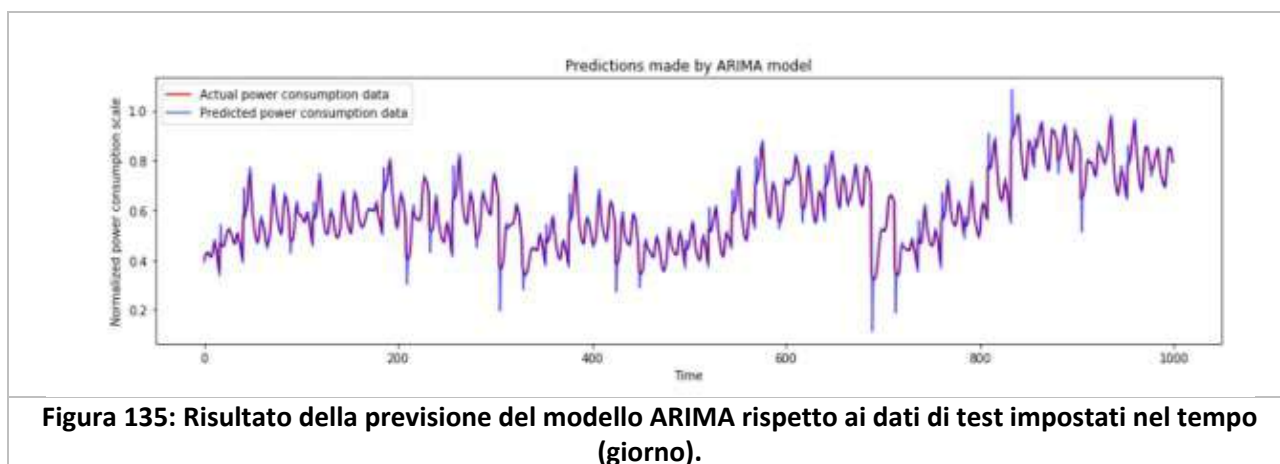
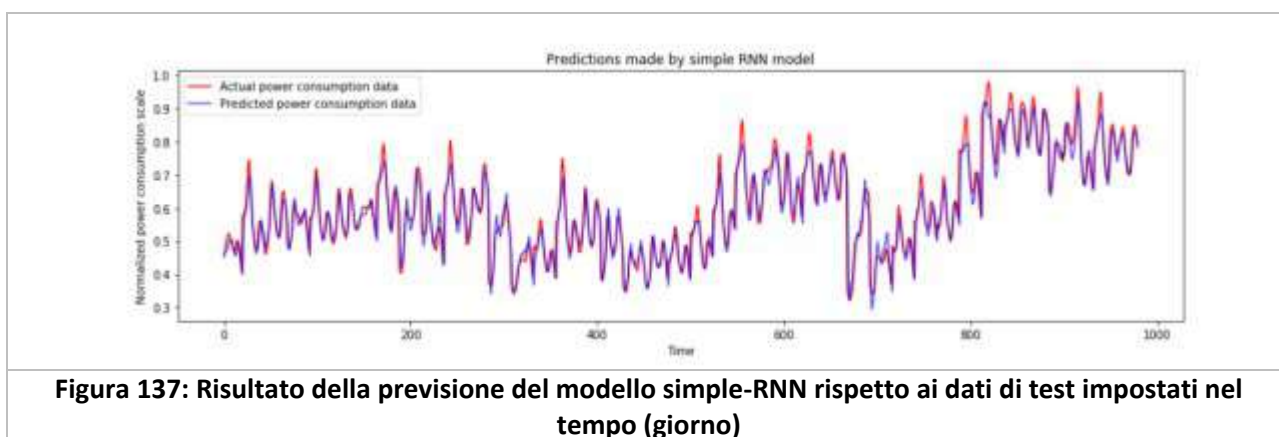
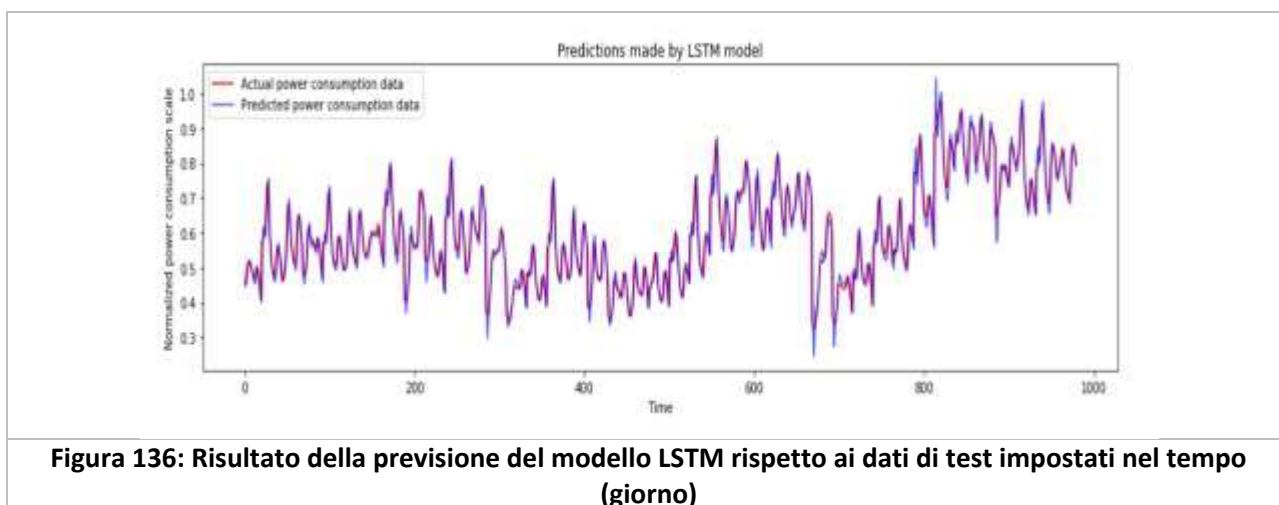


Figura 135: Risultato della previsione del modello ARIMA rispetto ai dati di test impostati nel tempo (giorno).



I risultati mostrano che il modello di previsione proposto ha una buona capacità predittiva e migliora significativamente la precisione rispetto ai metodi di previsione tradizionali. Inoltre, il metodo di rilevamento del furto di energia elettrica ha dimostrato di essere efficace nell'identificazione dei casi di furto di energia elettrica. Questi risultati indicano che il modello proposto e il metodo di rilevamento del furto di energia elettrica possono essere utili per migliorare la gestione della domanda di energia elettrica e ridurre le perdite di reddito dovute al furto di energia elettrica.

Il documento [77] presenta un metodo per la previsione della domanda di energia elettrica a lungo termine nel contesto di una catena di approvvigionamento di energia elettrica in una micro-rete di un complesso siderurgico.

Il metodo utilizza un approccio accoppiato che combina una tecnica di analisi di serie temporali con un modello di **apprendimento automatico basato su Artificial Neural Network (ANN)**. In particolare, il modello utilizza le previsioni a breve termine e le informazioni storiche sulla domanda di energia elettrica, insieme ad altre variabili come la temperatura e l'occupazione degli impianti, per prevedere la domanda di energia elettrica a lungo termine in un'ottica di catena di approvvigionamento.

Il metodo è stato applicato a un caso studio di un complesso siderurgico, e i risultati mostrano che il modello proposto ha una buona capacità predittiva rispetto ai dati reali della domanda di energia elettrica. Inoltre, il modello accoppiato ha dimostrato di essere efficace nel prevedere la domanda di energia elettrica in una catena di approvvigionamento di energia elettrica in una micro-rete, migliorando la pianificazione e la gestione della produzione di energia elettrica.

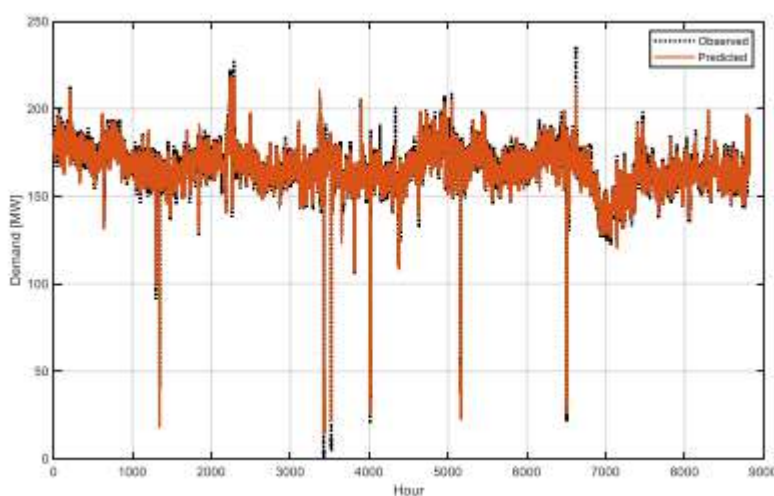


Figura 138: Domanda prevista e previsione dell'ESC utilizzando il metodo proposto

Method	Optimization Method	MAE	MAPE	MSE	RMSE	Standard Deviation
W-LSTM	ELATLBO	3.023	1.8161	39.515	6.2861	12.9090
Gaussian SVM	BO	3.2191	1.9339	53.895	7.3413	13.4734
Decision Tree	BO	3.1224	1.8758	49.847	7.0603	13.5507
Boosted Tree	BO	3.2901	1.9765	48.299	6.9498	13.8292
Random Forest	BO	3.0611	1.839	48.513	6.9651	13.6558

Figura 139: Risultati del test

La previsione accurata della domanda elettrica delle industrie metallurgiche di base per pianificare la generazione elettrica è il primo passo per bilanciare la filiera elettrica. Un metodo ibrido è stato proposto in questo lavoro per prevedere le serie storiche della domanda delle industrie metallurgiche. L'approccio accoppiato consiste **nella decomposizione wavelet** nella prima fase. Poi i vettori di output della decomposizione wavelet sono divisi in tre parti. L'allenamento dei sottoinsiemi e dei sottoinsiemi di convalida viene utilizzato nell'addestramento e nella messa a punto del modello **Long Short-Term Memory (LSTM)** utilizzando il metodo ELATLBO. Il set di dati ESC utilizzato in questo studio per la previsione della domanda di elettricità include 24 ore al giorno per 40 mesi dal 21 marzo 2017 al 21 giugno 2020. Per la valutazione dei risultati ottenuti è stato effettuato un confronto con altri metodi, inclusi **Support Vector Machine (SVM)**, **Decision Trees (DT)**, **Boosted Tree (BT)** e **Random Forest (RF)** che sono ottimizzati utilizzando il metodo di Bayesian optimization (BO). I risultati ottenuti con i dati minimi disponibili (solo le serie temporali della domanda) mostrano che le prestazioni dell'approccio accoppiato sono buone e i parametri impostati nel metodo proposto sono molto inferiori rispetto ad approcci simili. Nel seguito della ricerca svolta in questo studio, è possibile considerare e studiare un algoritmo di ottimizzazione con elevata adattabilità e minimizzazione dei parametri sintonizzabili esclusivamente per l'ottimizzazione dei metodi di regressione.

L'articolo [78] propone un nuovo **Adaptive Backpropagation Algorithm (ABPA)** per la previsione a lungo termine della domanda di energia elettrica. L'algoritmo proposto utilizza un approccio basato su **Artificial Neural Network (ANN)**, in cui le informazioni storiche sulla domanda di energia elettrica e altre variabili pertinenti vengono utilizzate come input per la rete neurale. L'algoritmo adattivo consente di adattare automaticamente i pesi delle connessioni della rete neurale per migliorare la precisione delle previsioni. Il paper confronta l'algoritmo adattivo con altri algoritmi di previsione

della domanda di energia elettrica, come il **Recurrent Neural Network (RNN)**, utilizzando i dati di domanda di energia elettrica del mercato energetico dell'Irlanda. In sintesi, il documento presenta un nuovo algoritmo di retropropagazione adattiva basato su reti neurali artificiali per la previsione a lungo termine della domanda di energia elettrica. L'algoritmo si è dimostrato efficace per migliorare l'accuratezza delle previsioni della domanda di energia elettrica rispetto ad altri algoritmi di previsione della domanda di energia elettrica.

Year 2020	Actual load	REG	BPA	ABPA	RBFN	RNNs-LSTM
Jan	23,849	18,324	19,484	19,784	25,344	20,924
Feb	21,988	14,100	12,510	12,138	23,212	16,126
Mar	25,615	17,117	17,276	19,080	17,812	19,858
Apr	21,923	18,707	19,529	20,394	15,960	23,365
May	26,592	19,653	20,494	22,047	21,271	25,437
June	28,015	22,830	22,319	24,235	25,539	28,868
July	28,354	21,903	20,558	20,847	27,043	23,375
Aug	29,792	23,829	22,243	23,326	27,329	27,746
Sept	29,060	21,355	20,445	20,436	26,513	24,679
Oct	27,615	19,997	19,753	20,376	21,052	27,327
MAPE	–	0.25	0.26	0.045	0.144	0.116
MSE	–	44.495.892	50.759.957	1.195.650	19.197.960	12.845.733

Figura 140: Previsione di quattro metodi con dati reali



Figura 141: Il confronto di previsione accurata di più approcci

I risultati evidenziano l'efficacia dell'algoritmo di retropropagazione adattiva proposto nella previsione a lungo termine della domanda di energia elettrica. In particolare, l'algoritmo adattivo ha dimostrato di superare gli altri algoritmi di previsione della domanda di energia elettrica confrontati nel documento, in termini di accuratezza delle previsioni. Inoltre, l'uso di **Artificial Neural Network (ANN)**, insieme all'algoritmo di retropropagazione adattiva, ha permesso di incorporare in modo efficace l'informazione storica sulla domanda di energia elettrica e altre variabili rilevanti nella previsione a lungo termine. Inoltre, per finire, si suggerisce che l'algoritmo adattivo può essere utilizzato per migliorare la pianificazione della produzione di energia elettrica, ridurre i costi e migliorare l'efficienza del sistema energetico.

L'articolo [79] di ricerca descrive come gli algoritmi di machine learning possono essere utilizzati per prevedere la domanda di energia in Corea. L'articolo inizia fornendo una panoramica del sistema energetico coreano e delle sfide che il paese sta affrontando nel gestire la crescente domanda di energia. Successivamente, viene descritto il modello di previsione utilizzato, basato su algoritmi di apprendimento automatico come **Linear Regression (LR)**, **Support Vector Machine (SVM)** e **Neural Network (NN)**. I risultati della previsione mostrano che i modelli di machine learning possono prevedere con successo la domanda di energia in Corea, ottenendo un'accuratezza maggiore rispetto ai metodi di previsione tradizionali. Inoltre, viene evidenziata la capacità dei modelli di identificare e considerare le variabili che influenzano la domanda di energia, come il clima e la temperatura.

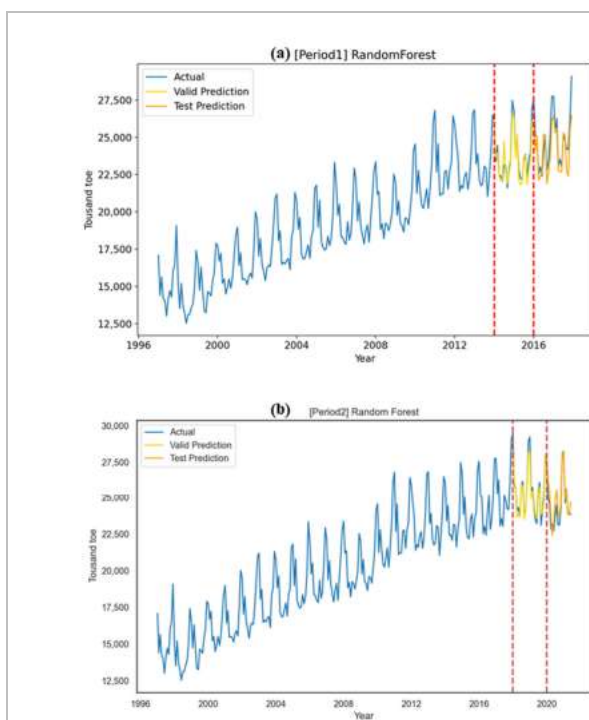


Figura 142: Previsione Random forest

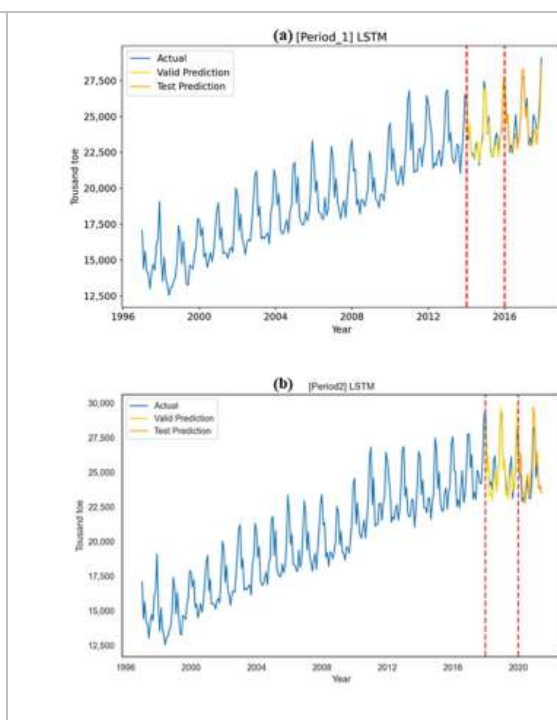


Figura 143: Previsione LSTM

Per riassumere i risultati della ricerca, in primo luogo, il periodo 1 è stato previsto con la massima precisione dal modello **Long Short-Term Memory (LSTM)**. In secondo luogo, il modello RF tendeva a produrre RMSE e MAPE più bassi nel periodo 2. I seguenti punti sono implicazioni derivate dai risultati dello studio. LSTM, che potrebbe tenere conto dei movimenti periodici, ha mostrato un significativo valore di prestazioni predittive relative ai diversi metodi di apprendimento automatico quando la tendenza del mercato era coerente. LSTM ha molti vantaggi rispetto ad altri NN feedforward. Tuttavia, nella modellazione di sistemi non lineari, il normale LSTM non funziona molto bene. Quando il comportamento del mercato è cambiato da una tendenza all'altra, RF, con la sua capacità di modellazione non lineare, ha mostrato i risultati predittivi più efficaci.

L'articolo di ricerca [80] descrive come le **tecniche di deep learning** possono essere utilizzate per prevedere la domanda di energia in Francia. L'articolo inizia fornendo una panoramica del sistema energetico francese e delle sfide che il paese sta affrontando nel gestire la crescente domanda di energia, come la transizione verso fonti di energia rinnovabile. Successivamente, viene descritto il

modello di previsione utilizzato, basato su una **Convolutional Neural Network (CNN)**, una delle tecniche di deep learning più utilizzate.

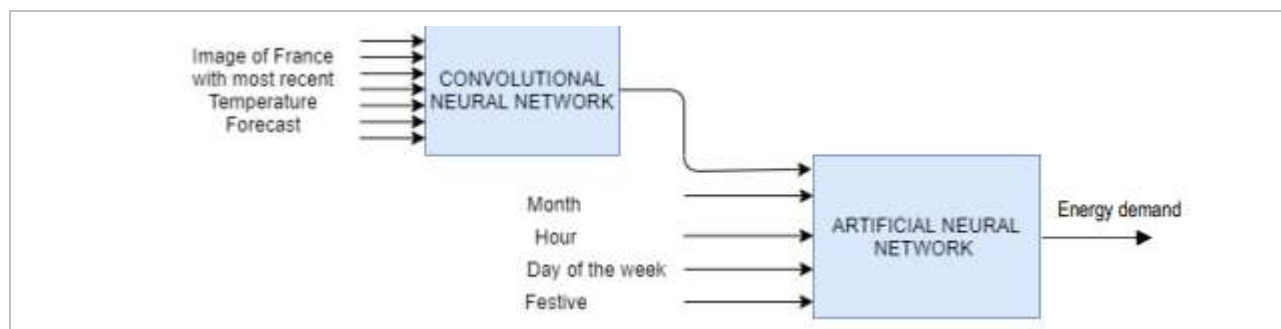


Figura 144: Struttura di deep learning composta da una rete neurale convoluzionale seguita da una rete neurale artificiale e adattata al problema di previsione della domanda di energia

Month	MAPE (%) provided by the proposed deep learning structure	MAPE(%) provided by the RTE subscription-based service
January	1.3542	1.5249
February	1.2424	1.5078
March	1.4667	1.8752
April	1.9774	1.4599
May	1.4718	1.5619
June	1.2693	1.3278
July	1.2318	1.3025
August	1.3592	1.4180
September	1.3575	1.2333
October	1.3964	1.4759
November	1.7511	1.5558
December	2.0131	1.6306

Figura 145: Confronto delle prestazioni delle metriche MAPE mensili

In questo documento, gli autori presentano comunemente l'adattamento di una struttura di rete neurale profonda utilizzata per la classificazione delle immagini applicata alla previsione del fabbisogno energetico. In particolare, la struttura è stata addestrata per la rete energetica francese. I risultati mostrano che le prestazioni della struttura proposta competono con i risultati forniti dal servizio di riferimento in abbonamento RTE. Nello specifico, la metrica MAPE complessiva dell'approccio proposto fornisce un errore dell'1,4933%, che è leggermente migliore rispetto alla

metrica dell'1,4940% ottenuta con i dati previsionali RTE. Se analizzati su base mensile, gli errori sono distribuiti uniformemente durante l'anno nonostante ci siano notevoli incrementi durante il tardo autunno e la stagione invernale. Anche questo fatto è conforme con i dati previsionali RTE di riferimento e può essere dovuto all'intermittenza del profilo di consumo energetico osservato quando le temperature, in Francia, sono basse.

Il documento di ricerca [81] descrive come le tecniche di intelligenza artificiale possono essere utilizzate per prevedere la domanda di energia elettrica in Turchia a lungo termine. L'articolo inizia fornendo una panoramica del sistema energetico turco e delle sfide che il paese sta affrontando nel gestire la crescente domanda di energia, tra cui la necessità di un'offerta energetica sostenibile e l'aumento del consumo di energia. Successivamente, viene descritto il modello di previsione utilizzato, basato su diverse tecniche di intelligenza artificiale, tra cui **Artificial Neural Network (ANN)**, **Decision Trees (DT)** e **Linear Regression (LR)**.

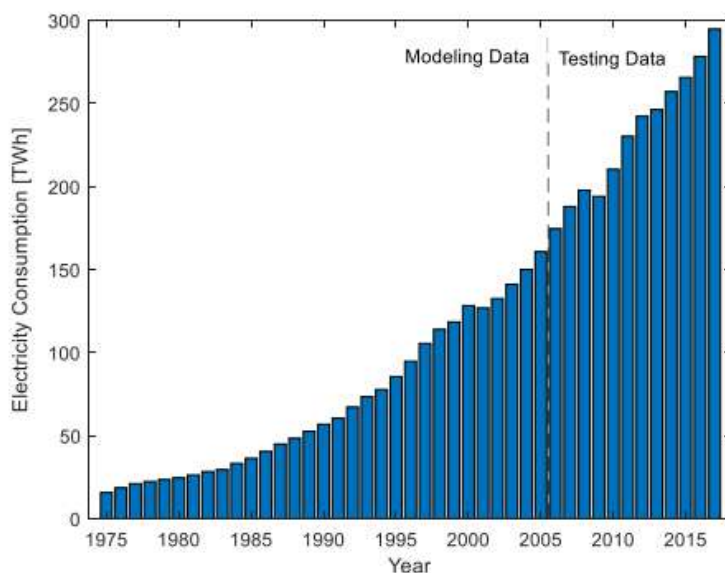
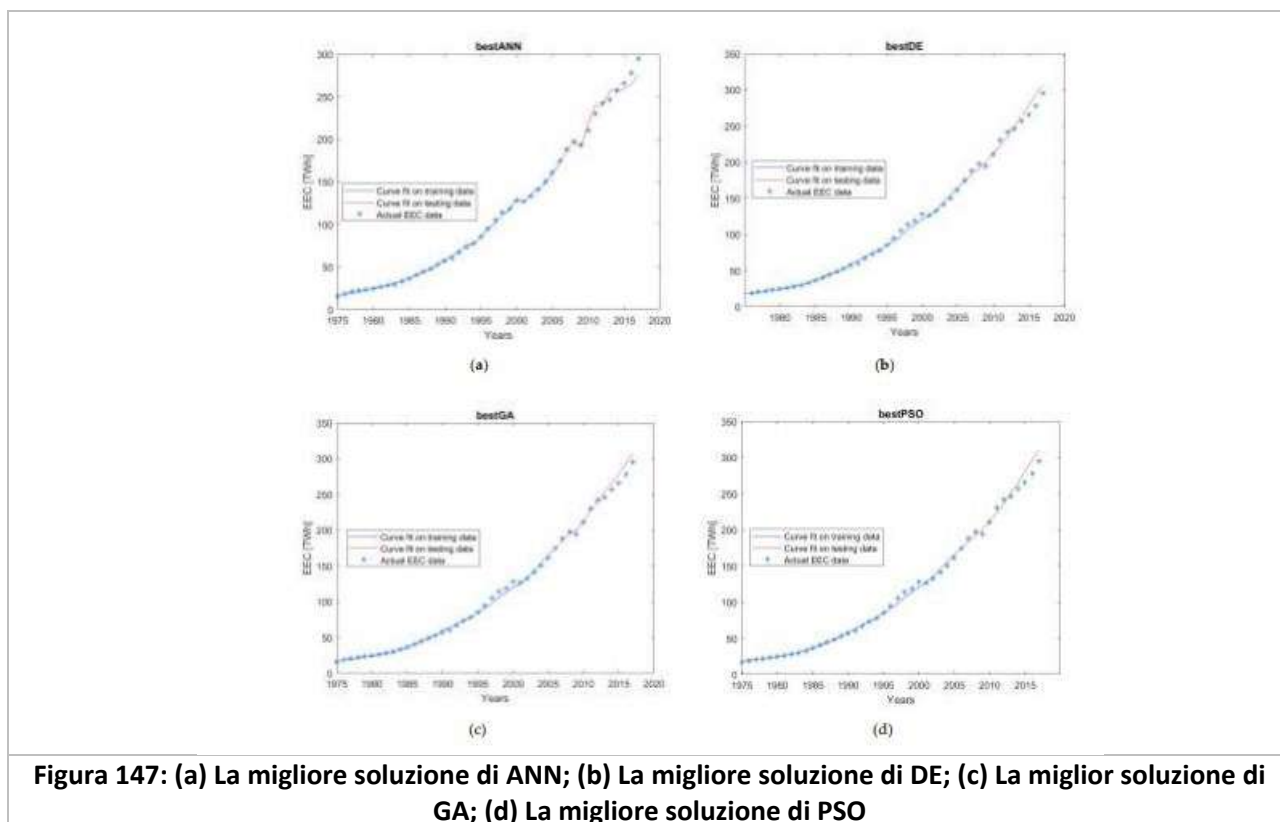


Figura 146: Domanda di elettricità della Turchia tra il 1975 e il 2017



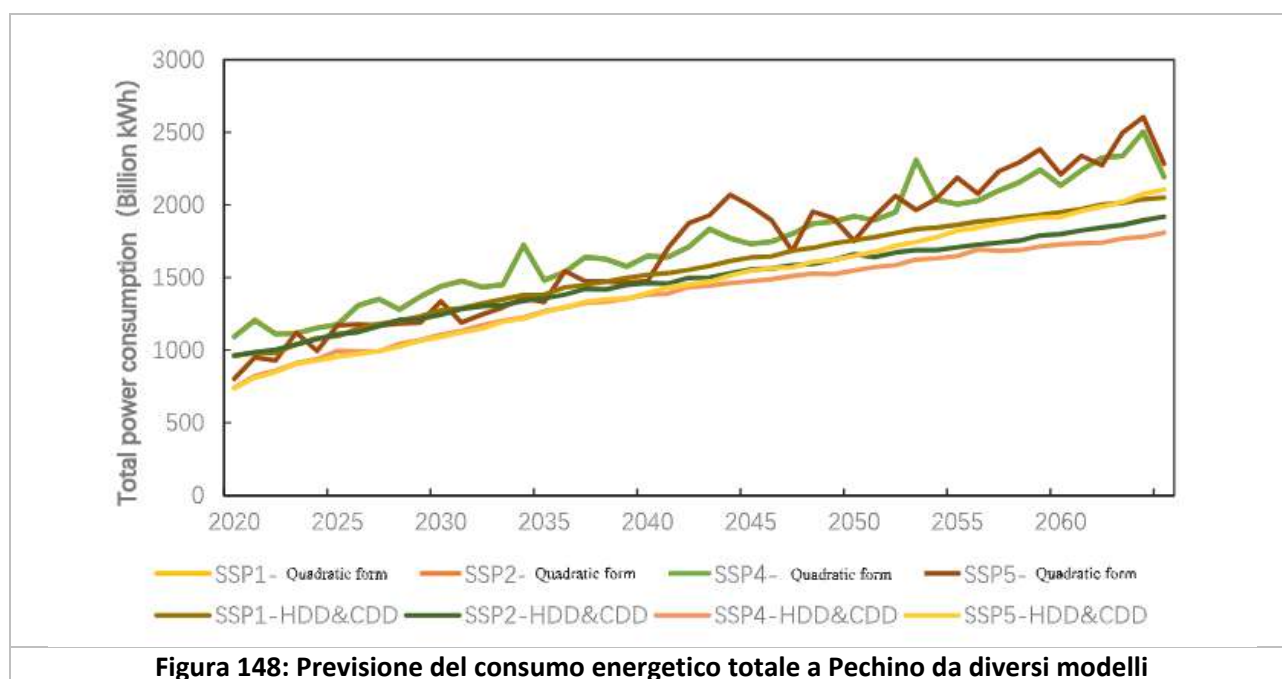
I risultati della previsione mostrano che il modello di intelligenza artificiale può prevedere con successo la domanda di energia elettrica in Turchia a lungo termine, ottenendo un'accuratezza maggiore rispetto ai metodi di previsione tradizionali. Inoltre, viene evidenziata la capacità del modello di identificare e considerare le variabili che influenzano la domanda di energia, come la temperatura, la popolazione e il PIL.

Infine, l'articolo conclude che l'utilizzo di tecniche di intelligenza artificiale può essere un'importante risorsa per i responsabili delle politiche energetiche, poiché può fornire previsioni più accurate e dettagliate sulla domanda di energia elettrica, aiutando a sviluppare strategie di gestione dell'offerta e della domanda di energia più efficienti e sostenibili in Turchia.

Infine, l'articolo [82] affronta la sfida di prevedere la domanda di energia elettrica a medio e lungo termine in un contesto di cambiamento climatico, con focus sulla città di Pechino.

L'articolo inizia fornendo una panoramica degli impatti del cambiamento climatico sulla domanda di energia elettrica, come le variazioni stagionali e l'aumento della domanda di raffreddamento durante i periodi di caldo estivo.

Successivamente, viene descritto il modello di previsione utilizzato, basato su una combinazione di tecniche statistiche e di intelligenza artificiale, tra cui la **Linear Regression (LR)**, le **Artificial Neural Network (ANN)** e i modelli di autocorrelazione.



Questo documento valuta innanzitutto la funzione di risposta temperatura-consumo energetico a Pechino e misura quantitativamente l'impatto della temperatura sul consumo energetico regionale; inoltre, sulla base della funzione dose-risposta è previsto il consumo energetico di Pechino nei prossimi 40 anni nell'ambito di diversi percorsi socioeconomici condivisi (SSPS). In questo processo, le principali conclusioni sono le seguenti: (1) utilizzando diversi modelli per la misurazione, la relazione tra temperatura e consumo energetico a Pechino presenta una curva a forma di U. Prendendo il modello lineare come, ad esempio, se il CDD aumenta di 1 grado giorno, il consumo energetico a Pechino aumenterà di 5,3 milioni kWh, mentre se l'HDD aumenta di 1 grado al giorno,

il consumo energetico a Pechino aumenterà di 900mila kWh. (2) Il futuro livello di consumo energetico di Pechino mostrerà una tendenza al rialzo fluttuante. Nel 2060, il potere di consumo di Pechino raggiungerà 219,3 ~ 290,4 miliardi di kWh. Sulla base di diversi percorsi di economia sociale condivisi (SSPS), questo documento prevede l'evoluzione del consumo energetico a Pechino nei prossimi 60 anni.

3.2 Previsione della produzione di energia

Poiché le energie rinnovabili intermittenti stanno diventando sempre più onnipresenti, le previsioni stanno diventando sempre più importanti per garantire che la rete rimanga equilibrata. La produzione di energia solare ed eolica è influenzata dalle condizioni meteorologiche; quindi, avere una previsione accurata di quanta energia sarà prodotta da queste risorse in qualsiasi momento è la chiave per una rete ottimizzata ed equilibrata. I modelli di machine learning sono perfettamente adatti a questo compito, in quanto possono apprendere dalle tendenze dei dati del passato per prevedere il prossimo futuro. L'intelligenza artificiale può aiutare a migliorare il sistema energetico riducendo i costi, ottimizzando le risorse, prevedendo guasti e interruzioni, gestendo la domanda di energia, aiutando a scoprire nuovi materiali e guidando l'innovazione. Può anche aiutare a ottimizzare il processo decisionale e le operazioni della rete elettrica integrando autonomamente la fornitura di energia, la domanda e le fonti rinnovabili.

3.2.1 Fonti rinnovabili

Tenendo conto delle ipotesi del Green Deal europeo relative a una riduzione significativa delle emissioni di serra, la produzione di elettricità da fonti di energia rinnovabile (FER) è maggiore e più importante al giorno d'oggi. Oltre a questo, previsioni accurate e affidabili sulla produzione di energia elettrica da RES sono necessarie per la pianificazione della capacità, la programmazione, la gestione dell'inerzia e la risposta in frequenza durante gli eventi imprevisti. Gli ultimi tre anni hanno dimostrato che i modelli di **Machine Learning (ML)** sono una soluzione promettente per la previsione della produzione di elettricità da FER. Nella relazione [83] si evidenzia come:

- 1) **Extreme Learning Machine e metodi Ensemble** siano i metodi più popolari utilizzati per la previsione della produzione di elettricità da RES in negli ultimi tre anni (2020–2022)
- 2) la maggior parte della ricerca è stata svolta per i sistemi eolici,
- 3) i modelli ibridi rappresentano circa un terzo dei lavori analizzati,
- 4) la maggior parte degli articoli interessati modelli a breve termine,
- 5) la maggior parte dei ricercatori proveniva dalla Cina,

3.2.1.1 *Energia eolica*

Il paper [84] utilizza algoritmi di machine learning per prevedere la produzione di energia eolica in un determinato sito. Il progetto si basa sulla raccolta di dati storici sulla produzione di energia eolica e sulle condizioni meteorologiche del sito. Questi dati vengono utilizzati per addestrare un modello di machine learning, che viene poi utilizzato per fare previsioni sulla produzione di energia futura. Il modello di machine learning utilizzato in questo progetto dipende dalle specifiche esigenze. Alcuni esempi includono l'utilizzo di **Neural Network (NN)**, **Decision Trees (DT)** e **Linear Regression (LR)**.

L'obiettivo principale dell'articolo è quello di aiutare le aziende che operano nell'industria dell'energia eolica a prendere decisioni informate sulla produzione di energia eolica in un determinato sito. Inoltre, il progetto mira anche a migliorare la comprensione della produzione di energia eolica attraverso l'utilizzo di dati storici e tecniche di machine learning.

	Mean	Std	Min	1 st quartile	Median	3 rd quartile	Max
Energy	8766.86	6846.41	135	3617	6944	11,800	33,626
v_1	2.94	1.19	0.75	2.07	2.72	3.58	7.93
v_2	4.3	1.61	1.35	3.14	3.95	5.11	11.16
v_{50m}	6.06	2.09	1.96	4.63	5.72	7.1	14.92
h_1	2.69	0.13	2.46	2.58	2.73	2.81	2.83
h_2	10.69	0.13	10.46	10.58	10.73	10.81	10.83
z_0	0.2	0.03	0.14	0.17	0.21	0.23	0.24
ρ	1.21	0.03	1.13	1.18	1.21	1.24	1.31
p	9.86E+04	8.58E+02	9.55E+04	9.81E+04	9.86E+04	9.91E+04	1.01E+05
T	9.22	7.58	-8.55	2.91	8.09	15.52	28.7

Figura 149: Descrizione del dataset

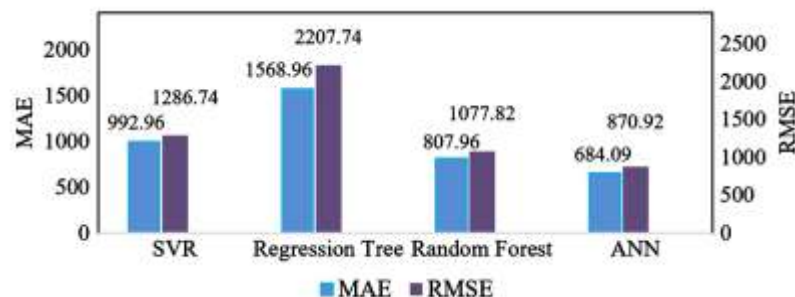


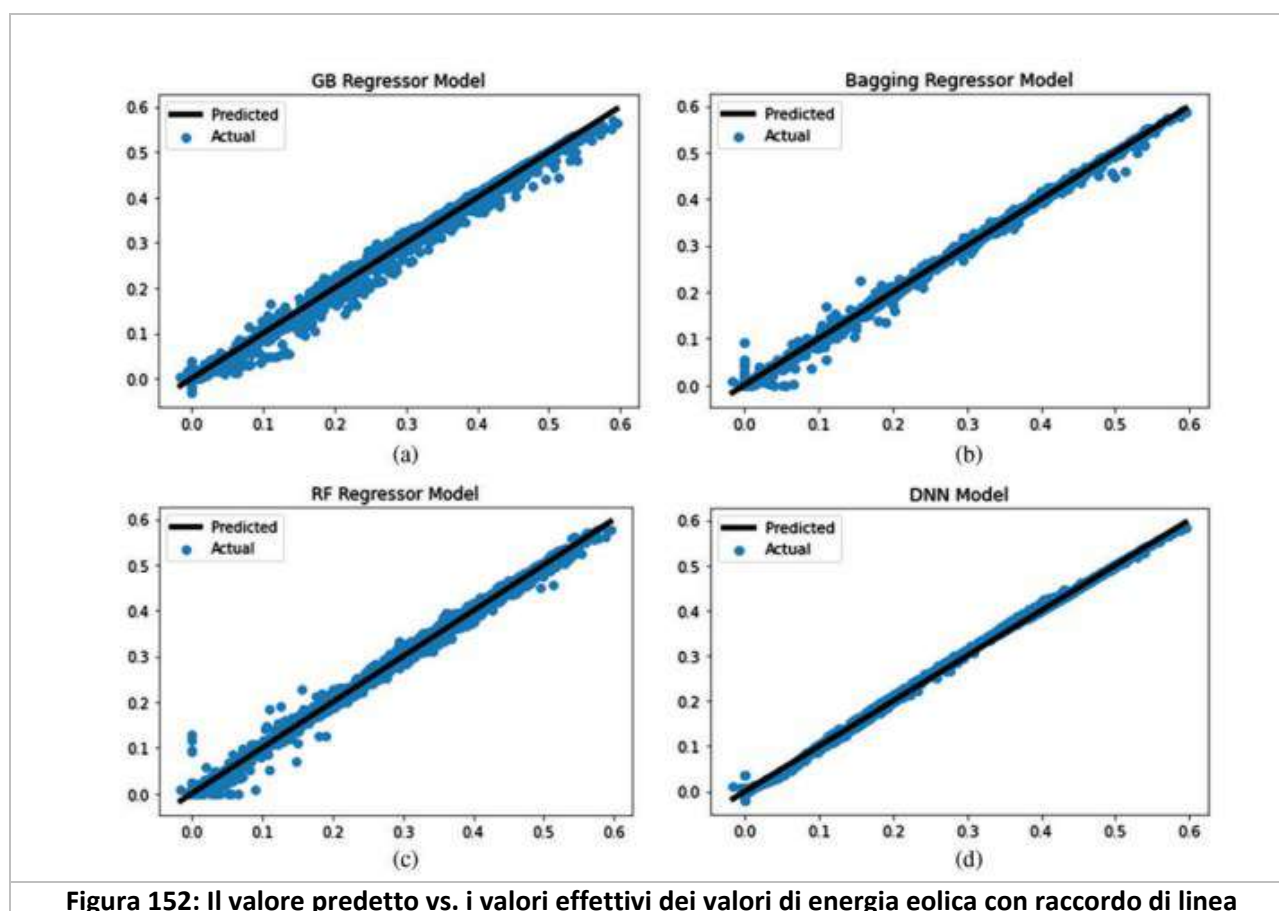
Figura 150: Confronto di MAE e RMSE per gli algoritmi selezionati

Questo approccio end-to-end basato sul framework CRISP-DM offre alcune linee guida per progetti futuri. **Le Artificial Neural Network (ANN)** ben sintonizzate possono fornire previsioni accurate per prevedere la produzione di energia eolica, ma soprattutto sostenere i costi associati alle fonti, al tempo e alla potenza di calcolo necessari per trovare le combinazioni di iperparametri. I modelli basati su alberi forniscono una trasparenza che neanche algoritmi black box come ANN danno. La **Support Vector Regression (SVR)** potrebbe essere stata la migliore soluzione proposta considerando le prestazioni proposte con il tempo di addestramento.

Il documento [85] propone algoritmi di **machine learning e deep learning** per prevedere la produzione di energia eolica. Il progetto si basa sulla raccolta di dati storici sulla produzione di energia eolica e sulle condizioni meteorologiche del sito. Questi dati vengono utilizzati per addestrare modelli di machine learning e deep learning, che vengono poi utilizzati per fare previsioni sulla produzione di energia futura. Tra i modelli di machine learning utilizzati in questo progetto ci sono la **Linear Regression (LR)**, la **regressione logistica** e i **K-nearest Neighbors algorithm (KNN)**, mentre tra i modelli di deep learning utilizzati ci sono **Convolutional Neural Network (CNN)** e le **Recurrent Neural Network (RNN)**. L'obiettivo principale è quello di migliorare la precisione delle previsioni sulla produzione di energia eolica e di aiutare le aziende che operano nell'industria dell'energia eolica a prendere decisioni informate sulla produzione di energia eolica in un determinato sito. Inoltre, il progetto mira anche a migliorare la comprensione della produzione di energia eolica attraverso l'utilizzo di dati storici e tecniche avanzate di machine learning e deep learning.

	Count	Mean	Std	Min	25%	50%	75%	Max
Power	50530	1307.68	1312.45	2.4714	50.677	825.83	2482.50	3618.73
Wind_speed	50530	7.55795	4.22716	0	4.2013	7.1045	10.3000	25.2060
Theor_power	50530	1492.17	1368.01	0	161.32	1063.77	2964.97	3600
Wind_dir	50530	123.687	93.4437	0	49.315	73.7129	201.696	359.997

Figura 151: Analisi statistica delle caratteristiche del set di dati sul vento.



I dati di previsione del vento vengono utilizzati in questo studio per testare i modelli di efficienza **Long Short-Term Memory (LSTM)** ottimizzati proposti. Una nuova tecnica di ottimizzazione viene utilizzata per ottimizzare i parametri della rete LSTM. Per migliorare le capacità di esplorazione e sfruttamento dell'ottimizzatore PSO, questa strategia di ottimizzazione utilizza sia l'ottimizzatore PSO che SFS per generare gruppi stocastici di agenti. Esperimenti sono stati fatti per valutare le prestazioni dell'approccio suggerito e confrontarlo con i risultati di sei altri modelli di apprendimento automatico per dimostrarne la fattibilità. Per valutare il successo del progetto, vengono utilizzati cinque criteri di valutazione. La stabilità e l'efficacia di questo approccio sono state ulteriormente dimostrate da grafici di analisi derivati dai dati ottenuti. È chiaro dai confronti con altri modelli che la strategia fornita è superiore.

L'articolo [86] utilizza tecniche di **deep learning** per prevedere la produzione di energia eolica in un parco eolico composto da più turbine. Ci si basa sulla raccolta di dati storici sulla produzione di energia eolica e sulle condizioni meteorologiche del sito, non solo per una singola turbina, ma per tutte le turbine all'interno del parco eolico. Questi dati vengono utilizzati per addestrare un modello di deep learning, che viene poi utilizzato per fare previsioni sulla produzione di energia futura per tutte le turbine del parco eolico. Il modello di deep learning utilizzato in questo progetto è una **Convolutional Neural Network (CNN)** con un'architettura specializzata per l'analisi di dati di serie temporali. L'obiettivo principale è quello di migliorare la precisione delle previsioni sulla produzione di energia eolica per tutte le turbine del parco eolico. Ciò può aiutare le aziende che operano nell'industria dell'energia eolica a massimizzare la produzione di energia eolica e a ridurre i costi operativi. Inoltre, il progetto mira anche a migliorare la comprensione della produzione di energia eolica in un ambiente complesso come un parco eolico.

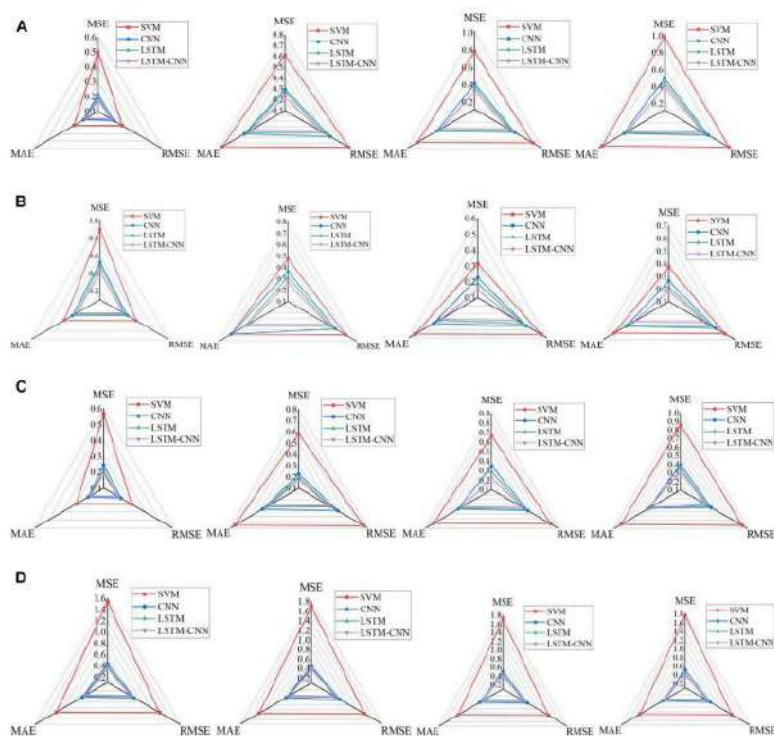


Figura 153: Risultati della previsione per le quattro turbine eoliche in diverse fasi di previsione: (A) fase 1; (B) fase 2; (C) fase 3; (D) fase 4.

Il modello di previsione proposto si basa sull'idea di modellazione di "catturare prima la dipendenza temporale e poi estrarre le caratteristiche spaziali." Composto da **Long Short-Term Memory (LSTM)** e **Convolutional Neural Network (CNN)**, il modello è stato addestrato end-to-end da una funzione di perdita unificata. Il modello addestrato può estrarre caratteristiche spazio-temporali di alto livello dai dati di potenza storici di turbine eoliche, in modo da raggiungere lo scopo, simultaneamente, di prevedere la potenza delle turbine eoliche in luoghi diversi. Sono stati utilizzati i dati misurati di un parco eolico offshore in Cina e, tramite esperimenti di simulazione, i valori reali sono stati confrontati con i valori previsti del modello. I risultati del confronto hanno mostrato che il metodo proposto ha prestazioni più eccellenti rispetto agli esistenti metodi di previsione, come **Long Short-Term Memory (LSTM)**, **Convolutional Neural Network (CNN)** e **Support Vector Machine (SVM)**. Con il modello proposto, è possibile prevedere con precisione la potenza eolica di più turbine all'interno di un parco eolico con il layout regolare. Su questa base, può essere eseguita la programmazione energetica accuratamente, che è la futura direzione della ricerca.

Lo studio proposto nell'articolo [87] compara diverse tecniche di machine learning per la previsione efficiente della produzione di energia eolica. La ricerca si basa sulla raccolta di dati storici sulla produzione di energia eolica e sulle condizioni meteorologiche della struttura. Questi dati vengono utilizzati per addestrare diversi modelli di machine learning, tra cui **Neural Network (NN)**, **Support Vector Machine (SVM)** e **Random Forest (RF)**.

Lo studio confronta le prestazioni dei modelli di machine learning in termini di accuratezza e velocità di previsione. In particolare, l'obiettivo è quello di trovare il modello di machine learning più efficiente per la previsione della produzione di energia eolica. In particolare, vengono utilizzati **l'Ottimizzazione bayesiana (BO)** per regolare in modo ottimale gli iperparametri della **Gaussian Process Regression (GPR)**, il **Support Vector Regression (SVR)** con diversi kernel e modelli di **Ensemble Learning (ES)** che studiano le prestazioni di previsione. In secondo luogo, le informazioni dinamiche sono state incorporate nella loro costruzione per migliorare ulteriormente la previsione delle prestazioni dei modelli studiati.

Models	France			Turkey			Kaggle		
	RMSE (kW)	MAE (kW)	R ²	RMSE (kW)	MAE (kW)	R ²	RMSE (kW)	MAE (kW)	R ²
SVR _L	185.50	126.07	0.83	111.48	70.01	0.97	123.43	74.88	0.95
SVR _Q	259.41	218.69	0.68	113.48	74.47	0.97	139.28	90.10	0.94
SVR _C	187.64	124.63	0.83	120.46	88.44	0.96	123.49	79.41	0.95
SVR _{FG}	186.16	123.82	0.83	131.25	98.08	0.95	124.30	80.81	0.95
SVR _{MG}	186.32	123.68	0.83	133.22	101.86	0.95	124.32	80.10	0.95
GPR _{RQ}	185.81	124.21	0.83	114.87	70.77	0.97	122.49	75.19	0.96
GPR _{SE}	184.40	123.98	0.84	112.89	69.96	0.97	122.62	75.15	0.96
GPR _{M52}	184.70	124.49	0.84	115.09	70.88	0.96	122.48	75.17	0.96
GPR _E	183.23	124.42	0.84	122.63	78.13	0.96	121.62	75.47	0.96
GPR _O	184.57	124.22	0.84	112.79	69.96	0.97	122.60	75.13	0.96
BS	183.66	124.33	0.84	114.34	71.78	0.97	129.71	82.38	0.95
BT	199.79	135.27	0.81	123.79	84.11	0.96	119.24	82.17	0.96
ELO	183.46	124.05	0.84	117.25	72.89	0.96	120.64	75.80	0.96
XGB	185.38	125.21	0.83	121.02	75.23	0.96	122.02	76.95	0.96
RF	223.40	148.95	0.76	185.13	115.01	0.91	132.52	91.95	0.95

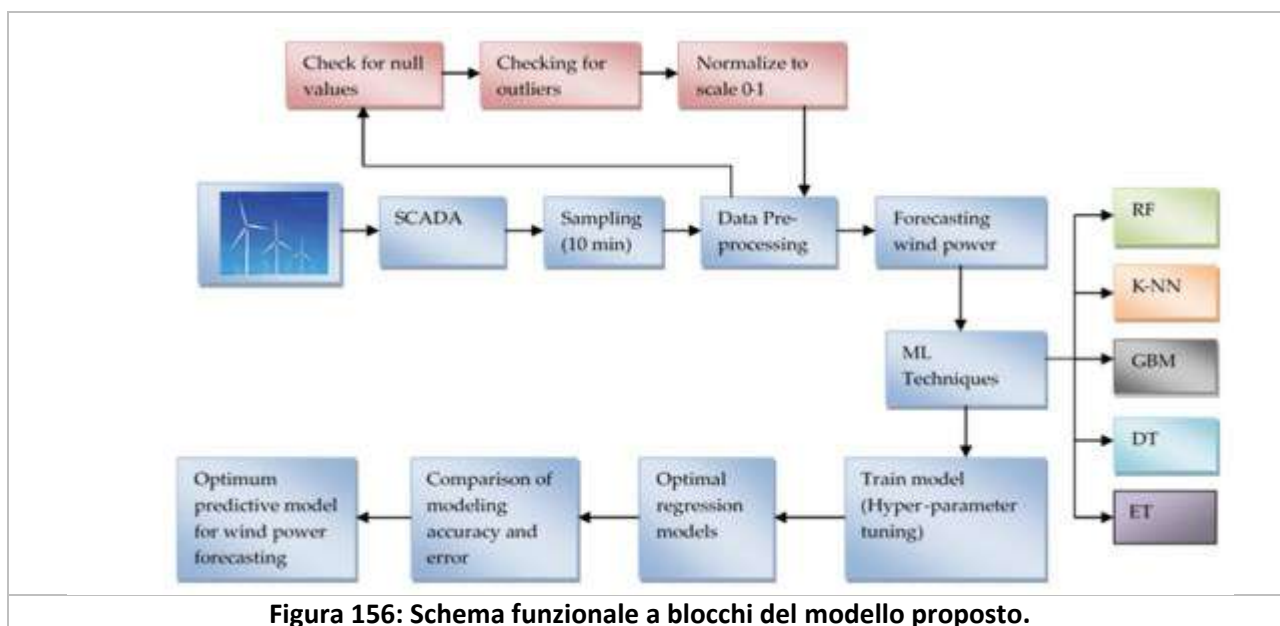
Figura 154: Risultati di previsione utilizzando modelli statici.

Models	France			Turkey			Kaggle		
	RMSE (kW)	MAE (kW)	R ²	RMSE (kW)	MAE (kW)	R ²	RMSE (kW)	MAE (kW)	R ²
SVR _L	130.89	85.36	0.91	111.98	73.04	0.97	130.89	85.36	0.91
SVR _Q	130.77	85.06	0.91	171.45	133.97	0.92	130.77	85.06	0.91
SVR _C	131.71	86.98	0.91	122.53	90.98	0.96	131.71	86.98	0.91
SVR _{FG}	149.47	97.84	0.89	133.39	100.31	0.95	149.47	97.84	0.89
SVR _{MG}	132.07	85.76	0.91	124.17	91.53	0.96	132.07	85.76	0.91
GPR _{RQ}	165.33	112.81	0.86	143.97	93.44	0.94	165.33	112.81	0.86
GPR _{SE}	199.63	125.11	0.80	176.64	101.93	0.92	199.63	125.11	0.80
GPR _{M52}	170.63	110.82	0.85	148.86	94.90	0.94	170.63	110.82	0.85
GPR _E	160.65	108.14	0.87	133.34	85.20	0.95	160.65	108.14	0.87
GPR _O	129.06	84.94	0.92	109.62	68.87	0.97	129.06	84.94	0.92
BS	129.26	87.08	0.92	114.71	71.94	0.96	129.26	87.08	0.92
BT	129.62	88.27	0.91	112.21	75.21	0.97	129.62	88.27	0.91
ELO	127.37	86.15	0.92	114.86	72.74	0.96	127.37	86.15	0.92
XGB	129.68	88.16	0.91	117.28	73.17	0.96	129.68	88.16	0.91
RF	133.70	90.41	0.91	118.87	80.09	0.96	133.70	90.41	0.91

Figura 155: Risultati di previsione utilizzando modelli dinamici.

Questo studio ha prima applicato e confrontato diversi approcci di macchine di apprendimento per modellare le dinamiche non lineari dell'energia eolica e prevedere le tendenze future dell'energia eolica. In particolare, **modelli di machine learning basati su kernel (Support Vector Regression (SVR) e Gaussian Process Regression (GPR))** e **i modelli di apprendimento d'insieme (Boosting, Bagging, Extreme Gradient Boosting (XGBoost) e Random Forest (RF))** sono considerati in questo studio comparativo. Per la valutazione vengono utilizzati i dati sull'energia eolica di tre turbine eoliche. I risultati hanno rivelato che utilizzando modelli dinamici che considerano le informazioni rilevanti dalle osservazioni passate del vento, le serie di potenze hanno migliorato i risultati della previsione rispetto ai modelli statici. Inoltre, questo studio ha dimostrato che l'inclusione di informazioni provenienti da variabili di input (ad es. velocità del vento e direzione del vento) migliora ulteriormente i risultati della previsione. Si osserva che non esiste un singolo approccio che domina gli altri per tutti gli esperimenti, ma in media, i modelli GPR e i modelli di ensemble hanno raggiunto prestazioni di previsione soddisfacenti, con un R^2 di circa 0,95.

Nell'articolo [88] si propongono e confrontano cinque macchine di regressione robuste ottimizzate, vale a dire **Random Forest (RF)**, **Gradient Boosting Machine (GBM)**, **K-nearest Neighbors algorithm (KNN)**, **Decision Trees (DT)** e **Extra Trees (XT)**, che vengono applicate per migliorare l'accuratezza della previsione della produzione di energia eolica a breve termine nei parchi eolici turchi, situati nella parte occidentale della Turchia, sulla base di dati storici di velocità e direzione del vento. I diagrammi polari vengono tracciati e l'impatto di variabili di input come la velocità e la direzione del vento sulla produzione di energia eolica vengono esaminati. Curve di dispersione che descrivono le relazioni tra la velocità del vento e quella della potenza prodotta dalla turbina viene tracciata per tutti i metodi e la potenza eolica media prevista viene confrontata con la potenza media reale della turbina con l'aiuto delle curve di errore tracciate.



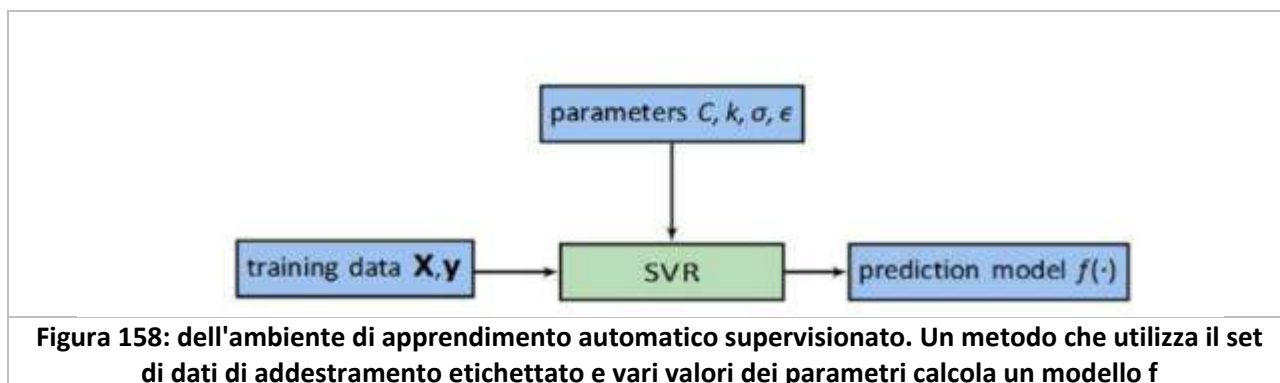
Regression Models	Performance Evaluation on Training Dataset					Performance Evaluation on Testing Dataset					Training Time (s)
	MAE	MAPE	RMSE	MSE	R ²	MAE	MAPE	RMSE	MSE	R ²	
Random Forest	0.0186	0.2966	0.0588	0.0040	0.9888	0.0277	0.3310	0.0672	0.0045	0.9651	11.9
K-NN	0.0278	0.2960	0.0580	0.0036	0.9742	0.0286	0.3248	0.0667	0.0044	0.9656	0.08
GBM	0.0260	0.0555	0.0228	0.0031	0.9897	0.0264	0.3012	0.0634	0.0040	0.9690	5.83
Decision Tree	0.0325	0.3213	0.0592	0.0055	0.9660	0.0336	0.3349	0.0884	0.0078	0.9497	0.22
Extra Tree	0.0274	0.2915	0.0522	0.0036	0.9782	0.0276	0.3243	0.0655	0.0041	0.9678	3.05

Figure 157: Prestazioni del modello basate sulle metriche MAE, MAPE, RMSE, MSE e R2.

In questo studio, l'analisi comparativa di vari metodi di machine learning è stata effettuata per prevedere l'energia eolica in base alla velocità del vento e ai dati sulla direzione del vento. Per raggiungere questo obiettivo, il parco eolico di Yalova, situato a ovest della Turchia, è stato utilizzato come caso di studio. Per la raccolta degli esperimenti è stato utilizzato un sistema SCADA di dati nel periodo da gennaio 2018 a dicembre 2018 a una frequenza di campionamento di 10 min per l'addestramento e il test dei modelli ML. I risultati mostrano che il **Random Forest (RF)**, **Gradient Boosting Machine (GBM)**, **K-nearest Neighbors algorithm (KNN)**, **Decision Trees (DT)** e **Extra Trees (XT)** sono potenti tecniche di previsione di energia eolica a breve termine. Tra questi algoritmi, la capacità dell'algoritmo **Gradient Boosting Machine (GBM)** con un valore MAE di 0,0277, un valore MAPE di 0,3310, un valore RMSE di 0,0672, un valore MSE di 0,0045 e un valore R2 di 0,9651 per la previsione di energia eolica, è stato verificato con maggiore accuratezza rispetto a RF, k-NN,

algoritmi DT ed ET. Le prestazioni dell'algoritmo DT non sono state soddisfacenti, con un MAE di 0,0336, MAPE di 0,3309, RMSE di 0,0884 e MSE di 0,0078.

Nell'articolo [89] si utilizzano i dati del National Renewable Energy Laboratory (NREL) e ci si concentra sulla proiezione sintetica dell'energia eolica. Riprendendo il modello di regressione come problema, si esplorano diverse strategie, tra cui modelli di regressione **K-nearest Neighbors algorithm (KNN)**, e **Support Vector Machine (SVM)**. Negli esperimenti, si guardano previsioni sia per singole turbine eoliche che per interi parchi eolici, dimostrando la fattibilità di una tecnica di apprendimento automatico per la proiezione dell'energia eolica a breve termine.



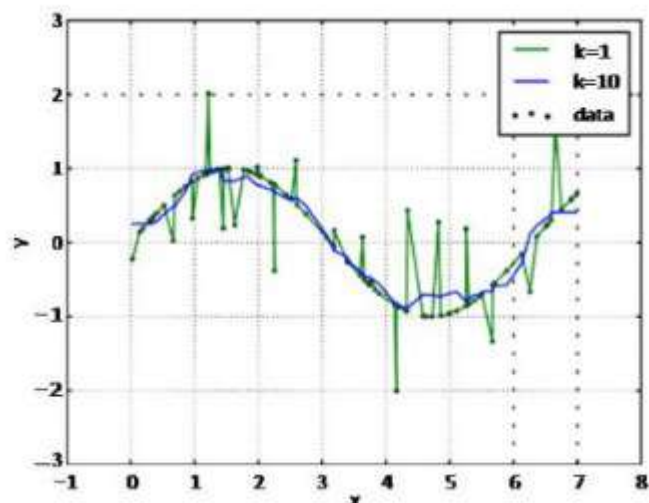


Figura 159: *Regressione con $k = 1$ e $k = 10$.*

Sono state utilizzate tecniche di classificazione per prevedere i valori delle osservazioni di velocità del vento fornite. Il vento medio giornaliero e i valori di velocità sono stati calcolati utilizzando i dati di velocità del vento orari e, la deviazione standard è stata utilizzata per ottenere la potenza eolica totale giornaliera.

Il metodo proposto è stato provato in diversi punti per vedere se gli algoritmi potevano produrre risultati adatti al sito di addestramento. Le previsioni sull'energia eolica totale giornaliera a lungo termine si sono dimostrate accurate utilizzando gli algoritmi **Extreme Gradient Boosting (XGBoost)**, **Support Vector Regression (SVR)** e **Random Forest (RF)**. RF si è dimostrato il più efficace, con un punteggio R^2 di 0,995 e un MAE di 7,048. L'algoritmo LASSO lineare è di gran lunga quello più debole. D'altra parte, il valore LASSO R^2 è 0,862. Il risultato chiave della ricerca è che, utilizzando il modello di energia eolica per un sito di base, è possibile utilizzare algoritmi di apprendimento automatico per decidere se ha senso costruire parchi eolici in regioni inesplorate.

Lo studio dell'articolo [90] propone un modello di previsione per risolvere un problema reale nel settore delle energie rinnovabili stimando con precisione la quantità di energia eolica prodotta all'ora nelle prossime 24 ore, applicando tecniche di apprendimento automatico (ML) e utilizzando

dati storici sulla generazione di energia eolica e bollettini meteorologici. Nell'approccio proposto, in primo luogo, un metodo ML non supervisionato (ovvero, **l'algoritmo di clustering K-Means**) viene applicato per raggruppare i dati in cluster significativi; quindi, questi cluster vengono accettati come nuovi valori di funzionalità e aggiunti al set di dati per ingrandirlo; infine, viene eseguito il metodo ML supervisionato per la previsione. Questo studio mette a confronto nove algoritmi di apprendimento supervisionato: **K-nearest Neighbors algorithm (KNN)**, **Support Vector Regression (SVR)**, **Random Forest (RF)**, **Extra Trees (XT)**, **Gradient Boosting Machine (GBM)**, **Ridge Regression (RR)**, **Least Absolute Shrinkage and Selection Operator**, **Decision Trees (DT)** e **Convolutional Neural Network (CNN)**. Lo scopo di questo studio è quello di indagare il successo di diversi algoritmi ML sui dati nel mondo reale delle turbine eoliche e propone una metodologia per confrontare vari algoritmi di machine learning per scegliere il modello finale più accurato per la generazione della previsione di energia eolica.

Algorithm	Without feature generation process (MAE)	With feature generation process (MAE) (Proposed Approach)
KNN	0.7720	0.7719
SVR	0.7283	0.7289
Random Forest	0.7049	0.7015
Extra Trees	0.7290	0.7269
Gradient Boosting	0.7742	0.7671
Ridge Regression	0.9363	0.9336
LASSO	1.0632	1.0632
Decision Tree	0.8329	0.8260
CNN	0.9433	0.9660

Figura 160: Confronto dei risultati MAE con e senza processo di generazione delle funzionalità

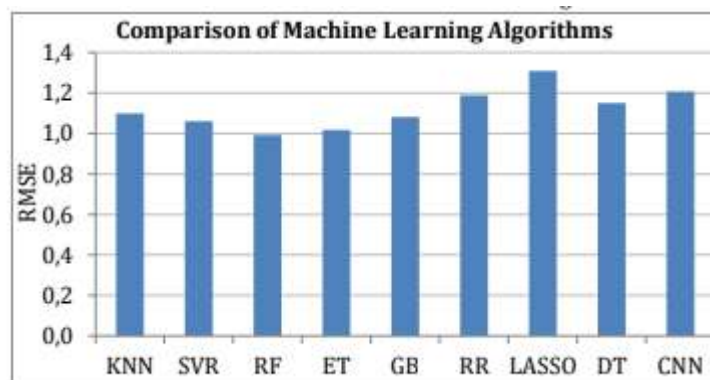


Figura 161: Confronto di algoritmi di apprendimento automatico per la previsione della generazione di energia eolica in termini di valori RMSE

Il paper ha indagato sugli algoritmi di apprendimento di previsione della produzione di energia. L'obiettivo di questo studio è stato quello di risolvere un problema della vita reale nel settore delle energie rinnovabili, cioè stimare accuratamente la quantità di produzione oraria di energia eolica nelle successive 24 ore. Per questo motivo, questo studio indaga e confronta molti algoritmi di apprendimento automatico per la previsione della produzione di energia eolica, inclusi KNN, SVR, RF, ET, GB, RR, LASSO, DT, e CNN.

L'articolo [91] utilizza un metodo per prevedere la produzione di energia eolica utilizzando tecniche di machine learning. In particolare, viene utilizzato un insieme di **modelli di previsione (Ensemble Learning)** per migliorare l'accuratezza della previsione. L'obiettivo è quello di prevedere la produzione di energia eolica per il giorno successivo, al fine di consentire ai trader di acquistare o vendere energia sul mercato in base alle previsioni. Questo approccio utilizza dati storici sulla produzione di energia eolica, condizioni meteorologiche, geografiche e altre variabili per addestrare i modelli di previsione. L'ensemble learning combina i risultati di diversi modelli di previsione per ottenere una previsione finale più accurata. In questo modo, il modello di previsione è in grado di considerare una vasta gamma di variabili e fornire una previsione affidabile sulla produzione di energia eolica.

Algorithm	SW1	SW2	SW3	SW4	MW1	MW2	NW1	NI1
Gradient Boosting	0.1451	0.1442	0.1676	0.1309	0.1197	0.1240	0.1307	0.1255
Linear Regression	0.1485	0.1477	0.1754	0.1403	0.1254	0.1409	0.1383	0.1387
Bayesian Ridge	0.1485	0.1476	0.1746	0.1402	0.1254	0.1408	0.1382	0.1387
SVR	0.1461	0.1424	0.1498	0.1332	0.1248	0.1371	0.1355	0.1364
LSTM	0.1516	0.1521	0.1556	0.1356	0.1238	0.1325	0.1395	0.1407
Dense LSTM	0.1490	0.1515	0.1549	0.1339	0.1249	0.1329	0.1382	0.1362
LSTM Ens. [G]	0.1433	0.1345	0.1174	0.1594	0.1207	0.1359	0.1490	0.1215
Neural Net	0.1488	0.1529	0.1578	0.1315	0.1166	0.1264	0.1361	0.1313
NN Ens. [M]	0.1499	0.1521	0.1595	0.1320	0.1168	0.1307	0.1327	0.1315
NN Ens. [G]	0.1411	0.1353	0.1168	0.1611	0.1238	0.1324	0.1514	0.1223
Boosting Ens. [G]	0.1433	0.1353	0.1182	0.1486	0.1245	0.1311	0.1502	0.1227
Boosting Ens. [M]	0.1460	0.1446	0.1616	0.1317	0.1188	0.1233	0.1308	0.1309
Boosting Ens. [G]	0.1433	0.1353	0.1182	0.1486	0.1245	0.1311	0.1502	0.1227
xGBoost	0.1454	0.1431	0.1675	0.1305	0.1190	0.1242	0.1309	0.1247
Reference Curve	0.2360	0.2279	0.1317	0.2051	0.1444	0.1701	0.2661	0.1922

Figura 162: Risultati medi per modello su 20 run per ogni posizione.

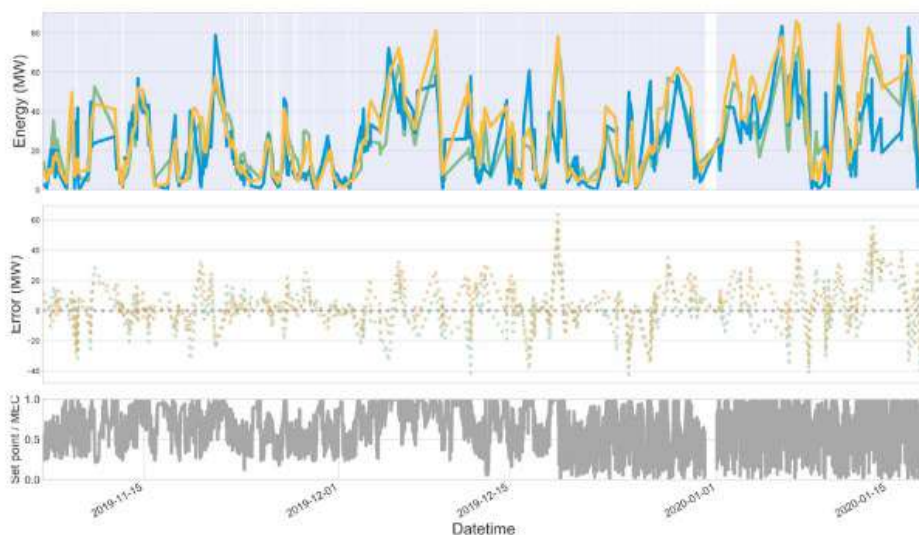


Figura 163: Previsioni del tempo in una singola run in Irlanda del Nord

Le ricerche hanno dimostrato che questo approccio può migliorare significativamente l'accuratezza della previsione della produzione di energia eolica rispetto ai modelli di previsione tradizionali. In particolare, l'utilizzo dell'ensemble learning ha permesso di ottenere una previsione più affidabile grazie alla combinazione di più modelli di previsione. Inoltre, l'ensemble learning ha dimostrato di essere in grado di selezionare automaticamente i modelli di previsione più accurati per una specifica situazione.

Le ricerche hanno inoltre evidenziato che l'accuratezza della previsione può essere influenzata da diversi fattori, come la disponibilità di dati storici sulla produzione di energia eolica, le condizioni meteorologiche e geografiche, e la qualità dei modelli di previsione utilizzati. In generale, gli studi

hanno mostrato un approccio promettente per migliorare la gestione dell'energia eolica e la sua integrazione nel sistema energetico, consentendo una migliore pianificazione della produzione di energia eolica e una gestione più efficiente del mercato energetico.

L'articolo [92] discute sulla ricerca di **Deep Learning** e le applicazioni dell'energia del vento e delle onde ed evidenzia il potenziale dell'utilizzo delle tecniche di deep learning nel settore delle energie rinnovabili. Il paper discute vari modelli di deep learning, come le **Convolutional Neural Network (CNN)**, le **Long Short-Term Memory (LSTM)** e gli **autoencoder**, che sono stati applicati ai dati dell'energia del vento e delle onde.

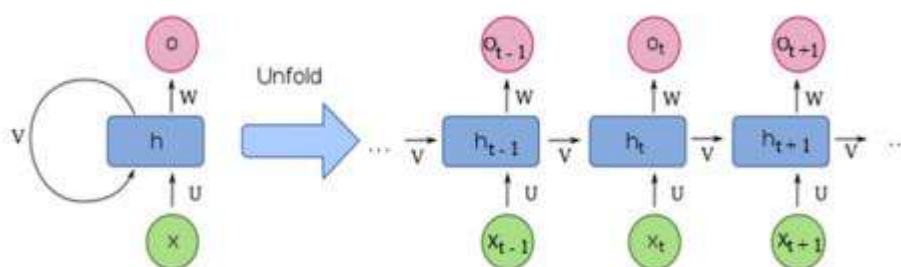


Figura 164: Semplice struttura di una rete neurale ricorrente

Forecasting Application	MAE	RMSE	R ²	MAPE	SI	Models	Lead Time
Wind speed [59]	0.28	0.49		0.257		SC-LSTM	
Wind speed [66]	0.3	0.3925		0.4285		CNN-LSTM	
Wind speed [63]	0.07988	0.1052	0.9992			LSTM	
Wind speed [72]	0.13–0.27	0.14–0.19				LSTM	6 h
Wind speed [81]	0.3509	0.5193	0.981			CNN + LSTM	1–3 h
Wave height and wind speed [37]		0.844				CNN + GRU	1–6 h
Wave condition and power [54]		0.49				LSTM	
Wave height [56]	0.78	0.8638					3–24 h
Wave height [61]	0.425	0.64			0.109		
Wind power [49]	0.3576	0.5058	-	0.2173			

Figura 165: Riepilogo delle prestazioni di diversi modelli di deep learning

Il paper ha esaminato studi recenti sulle applicazioni del vento e delle onde nel forecasting, nell'ottimizzazione e nel riconoscimento di pattern, che sono i principali focus di questi due tipi di energie rinnovabili. Sono state elencate le differenze nei dataset e i relativi metodi di preprocessing dei dati. La maggior parte dei metodi di estrazione delle caratteristiche si sono concentrati sulle CNN, che sono flessibili e consistono in strutture supervisionate e non supervisionate, conosciute insieme come “fine tuning”. Successivamente, il paper ha esaminato i modelli di previsione con il dataset gestito utilizzando il preprocessing dei dati, tra cui **Recurrent Neural Network (RNN)**, **Long Short-Term Memory (LSTM)**, **Gated Recurrent Unit (GRU)**, **Direction Basis Function Neural Networks (DBF-NN)** e **modelli ibridi**. Vengono elencati i vantaggi e gli svantaggi per aiutare a scegliere un modello appropriato per la ricerca futura. Gli indicatori di stima includono RMSE, MAE, MAPE, indice di dispersione e l'efficienza di correzione per valutare le prestazioni dei modelli. L'ottimizzazione dell'applicazione del vento e delle onde potrebbe estendersi a molte aree dell'ottimizzazione dei dispositivi, dell'ottimizzazione della disposizione e della gestione dell'energia.

Nell'articolo [93] è stato creato un algoritmo di ottimizzazione ibrido, che combina la scomposizione modale variazionale (VMD), l'algoritmo di massima rilevanza e minima ridondanza (mRMR), **la Long Short-Term Memory (LSTM)** e il **Firefly algorithm (FA)** per la previsione accurata della potenza del vento. In primo luogo, la sequenza originale di potenza del vento storica viene scomposta in diverse

funzioni di modello caratteristiche con VMD. Successivamente, mRMR viene applicato per ottenere il miglior insieme di caratteristiche analizzando la correlazione tra ciascuna componente. Infine, FA viene utilizzato per ottimizzare i vari parametri LSTM. L'aggiunta dei risultati di previsione di tutte le sottosequenze produce il risultato di previsione. Si è scoperto che il metodo ibrido proposto è superiore agli altri sei algoritmi confrontati. Allo stesso tempo, viene fornito un caso aggiuntivo per verificare ulteriormente l'adattabilità e la stabilità del modello ibrido proposto.

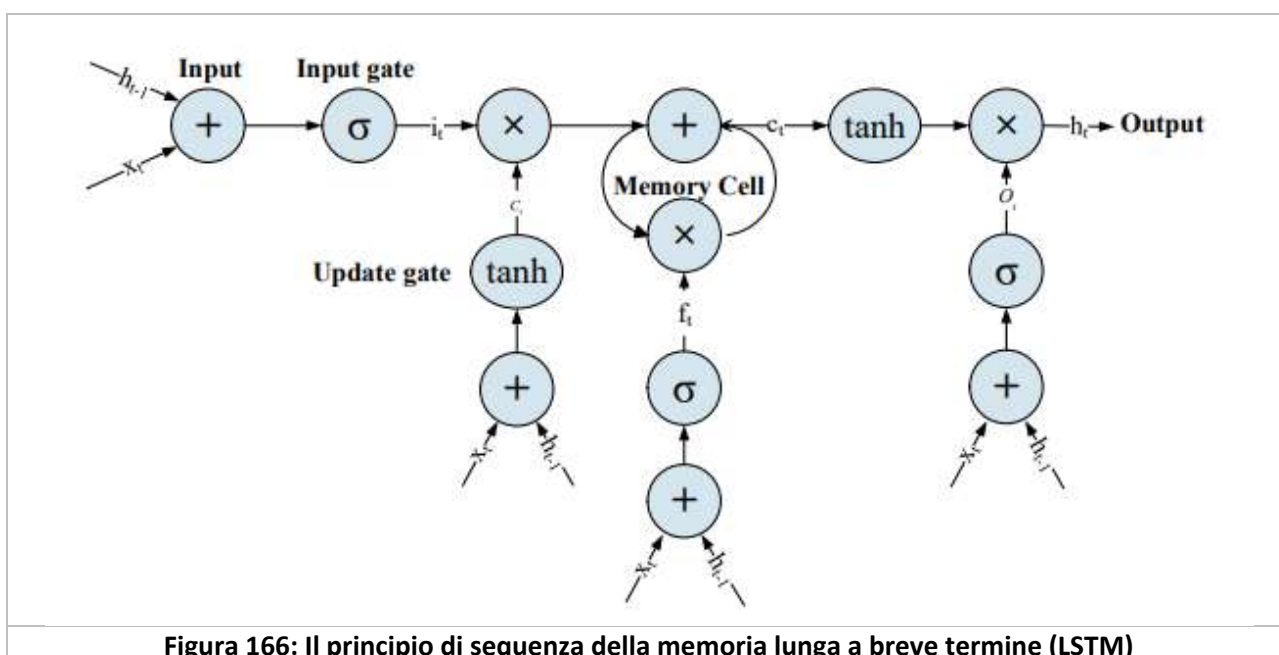


Figura 166: Il principio di sequenza della memoria lunga a breve termine (LSTM)

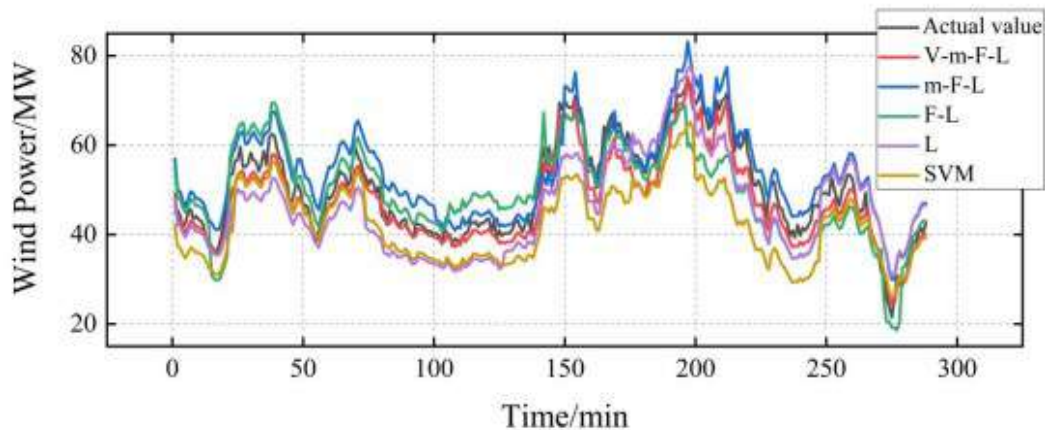


Figura 167: Risultati finali di previsione

Questa ricerca si propone di affrontare l'instabilità della previsione di potenza eolica a breve termine attraverso la proposta di un modello ibrido, denominato VMD-mRMR-FA-LSTM. Il modello utilizza la decomposizione VMD per scomporre il carico originale; quindi, l'algoritmo mRMR viene utilizzato per selezionare le migliori caratteristiche da utilizzare, tra cui temperatura, velocità del vento, direzione del vento, ecc. Successivamente, vengono costruiti diversi modelli di previsione LSTM per ogni nuova sequenza sulla base del risultato della selezione mRMR. Infine, l'algoritmo FA viene utilizzato per ottimizzare i parametri di LSTM. Il caso di studio del modello ibrido proposto dimostra che:

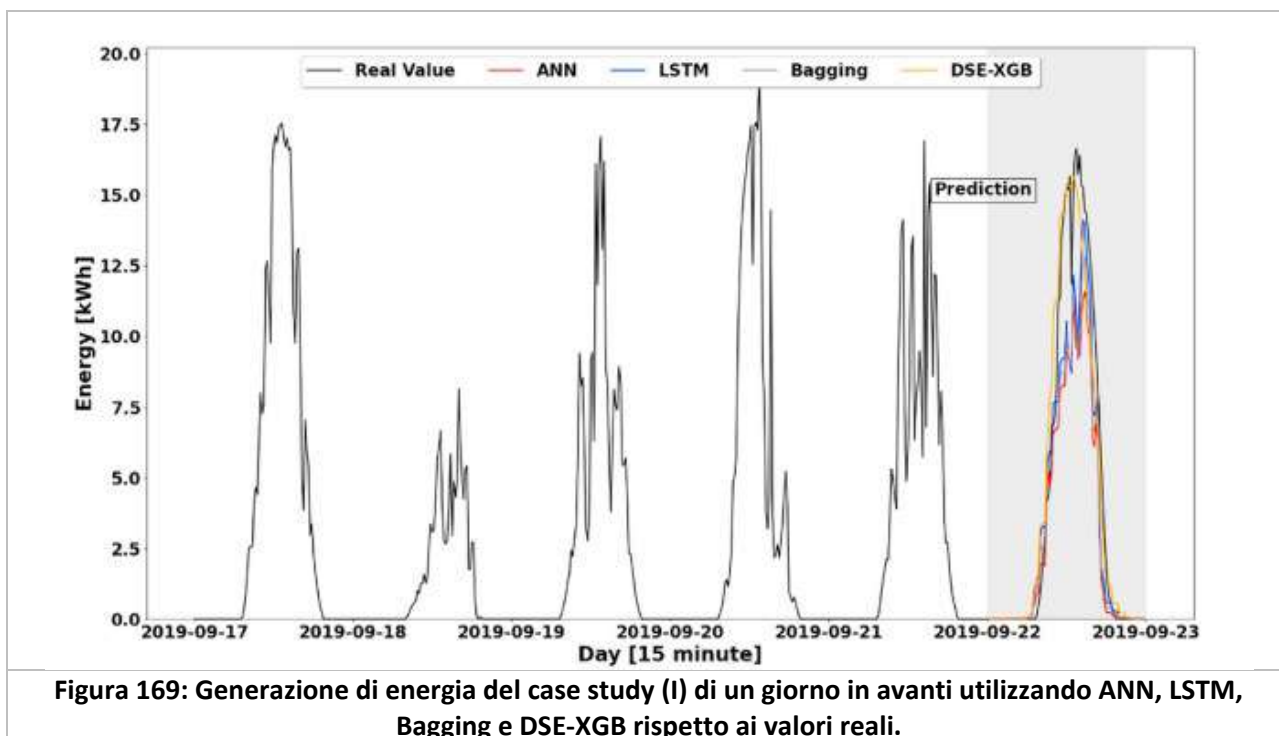
- Il modello ibrido ha una maggiore precisione di previsione rispetto ai modelli di benchmarking, e ha ampie prospettive di applicazione nella previsione di potenza eolica a breve termine.
- Rispetto al singolo modello LSTM, l'algoritmo FA può ottimizzare i parametri e la funzione di LSTM per ottenere una maggiore precisione di previsione, il che indica che il modello FA-LSTM ha una maggiore capacità di ricerca globale e una maggiore stabilità nella prestazione di previsione.
- Rispetto ad altre strategie di preprocessing dei dati, VMD-mRMR ha una migliore prestazione e migliora efficacemente la precisione di previsione. Il modello ibrido proposto in questo articolo può essere applicato con successo alla previsione di potenza eolica a breve termine.

3.2.1.2 Energia solare

Nel documento [94], viene proposto un algoritmo stacked ensemble migliorato generalmente applicabile (DSE XGB), che utilizza due algoritmi di **deep learning**, vale a dire **Artificial Neural Network (ANN)** e **Long Short-Term Memory (LSTM)** come modelli di base per la previsione dell'energia solare. I modelli di previsione sono integrati utilizzando un algoritmo di potenziamento del gradiente estremo per migliorare la precisione della previsione della generazione fotovoltaica. Il modello proposto è stato valutato su quattro diversi dataset di generazione solare fornendo una valutazione completa. Inoltre, il quadro esplicativo è utilizzato in questo studio per fornire una visione più approfondita del meccanismo di apprendimento dell'algoritmo. Le prestazioni del modello proposto sono state valutate confrontando i risultati della previsione con i modelli individuali ANN e LSTM.

Case Study (I)				
Error Metrics	Model			
	ANN	LSTM	Bagging	DSE-XGB
R-squared	0.92	0.91	0.92	0.99
RMSE (kWh)	0.87	0.90	0.87	0.35
MAE (kWh)	0.45	0.49	0.45	0.22
R-squared	0.92	0.92	0.92	0.98
RMSE (kWh)	0.60	0.60	0.60	0.26
MAE (kWh)	0.31	0.32	0.31	0.15
R-squared	0.92	0.92	0.92	0.99
RMSE (kWh)	0.31	0.31	0.31	0.14
MAE (kWh)	0.17	0.17	0.17	0.09
R-squared	0.93	0.92	0.92	0.98
RMSE (kWh)	1.73	1.83	1.74	0.74
MAE (kWh)	0.82	1.01	0.88	0.47
Case study (II)				
R-squared	0.90	0.91	0.89	0.96
RMSE (kWh)	1.08	0.93	0.89	0.78
MAE (kWh)	1.02	0.89	0.82	0.59

Figura 168: Metriche delle prestazioni R-quadrato, RMSE (kWh) e MAE (kWh) di tutti i modelli



In questo studio viene proposto un algoritmo di deep learning migliorato combinando **Artificial Neural Network (ANN)**, **Long Short-Term Memory (LSTM)** e **Extreme Gradient Boosting (XGBoost)**. La proposta del metodo DSE-XGB ha superato i singoli algoritmi di deep learning. Il modello ANN ha catturato in modo esplicito la dipendenza di previsione della generazione solare fotovoltaica mentre LSTM ha estratto le tendenze ripetitive dai dati. Questo studio ha integrato il machine learning interpretabile con la modellazione per capire come il meta-discente apprende dalle previsioni del modello di base.

Lo scopo principale dell'articolo [95] è quello di esplorare la relazione tra numerosi parametri di input e l'energia solare fotovoltaica (FV) utilizzando modelli di machine learning (ML). Due diversi approcci ML, come la **Support Vector Machine (SVM)** e la **Gaussian Process Regression (GPR)** sono state considerate e confrontate. I parametri di input di base, incluso il fotovoltaico solare, la temperatura del pannello, la temperatura ambiente, il flusso solare, l'ora del giorno e l'umidità

relativa sono state considerate per la previsione dell'energia solare fotovoltaica. I risultati hanno mostrato che tra gli approcci ML proposti, l'algoritmo Matern 5/2 GPR ha fornito la prestazione ottimale; mentre l'SVM cubico ha avuto le prestazioni peggiori. Inoltre, i risultati di output previsti sono in buon accordo con i valori sperimentali, indicando che gli approcci ML proposti sono appropriati per l'uso nella previsione della potenza di un diverso pannello solare fotovoltaico. Inoltre, per mostrare l'efficacia e l'accuratezza dei modelli SVM e GPR nelle previsioni, i risultati di questi modelli vengono confrontati utilizzando l'errore quadratico medio (RMSE) e l'errore assoluto medio (MAE).

Algorithm	RMSE	MAE	R ²
Linear SVM	13.115	10.244	0.96
Cubic SVM	21.72	15.667	0.88
Quadratic SVM	14.984	11.914	0.94
Rational quadratic GPR	9.052	6.453	0.98
Squared exponential GPR	9.0658	6.624	0.98
Matern 5/2 GPR	7.967	5.3025	0.98

Figura 170: Risultati degli algoritmi ML per la potenza dei pannelli fotovoltaici.

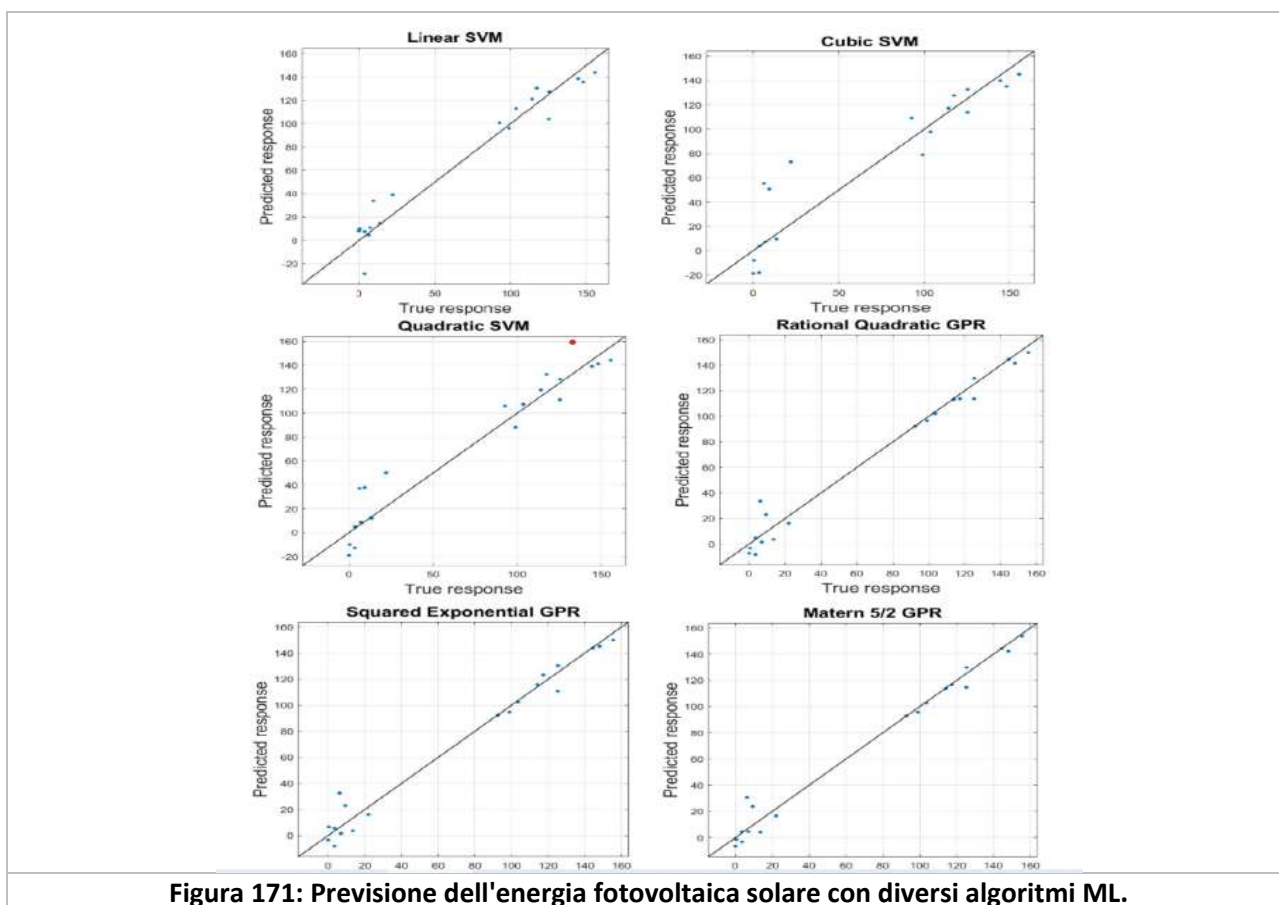


Figura 171: Previsione dell'energia fotovoltaica solare con diversi algoritmi ML.

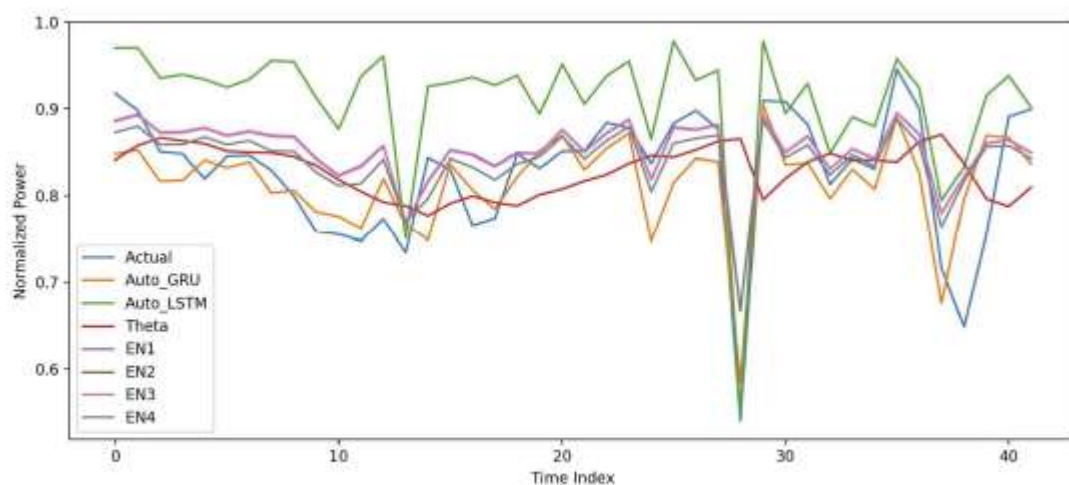
Questo studio ha rivelato che i modelli ML potrebbero essere impiegati come strumento rapido per prevedere le prestazioni adeguate della potenza di qualsiasi pannello solare fotovoltaico. Indipendentemente dalle ottime prestazioni degli approcci ML suggeriti, ulteriori parametri, inclusa la dimensione del file del pannello solare fotovoltaico, polvere e vento possono anche influenzare notevolmente le potenze solari fotovoltaiche. La conoscenza di questo effetto è ancora incompleta in termini di indagine sperimentale e analisi teorica.

Nella ricerca dell'articolo [96] si propone un modello ibrido che combina metodi di apprendimento automatico con il metodo statistico Theta per una previsione più accurata della futura generazione di energia solare da impianti di energia rinnovabile. I modelli di machine learning includono **la Long Short-Term Memory (LSTM), la Gated Recurrent Unit (GRU)), l'AutoEncoder LSTM (Auto-LSTM) e una Auto-GRU appena proposta**. Per migliorare la precisione del machine learning proposto e dello

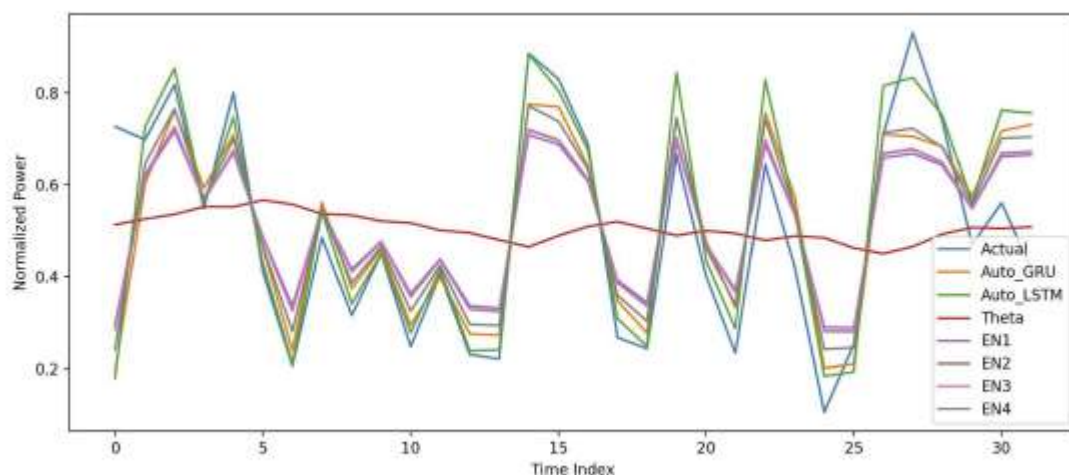
Statistical Hybrid Model (MLSHM), si utilizzano due tecniche di diversità, ovvero la diversità strutturale e la diversità dei dati. Per combinare la previsione dei membri dell'ensemble nel MLSHM proposto, si sfruttano quattro metodi di combinazione: approccio di media semplice, media ponderata con approccio lineare, approccio non lineare e combinazione attraverso la varianza utilizzando l'approccio inverso. Lo schema proposto MLSHM è stato convalidato su due serie di dati in tempo reale, ovvero Shagaya in Kuwait e Cocoa negli Stati Uniti. Gli esperimenti mostrano che l'MLSHM proposto, utilizzando tutti i metodi di combinazione, ha raggiunto una maggiore precisione rispetto alla previsione dei tradizionali modelli individuali.

Dataset	Shagaya Poly-SI [33]	Shagaya TFSC [33]	Cocoa Poly-SI [34]
System Size	5 MW	5 MW	166 W
No. of Days Before Preprocessing	420	420	316
No. of Days After Preprocessing	376	409	316
Training Set	304	330	255
Validating Set	34	37	29
Testing Set	38	42	32
Weather Parameters	GHI, GTI, wind speed, wind direction, air & panel's temperature	GHI, GTI, wind speed, wind direction, air & panel's temperature	Solar irradiance, ambient temperature, humidity & precipitation

Figura 172: Sommario di dataset di serie temporali.



(b) Shagaya TFCS Dataset.



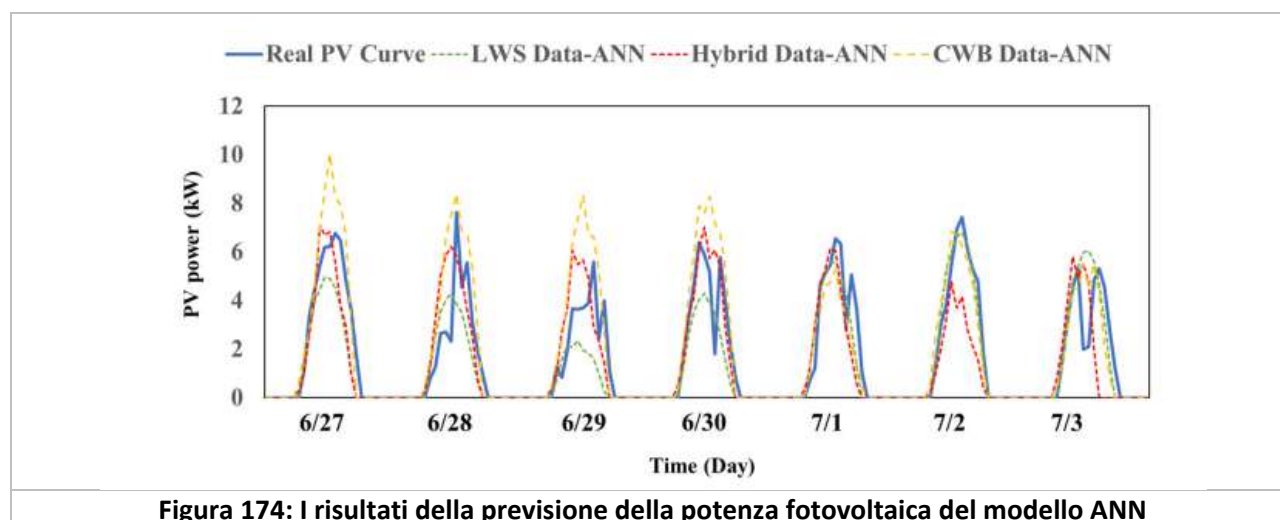
(c) Cocoa Dataset.

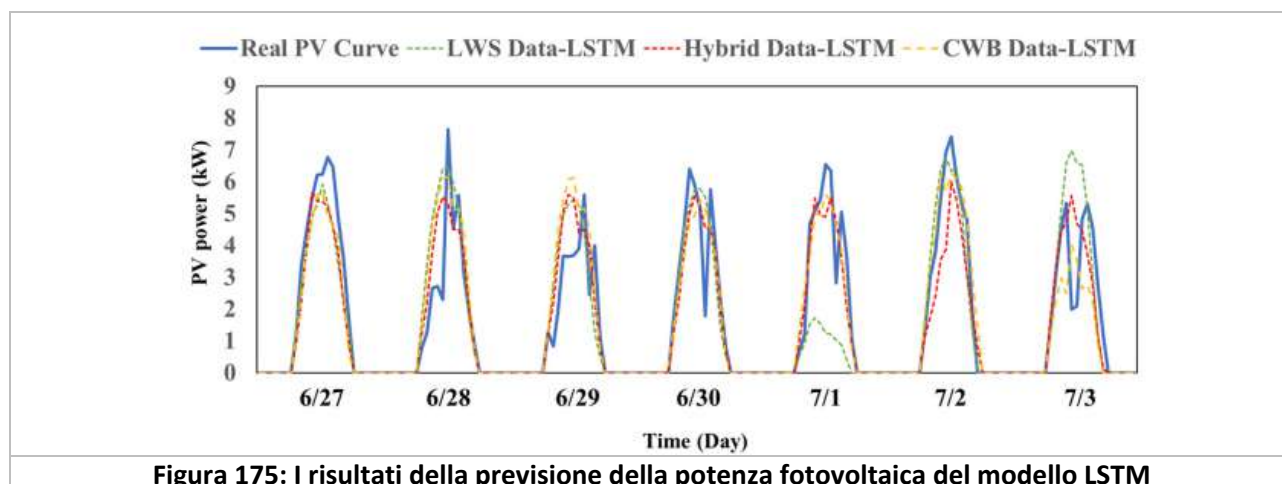
Figura 173: Effettiva potenza fotovoltaica solare prevista dei due set di dati.

In questa ricerca, si è proposto un modello ibrido (MLSHM) che combina i risultati di previsione di entrambi i modelli ML e del metodo statistico. Per lo studio si è sviluppato un nuovo modello ML, Auto-GRU, che apprende da dati di serie temporali storiche per prevedere la potenza fotovoltaica solare desiderata. Al fine di potenziare il modello ibrido, in questo studio sono state condotte due tecniche di diversità, ovvero la diversità strutturale tra i membri dell'ensemble e la diversità dei dati tra i training set del modello ML. Il modello ibrido proposto e il modello Auto-GRU sono stati testati su set di dati di serie in tempo reale di energia solare fotovoltaica e dati meteorologici raccolti da Shagaya situati in Kuwait e Cocoa, Florida, USA. Gli esperimenti permettono di concludere che un

modello ibrido che combina la previsione dei modelli ML e il metodo statistico ottiene maggiore precisione rispetto a un modello ibrido che combina la previsione dei modelli ML senza metodo statistico.

L'articolo scientifico [97] mette a confronto i risultati di previsione dell'energia fotovoltaica con un giorno di anticipo su tre modelli abbinati a tre gruppi di dati meteorologici. Dal momento che il numero, la perdita e il problema della corrispondenza dei dati meteorologici influenzeranno tutti i risultati dell'addestramento del modello, un framework di dati di pre-elaborazione viene proposto per risolvere il problema in questo caso di studio. I modelli utilizzati sono basati su **algoritmi di deep learning come Artificial Neural Network (ANN), Long Short-Term Memory (LSTM) e Gated Recurrent Unit (GRU)**. I gruppi di dati meteorologici sono il Central Weather Bureau (CWB), la stazione meteorologica locale (LWS) e i dati ibridi (la combinazione di dati CWB e LWS).





Rispetto agli altri due gruppi, i dati ibridi hanno mostrato un miglioramento del 5-8% nelle misurazioni. Inoltre, quando si tratta di condizioni meteorologiche diverse, sono stati evidenziati i vantaggi del modello LSTM. Dopo ulteriori analisi, il modello LSTM combinato con i dati ibridi ha mostrato le misurazioni più accurate, il che è stato dimostrato attraverso la previsione dei risultati per un mese. Infine, i risultati indicano che quando la quantità di dati sono limitati, l'utilizzo di dati ibridi e delle cinque funzioni meteorologiche è utile per addestrare il modello. Di conseguenza, il modello proposto mostra una migliore previsione FV con un giorno in avanti.

A causa dell'imprevedibilità nella generazione fotovoltaica, è fondamentale pianificare in avanti la generazione di energia solare, come viene richiesto nella previsione. La produzione di energia solare è dipendente dal tempo e imprevedibile, infatti questa previsione è complessa. Nell'articolo [98] vengono discusse le condizioni ambientali sull'uscita di un sistema fotovoltaico. Algoritmi di Machine Learning (ML) hanno mostrato ottimi risultati nella previsione di serie temporali e quindi possono essere utilizzati per anticipare la potenza con il tempo come input del modello. L'uso di più machine learning e deep learning vengono usate per eseguire previsioni sull'energia solare. Qui vengono usate tecniche di apprendimento come il **Support Vector Machine (SVM)**, il **Random Forest (RF)** e **Linear Regression (LR)** tra cui, il random forest regressor ha battuto gli altri due modelli di regressione con ottima precisione.

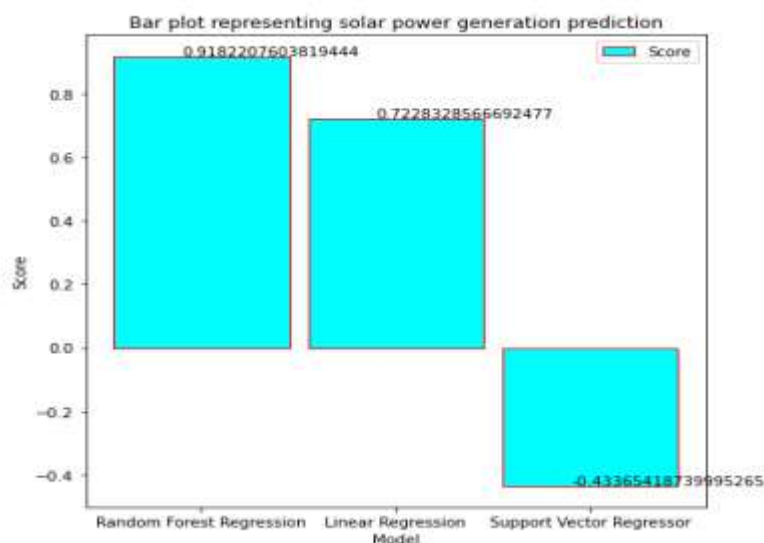


Figura 176: Punteggio di tutti e tre i modelli utilizzati nell'analisi dell'energia solare

Model	RMSE	MAE	MSE	Accuracy
Support vector machine regressor	131.44	77.16	172.76	-44.92
Linear regression	58.57	48.39	343.08	72.4
Random forest regressor	27.32	12.45	746.48	94.01

Figura 177: Modello di previsione di RMSE, MAE, MAPE e MBE e Precisione

È stato presentato un approccio basato sull'apprendimento automatico per analisi della generazione di energia solare, che prevede con precisione l'energia generata negli stati dell'India sulla base di dati ambientali. Soprattutto, la metodologia è andata oltre la previsione fornendo risultati chiave che hanno aiutato nella comprensione dell'analisi dell'energia solare (importanza variabile per periodo). Da un largo margine, il metodo proposto ha sovraperformato gli altri metodi popolari, come **Random Forest (RF)**. I modelli proposti sono **Support Vector Regression (SVR)**, **Linear Regression (LR)** e **Random Forest (RF)** rispetto alla temperatura con i dati forniti. La

temperatura <46F è una media della temperatura del 17%. Come visto nei risultati della precedente figura, si può vedere che il modello Random Forest Regressor sta ottenendo risultati migliori con il 94,01% di precisione e quindi quel modello è preferito per la distribuzione.

Con il miglioramento dell'integrazione della generazione di energia solare, l'energia fotovoltaica, le previsioni svolgono un ruolo significativo nel garantire la sicurezza operativa. Pertanto, nell'articolo [99], in base alla potenza FV dell'impianto a Lijiang, considerando i fattori correlati che influenzano la produzione fotovoltaica come l'irraggiamento solare, temperatura ambiente, pressione atmosferica, velocità del vento, direzione del vento e dati di generazione storici della centrale fotovoltaica, vengono utilizzati tre algoritmi di reti neurali, cioè **Back Propagation (BP)**, **Genetic Algorithm optimized Back Propagation (GA-BP)** e **Particle Swarm Optimization Back Propagation (PSO-BP)**, per costruire un modello di previsione a breve termine della produzione fotovoltaica.

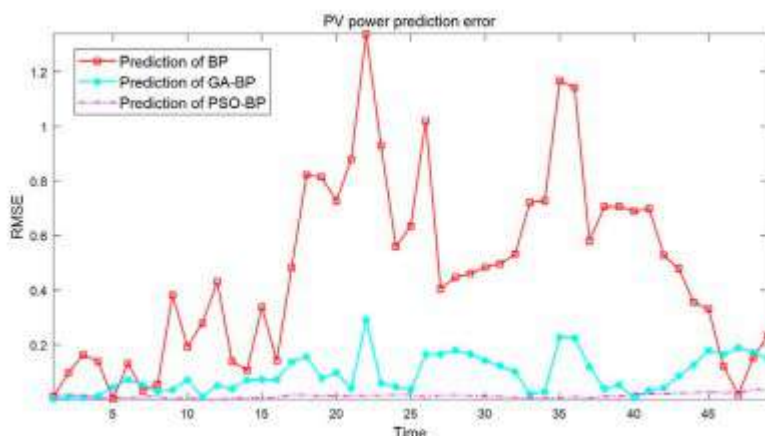


Figura 178: MAPE di tre modelli di previsione della potenza fotovoltaica

Neural network algorithm	BP	GA-BP	PSO-BP
RMSE/MW	0.57432	0.11499	0.015805
MAPE	11.2197	3.5884	0.94797

Figura 179: Dati dell'errore di previsione di ciascun algoritmo di rete neurale

I risultati della simulazione mostrano che GA-BP e i modelli di previsione della rete PSO-BP ottengono entrambi un'elevata precisione di previsione, che indica che i metodi GA e PSO possono ridurre efficacemente gli errori di previsione al contrario del modello BP originale. In particolare, PSO possiede una migliore applicabilità rispetto a GA, che può ridurre ulteriormente gli errori del modello di previsione della potenza fotovoltaica.

Studi esistenti hanno stabilito la superiorità degli algoritmi della rete neurale artificiale (ANN) e del random forest (RF) in questo campo. Tuttavia, studi più recenti hanno dimostrato che l'energia solare fotovoltaica ha promettenti prestazioni di previsione dell'output, se gestita mediante il **Decision Trees (DT), Extreme Gradient Boosting (XGBoost) e le Long Short-Term Memory (LSTM)**. Pertanto, lo studio del documento [100] si propone di affrontare una ricerca in questo campo determinando il miglior esecutore tra questi 5 algoritmi. Un set di dati da National Renewable Energy Laboratory (NREL) degli Stati Uniti costituito da parametri meteorologici, è stato utilizzato per i dati sulla produzione di energia solare per un modulo fotovoltaico in silicio monocristallino a Cocoa, in Florida.

Algorithm	Best model	MAE (rank)	MAPE (rank)	RMSE (rank)	R^2 (rank)	RM
RF	10	0.4852 (2)	0.0314 (1)	0.9209 (2)	0.9987 (2)	1.75
DT	10	0.6183 (3)	0.0388 (3)	1.0722 (4)	0.9983 (3)	3.25
XGB	7	0.6310 (4)	0.0403 (4)	1.0667 (3)	0.9983 (3)	3.5
ANN	16	0.4693 (1)	0.0338 (2)	0.8816 (1)	0.9988 (1)	1.25
LSTM	8	5.4641 (5)	0.3115 (5)	10.0652 (5)	0.8474 (5)	5

Figura 180: Confronto del punteggio delle prestazioni dei migliori modelli per ciascun algoritmo

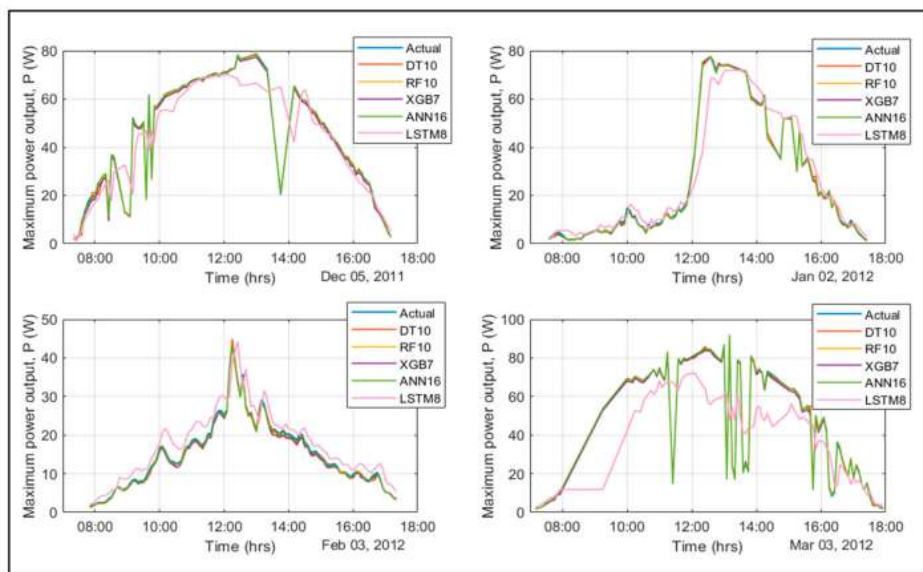
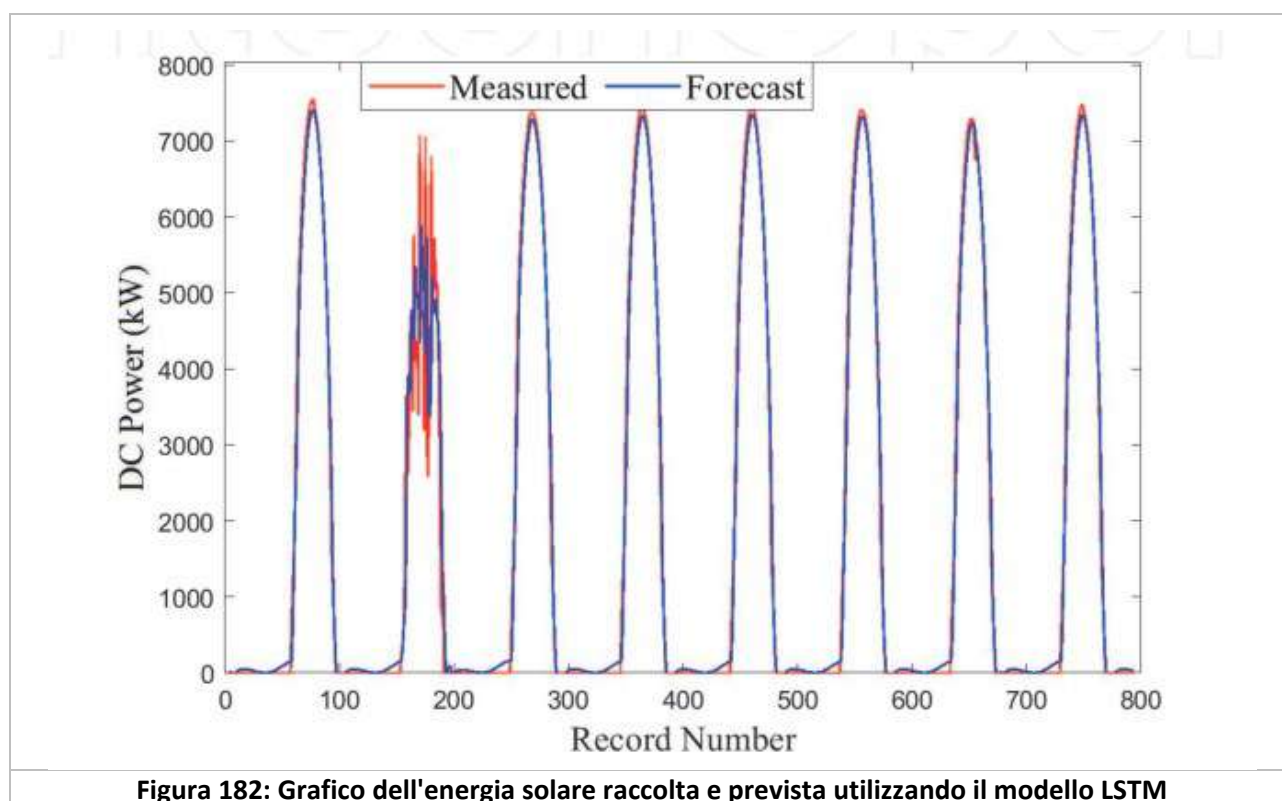


Figura 181: Produzione di energia solare fotovoltaica effettiva e prevista rispetto al tempo per i migliori modelli

I confronti dei punteggi di previsione mostrano che l'algoritmo ANN è superiore a quanto prodotto dal modello ANN16, che produce il miglior errore medio assoluto (MAE), il miglior errore quadratico medio (RMSE) e il miglior coefficiente di determinazione (R^2) con valori rispettivamente di 0,4693, 0,8816 W e 0,9988. Si conclude che, per questo studio in particolare, ANN è l'algoritmo più affidabile e applicabile per la previsione della produzione di energia solare fotovoltaica.

La produzione di energia basata su sistemi solari fotovoltaici (FV) ha guadagnato molta attenzione da parte di ricercatori e professionisti di recente, grazie alle sue caratteristiche. Tuttavia, la principale difficoltà nella produzione di energia solare è la volatilità intermittente della generazione di energia del sistema fotovoltaico, che è dovuta principalmente alle condizioni meteorologiche. Per i parchi solari su larga scala, lo squilibrio di potenza del sistema fotovoltaico può causare una perdita significativa del loro profitto economico. Per questo, la previsione accurata della potenza di produzione di impianti fotovoltaici a breve termine è di grande importanza per la gestione efficiente giornaliera/oraria della produzione, consegna e stoccaggio della rete elettrica, nonché per affrontare decisioni sul mercato dell'energia. Lo scopo dell'articolo [101] è fornire dati affidabili sulla

previsione a breve termine della produzione di energia dei sistemi solari fotovoltaici. Nello specifico, il documento presenta un **approccio di deep learning basato sulla Long Short-Term Memory (LSTM)**, per la previsione della produzione di energia di un impianto fotovoltaico. Questo è motivato dalle caratteristiche di LSTM, per descrivere le dipendenze nei dati delle serie temporali. Le prestazioni dell'algoritmo vengono valutate utilizzando i dati di un impianto connesso alla rete da 9 MWp. I risultati mostrano promettenti risultati di previsione sulla potenza di LSTM.



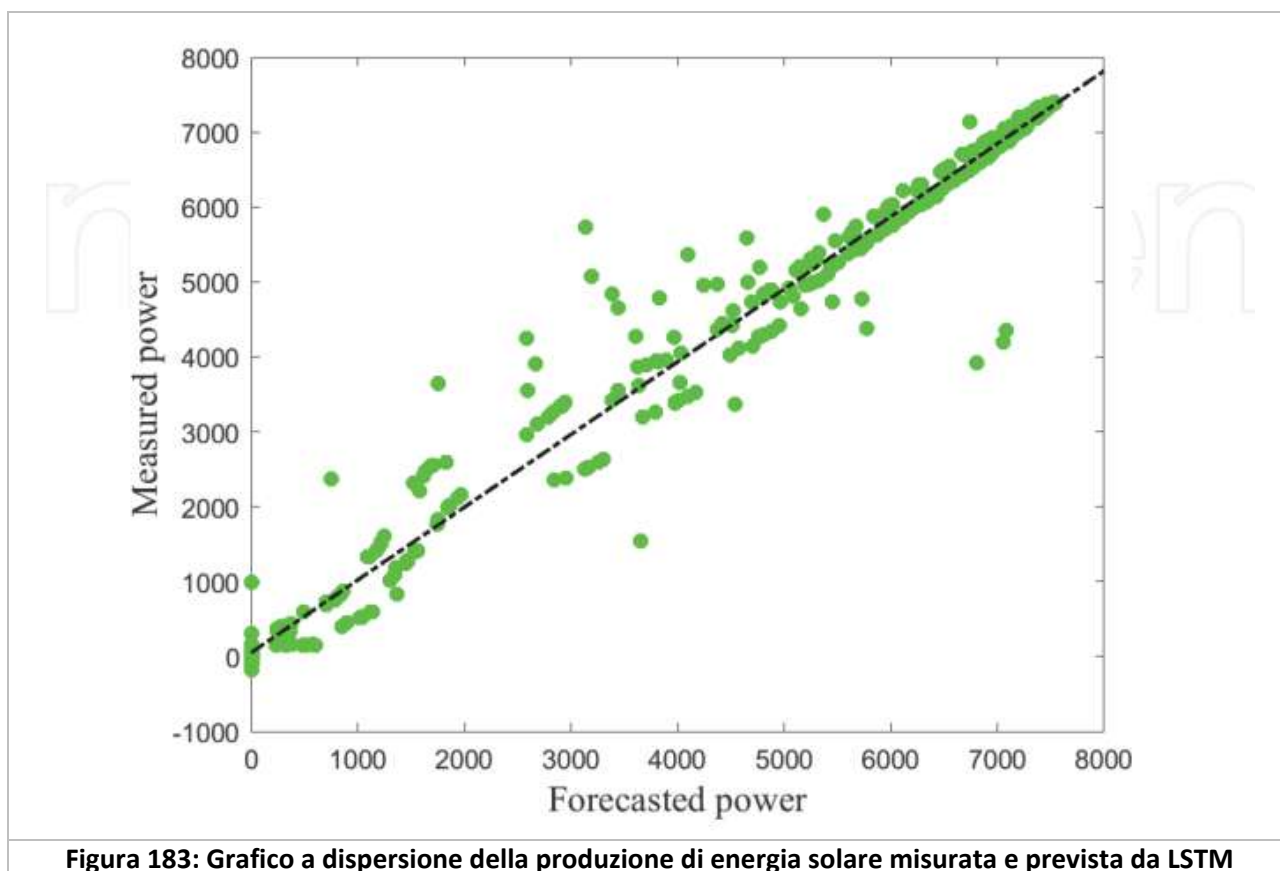


Figura 183: Grafico a dispersione della produzione di energia solare misurata e prevista da LSTM

La principale sfida nella generazione di energia solare è la volatilità intermittente di produzione di energia da impianto fotovoltaico dovuta principalmente alle condizioni atmosferiche. Così, sta diventando indispensabile una previsione accurata della produzione di energia fotovoltaica per ridurre l'effetto dell'incertezza e dei costi energetici, così da consentire un'adeguata integrazione dei sistemi fotovoltaici in una rete intelligente. Questo paper ha utilizzato un modello **Long Short-Term Memory (LSTM)** per prevedere con precisione l'energia solare fotovoltaica a breve termine. Questo approccio sfrutta le proprietà desiderabili di LSTM, che è un potente strumento per modellare la dipendenza nei dati. La qualità previsionale di questo approccio è stata verificata utilizzando i dati da gennaio 2018 a dicembre 2018 raccolti da un impianto connesso alla rete 9MWp. Risultati promettenti sono stati raggiunti dall'approccio proposto basato su LSTM, sulla previsione a breve termine della produzione di energia solare fotovoltaica.

3.2.2 Fonti non rinnovabili

Il passaggio dal fossile al rinnovabile, punto chiave nella lotta al cambiamento climatico e nella direzione della sostenibilità, rappresenta un cambio di paradigma. Da un modello di generazione di energia del tutto programmabile, si va verso uno scenario in cui la caratteristica intrinseca è la non programmabilità. Un percorso, quindi, che pone delle sfide tecniche e d'infrastruttura, anche perché non ci si può permettere di destabilizzare le reti, né di causare blackout o interruzioni del servizio.

Se proviamo a immaginare come dovrà essere la gestione energetica del futuro, è certo che servirà flessibilità. Alterazioni improvvise dell'equilibrio tra domanda e offerta di energia, stress di rete e situazioni eccezionali imporranno – e già ora impongono – una gestione con impianti capaci di anticipare e tollerare le situazioni critiche, affrontandole in tempo reale e poi ritornando in condizioni di normalità.

La sfida più grande da affrontare sta nel trovare un modo di gestire le differenze quotidiane tra domanda e offerta. Impianti eolici e fotovoltaici, infatti, generano un disallineamento tra la produzione di energia e il suo consumo, in parte prevedibile e in parte dovuto alle condizioni meteorologiche. La risposta non può che orientarsi in due principali direzioni. Anzitutto il potenziamento dei sistemi di accumulo di energia (storage) per differire l'erogazione di energia rispetto all'effettiva domanda. E poi, in una fase temporanea, la sostituzione del carbone con altre fonti che siano sì meno inquinanti, ma anche capaci di garantire una fornitura di energia programmabile. In questa ottica, oggi il gas naturale rappresenta un'alternativa promettente ed efficace e un ottimo alleato della transizione energetica in atto.

Nell'articolo [102] si propone di utilizzare algoritmi di apprendimento automatico per prevedere la potenza elettrica generata in base all'ora da una Combined Cycle Power Plant (CCPP). A tal fine, la potenza generata è considerata una funzione di quattro parametri fondamentali: umidità relativa, pressione atmosferica, temperatura dell'ambiente e vuoto di scarico. Le misurazioni di questi parametri e la potenza prodotta sono utilizzate per addestrare e testare i modelli di apprendimento automatico. Il dataset per la ricerca proposta è stato raccolto in un periodo di sei anni ed è stato

preso da un archivio di apprendimento automatico standard e pubblicamente disponibile. Gli algoritmi di apprendimento automatico utilizzati sono: **K-nearest neighbors (KNN)**, **Gradient-Boosted Regression Tree (GBRT)**, **Linear Regression (LR)**, **Artificial Neural Network (ANN)** e **Deep Neural Network (DNN)**. Si riporta una performance all'avanguardia in cui il GBRT supera non solo gli algoritmi utilizzati, ma anche tutti i metodi precedenti sul dataset CCPP fornito. Raggiunge i valori minimi di errore quadratico medio (RMSE) di 2,58 e di errore assoluto (AE) di 1,85.

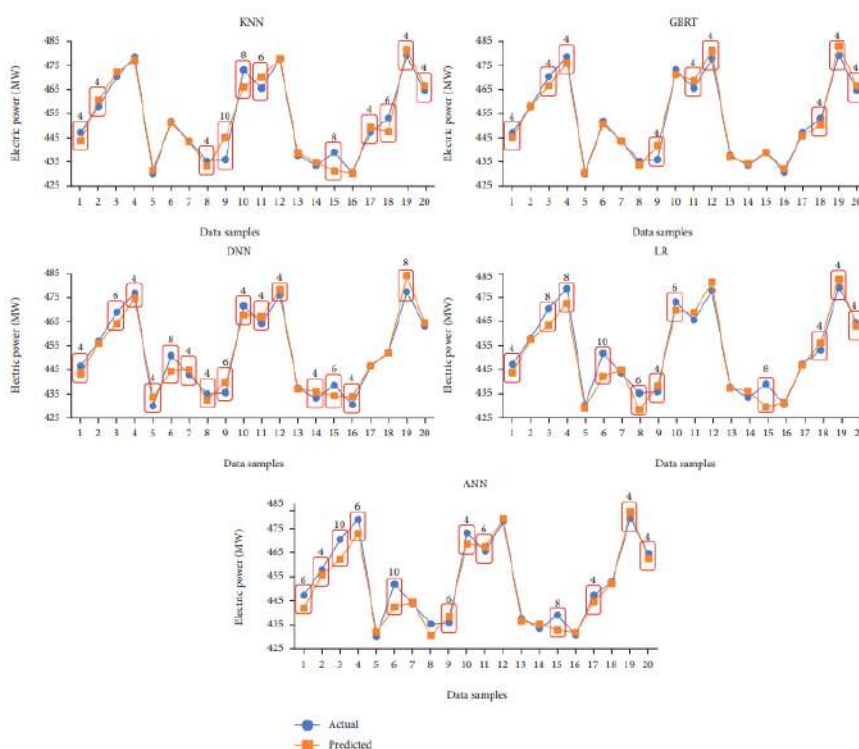


Figura 184: Potenza effettiva rispetto a quella prevista per 20 campioni casuali del set di dati utilizzando cinque algoritmi ML

In questo studio viene effettuata la previsione della potenza di una centrale elettrica a ciclo combinato su base oraria utilizzando il paradigma del machine learning. A tal fine, viene utilizzato un dataset pubblicamente disponibile, raccolto nel corso di 6 anni, durante i quali la centrale funziona a pieno carico. Il dataset misura la potenza in uscita come funzione di quattro parametri di input: temperatura, umidità, pressione e vuoto. Mantenendo i parametri predefiniti, vengono

valutati i parametri più cruciali di ciascun algoritmo per trovare quello che raggiunge il minimo RMSE e AE. Inoltre, viene valutato l'effetto della dimensione del set di addestramento e del numero di caratteristiche sui risultati ottenuti. GBRT si dimostra il migliore tra gli algoritmi, raggiungendo il minimo RMSE e AE con 450 alberi di decisione addestrando il modello sul 90% del dataset. Interessantemente, GBRT supera anche tutti i metodi precedentemente proposti sullo stesso dataset raggiungendo il minimo RMSE e AE.

Il lavoro presente nel paper [103] analizza un insieme di modelli per la previsione della generazione di energia. Inoltre, la generazione di energia è intrinsecamente influenzata dallo stato dei componenti della centrale elettrica, e quindi lo sviluppo del modello incorpora anche variabili aggiuntive del processo della centrale elettrica durante la previsione della generazione di energia.

Si presentano e confrontano diversi metodi di previsione all'avanguardia basati su dati per la generazione di energia per determinare il modello più adeguato e preciso. Si sviluppa anche un modello interpretabile basato sull'attenzione del trasformatore per spiegare l'importanza delle variabili del processo durante l'addestramento e la previsione. **Il modello di Deep Neural Network (DNN) basato su Long Short-Term Memory (LSTM)** addestrato ha una buona accuratezza nella previsione della generazione lorda di potenza in diversi orizzonti di previsione con/senza eventi di ciclismo e supera gli altri modelli per la previsione a lungo termine. I modelli basati sulla memoria DNN mostrano una superiorità significativa rispetto ad altri modelli di apprendimento automatico all'avanguardia per le previsioni a breve, medio e lungo termine.

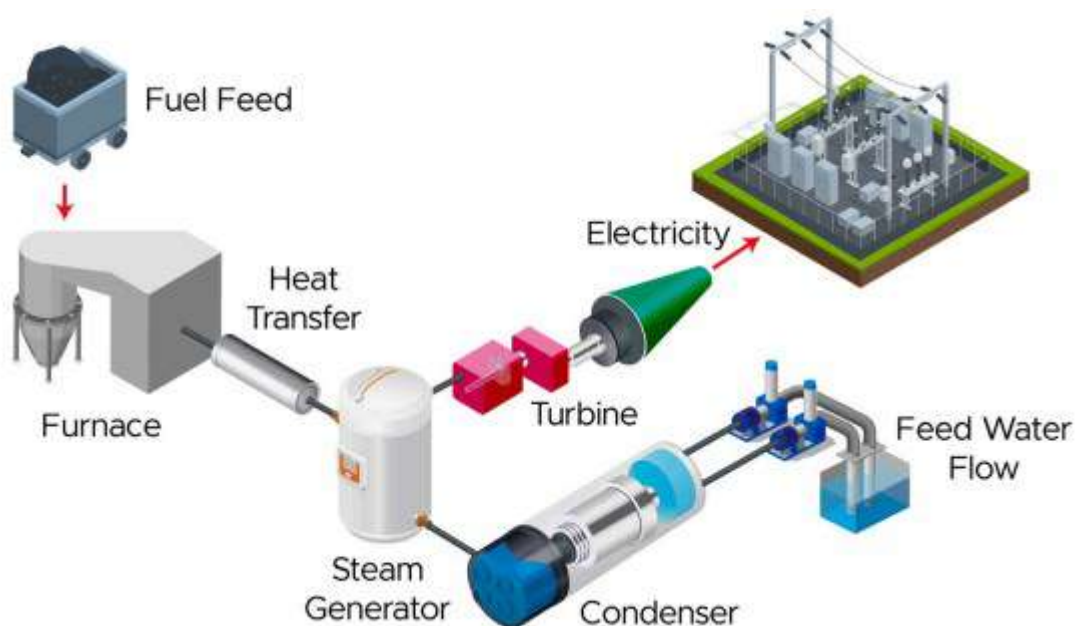


Figura 185: Uno schema di funzionamento semplificato di un impianto di generazione di energia a carbone.

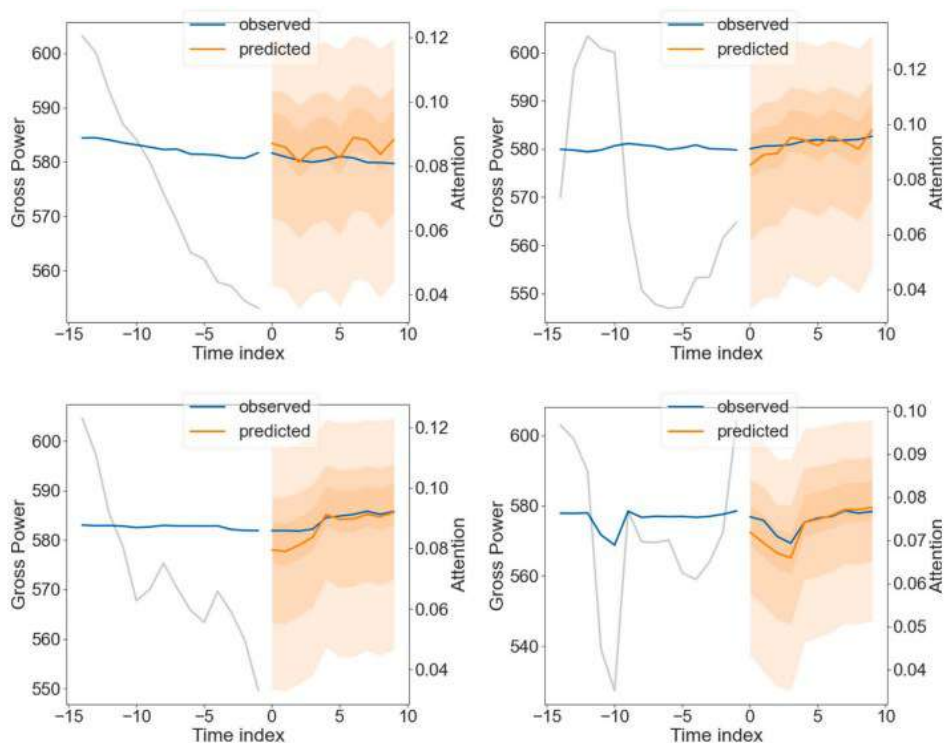


Figura 186: La sequenza di quattro campioni utilizzata per valutare le previsioni multi-step del modello per 10 ore

In questo lavoro viene presentato un modello DNN LSTM per la previsione della generazione lorda di energia. Il modello data-driven è sviluppato incorporando il profilo di generazione lorda ciclica e i dati operativi dei componenti della centrale elettrica, come il tasso di alimentazione del combustibile, il tasso di flusso di alimentazione dell'acqua, il tasso di flusso del vapore principale, la pressione del vapore principale, il tasso di flusso dell'aria primaria e il tasso di flusso dell'aria secondaria. Inoltre, il modello encoder-decoder LSTM è confrontato con altri modelli di machine learning competitivi per stabilire l'efficacia dell'architettura basata su reti neurali per considerare le informazioni di ciclaggio della generazione lorda per garantire una previsione accurata in avanti.

L'articolo [104] confronta l'efficacia dell'**Adaptive Neuro-Fuzzy Inference System (ANFIS)** con altri algoritmi di machine learning per la previsione della produzione di energia elettrica. Vengono utilizzati dati storici sulla produzione di energia e sulle condizioni atmosferiche per addestrare e testare i modelli. L'ANFIS viene confrontato con il **Support Vector Regression (SVR)**, la **Regression Tree (RT)**, il **Gradient Boosting Regression Tree (GBRT)** e la **Artificial Neural Network (ANN)**. I risultati mostrano che l'ANFIS ottiene prestazioni migliori rispetto ad altri algoritmi di machine learning in termini di accuratezza e di riduzione degli errori di previsione.

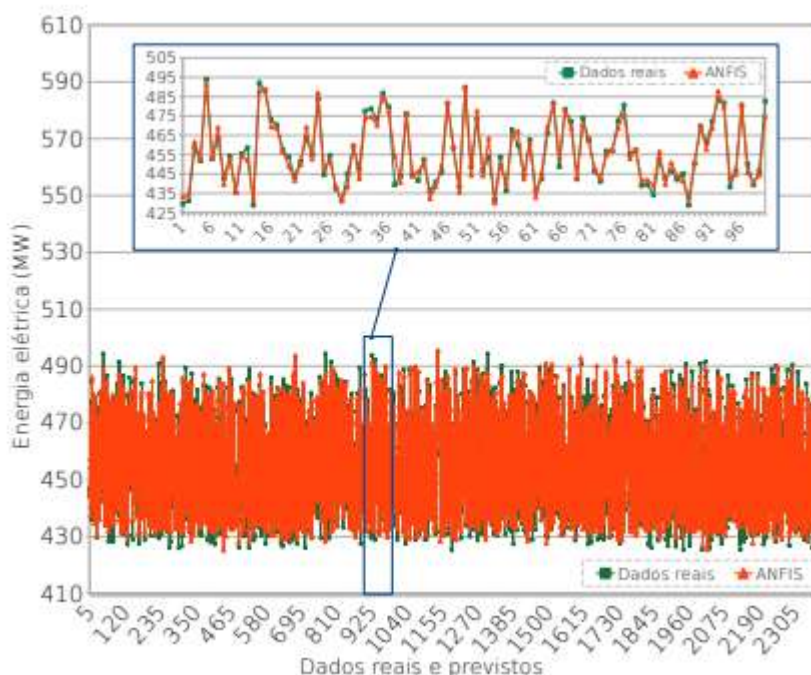


Figura 187: Grafico a linea di valori reali e previsti con ANFIS.

In questo lavoro è stata presentata un'alternativa per stimare l'energia elettrica prodotta per ora da una CCPP, utilizzando variabili termodinamiche ambientali e applicando tecniche di ML e ANFIS, per ottenere risultati affidabili nel minor tempo possibile.

I risultati hanno dimostrato che il modello sviluppato con ANFIS si è dimostrato il migliore rispetto agli altri modelli utilizzati in questo lavoro, anche rispetto ad altri modelli testati nella letteratura. I valori ottenuti dalle metriche di performance dell'ANFIS sono stati $R^2 = 0,95$, $r = 0,98$, $EMA = 2,91$, $REQM = 3,72$, $ERA = 0,20$ e $EQR = 0,22$.

Lo scopo del lavoro dell'articolo [105] è quello di applicare e sperimentare varie opzioni per gli effetti sulla **feed-forward Artificial Neural Network (ANN)** utilizzata per ottenere un modello di regressione che prevede la potenza di uscita elettrica (EP) della centrale elettrica a ciclo combinato basata su 4 input. I dati sono ottenuti da una fonte online aperta. Il lavoro mostra e spiega il comportamento stocastico della rete neurale di regressione, sperimenta l'effetto del numero di neuroni dei livelli nascosti. Mostra anche una prestazione più elevata per una maggiore dimensione del dataset di addestramento; d'altra parte, mostra un effetto diverso di un maggior numero di variabili in ingresso. Inoltre, vengono applicate e confrontate due diverse funzioni di addestramento. Infine, viene condotto uno studio statistico semplice sull'errore tra i valori reali e i valori stimati utilizzando ANN, che mostra l'affidabilità del modello. Questo lavoro fornisce un riferimento rapido sugli effetti dei principali parametri delle reti neurali di regressione.

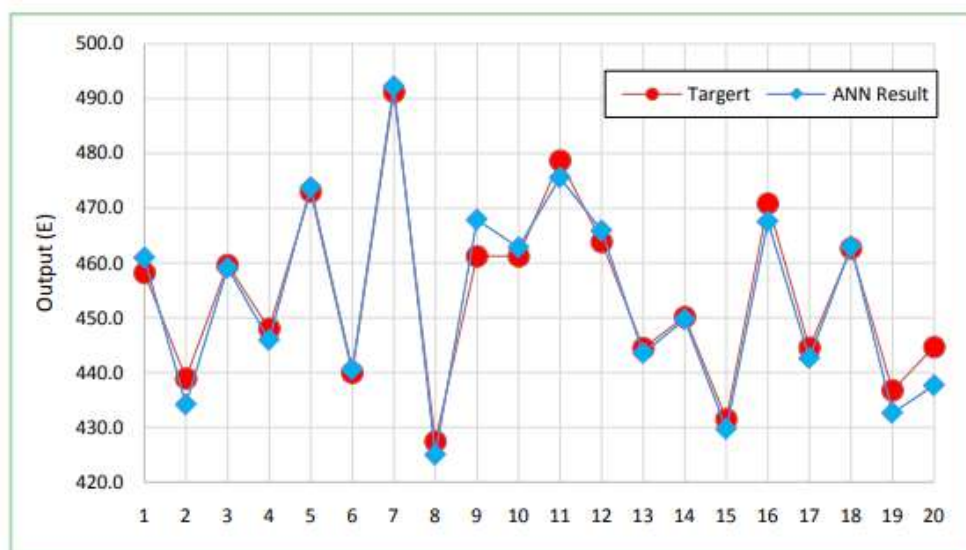


Figura 188: Potenza di uscita prevista (MW) ed errore (MW) a 20 punti dati casuali dai risultati dei test.

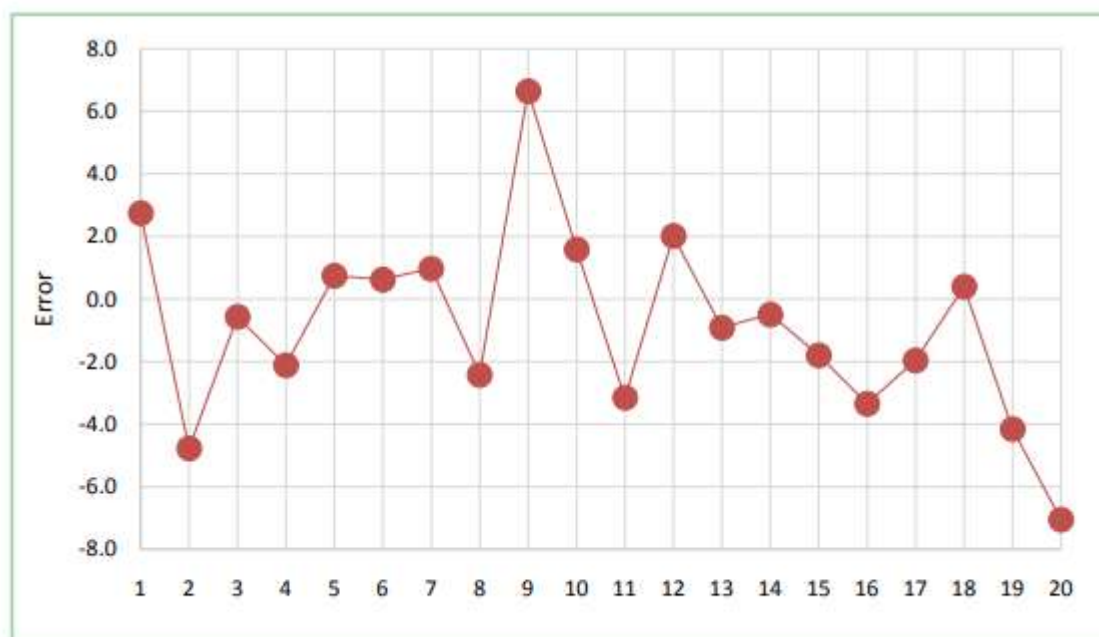
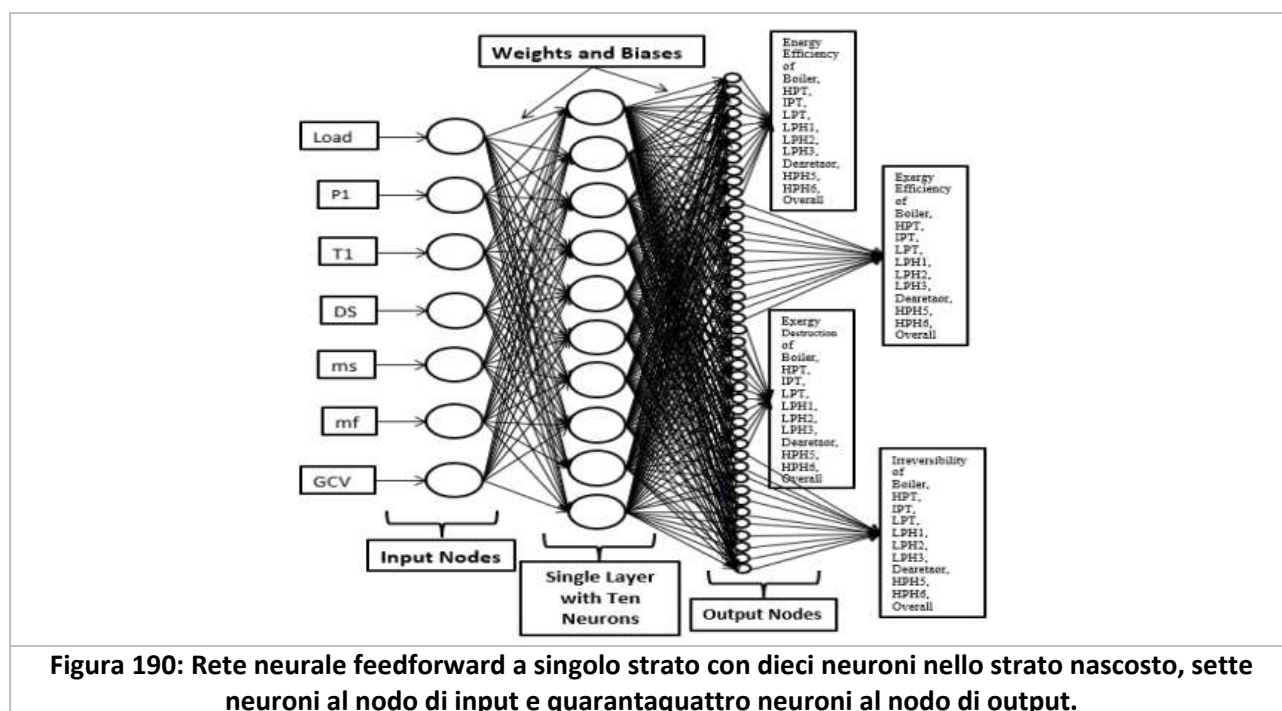


Figura 189: Errore di previsione (MW) a 20 punti dati casuali dai risultati del test.

I risultati mostrano la casualità della performance del modello ANN per ogni volta che viene addestrato, a causa della casualità dei valori iniziali dei pesi e dei bias. Si osserva inoltre che l'aumento del numero di neuroni nel layer nascosto non porta necessariamente ad un

miglioramento della qualità del modello; infatti, il numero di neuroni ha un effetto oscillatorio sulla performance del modello. L'aumento della dimensione del dataset (più punti dati per le stesse variabili) fornisce reti migliori fino ad un certo punto. L'aumento del numero di variabili di input non sempre porta ad una migliore qualità della rete; alcune variabili, quando introdotte, riducono la qualità del modello, altre la aumentano. È necessario studiarlo attraverso la correlazione tra le variabili stesse e tra le variabili e l'output. Inoltre, diverse funzioni di addestramento sono confrontate per le stesse impostazioni e dataset; in questo lavoro la **regolarizzazione bayesiana** ha funzionato meglio dell'algoritmo di Levenberg-Marquardt. Sono stati sperimentati anche metodi di normalizzazione del dataset forniti dalla toolbox. Infine, i risultati sono confrontati con i valori target di output per il gruppo di addestramento e di validazione e per il gruppo di test, ovvero l'intero gruppo di dataset. Il confronto mostra che i risultati sono molto vicini agli output target per entrambi i gruppi. Inoltre, mostra la distribuzione normale dell'errore tra ogni gruppo con valore medio di zero. Le deviazioni standard dell'errore dei due gruppi sono quasi uguali.

L'analisi energetica ed exergetica sono le tecniche più frequentemente utilizzate per valutare le prestazioni termiche di una centrale elettrica. Poiché le centrali elettriche sono operate 24 ore al giorno, è essenziale conoscere le prestazioni energetiche ed exergetiche delle centrali a diversi carichi e comprendere l'effetto dei parametri critici su di esse. L'efficienza energetica ed exergetica istantanea della centrale può essere prevista combinando la valutazione teorica con una **Artificial Neural Network (ANN)**. Nel documento [106] si parla di sessantacinque modelli di ANN che sono stati preparati e addestrati utilizzando quaranta set di dati empirici delle centrali termiche situate nella regione occidentale dell'India. Sono stati considerati sette parametri come neuroni di input e quarantaquattro parametri come neuroni di output. I modelli di ANN addestrati sono stati impiegati per prevedere le prestazioni di un'altra centrale elettrica con la stessa capacità e progettazione utilizzando set di dati separati. Gli errori nella previsione sono stati valutati in termini di MSE, NMSE, MAE, MARD e MRE. Confrontando gli errori nella previsione, si è scoperto che il modello ANN E1 (rete di regressione generalizzata con costante di diffusione di uno) produce il valore minimo di errore. Pertanto, è stato proposto come modello adatto per prevedere le prestazioni termiche della centrale KWU progettata da 210 MW cambiando il carico operativo e decidendo i parametri al momento richiesto.



Gli errori nelle previsioni per i vari modelli di reti neurali sono stati valutati confrontando il valore previsto con il valore valutato analiticamente in termini di MSE, NMSE, MAE, MARD e MRE. Confrontando gli errori nella previsione, si è scoperto che il modello di rete neurale E1 (rete di regressione generalizzata con costante di propagazione pari a uno) fornisce il valore minimo di errore nella previsione per entrambi i casi, ovvero quando vengono utilizzati i set di dati della stessa centrale e di un'altra centrale. Basandosi su questo studio, si conclude che le prestazioni energetiche ed exergetiche di una centrale elettrica con design KWU e capacità di 210 MW possono essere previste utilizzando il modello proposto. Il modello proposto aiuta a decidere i valori ottimali di vari parametri vitali durante il cambiamento del carico operativo al fine di ottenere le prestazioni ottimali della centrale elettrica.

L'articolo [107] propone un metodo di previsione della generazione di energia basato su uno degli algoritmi di apprendimento automatico più popolari, ovvero il **Support Vector Machine (SVM)**, per prevedere sia la generazione complessiva di energia che alcuni tipi speciali di generazione di energia. La relazione non lineare della generazione di energia elettrica netta viene esplorata dai dati storici

registrati mensilmente, questa relazione può aiutare la previsione della generazione elettrica netta per il mese successivo.

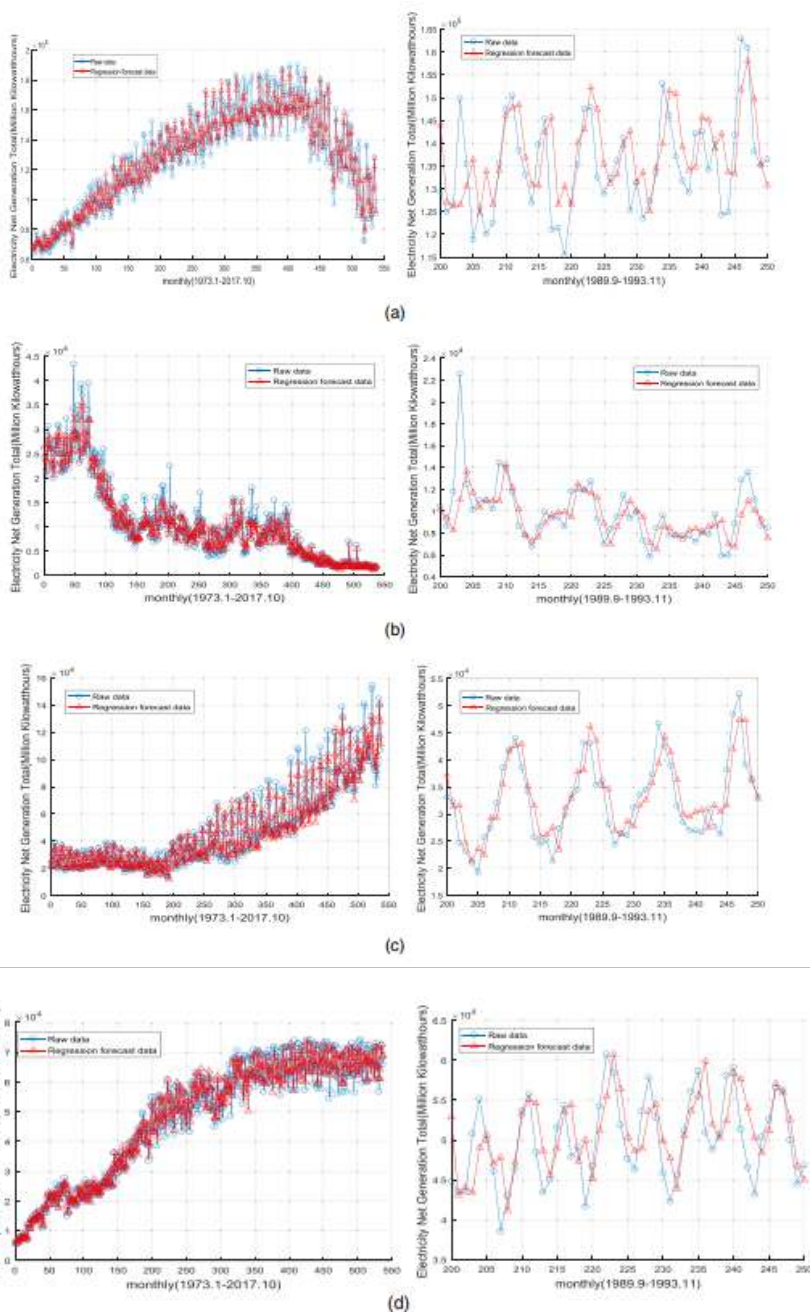


Figura 191: Confronto di previsione tra diverse fonti di dati: (a) curve di regressione della generazione di energia elettrica netta dal carbone, (b) curve di regressione della generazione di energia elettrica netta dal petrolio, (c) curve di regressione della generazione di energia elettrica netta dal gas naturale, (d) curve di regressione della generazione di energia elettrica netta dalle centrali nucleari.

Dal gennaio 1973 all'ottobre 2017, la generazione complessiva di energia elettrica netta, la generazione a carbone, la generazione a petrolio, la generazione a gas naturale e la generazione di energia elettrica nucleare sono state utilizzate per prevedere la generazione di energia elettrica netta di ciascun settore per il mese successivo tramite SVM in questo studio. Rispetto ai lavori precedenti in questo campo, questo lavoro può ottenere un risultato di previsione più accurato rispetto ai dati grezzi reali presenti nel record. In futuro, si potranno esplorare i vantaggi di un minor numero di variabili indipendenti e di un calcolo più semplice con SVM avanzati o altri metodi di previsione superiori.

La ricerca presente nell'articolo [108] tenta di costruire la curva di potenza basata sui dati del generatore installato in una centrale elettrica da 660 MW incorporando strumenti di modellizzazione basati sull'intelligenza artificiale (AI). La potenza prodotta dal generatore è modellata da una **Artificial Neural Network (ANN)** - una tecnica affidabile di analisi dei dati di apprendimento profondo. Allo stesso modo, viene impiegata l'applicazione R2.ai, che appartiene alla piattaforma di apprendimento automatico automatizzato (AutoML), per mostrare i metodi di modellizzazione alternativi nell'utilizzo dell'approccio AI. In confronto, l'ANN si è comportata bene nel test di convalida esterna ed è stata utilizzata per costruire la curva di potenza del generatore. Sono stati simulati esperimenti di Monte Carlo che comprendono i parametri operativi termoelettrici dell'impianto elettrico e il rumore gaussiano con l'ANN, e quindi la curva di potenza del generatore è stata costruita con un intervallo di confidenza del 95%. Le curve di prestazione dei sistemi industriali e delle macchine basate sui loro dati operativi possono essere costruite utilizzando le ANN, e le decisioni basate su queste curve di prestazione potrebbero contribuire alla visione dell'Industria 4.0 di una gestione operativa efficace.

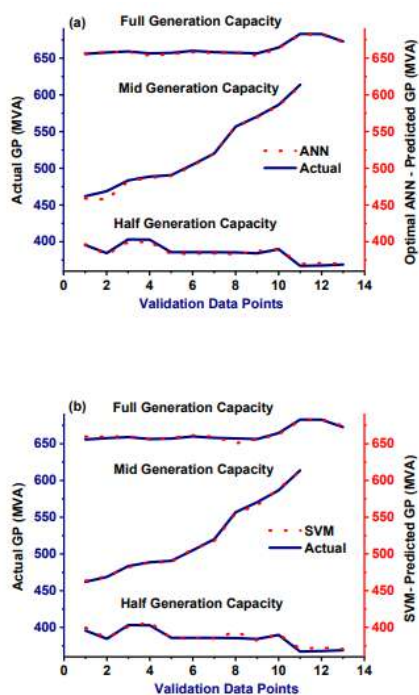


Figura 192: Grafici dei dati di validazione esterna: (a) RNA ottimale; (b) macchina vettoriale di supporto (SVM).

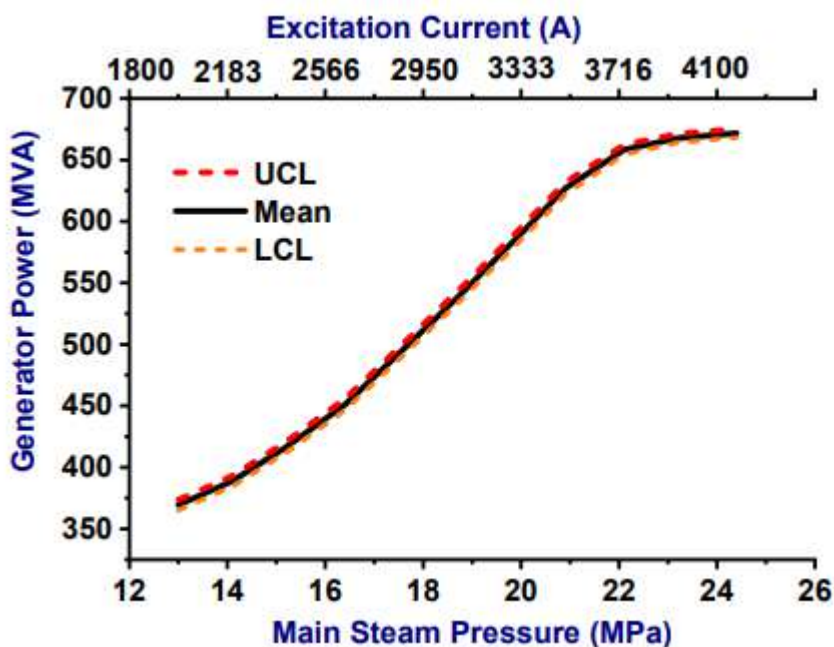
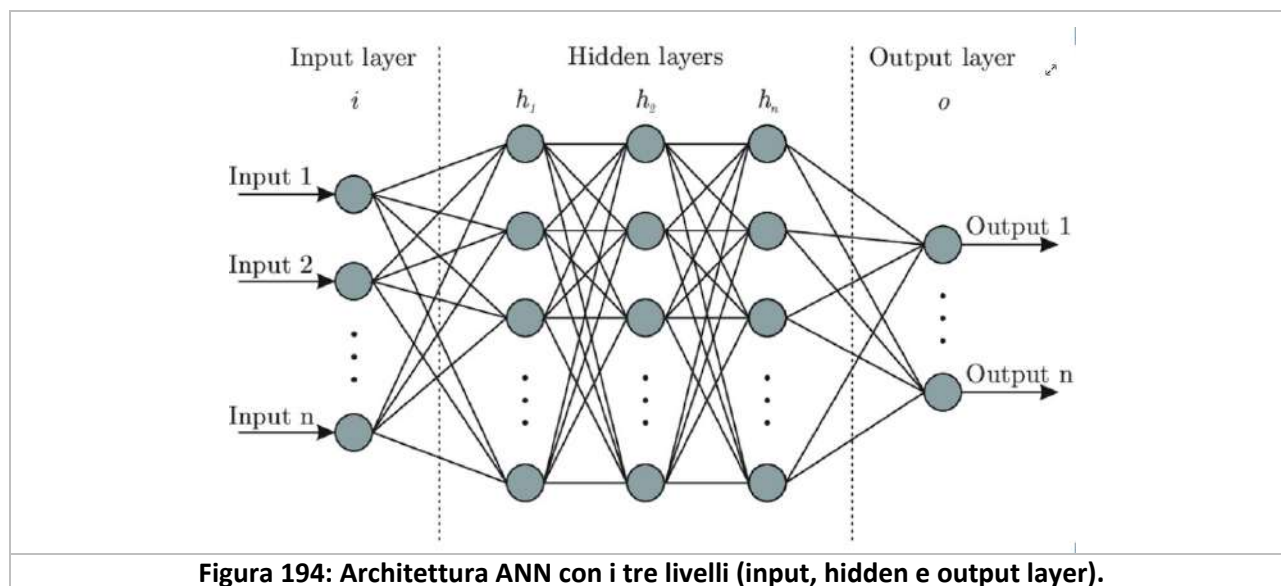


Figura 193: L'effetto combinato dei parametri operativi termoelettrici sul generatore di potenza

In questo lavoro, le tecniche di modellizzazione matematica per la costruzione delle curve di potenza sono limitate nelle loro applicazioni ai sistemi industriali complessi e di grande scala. In questo studio, sono state utilizzate due tecniche di modellizzazione basate sui dati delle AI, ovvero ANN e AutoML, per costruire la curva di potenza di un generatore in relazione a ventiquattro parametri operativi termoelettrici di una centrale elettrica a carbone supercritico da 660 MW. I dati operativi completi della centrale contenenti tutte le possibili modalità operative della potenza del generatore sono stati presi dal SIS. Tecniche di elaborazione e visualizzazione dati sono state utilizzate per eliminare le osservazioni difettose presenti nei dati grezzi. I dati filtrati sono stati quindi utilizzati per costruire i modelli ANN e AutoML per la potenza del generatore della centrale. Tuttavia, nel test di convalida esterna, l'ANN ha superato il modello SVM basato su AutoML e quindi è stato selezionato. Sono stati effettuati esperimenti di Monte Carlo che comprendono i parametri operativi termoelettrici della centrale elettrica e il rumore gaussiano, simulati dall'ANN e impiegati per costruire la curva di potenza del generatore con un intervallo di confidenza del 95%.

L'articolo [109] sviluppa un metodo basato sull'apprendimento automatico per prevedere le prestazioni delle turbine a gas per la generazione di energia. Vengono sviluppati due modelli surrogati basati sulla rappresentazione del modello ad alta dimensionalità (HDMR) e **sulla Artificial Neural Network (ANN)** a partire da dati operativi reali per prevedere le caratteristiche operative del compressore d'aria e della turbina. Entrambi i modelli catturano bene le caratteristiche operative con errori medi inferiori all'1,0%. Inoltre, vengono sviluppati altri quattro modelli più olistici per catturare le prestazioni a carico parziale e a pieno carico della turbina a gas. I modelli per il compressore d'aria e la turbina vengono quindi incorporati in un programma di simulazione di turbine a gas, e tutti i modelli surrogati vengono validati utilizzando un set di dati separato. Si dimostra che la potenza in uscita, il rapporto di pressione, il flusso di combustibile e la temperatura di scarico della turbina da questi modelli corrispondono bene ai loro valori misurati con errori medi e massimi inferiori al 2,0% e al 4,3%, rispettivamente. Poiché i modelli olistici ANN hanno una complessità inferiore e una maggiore precisione, il modello ANN per prevedere le prestazioni a pieno carico viene utilizzato per costruire le curve di correzione delle prestazioni della turbina a gas. Le curve di correzione insieme al modello ANN per prevedere le prestazioni a carico parziale offrono

una base eccellente per il monitoraggio continuo della salute e la diagnosi dei guasti. La metodologia proposta è applicabile a qualsiasi turbina a gas e può aiutare le centrali elettriche a studiare e quantificare il degrado delle prestazioni nel tempo.



Questo studio ha proposto un metodo basato sull'apprendimento supervisionato tramite dati storici per predire le prestazioni di una turbina a gas, utilizzando modelli surrogati HDMM e ANN. Sono stati sviluppati modelli HDMM e ANN per prevedere le caratteristiche operative del compressore e della turbina, quindi incorporati in un programma di simulazione per predire le prestazioni della turbina a gas. Inoltre, sono stati sviluppati modelli surrogati HDMM e ANN olistici per catturare le prestazioni a carico parziale e a carico completo della turbina a gas. Si è riscontrato che tutti i modelli surrogati catturano le prestazioni complessive della turbina a gas.

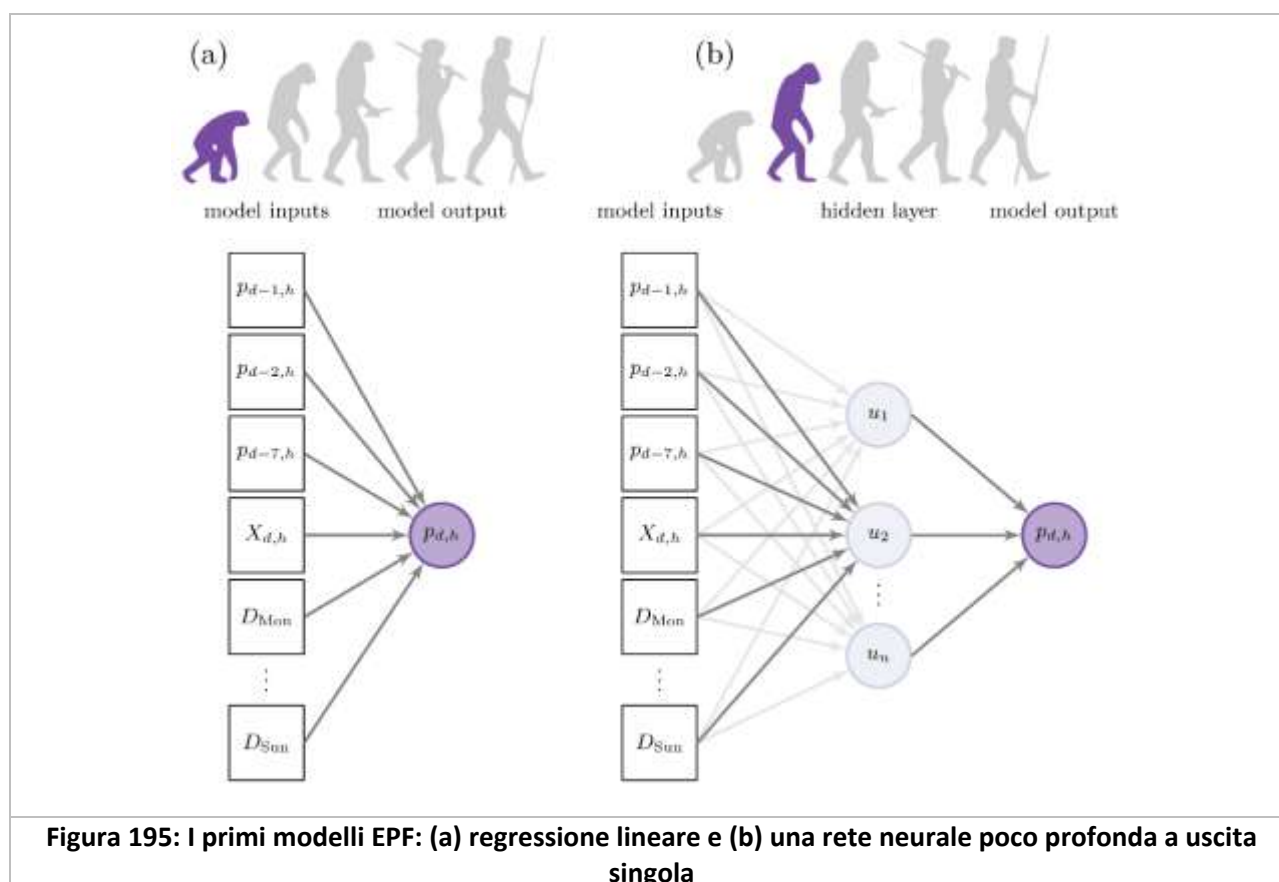
3.3 Previsione del costo dell'approvvigionamento da reti di fornitura

La previsione dei prezzi dell'elettricità (Electricity price forecasting o EPF) è una branca delle previsioni sull'interfaccia tra ingegneria elettrica, statistica, informatica e finanza, che si concentra sulla previsione dei prezzi nei mercati all'ingrosso dell'elettricità per un intero spettro di orizzonti.

Questi vanno da pochi minuti (aste in tempo reale/intraday e negoziazione continua), a giorni (aste day-ahead), a settimane, mesi o persino anni (futures e contratti a termine negoziati in borsa e over-the-counter).

I mercati del giorno prima (Day-ahead o DA) sono il cavallo di battaglia del commercio di energia, in particolare in Europa, e un proxy comunemente usato per "il prezzo dell'elettricità". La stragrande maggioranza - fino al 90% - della letteratura EPF si è concentrata sulla previsione dei prezzi DA. Questi ultimi vengono tipicamente determinati intorno a mezzogiorno durante 24 aste a prezzo uniforme, una per ogni ora del giorno successivo. Ciò ha implicazioni dirette per il modo in cui vengono costruiti i modelli EPF. Negli ultimi 25 anni, vari metodi e strumenti computazionali sono stati applicati all'EPF intraday e day-ahead [110]. Fino all'inizio degli anni 2010, il campo era dominato da modelli di regressione lineare relativamente piccoli e reti neurali (artificiali), tipicamente con non più di due dozzine di input. Col passare del tempo, sono diventati disponibili più dati e più potenza di calcolo. I modelli sono diventati più grandi nella misura in cui la conoscenza esperta non era più sufficiente per gestire le strutture complesse. Questo, a sua volta, ha portato all'introduzione di tecniche di apprendimento automatico (ML) in questo settore affascinante e in rapido sviluppo.

Gli inizi dell'EPF risalgono agli anni '90. Il primo tentativo di prevedere l'elettricità e la dinamica dei prezzi utilizzava tecniche di regressione lineare. Ricordiamo che un tale modello assume una lineare relazione tra la variabile prevista (ad esempio, il prezzo dell'elettricità oggi alle 18:00) e l'input (ad esempio, prezzi dell'elettricità passati, previsione del carico per oggi). Gli input sono anche chiamati caratteristiche, variabili esplicative, regressori o predittori. Più formalmente, la variabile prevista è rappresentata come somma ponderata degli input.



I primi modelli di regressione sono stati costruiti sulla conoscenza degli esperti. Gli input includevano i prezzi passati (tipicamente ieri, due e sette giorni fa), variabili esogene e componenti stagionali. La stagionalità è stata catturata da una funzione sinusoidale o, più comunemente, un insieme di cosiddette variabili fittizie che assumono il valore 1 in un dato giorno della settimana e 0 altrimenti. Un modello EPF di regressione di esempio è illustrato nella figura precedente (a). La variabile esogena più rilevante e spesso l'unica utilizzata, anche nei modelli più recenti, è la previsione del carico del giorno prima. È interessante notare che il carico effettivo non è il predittore migliore. Il motivo è che le offerte nel mercato DA vengono invece piazzate sulla base delle previsioni di carico del giorno prima dei carichi effettivi osservati il giorno successivo. Poiché il rapporto prezzo-carico dell'elettricità è non lineare, per effettuare le previsioni si sono applicate anche le **Neural Network (NN)**. I modelli NN degli anni '90 e 2000 erano tipicamente strutture poco profonde con un solo strato nascosto, un neurone di output (ad esempio, il prezzo dell'elettricità oggi alle 18:00) e un'architettura feed-forward. Quest'ultima significa che le informazioni sono state propagate in una

sola direzione, dagli input all'output. Contrariamente alla regressione lineare in cui l'output è solo una somma ponderata di tutti gli input, nei modelli di rete neurale è possibile applicare ulteriori trasformazioni non lineari a questa somma in ciascun nodo.

Scoprire la relazione non lineare tra le variabili è una sfida per l'approccio basato sulla regressione lineare. Un'altra sfida è la gestione di un gran numero di input. Anche poche dozzine di variabili esplicative possono portare a stime e previsioni inaffidabili quando si utilizza il modello di regressione classico e i minimi quadrati ordinari per minimizzare la somma degli errori al quadrato tra i valori previsti e quelli osservati. Un rimedio funzionante fornisce i cosiddetti algoritmi di regolarizzazione. Riducono (da qui il nome alternativo - algoritmi di restringimento) i coefficienti o i pesi degli input meno importanti verso lo zero aggiungendo una penalità per i coefficienti grandi alla somma degli errori al quadrato.

Uno degli algoritmi di regolarizzazione più utilizzati è l'operatore di restringimento e selezione meno assoluto (**least absolute shrinkage and selection operator o LASSO**), reso popolare da Robert Tibshirani a metà degli anni '90. Questa tecnica di apprendimento automatico (chiamata anche apprendimento statistico) riduce alcuni coefficienti a zero e quindi esegue una selezione di variabili semiautomatica basata sui dati. I modelli EPF, che inizialmente hanno centinaia di input, possono essere ridotti a strutture con solo una dozzina o due regressori rilevanti. La selezione è "semi-automatica" perché LASSO richiede l'impostazione dell'iperparametro di sintonia che regola quali caratteristiche vengono eliminate. Questo iperparametro può essere impostato utilizzando la convalida incrociata, un metodo di ricampionamento che utilizza parti diverse dei dati per testare e addestrare un modello.

I modelli EPF stimati da LASSO sono apparsi a metà degli anni 2010 e hanno rapidamente rivoluzionato il campo. Nonostante la loro relativa semplicità e l'incapacità di gestire le dipendenze non lineari, non è facile superarle.

Un esempio più recente è il modello **LASSO-Estimated AutoRegressive (LEAR)**, un metodo open source proposto nel 2021 come uno dei due benchmark impegnativi per la previsione contemporanea dei prezzi dell'elettricità. Il modello LEAR è una struttura autoregressiva ricca di

parametri con circa 250 variabili esplicative. In altre parole, per prevedere il prezzo per un'ora, il modello utilizza i prezzi orari di ieri, due, tre e sette giorni fa, previsioni orarie day-ahead di due variabili esogene (ad esempio, carico del sistema e generazione di energia eolica) per oggi, ieri, e una settimana fa, e manichini nei giorni feriali.

In passato, la sfida dell'utilizzo delle reti neurali era la complessità computazionale. Nell'ultimo decennio, i progressi nelle risorse computazionali (ad esempio, l'uso massiccio di unità di elaborazione grafica, GPU e algoritmi di ottimizzazione) hanno reso possibile addestrare in modo efficiente strutture complesse, comprese le reti neurali con centinaia di input, output multipli, più strati nascosti, e collegamenti a livelli precedenti. Poiché le reti la cui profondità (ovvero il numero di livelli) non era limitata a un singolo livello nascosto hanno mostrato sistematicamente risultati migliori e capacità di generalizzazione, il campo e i modelli sono stati denominati **Deep Learning (DL)** e **Deep Neural Network (DNN)** per sottolineare l'importanza della profondità nei miglioramenti ottenuti.

Sebbene questo successo dei modelli DL sia iniziato nelle applicazioni informatiche, ad esempio il riconoscimento di immagini, il riconoscimento vocale o la traduzione automatica, i vantaggi di DL si sono diffusi anche alla fine degli anni 2010 alle applicazioni relative all'energia, ad esempio il prezzo dell'elettricità o la previsione dell'energia eolica. Da allora, i modelli profondi sono stati ampiamente utilizzati per sfruttare e modellare meglio le relazioni non lineari tra le quantità basate sull'energia (prezzi, carico, generazione) e i loro driver (ad esempio, comportamento umano, effetti del calendario).

La prima ondata di modelli DL è stata seguita negli ultimi anni da architetture ibride che utilizzano la cosiddetta **Long Short-Term Memory (LSTM)** e/o **Convolutional Neural Network (CNN)**. A differenza delle reti feed-forward, l'architettura della memoria a lungo-breve termine ha più connessioni di feedback e può elaborare non solo singoli punti dati ma anche intere serie temporali. Le reti neurali convoluzionali sono versioni regolarizzate, per evitare l'overfitting, di reti feed-forward multistrato che utilizzano la convoluzione (invece della moltiplicazione di matrici) in almeno uno dei loro strati. Tuttavia, nonostante il recente lavoro sui modelli DL, non è chiaro se tutta la complessità aggiuntiva apporti miglioramenti nell'accuratezza delle previsioni. Le prove esistenti

indicano che, sebbene i modelli DL semplici (come DNN) possano in media migliorare rispetto ai modelli di regressione stimati con LASSO (come LEAR), le prestazioni di modelli DL più complessi sono generalmente sconosciute. Un modello comune condiviso da molti lavori in DL è l'incapacità di confrontare i modelli proposti con metodi statistici all'avanguardia e/o di impiegare set di dati sufficientemente lunghi per trarre conclusioni statisticamente valide.

In [111] si propone un modello integrato di **rete convoluzionale ricorrente a lungo termine (Integrated Long-Term Recurrent Convolutional Network o ILRCN)** per prevedere i prezzi dell'elettricità considerando come input la maggior parte degli attributi che contribuiscono al prezzo di mercato. Il modello ILRCN proposto combina le funzionalità di una **Convolutional Neural Network (CNN)** e della **Long Short-Term Memory (LSTM)** insieme al nuovo termine di correzione dell'errore condizionale proposto. Il modello ILRCN combinato può identificare il comportamento lineare e non lineare all'interno dei dati di input. I dati dei prezzi di mercato ERCOT all'ingrosso insieme al profilo di carico, temperatura e altri fattori per la regione di Houston sono stati utilizzati per illustrare il modello proposto.

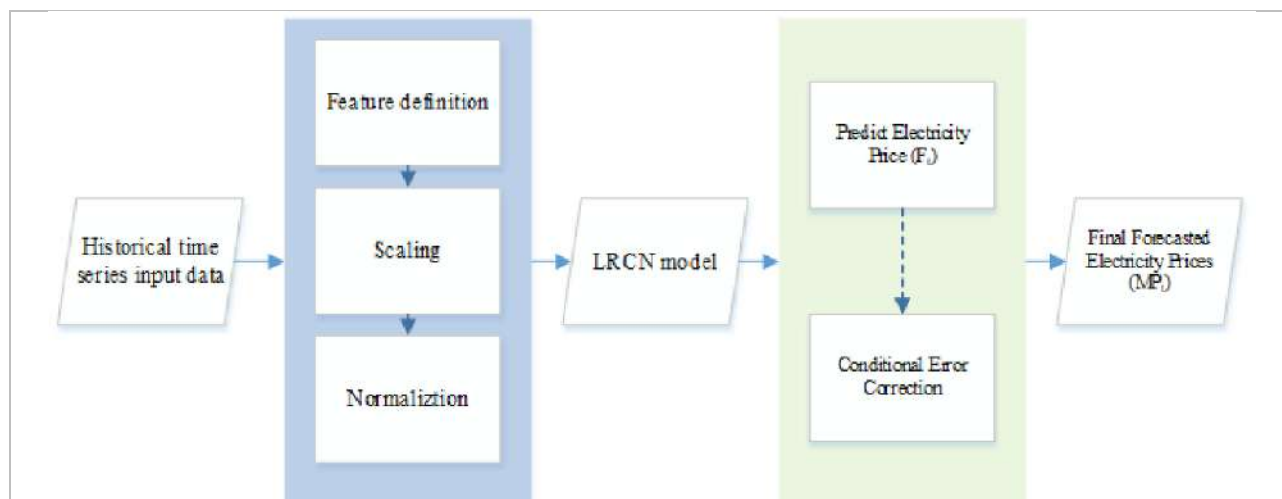
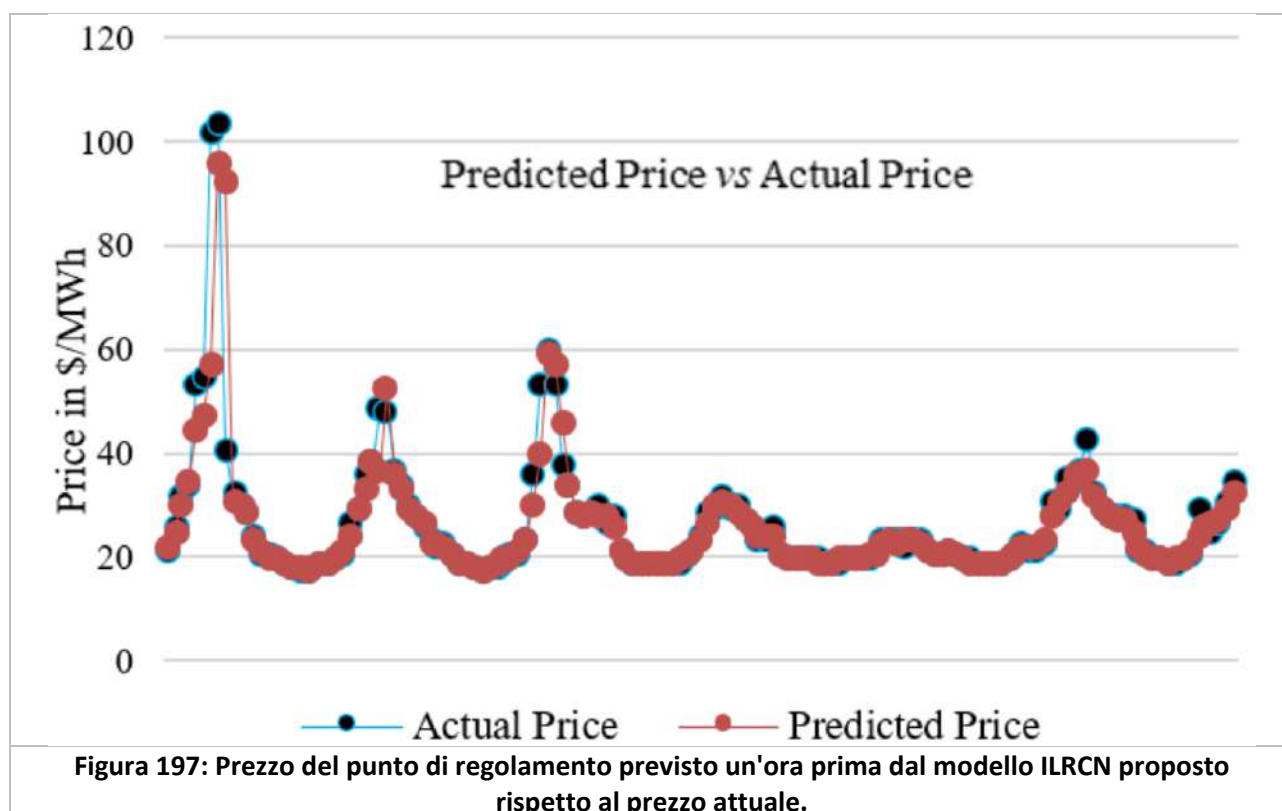


Figura 196: Diagramma di flusso della previsione dei prezzi dell'energia elettrica a breve termine con ILRCN

Il modello ILRCN proposto per il mercato all'ingrosso la previsione dei prezzi dell'elettricità supera il modello **Support Vector Regression (SVR)**, il modello **Fully Convolutional Neural Network (FCNN)**,

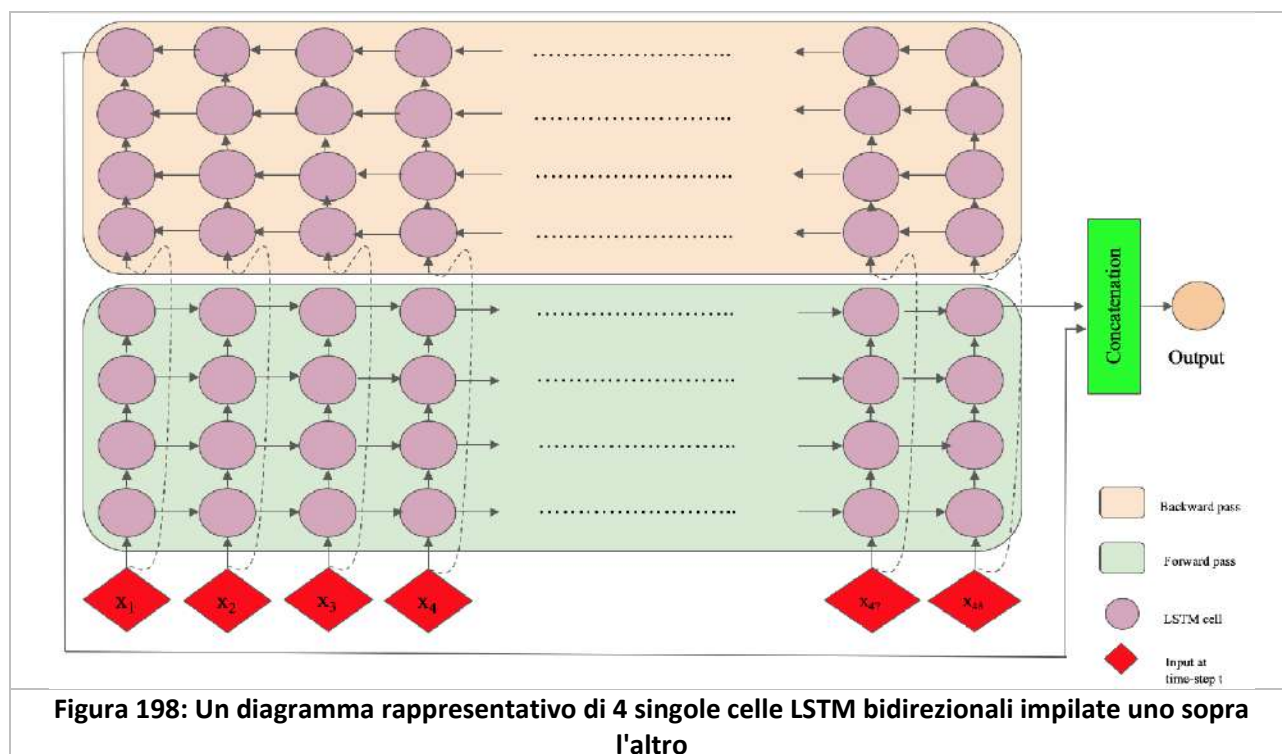
il modello Long Short-Term Memory (LSTM) e il modello Long-term Recurrent Convolutional Network (LRCN), così come il metodo naive in termini di varie metriche, tra cui la formazione efficienza, accuratezza, MSE e MAE, come dimostrato nella sezione dei casi di studio. Il modello ILRCN proposto prevede il prezzo dell'elettricità con alta e bassa precisione d'errore sia dell'ora che del giorno prima rispetto al metodo naive, SVR, FCNN, LSTM, e modelli LRCN.

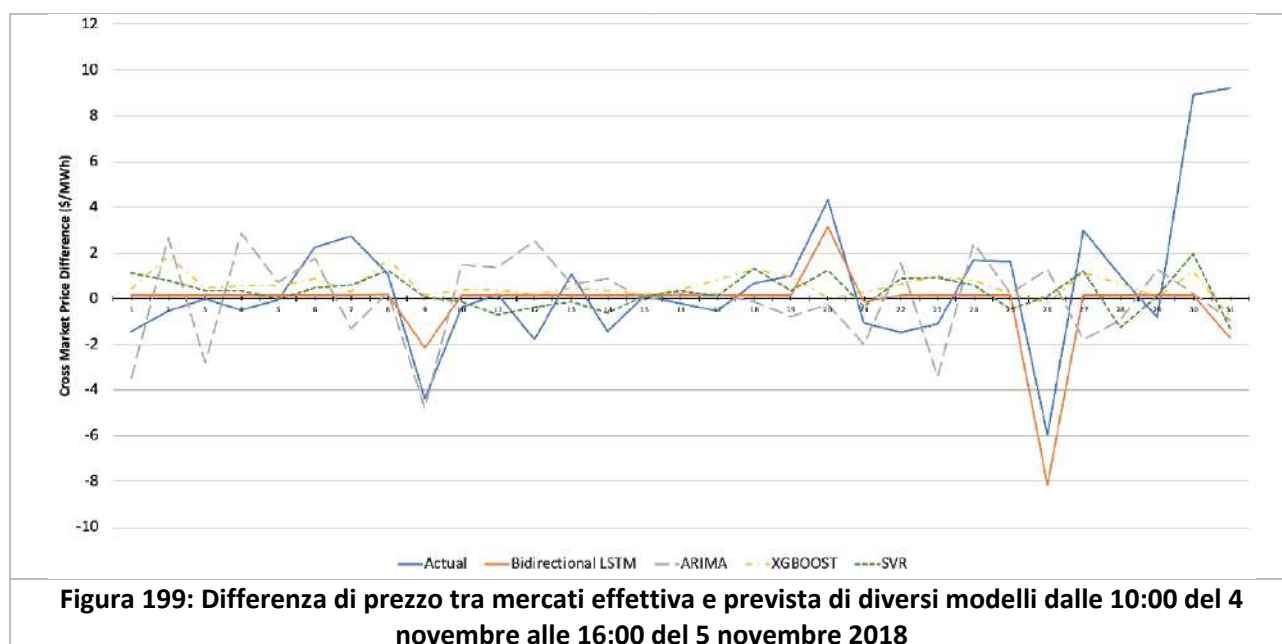


Il termine di correzione dell'errore condizionale proposto può essere ulteriormente migliorato per migliorare le prestazioni di previsione del modello ILRCN. In sintesi, il modello ILRCN proposto supera sia i modelli tradizionali che altri modelli di rete neurale, e si è dimostrato essere un modello accurato ed efficiente nella previsione dei prezzi al punto di regolamento. La praticità e la fattibilità del modello ILRCN proposto sono confermate dalle metriche di performance.

In [112] si effettua il primo tentativo di impiegare una nuova architettura di deep learning **Bidirectional Long-Short Term Memory (Bi-LSTM)** per prevedere la differenza di prezzo tra i mercati del giorno prima e quelli in tempo reale per lo stesso nodo. I dati grezzi vengono raccolti dal mercato PJM, elaborati e immessi nella rete proposta. Più nel dettaglio l'architettura si basa sui seguenti concetti

- 1) Memoria a lungo-breve termine: la memoria a lungo-breve termine è una specie di Rete Neurale Ricorrente che può catturare dipendenze sia a lungo che a breve termine nei dati, rendendolo adatto per previsioni, riconoscimento vocale, ecc.
- 2) LSTM bidirezionale: è simile a una rete LSTM. Una rete bidirezionale addestra 2 reti LSTM sulla sequenza di input, invece di una sola. Il primo è la sequenza di input stessa e la seconda è la copia invertita della sequenza di input.





Le performance previsionali del modello sono misurate utilizzando i valori Test RMSE, che mostrano che il modello ha una certa capacità di generalizzazione. La capacità del modello di prevedere correttamente i picchi della differenza di prezzo al momento giusto viene anche dimostrata. Il modello LSTM bidirezionale viene confrontato con **Autoregressive Integrated Moving Average (ARIMA)**, **Extreme Gradient Boosting (XGBoost)** e metodi **Support Vector Regression (SVR)** e li supera tutti sia in RMSE che nel prevedere correttamente la direzione delle differenze di prezzo. Il documento si occupa però solo di un modello di serie temporali univariate, dove si guarda al prezzo, o meglio alle differenze di prezzo a diversi intervalli di tempo. Tuttavia, la previsione potrebbe essere ulteriormente migliorata tenendo conto di altri fattori di influenza. Un altro problema con le tecniche di deep learning come gli LSTM è che impiegano una notevole quantità di tempo per addestrarsi.

In [113] si utilizzano modelli di **Machine Learning (ML) distribuiti** per prevedere i prezzi delle azioni energetiche utilizzando indici di incertezza come proxy per prevedere la volatilità del mercato energetico. Si esaminano empiricamente l'efficacia comparativa dei modelli prevalenti di ML e approcci convenzionali (regressione) per prevedere l'equità energetica dei prezzi utilizzando i dati giornalieri dal 1/6/2011 al 18/1/2022 per quattro indici di incertezze statunitensi e otto indici azionari energetici.

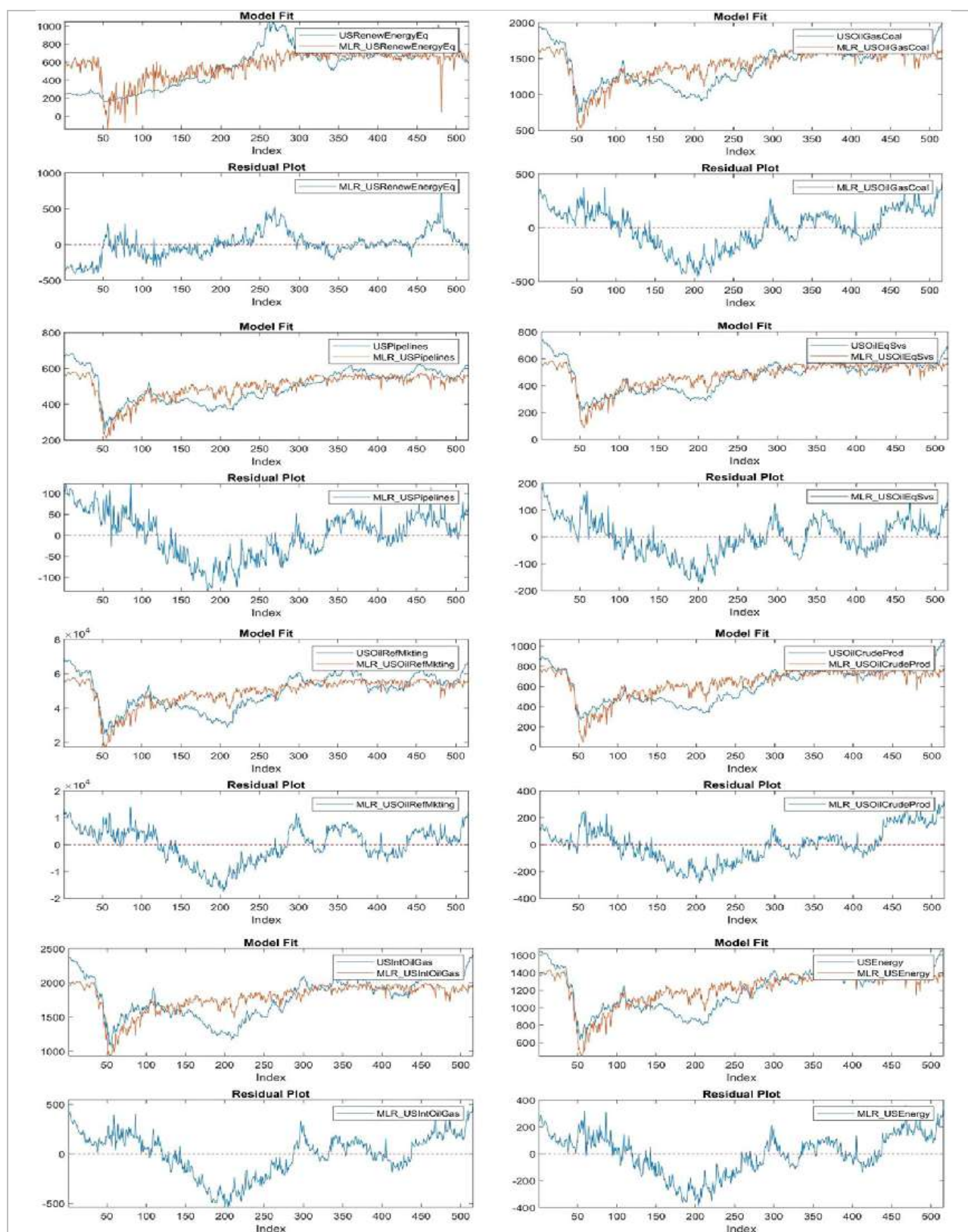


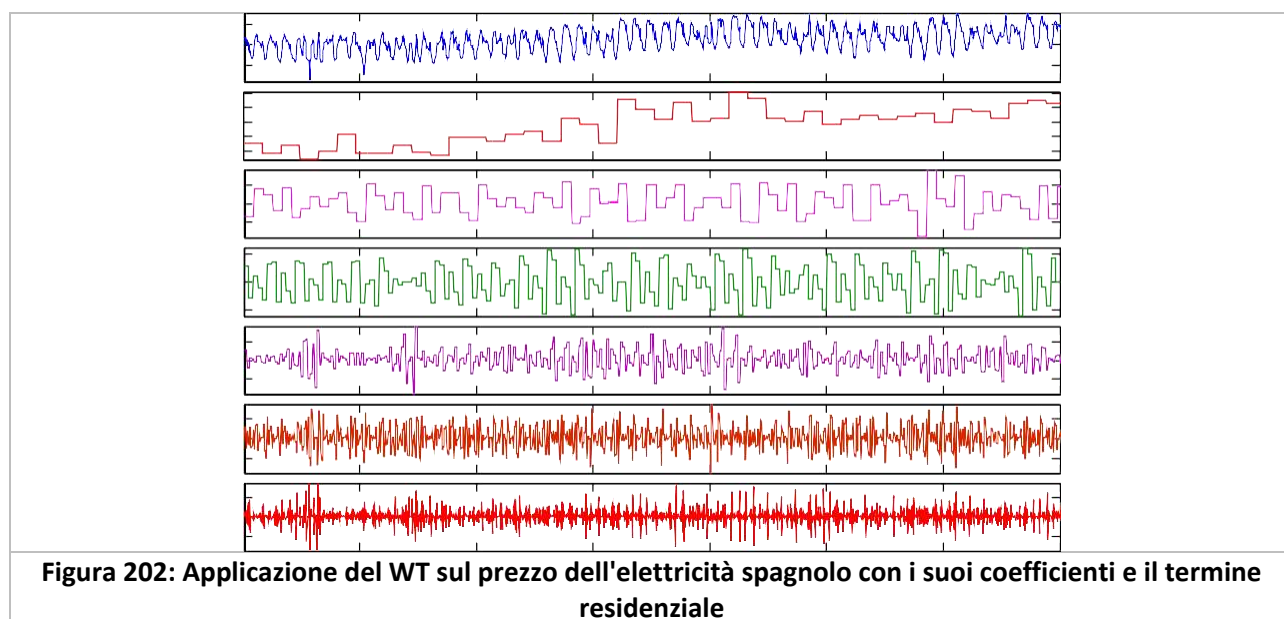
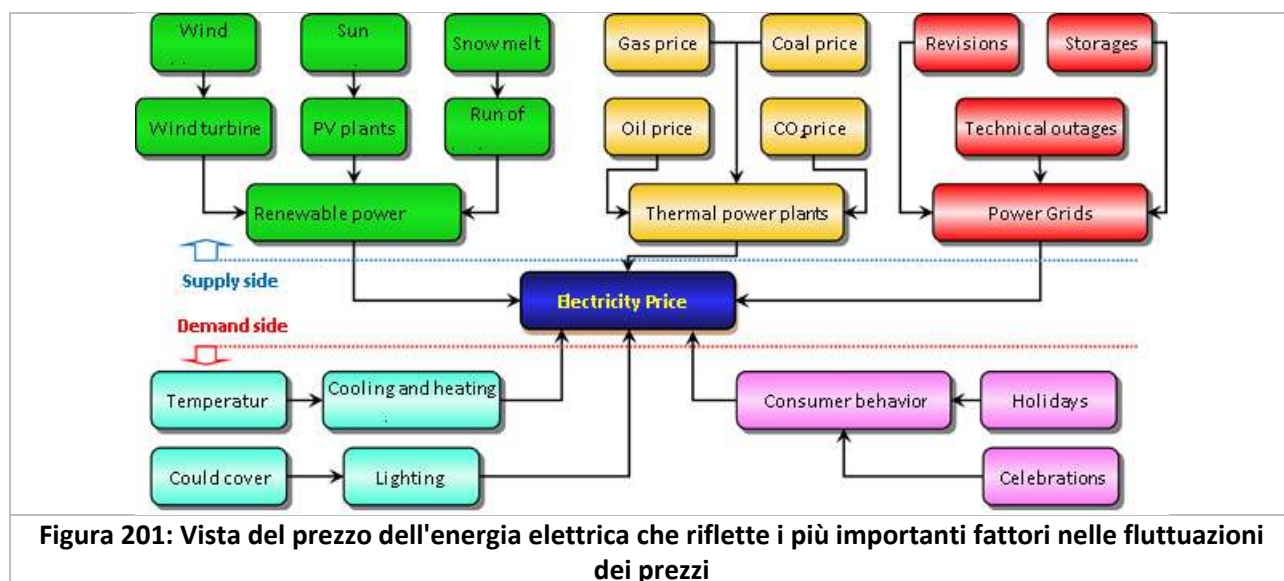
Figura 200: Prezzi effettivi e previsti con i residui del modello MLR: periodo pandemia COVID19. Note: Il la linea blu presenta i prezzi effettivi mentre l'arancione le previsioni del modello MLR

Questo studio utilizza ML e approcci tradizionali (ad es. regressione) per prevedere i prezzi di indici azionari energetici. Il Root Mean Square Error (RMSE) viene utilizzato per esaminare l'implementazione dei modelli insieme all'accuratezza della prevedibilità. In particolare, si sono applicati 25 approcci ML, inclusi **NN**, **Support Vector Regression (SVR)**, **Regression Trees (RT)** e **modelli Gaussian Process Regression (GPR)**.

Si è scoperto che i modelli ML hanno sovraperformato l'approccio MLR in tutti i casi. Il **Nonlinear Autoregressive with External (Exogenous) parameters (NARX)** era superiore in quanto significativamente migliore in precisione. L'uso simultaneo dell'algoritmo di Levenberg-Marquardt e l'autoregressione con parametri di input esogeni ha reso il risultato della risposta del tutto accurato e ridotto al minimo l'RMSE rispetto ad altri modelli di machine learning NN. Insomma, i modelli di Intelligenza Artificiale (AI)" e Machine Learning (ML) hanno consentito una maggiore precisione di previsioni rispetto ai tradizionali modelli di regressione lineare multipla. Le Reti Neurali appaiono come il modello migliore.

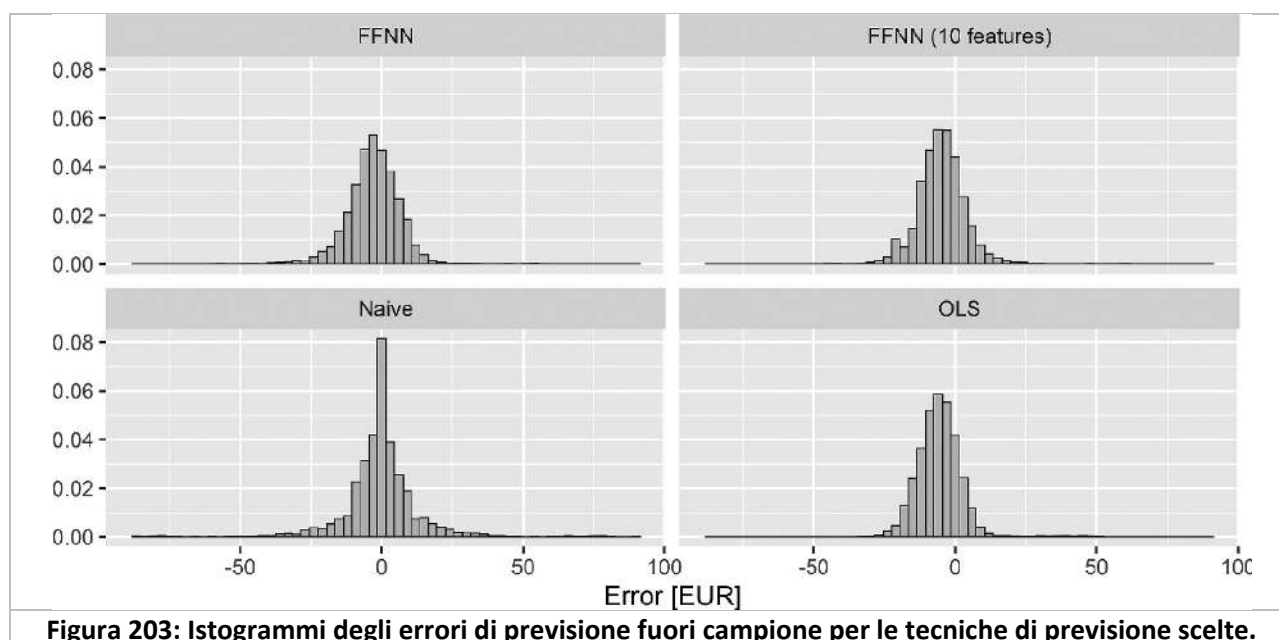
Sebbene lo studio abbia contribuito a importanti scoperte, ha sofferto di alcune limitazioni. In primo luogo, lo studio è stato limitato al contesto statunitense, anche se sarebbe utile eseguirlo a livello globale. Ciò è dovuto alla disponibilità di dati statunitensi relativi ai prezzi delle azioni energetiche e indici di incertezza. In secondo luogo, i dati ad alta frequenza non sono facilmente accessibili.

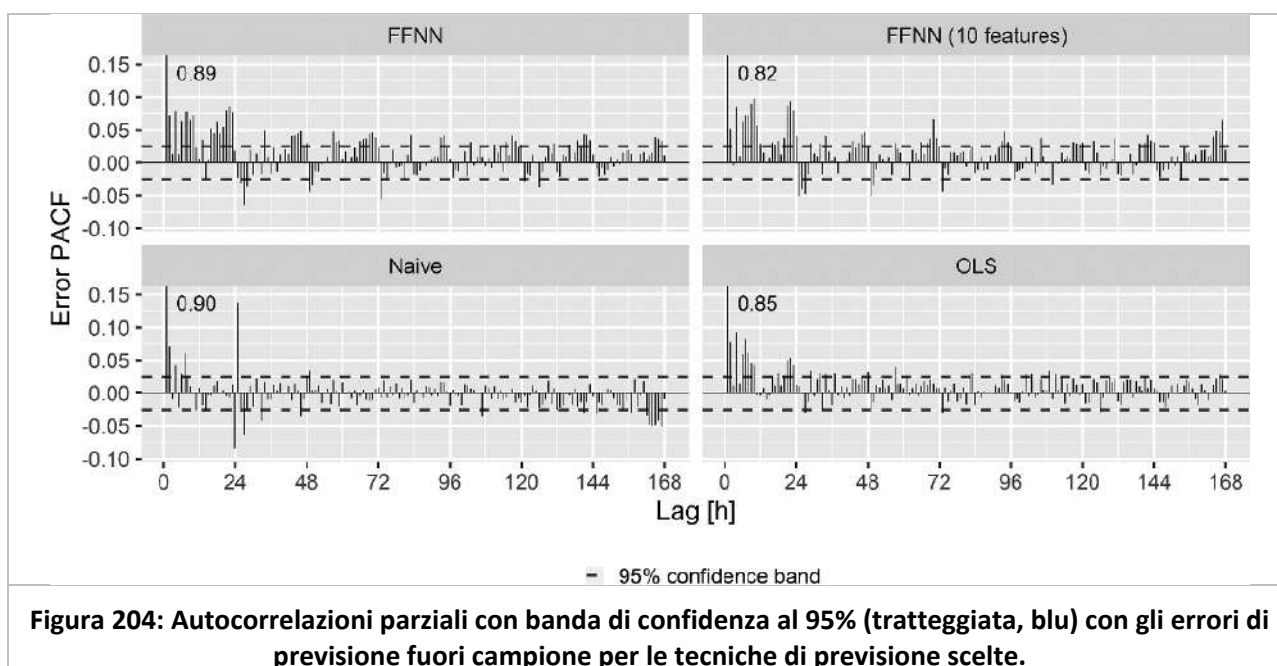
In [114] si propone un algoritmo per la predizione del prezzo dell'energia elettrica. Per coprire questo obiettivo, si utilizza un algoritmo diviso in tre passaggi, vale a dire pre-elaborazione, apprendimento e messa a punto. La parte di pre-elaborazione consiste in **Wavelet Packet Transform (WPT)** per analizzare il segnale di prezzo per le sottoserie ad alta e bassa frequenza e l'informazione mutua variazionale (VMI) per selezionare preziosi dati di input al fine di aiutare la parte di apprendimento e diminuire il calcolo fardello. Grazie alla parte di apprendimento, **un nuovo kernel fuzzy autoadattativo basato su Least Squares Support Vector Machine (LSSVM-SFK)** viene proposto per estrarre il miglior modello di mappa dai dati di input. Un nuovo HBMO viene introdotto per impostare in modo ottimale le variabili LSSVM-SFK come bias, peso, ecc. Per migliorare le prestazioni del honey bee mating optimization (HBMO), vengono proposte due modifiche che hanno un'elevata stabilità.



Il risultato numerico basato sugli indici di errore di previsione dimostra che il quadro di previsione proposto migliora notevolmente l'accuratezza delle previsioni in tutti i mercati dell'energia elettrica.

In [115] si utilizzano algoritmi di apprendimento automatico per prevedere i prezzi del mercato spot dell'elettricità in Germania. Le previsioni utilizzano in particolare i dati del portafoglio ordini bid e ask del mercato spot, ma anche dati di mercato fondamentali come l'alimentazione rinnovabile e la domanda totale attesa. L'estrazione di caratteristiche appropriate per i dati del portafoglio ordini, viene sviluppata procedendo dalla letteratura esistente. Utilizzando la convalida incrociata per ottimizzare gli iperparametri, le reti neurali e i random forest si adattano ai dati. Le loro prestazioni all'interno e all'esterno del campione vengono confrontate con i modelli di riferimento statistici.

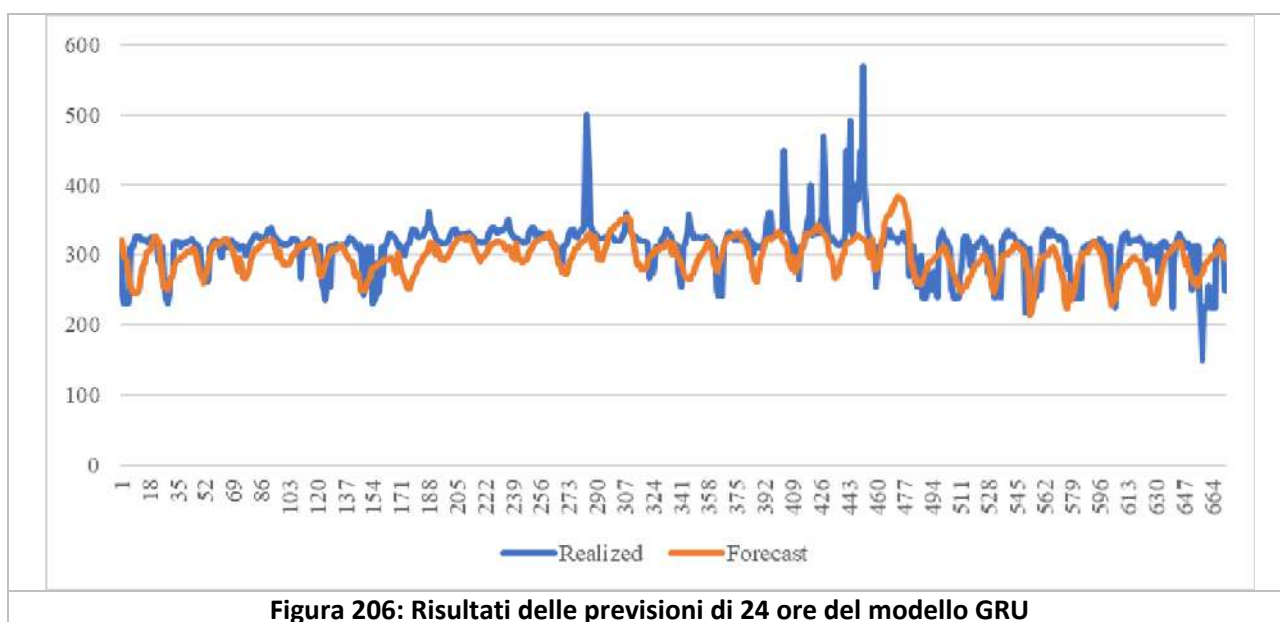
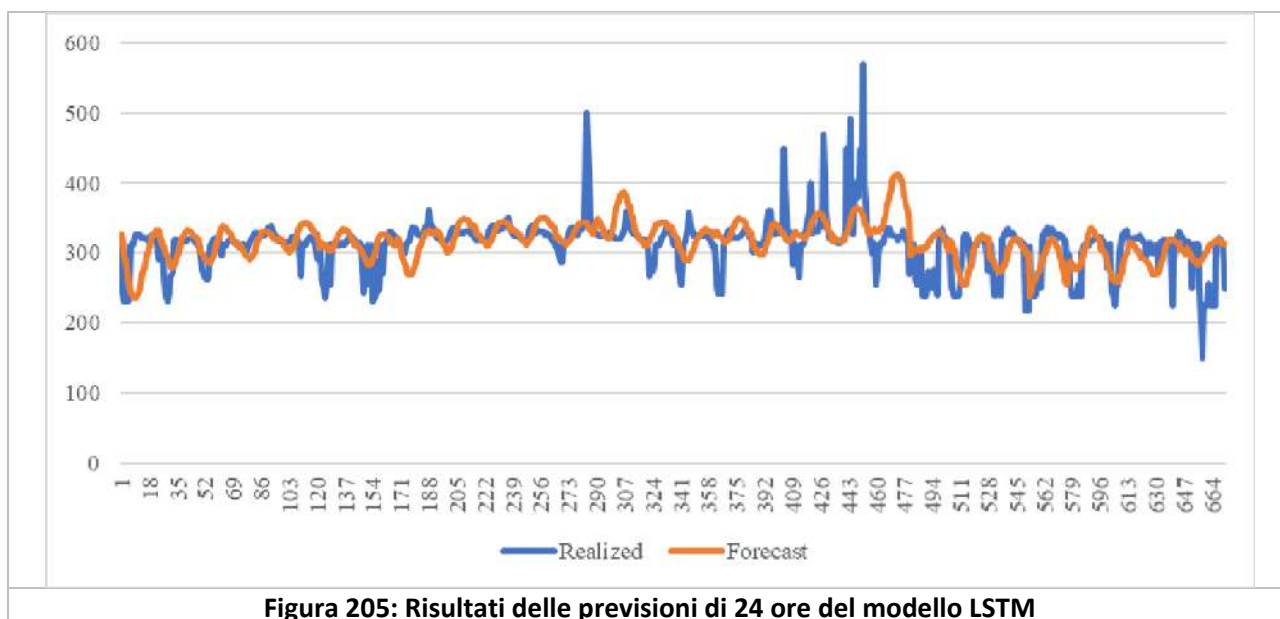




I risultati mostrano che le **Neural Network (NN)** possono effettivamente fornire un prezzo basato sul libro degli ordini di previsioni con risultati competitivi. Tuttavia, non hanno prestazioni significativamente migliori rispetto a metodi più semplici come la normale regressione lineare. Sebbene la classica tecnica di previsione basata sul libro degli ordini richieda molte analisi statistiche, l'ottimizzazione dell'architettura di rete richiede anche risorse significative. Si è anche scoperto che riducendo il numero di funzioni, generalmente si migliorano i risultati. Per quanto riguarda l'RMSE, si è trovato che una rete neurale feed-forward con solo 10 funzionalità selezionate dalla foresta casuale si comporta meglio. Considerando il MAE (misura direttamente collegata ai ricavi da trading finanziario), la rete neurale feed-forward senza selezione delle funzionalità è la più performante. Tuttavia, anche il modello ingenuo mostra buoni risultati, a sostegno del fatto che questa euristica tradizionale è spesso applicata in economia energetica. Le architetture di rete neurale dalla letteratura mostrano risultati competitivi nel campione, ma le loro prestazioni diminuiscono in modo significativo in un'analisi fuori campione. Ciò indica che queste architetture di modello sono overfitting sul set di dati.

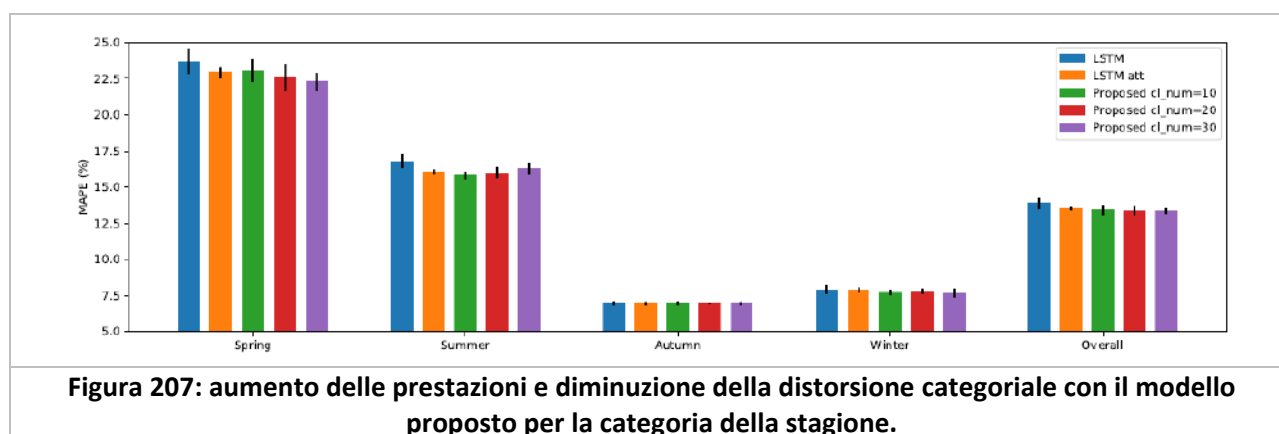
Lo studio in [116] mira a prevedere il Market Clearing Price (il prezzo al quale la domanda di un bene da parte dei consumatori è uguale al numero di beni che possono essere prodotti a quel prezzo) utilizzando tecniche di deep learning. In questo contesto, i prezzi di compensazione del mercato 24 ore su 24 sono stati previsti con MLP, CNN, LSTM e GRU. In particolare:

- **Multilayer Perceptrons (MLP):** le reti neurali artificiali sono metodi di intelligenza artificiale che cercano di apprendere imitando l'acquisizione e l'elaborazione delle informazioni del cervello umano. Il metodo di rete neurale artificiale più utilizzato è il modello di rete percettiva con più livelli, chiamati percettroni multistrato,
- **Convolutional Neural Networks (CNN):** le CNN utilizzano una speciale operazione matematica lineare, chiamata convoluzione, in almeno uno dei suoi livelli per elaborare i dati. Le reti neurali convoluzionali sono costituite da neuroni che si auto-ottimizzano attraverso l'apprendimento.
- **Long Short Term Memory (LSTM):** le reti neurali ricorrenti sono un tipo di rete neurale con loop che consentono l'uso permanente di informazioni del passato nella struttura della rete. In RNN, il gradiente dei pesi diventa troppo piccolo o troppo grande man mano che il passo temporale si allunga. LSTM, un tipo speciale di rete neurale ripetitiva con celle di memoria, ricorda le informazioni per lungo tempo con l'aiuto delle celle di memoria.
- **Gated Recurrent Unit (GRU):** come per le LSTM, il modello GRU è progettato per ricordare i dati della fase temporale precedente. Mentre LSTM utilizza tre porte, il GRU utilizza due porte come porta di ripristino e porta di aggiornamento. Il cancello di ripristino viene utilizzato per decidere quante informazioni passate devono essere dimenticate o ricordate.



La LSTM ha avuto la media migliore di prestazione previsionale con un valore MAPE di 8,15, in base ai risultati ottenuti. MLP ha seguito LSTM con 8.44 MAPE, GRU con 8.72 MAPE e CNN con 9.27 MAPE. Nello studio, le province dove ci sono molti impianti che producono con risorse rinnovabili sono state selezionate per l'analisi delle variabili meteorologiche.

In [117] si presenta un nuovo approccio alla previsione dei prezzi orari dell'elettricità del giorno prima. L'approccio presentato in questo lavoro si concentra sul problema della costruzione di un modello equo e imparziale. In questo caso, imparziale significa che il modello può aumentare l'accuratezza della previsione e ridurre la distorsione categoriale tra diversi cluster di dati. A tale scopo, viene creato un modello che combina tecniche come la **Long Short-Term Memory (LSTM)**, la **Recurrent Neural Network (RNN)**, il **meccanismo di attenzione** e il **clustering**. La caratteristica principale del modello proposto è che i pesi di attenzione per gli stati nascosti LSTM sono calcolati considerando un vettore di contesto dato individualmente per ogni campione come centro del cluster a cui appartiene il campione. In modalità di addestramento, i campioni vengono ripetuti in modo iterativo (una volta per epoca) raggruppati in base a vettori di rappresentazione dati dal meccanismo dell'attenzione. Nello studio empirico, il modello proposto è stato applicato e valutato sui dati di mercato di Nord Pool. Per confermare che il modello diminuisce il bias categorico, i risultati ottenuti sono stati confrontati con i risultati di modelli LSTM simili ma senza il proposto meccanismo di attenzione.



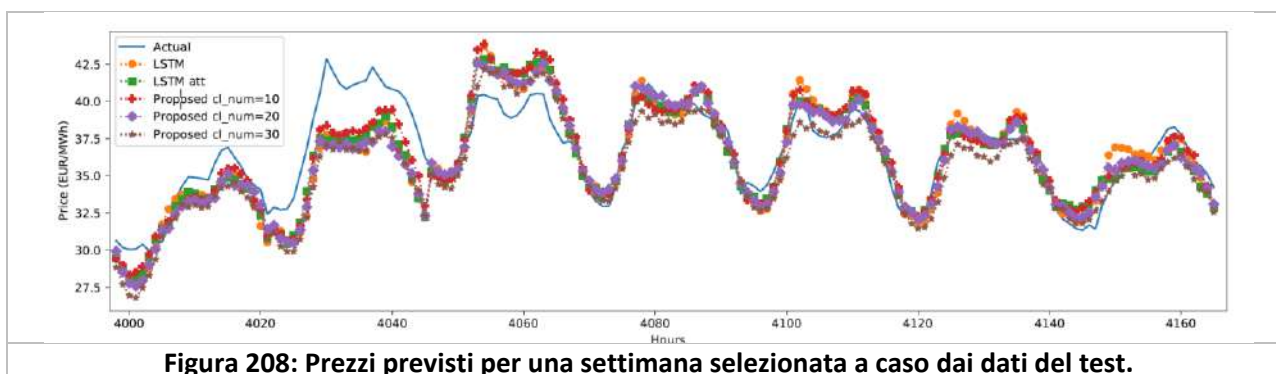


Figura 208: Prezzi previsti per una settimana selezionata a caso dai dati del test.

Per confrontare e analizzare l'accuratezza predittiva dei vari predittori, è stato calcolato il loro MAPE sul set di test. I modelli sono stati riaddestrati da zero cinque volte ciascuno per una maggiore solidità statistica dei risultati. È importante da notare che i predittori non vengono nuovamente stimati quando sono disponibili nuovi dati, ovvero i modelli sono stati addestrati con dati dal 2013 al 2017, mentre i dati dei test coprono il 2019. I risultati dimostrano che l'aggiunta di un meccanismo di attenzione ha migliorato le previsioni e l'aggiunta del clustering li ha migliorati ancora di più. C'è anche una dipendenza tra l'aumento del numero di cluster e il miglioramento della previsione. Inoltre, per quasi tutte le divisioni, il modello proposto si è rivelato il predittore più imparziale. La deviazione standard è quindi la più piccola e, nella maggior parte dei casi, diminuisce con il numero di cluster che aumenta. L'unica eccezione è la suddivisione in ore, che può derivare dal fatto che si effettuano previsioni per l'intera giornata, non individualmente per ogni ora. Vale però la pena di notare che le differenze nella qualità delle previsioni tra le categorie sono significative, confermando la necessità di sviluppare metodi per eliminare questi problemi.

Nello studio di [118], si è cercato di migliorare l'accuratezza della previsione dei prezzi dell'elettricità attraverso l'uso di una tecnica di machine learning, ovvero un nuovo approccio di programmazione genetica. Sulla base dei dati empirici dei più grandi mercati dell'energia dell'UE, si è proposto un modello di previsione che considera le variabili relative a condizioni meteorologiche, prezzi del petrolio e tagliandi di CO₂ e prevede i prezzi dell'energia con 24 ore di anticipo. Più specificamente, questo studio ha presentato un nuovo sistema di **programmazione genetica Liquid State Genetic Programming (LSGP)** che include la consapevolezza semantica nel processo di ricerca, così come un ottimizzatore di ricerca locale per accelerare la convergenza del processo di apprendimento. Questo

nuovo sistema è stato testato contro diverse esistenti tecniche all'avanguardia considerate in studi precedenti per affrontare lo stesso problema di ottimizzazione.

Variable	Description
X ₀	Crude oil spot price (Euro)
X ₁	Settlement price EUR/t CO ₂ (Euro)
X ₂	Air temperature (degrees Celsius)
X ₃	Air density (Kg/m ³)
X ₄	Ground temperature (degrees Celsius)
X ₅	Air pressure (Pascal)
X ₆	Relative humidity (expressed as a percentage)
X ₇	Wind speed (meters/second)
X ₈	Maximum intraday air temperature (degrees Celsius)
X ₉	Minimum intraday air temperature (degrees Celsius)
X ₁₀	Minimum intraday air temperature at ground level (degrees Celsius)
X ₁₁	Maximum intraday wind speed (meters/second)
X ₁₂	Next-day rain precipitation forecast to cover a horizontal ground surface (millimeters)
X ₁₃	Next-day snow precipitation forecast to cover a horizontal ground surface (centimeters)
X ₁₄	Sunshine duration (hours)
X ₁₅	Snow depth (centimeters)
TARGET	Electricity price in EUR in the following day (24 hour prediction)

Figura 209: Variabili del modello disponibili nel dataset considerato

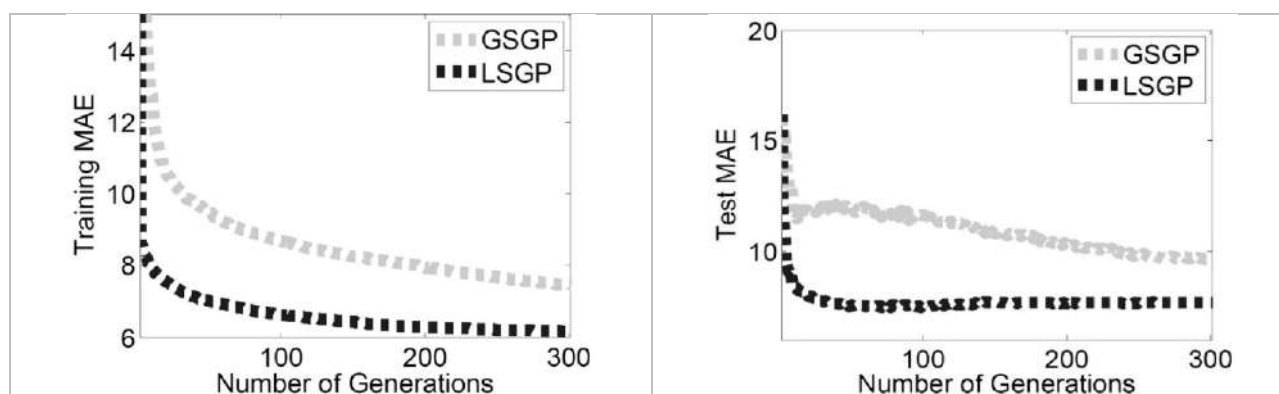


Figura 210: MAE di addestramento e test

I risultati mostrano che l'approccio proposto è molto competitivo; infatti, l'LSQP ha superato tutti le altre alternative per il problema in questione, fornendo un robusto software standard che può essere utilizzati dagli operatori del mercato dell'energia per ottenere previsioni accurate e precise dei prezzi dell'energia. LSQP ci consente di contrastare parzialmente uno degli svantaggi di GSGP, ovvero la creazione di soluzioni che, pur rappresentando espressioni matematiche, sono molto

difficili da comprendere a causa della loro complessità. LSGP può infatti produrre un modello finale rappresentato da un'espressione matematica che contiene pochi operatori e variabili. Nel documento, ci si è concentrati sui mercati dell'elettricità dell'UE, poiché la borsa dell'energia tedesca è la principale nello scambio di energia in Europa centrale. Questo rappresenta un limite ma fornisce anche indicazioni per ulteriori ricerche.

Lo studio di [119] si concentra su cinque variabili LOLP (Loss of Load Probability) con diverse stime del tempo in anticipo e altri parametri quasi deterministici, per spiegare il comportamento dei prezzi di un modello di regressione multivariabile. Questi includono la produzione di base, il carico del sistema, la generazione solare ed eolica, la stagionalità, il prezzo del giorno prima e lo squilibrio sui contributi di volume. Questo studio presenta lo sviluppo di un modello multi-regressione, testando tre algoritmi di machine learning, **Gradient Boosting (GB)**, **Random Forest (RF)** ed **Extreme Gradient Boosting (XGBoost)**, presentando un approccio combinato di diverse categorie. I dati storici di mercato vengono utilizzati per la generazione, l'offerta, il carico, l'effetto temporale come il periodo di liquidazione, giorno, mese, festività, stagione e incertezze non strategiche, come il carico previsto e la probabilità di riserve per soddisfare la domanda. Per quest'ultima variabile, la probabilità di perdita di carico (LOLP) viene utilizzata con diversi orizzonti temporali per catturare questa incertezza. Il modello è uno strumento per la previsione a breve termine, utilizzabile dalle 12 h avanti fino a 1 h prima della chiusura del cancello. Per condurre l'analisi, viene preso in considerazione il mercato dell'energia di bilanciamento ELEXON nel Regno Unito.

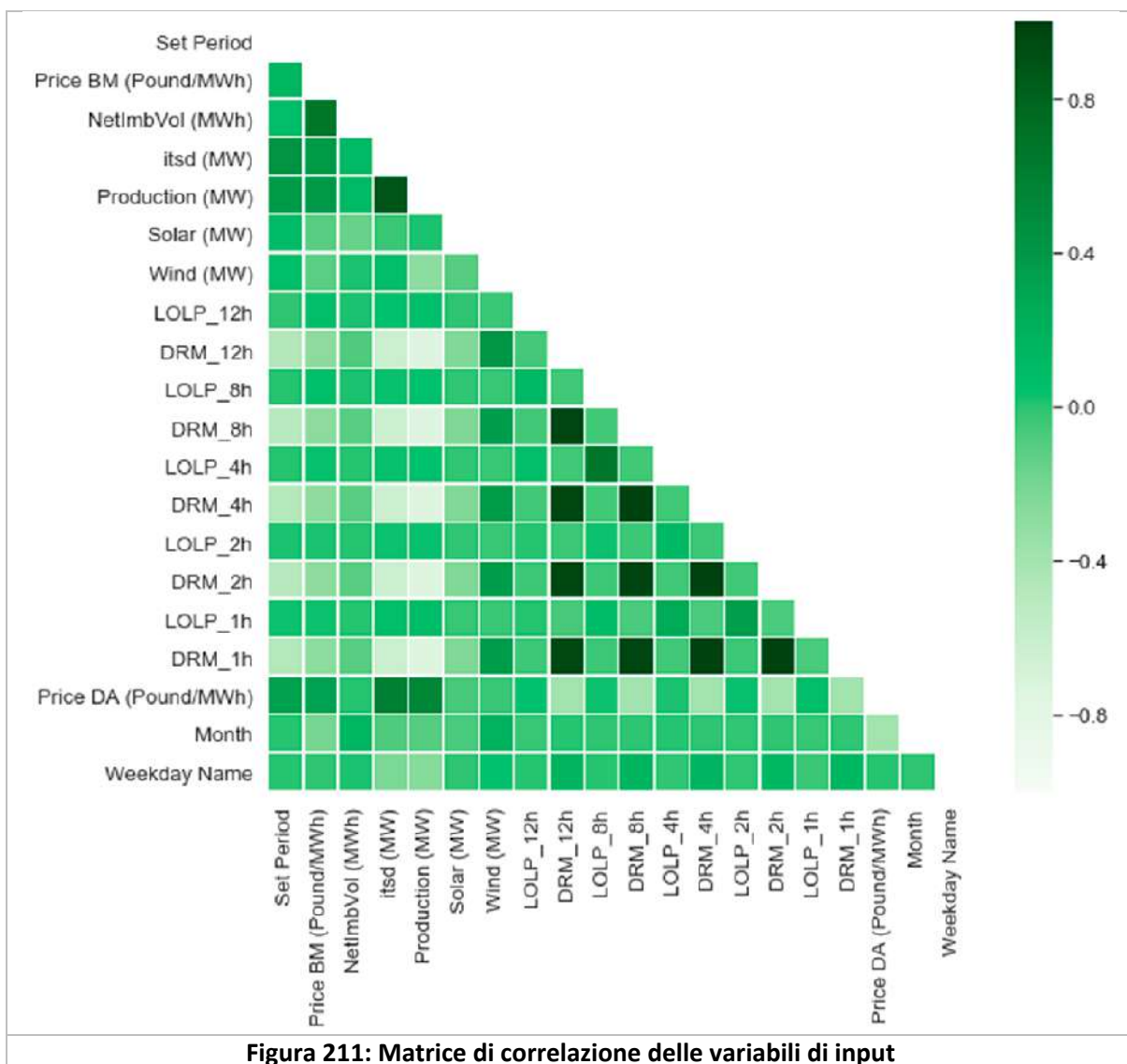


Figura 211: Matrice di correlazione delle variabili di input

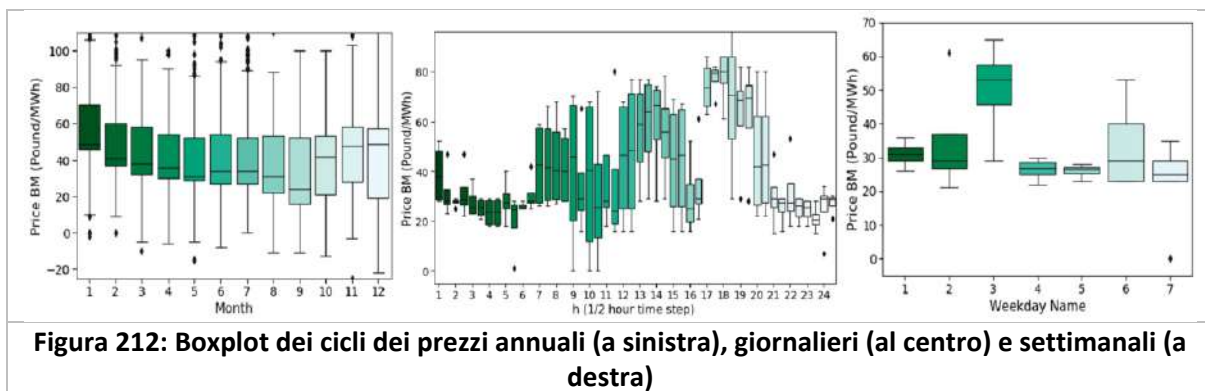


Figura 212: Boxplot dei cicli dei prezzi annuali (a sinistra), giornalieri (al centro) e settimanali (a destra)

In termini di importanza delle caratteristiche, il Net Imbalance Volume, il LOLP (aggregato), i margini declassati (aggregati) e le variabili mensili hanno ottenuto il punteggio più alto, con rispettivamente il 28,6%, il 27,5%, il 14,0% e l'8,9% del peso sull'importanza delle caratteristiche. Il modello ha un MAE di 7,89 £/MWh, un punteggio R2 del 76,8% e un MSE di 124,74, che è accettabile per il problema affrontato. Lo studio mostra che i LOLP sono predittori importanti da considerare, mentre l'incertezza relativa alla variabile NetImbVol può essere mitigata con una regressione quantile. Tuttavia, rimane un predittore che merita ulteriori indagini. Nell'esempio reale fornito, i picchi del quotidiano e le fluttuazioni dei prezzi sono stati ben previsti dal modello e confermate dall'analisi statistica e quindi si può presumere che il PS corretto, potrebbe essere potenzialmente ben individuato per allocare la Flessibilità DR. Inoltre, l'ampiezza del prezzo è stata prevista con una media assoluta d'errore accettabile.

In [120] si esaminano le prestazioni di diversi modelli di deep learning per il mercato elettrico finlandese. L'indagine comprende diversi tipi di architetture, schemi di aggregazione dei dati e metodo di preformazione. I contributi di questo scritto sono tre:

- Proporre nuovi modelli previsionali per i Prezzi Day-Ahead dell'elettricità.
- Analizzare l'influenza dei mercati vicini in previsione dei prezzi.
- Studio di diversi schemi di aggregazione dei dati.

Per questo studio sono stati presi in considerazione tre anni di dati, dal 01/01/2016 al 31/12/2018, relativi al Nordic electricity market, noto anche come Nord Pool. Inoltre, è stata utilizzata una risoluzione oraria, dal day-ahead i prezzi sono commercializzati in questa risoluzione. Infine, sono stati utilizzati i dati sui prezzi dei due giorni precedenti.

Sono state esplorate tre diverse architetture, vale a dire **Feedforward Neural Network (FNN)**, **Long Short-Term Memory (LSTM)** e **Convolution Neural Networks (CNN)**. Per i tre tipi di architettura, le informazioni sui prezzi di un mercato esterno sono state incluse come un modo per esaminare l'influenza di un mercato non direttamente connesso a quello in analisi. Inoltre, per le architetture

LSTM e CNN sono stati inoltre implementati due diversi schemi di concatenazione e un processo di pre-addestramento.

Model	Architecture	Concatenation type	Inputs	Pre training
M1	ANN	Input level	P	No
M2			P, G, C	No
M3			P, G, C, P_{UK}	No
M4	LSTM	Input level	P, G, C	No
M5			P, G, C, P_{UK}	No
M6			P, G, C, P_{UK}	Yes
M7		Hidden layer level	P, G, C	No
M8			P, G, C, P_{UK}	No
M9			P, G, C, P_{UK}	Yes
M10	CNN	Input level	P, G, C	No
M11			P, G, C, P_{UK}	No
M12			P, G, C, P_{UK}	Yes
M13		Hidden layer level	P, G, C	No
M14			P, G, C, P_{UK}	No
M15			P, G, C, P_{UK}	Yes

Figura 213: Riepilogo dei modelli

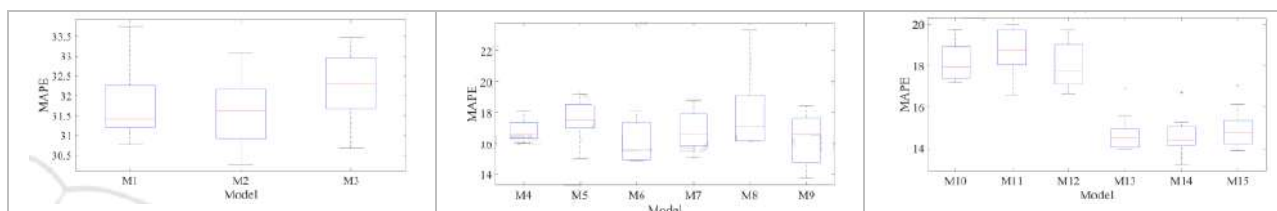


Figura 214: Box-plot dell'errore percentuale assoluto medio: (a) modelli basati su FNN, (b) modelli basati su LSTM e (c) modelli basati su CNN.

Complessivamente, 15 modelli sono stati testati e i risultati hanno indicato schemi di architettura promettenti anche per la previsione dei prezzi, come l'importanza di sviluppare architetture più complesse quando si ha a che fare con informazioni così volatili. La CNN unidimensionale ha mostrato i migliori risultati tra tutti i modelli e, pertanto, è l'architettura consigliata dagli autori per ulteriori ricerche.

4 Conclusioni

In questo documento è stato introdotto il concetto di modellazione di asset e funzionalità tramite digital twin e confrontati i principali due metodi per ottenere tale obiettivo: Machine learning based e Physics based. Si sono quindi spiegate le ragioni per le quali il primo approccio è preferibile al primo. Partendo da tale assunto, si è poi proceduto ad analizzare lo Stato dell'Arte (nella maggior parte dei casi si fa riferimento ad articoli scientifici di riviste e di conferenze apparsi dal 2020 in poi) dell'applicazione di tecniche di intelligenza artificiale (Machine learning) per la generazione dei modelli di carichi e generatori. In particolare, si è suddivisa tale analisi in base ai fini della modellazione stessa: nel capitolo 2 si è trattato della modellazione con obiettivo l'ottimizzazione di processo e la manutenzione, mentre nel capitolo 3 ci si è dedicati alle tecniche per la predizione dell'energia prodotta o richiesta da un sistema. Il capitolo 3 ha inoltre trattato l'argomento della modellazione dell'andamento del costo dell'energia stessa, coprendo il caso d'interesse in cui l'energia, non potendo essere prodotta dal sistema in esame, deve essere acquistata da enti terzi. Tutte le analisi prodotte sono state ulteriormente distinte in base al tipo di generazione dell'energia (fonti rinnovabili e non) e all'orizzonte temporale considerato nella predizione. Per quanto riguarda l'analisi dei consumi, si è dedicata una sezione allo studio dell'approccio NILM che, pur essendo una metodologia molto recente, potrebbe permettere di superare i limiti imposti dalla attuale scarsità di dati energetici provenienti da singoli asset (comparata alla disponibilità di dati aggregati).

Complessivamente questo documento ha dimostrato l'enorme crescita delle applicazioni dell'intelligenza artificiale nella modellazione di fenomeni energetici con talune metodologie (Artificial Neural Network, Convolutional Neural Network, Long Short-Term Memory algorithms, Decision Trees, Support Vector Machine) che si sono consolidate in ambito sia accademico che applicativo mentre altre (Deep Learning, Gated Recurrent Unit, Gradient Boosting Machine) che sono tuttora in ascesa. Per tale ragione, si dedicherà una particolare attenzione a tali algoritmi nell'ambito di progetto, mantenendo comunque un approccio aperto alle tecnologie emergenti.

Riferimenti

- [1] E. S. Ove, Merging Physics, Big Data Analytics and Simulation for the Next-Generation, Norwegian University of Science and Technology, September 2017.
- [2] D. Julianna, «Data Science/Machine Learning,» IBM Analytics, 12 March 2021. [Online].
- [3] H. H. Hosamo, P. R. Svennevig, K. Svidt, D. Han e H. K. Nielsen, «A Digital Twin predictive maintenance framework of air handling units based on automatic fault detection and diagnostics,» *Energy and Buildings*, vol. 261, n. 111988, 15 April 2022.
- [4] H. H. Hosamo, H. K. Nielsen, D. Kraniotis, P. R. Svennevig e K. Svidt, «Digital Twin framework for automated fault source detection and prediction for comfort performance evaluation of existing non-residential Norwegian buildings,» *Energy and Buildings*, vol. 281, n. 112732, 15 Febbraio 2022.
- [5] H.-A. Park, G. Byeon, W. Son, H.-C. Jo, J. Kim e S. Kim, «Digital Twin for Operation of Microgrid: Optimal Scheduling in Virtual Space of Digital Twin,» *MDPI Energies* 2020, vol. 13, n. 5504, 20 October 2020.
- [6] H. Hosam, M. H. Hosamo, H. K. Nielsen, P. R. Svennevig e K. Svidt, «Digital Twin of HVAC system (HVACDT) for multiobjective optimization of energy consumption and thermal comfort based on BIM framework with ANN-MOGA,» Taylor&Francis Online, 26 ottobre 2022. [Online]. Available: <https://www.tandfonline.com/doi/full/10.1080/17512549.2022.2136240>.
- [7] X. Xie, J. Merino, N. Moretti, P. Pauwels, J. Y. Chang e A. Parlikad, «Digital twin enabled fault detection and diagnosis process for building HVAC systems,» *Automation in Construction*, vol. 146, n. 104695, Febbraio 2023.
- [8] Y. Fathy, M. Jaber e Z. Nadeem, «Digital Twin-Driven Decision Making and Planning for Energy Consumption,» *MDPI J. Sens. Actuator Netw.*, vol. 10(2), 20 June 2021.
- [9] A. Rosato, S. S. Francesco Guarino, E. Entchev, M. Masullo e L. Maffei, «Healthy and Faulty Experimental Performance of a Typical HVAC System under Italian Climatic Conditions: Artificial Neural Network-Based Model and Fault Impact Assessment,» *Energies* 2021, vol. 14, n. 5362, 28 August 2021.
- [10] A. Rosato, F. Guarino, S. Sibilio, E. Entchev, M. Masullo e L. Maffei, «Healthy and Faulty Experimental Performance of a Typical HVAC System under Italian Climatic Conditions: Artificial Neural Network-Based Model and Fault Impact Assessment,» *Energies* 2021, vol. 14, n. 5362, 28 August 2021.
- [11] Y. Liu, J. Wang, J. Deng e W. S. P. Tan, «Non-Intrusive Load Monitoring Based on Unsupervised Optimization Enhanced Neural Network Deep Learning,» *Front. Energy Res.*, 30 September 2021, *Sec. Smart Grids*, vol. 9, 30 September 2021.
- [12] S. Naderian, «A Novel Hybrid Deep Learning Approach for Non-Intrusive Load Monitoring of Residential Appliance Based on Long Short Term Memory and Convolutional Neural Networks,» *Computer Science Machine Learning*, n. arXiv:2104.07809, 15 Apr 2021.
- [13] L. Ogrizek, B. Bertalanic, G. Cerar e M. Meza, «Designing a Machine Learning based Non-

- intrusive Load Monitoring Classifier,» in *IEEE ERK*, Ljubljana, Slovenia, 06 October 2021.
- [14] L. Zhang e H. Zhu, «A Simple Method of Non-Intrusive Load Monitoring Based on BP Neural Network,» *IOP science J. Phys.: Conf. Ser.* 2237 012010, 2022.
- [15] S. Athanasoulas, S. Sykiotis, M. Kaselimi, E. Protopapadakis e N. Ipiotis, «A First Approach using Graph Neural Networks on Non-Intrusive-Load-Monitoring,» in *PETRA '22: Proceedings of the 15th International Conference on Pervasive Technologies Related to Assistive Environments*, 11 July 2022.
- [16] M. U. Hadi, N. H. N. Suhaimi e A. Basit, «Efficient Supervised Machine Learning Network for Non-Intrusive Load Monitoring,» *MDPI Technologies* 2022, vol. 10(4), n. 85, 16 July 2022.
- [17] M. D’Incecco, S. Squartini e M. Zhong, «Transfer Learning for Non-Intrusive Load Monitoring,» *IEEE Transactions on Smart Grid*, pp. 11 (2). pp. 1419-1429 ISSN 1949-3053, 2020.
- [18] I. Wang, S. Mao e R. M. Nelms, «Transformer for Non-Intrusive Load Monitoring: Complexity Reduction and Transferability,» *IEEE Internet Of Things Journal*, 01 October 2022.
- [19] S. Chen, B. Zhao, M. Zhong, W. Luan e Y. Yu, «Non-intrusive Load Monitoring based on Selfsupervised Learning,» *Electrical Engineering and Systems Science, Signal Processing*, n. <https://doi.org/10.48550/arXiv.2210.04176>, 9 Oct 2022.
- [20] J. Jorgensen, M. Hodkiewicz, E. Cripps e G. M. Hassan, «Requirements for the application of the Digital Twin Paradigm to offshore wind turbine structures for uncertain fatigue analysis,» *Computers in Industry*, vol. 145, n. 103806, 28 November 2022.
- [21] M. Fahim, V. Sharma, T.-v. Cao, B. Canberk e T. Q. Duong, «Machine Learning-Based Digital Twin for Predictive Modeling in Wind Turbines,» *IEEE Access*, vol. 10, February 8, 2022.
- [22] J. Ma, Y. Yuan e P. Chen, «A fault prediction framework for Doubly-fed induction generator under time-varying operating conditions driven by digital twin,» *IET Electric Power Applications*, 15 December 2022.
- [23] J. Maron, D. Anagnostos, B. Brodbeck e A. Meyer, «Artificial intelligence-based condition monitoring and predictive maintenance framework for wind turbines,» *Journal of Physics: Conference Series*, January 2022.
- [24] E. E. o. Mammadov, «Predictive Maintenance of Wind Generators based on AI Techniques,» 2019-12-11.
- [25] A. P. Gonzalo, T. Benmessaoud, M. Entezami e F. P. GarcíaMarquez, «Optimal maintenance management of offshore wind turbines by minimizing the costs,» *Sustainable Energy Technologies and Assessments*, vol. 52C, n. 102230, 19 April 2022.
- [26] A. Amini, J. Kanfoud e T.-H. Gan, «An Artificial Intelligence Neural Network Predictive Model for Anomaly Detection and Monitoring of Wind Turbines Using SCADA Data,» *Applied Artificial Intelligence, An International Journal*, vol. 36, n. 1, 25 Mar 2022.
- [27] C. Martinez, F. A. Yeboah, M. B. Scott Herford e V. Puttagunta, «Predicting Wind Turbine Blade Erosion using Machine Learning,» *SMU Data Science Review*, vol. 2, n. 17, 2019.
- [28] Y.-Y. Hong e R. A. Pula, «Diagnosis of PV faults using digital twin and convolutional mixer with LoRa notification system,» *Energy Reports*, vol. 9, pp. 1963-1976, 3 January 2023.

- [29] A. A. Al-Katheri, E. A. Al-Ammar, M. A. Alotaibi, W. Ko, S. Park e H.-J. Choi, «Application of Artificial Intelligence in PV Fault Detection,» *MDPI Sustainability*, vol. 14, n. 21, 25 October 2022.
- [30] A. Ba, A. Ndiayee, E. H. M. Ndiaye e S. Mbodji, «Power optimization of a photovoltaic system with artificial intelligence algorithms over two seasons in tropical area,» *ScienceDirect MethodsX*, vol. 10, n. 101959, 2023.
- [31] B. Li, C. Delpha, A. Migan-Dubois e D. Diallo, «Fault diagnosis of photovoltaic panels using full I–V characteristics and machine learning techniques,» *HAL open science*, n. 03415367, 4 Nov 2021.
- [32] E. H. M. Ndiaye, A. Ndiaye e M. Faye, «Photovoltaic power optimization based on artificial intelligence method,» *IJSET - International Journal of Innovative Science, Engineering & Technology*, vol. 8, n. 5, May 2021.
- [33] a. J. Benavides, P. Arévalo-Cordero, L. G. Gonzalez, L. Hernández-Callejo, F. Jurado e J. A. Aguado, «Method of monitoring and detection of failures in PV system based on machine learning,» *Rev.fac.ing.univ. Antioquia no.102 Medellín*, Oct 08, 2021.
- [34] U. C. MURAT KUZLU, V. SHARMA e Ö. GÜLER, «Gaining Insight Into Solar PV Power Generation Forecasting Utilizing XAI Tools,» *IEEE Access*, vol. 8, October 26, 2020.
- [35] M. Chang, C.-C. Hsu, K. H. Chen, T. P. Hsu, K. Wei, K. Chuang e Y.-S. Chen, «PV O&M OPTIMIZATION BY AI PRACTICE,» in *36th European Photovoltaic Solar Energy Conference and Exhibition.*, 03 November 2019.
- [36] M. Zamen, A. Baghban, S. M. Pourkiaei e M. H. Ahmadi, «Optimization methods using artificial intelligence algorithms to estimate thermal efficiency of PV/T system,» *Energy Science & Engineering*, vol. 7, pp. 821-834, 12 April 2019.
- [37] M. A. Ebrahim, S. M. Ramadan, H. A. Attia, E. M. Saied, M. Lehtonen e a. H. A. Abdelhadi, «Improving the Performance of Photovoltaic by Using Artificial Intelligence Optimization Techniques,» *International Journal of renewable energy research*, vol. 11, n. 1, March, 2021.
- [38] A. M. Abed, L. F. Seddek e S. Elattar, «Building a Digital Twin Simulator Checking the Effectiveness of TEG-ICE Integration in Reducing Fuel Consumption Using Spatiotemporal Thermal Filming Handled by Neural Network Technique,» *Processes* 2022, 10(12), 2701; <https://doi.org/10.3390/pr10122701>, 14 December 2022.
- [39] A. Hasanzadeh, A. Chitsaz, A. Ghasemi e P. Mojaver, «Soft computing investigation of stand-alone gas turbine and hybrid gas,» *Energy Reports*, vol. 8, pp. 7537-7556, 13 June 2022.
- [40] R. Y. a. T. R. Biyanto, «Performance Optimization of Gas Turbine Generator Based on Operating Conditions Using ANN-GA at Saka Indonesia Pangkah Ltd,» *IPTEK Journal of Proceedings Series*, vol. 6, 2020.
- [41] H.-R. Kim e T.-S. Kim, «Optimization of Sizing and Operation Strategy of Distributed Generation System Based on a Gas Turbine and Renewable Energy,» *Energies* 2021, vol. 14, n. 8448, 14 December 2021.
- [42] A. V. T. Jayachandran, A. T. HH Omar e A. Krishnakumar, «Machine learning predictor for micro gas turbine performance evaluation,» *Aeron Aero Open Access J*, vol. 4, pp. 172-180, December 30, 2020.

-
- [43] A. Vafaei e M. A. Aliehyaei, «Optimization of micro gas turbine by economic, exergy and environment analysis using genetic bee colony and searching algorithms,» *Journal of Thermal Engineering*, vol. 6, n. 1, pp. 117-140, January 2020.
- [44] Y. JIN, Y. YING, J. LI e A. H. ZHOU, «Gas Path Fault Diagnosis of Gas Turbine Engine Based on Knowledge Data-Driven Artificial Intelligence Algorithm,» *IEEE Access*, vol. 9, July 30, 2021.
- [45] V. Mrzljak, N. Anđelić, I. Lorencin e S. B. Šegota, «The influence of various optimization algorithms on nuclear power plant steam turbine exergy efficiency and destruction,» *Scientific Journal of Maritime Research*, vol. 35, pp. 69-86, 2021.
- [46] A. A. M. Ali, O. A. M. Nasher e A. S. Al-Hegami, «A Framework for Total Quality Management of Diesel Generator Fuel Consumption using Machine Learning and Internet of Things (IoT),» *International Journal of Computer Applications*, vol. 176, 2020.
- [47] S. JAFARI e Y.-C. BYUN, «Prediction of the Battery State Using the Digital Twin Framework Based on the Battery Management System,» *IEEE Access*, vol. 10, n. 3225093, 01 December 2022.
- [48] D. Yang, Y. Cui, Q. Xia, F. Jiang, Y. Ren, B. Sun, Q. Feng, Z. Wang e C. Yang, «A Digital Twin-Driven Life Prediction Method of Lithium-Ion Batteries Based on Adaptive Model Evolution,» *Materials 2022*, vol. 15(9), n. 3331, 6 May 2022.
- [49] X. Hu, Y. Che, X. Lin e S. Onori, «Battery health prediction using fusion-based feature selection and machine learning,» *IEEE Transactions on Transportation Electrification*, 2020.
- [50] S. M. Hell e C. D. Kim, «Development of a Data-Driven Method for Online Battery Remaining-Useful-Life Prediction,» *Batteries 2022*, vol. 8, n. 192, 18 October 2022.
- [51] P. Khumprom e N. Yodo, «A Data-Driven Predictive Prognostic Model for Lithium-ion Batteries based on a Deep Learning Algorithm,» *Energies 2019*, vol. 12, n. 660, 18 February 2019.
- [52] Y. Gao e J. L. a. M. Hong, «Machine Learning Based Optimization Model for Energy Management of Energy Storage System for Large Industrial Park,» *Processes 2021*, vol. 9, n. 825, 9 May 2021.
- [53] T. Som, M. Dwivedi e A. C. Dubey, «Parametric Studies on Artificial Intelligence Techniques for Battery SOC Management and Optimization of Renewable Power,» *Procedia Computer Science*, vol. 167, pp. 353-362, 2020.
- [54] J. Vardakas, I. Zenginis e a. C. Verikoukis, «Machine Learning Methodologies for Electric-Vehicle Energy Management Strategies,» *Connected and Autonomous Vehicles in Smart Cities*, pp. 115-132, December 2020.
- [55] T. Jayakumara, N. M. Gowdab, R. Sujathac, S. N. Bhukyad, G. Padmapriyae, S. Radhikaf, V. Mohanavelg, M. S. h e R. Sathyamurthyj, «Machine Learning approach for Prediction of residual energy in batteries,» *Energy Reports*, vol. 8, pp. 756-764, 19 October 2022.
- [56] I. Ozer, S. B. Efe e H. Ozbay, «A combined deep learning application for short term load forecasting,» *Alexandria Engineering Journal (2021)*, vol. 60, pp. 3807-3818, 2021.
- [57] E. A. Madrid e N. Antonio, «Short-Term Electricity Load Forecasting with Machine Learning,» *Information 2021*, vol. 12, n. 50, 22 January 2021.
- [58] B. u. Islam e S. F. Ahmed, «Short-Term Electrical Load Demand Forecasting Based on LSTM

- and RNN Deep Neural Networks,» *Hindawi Mathematical Problems in Engineering*, vol. 2022, n. 2316474, 31 July 2022.
- [59] A. M. N. C. Ribeiro, P. R. X. d. Carmo, I. R. Rodrigues, D. Sadok e T. L. a. P. T. Endo, «Short-Term Firm-Level Energy Consumption Forecasting for Energy-Intensive Manufacturing: A Comparison of Machine Learning and Deep Learning Models,» *Algorithms 2020*, vol. 13, n. 274, 30 October 2020.
- [60] W. Chen, G. Han e H. Z. a. L. Liao, «Short-Term Load Forecasting with an Ensemble Model Based on 1D-UCNN and Bi-LSTM,» *Electronics 2022*, vol. 11, n. 3242, 9 October 2022.
- [61] A. M. N. C. Ribeiro, P. R. X. d. Carmo, P. T. Endo e P. R. a. T. Lynn, «Short- and Very Short-Term Firm-Level Load Forecasting for Warehouses: A Comparison of Machine Learning and Deep Learning Models,» *Energies 2022*, vol. 15, n. 750, 20 January 2022.
- [62] J. V. J. Melo, G. R. S. Lira, E. G. Costa e A. F. L. N. a. I. B. Oliveira, «Short-Term Load Forecasting on Individual Consumers,» *Energies 2022*, vol. 15, n. 5856, 12 August 2022.
- [63] B. I. L. Rabelo e E. G.-F. a. N. Clavijo-Buritica, «Machine Learning for Short-Term Load Forecasting in Smart Grids,» *Energies 2022*, 15, 8079, 31 October 2022.
- [64] C. Sankalpa, S. Kittipiyakul e a. S. Laitrakun, «Forecasting Short-Term Electricity Load Using Validated Ensemble Learning,» *Energies 2022*, vol. 15, n. 8567, 16 November 2022.
- [65] V. Tiwari e K. Pal, «Short-Term Load Forecasting for a Captive Power Plant Using Artificial Neural Network,» *International Journal of Information Retrieval Research*, vol. 12, 2022.
- [66] P. Matrenina, M. Safaralievb, S. Dmitrievb, S. Kokinb, . Ghulomzodac e S. Mitrofanov, «Medium-term load forecasting in isolated power systems based on ensemble machine learning models,» *Energy Reports*, Vol. %1 di %2Volume 8, Supplement 1, pp. Pages 612-618, April 2022.
- [67] B. N. Oreshkina, G. Dudekb, P. Petkab e E. Turkina, «N-BEATS neural network for mid-term electricity load forecasting,» *Applied Energy*, vol. 293, n. 116918, 1 July 2021,.
- [68] S.-R. YAN, M. TIAN, K. A. ALATTAS, ARDASHIR MOHAMMADZADEH, M. H. SABZALIAN e A. MOSAVI, «An experimental machine learning approach for mid-term energy demand forecasting,» *IEEE Access*, vol. 10, pp. 118926-118940, 10 November 2022.
- [69] N. Shirzadi, A. Nizami e M. K. a. M. Nik-Bakht, «Medium-Term Regional Electricity Load Forecasting through Machine Learning and Deep Learning,» *Designs 2021*, vol. 5, n. 27, 6 April 2021.
- [70] D. Zhang, W. G. J. Yang, H. Yu e W. X. a. T. Yu, «Medium—And Long-Term Load Forecasting Method for Group Objects Based on the Image Representation Learning,» *Frontiers in Energy Research August 2021*, vol. 9, n. 739993, 26 August 2021.
- [71] M. Askari e F. Keynia, «Mid-term electricity load forecasting by a new composite method based on optimal learning MLP algorithm,» *IET Generation, Transmission & Distribution*, 2020.
- [72] F. M. Butt, L. Hussain, S. H. M. Jafri, H. Mesfer Alshahrani, F. N. Al-Wesabi, K. J. Lone e E. M. T. E. D. & A. Duhayyim, «Intelligence based Accurate Medium and Long Term Load Forecasting System,» *Applied Artificial Intelligence*, vol. 36, 26 Jun 2022.
- [73] R. Mansor, B. J. Zaini e C. C. M. May, «Medium-term Electricity Load Demand Forecasting

- using Holt-Winter Exponential Smoothing and SARIMA in University Campus,» *International Journal of Engineering Trends and Technology*, vol. 69, n. 3, pp. 165-171, March 2021.
- [74] E. C. Ashigwuike, A. R. A. Aluya e J. E. C. E. a. S. A. B. Benson, «Medium term electrical load forecast of Abuja municipal area council using artificial neural network method,» *Nigerian Journal of Technology (NIJOTECH)*, vol. 39, n. 3, pp. 860-870, July 2020.
- [75] T. K. Lukong, D. N. Tanyu e T. T. T. & D. Schulz, «Long Term Electricity Load Forecast Based on Machine Learning for Cameroon's Power System,» *Energy and Environment Research*, vol. 12, n. 1, May 30, 2022.
- [76] S. K. Dash , M. Roccotelli, R. R. Khansama e M. P. F. A. M. Mangini, «Long Term Household Electricity Demand Forecasting Based on RNN-GBRT Model and a Novel Energy Theft Detection Method,» *Energy Theft Detection Method. Appl. Sci.* 2021, vol. 11, n. 8612, 16 September 2021.
- [77] S. Moalem , R. M. Ahari , G. Shahgholian e M. M. S. M. Kazemi, «Long-Term Electricity Demand Forecasting in the Steel Complex Micro-Grid Electricity Supply Chain—A Coupled Approach,» *Energies* 2022, vol. 15, n. 7972, 27 October 2022.
- [78] N. A. Mohammed e A. Al-Bazi, «An adaptive backpropagation algorithm for long-term electricity load forecasting,» *Neural Computing and Applications (2022)*, vol. 34, pp. 477-491, 11 August 2021.
- [79] S.-Y. Shin e H.-G. Woo, «Energy Consumption Forecasting in Korea Using Machine Learning Algorithms,» *Energies* 2022, vol. 15, n. 4880, 2 July 2022.
- [80] A. .. d. Real e F. D. a. J. Durán, «Energy Demand Forecasting Using Deep Learning: Application to the French Grid,» *Energies* 2020, vol. 13, n. 2242, 10 March 2020.
- [81] G. Ozdemir, «Long-term electrical energy demand forecasting by using artificial intelligence techniques – a case study of Turkey,» *Electrical engineering*, July 26th, 2022.
- [82] Z. Hongyua, L. Y. Chen Boa, G. Junweia, Z. W. Li Conga e Y. Haobo, «Research on medium- and long-term electricity demand forecasting under climate change,» *Energy Reports*, vol. 8, pp. 1585-1600, 10 March 2022.
- [83] A. Krechowicz, M. Krechowicz e K. Poczeta, «Machine Learning Approaches to Predict Electricity Production from Renewable Energy Sources,» *Energies*, vol. 15, n. 23, 2 December 2022.
- [84] A.-N. Buturache e S. Stancu, «Wind Energy Prediction Using Machine Learning,» *Low Carbon Economy*, vol. 12, pp. 1-21, January 27, 2021.
- [85] «Wind Power Prediction Based on Machine Learning and Deep Learning ModelsZahraa Tarek, Mahmoud Y. Shams2, Ahmed M. Elshewey, El-Sayed M. El-kenawy, Abdelhameed Ibrahim, Abdelaziz A. Abdelhamid and Mohamed A. El-dosuky,» *Computers, Materials & Continua*, vol. 74, n. 1, 2023.
- [86] X. Chen, X. Zhang, M. Dong, L. Huang e Y. G. a. S. He, «Deep Learning-Based Prediction of Wind Power for Multi-turbines in a Wind Farm,» *Frontiers in Energy Research*, 29 July 2021.
- [87] A. Alkesaiberi , F. Harrou e a. Y. Sun, «Efficient Wind Power Prediction Using Machine Learning Methods: A Comparative Study,» *Energies* 2022, vol. 15, n. 2327, 23 March 2022.
- [88] U. Singh, M. Rizwan e M. A. I. Alsaidan, «A Machine Learning-Based Gradient Boosting

- Regression Approach for Wind Power Production Forecasting: A Step towards Smart Grid Environments,» *Energies* 2021, vol. 14, n. 5196, 23 August 2021.
- [89] Z. Mushtaq, B. Kaur e A. Y. Malik, «Machine Learning for Wind Energy Analysis and Forecasting,» *International Journal for Research in Applied Science & Engineering Technology (IJRASET)*, 2020.
- [90] Ö. E. Yürek , D. Birant e İ. Yürek, «Wind Power Generation Prediction Using Machine Learning Algorithms,» *Deu Muhendislik Fakultesi Fen ve Muhendislik*, vol. 23(67), pp. 107-119.
- [91] A. L. S. arez-Cetrulo , L. Burnham-King, D. Haughton e R. S. Carbajo, «Wind power forecasting using ensemble learning for day-ahead energy trading,» *Renewable Energy* 191, n. 685e698, 9 April 2022.
- [92] C. G. a. H. Li, «Review on Deep Learning Research and Applications in Wind and Wave Energy,» *Energies* 2022, vol. 15, n. 1510, 17 February 2022.
- [93] G. Qin, Q. Yan, J. Zhu, C. Xu e a. D. M. Kammen, «Day-Ahead Wind Power Forecasting Based on Wind Load Data Using Hybrid Optimization Algorithm,» *Sustainability* 2021, vol. 13, n. 1164, 22 January 2021.
- [94] W. Khan, S. Walker e W. Zeiler, «Improved solar photovoltaic energy generation forecast using deep learning-based ensemble stacking approach,» *Energy* 2022, vol. 240, n. 122812, 4 December 2021.
- [95] B. Zazoum, «Solar photovoltaic power prediction using different machine learning methods,» *Energy Reports*, n. 8, pp. 19-25, 26 November 2021.
- [96] M. A. a. I. Ahmad, «Solar power generation forecasting using ensemble approach based on deep learning and statistical methods,» *Applied Computing and Informatics ahead-of-print*, 2020.
- [97] W.-C. Kuo , C.-H. Chen e S.-H. H. a. C.-C. Wang, «Assessment of Different Deep Learning Methods of Power Generation Forecasting for Solar PV System,» *Appl. Sci.* 2022, vol. 12, n. 7529, 27 July 2022.
- [98] K. Anuradha, D. Erlapally, V. S. G. Karuna e K. Adilakshmi, «Analysis Of Solar Power Generation Forecasting Using Machine Learning Techniques,» in *E3S Web of Conferences* 309, 01163, 2021.
- [99] Y. Li , L. Zhou, P. Gao, B. Yang e Y. H. C. Lian, «Short-Term Power Generation Forecasting of a Photovoltaic Plant Based on PSO-BP and GA-BP Neural Networks,» *Frontiers in Energy Research*, 19 January 2022.
- [100] Y. Essam, A. N. Ahmed, R. Ramli, K.-W. Chau, M. S. I. Ibrahim, M. Sherif e A. S. &. Ahmed El-Shafie, «Investigating photovoltaic solar power output forecasting using machine learning algorithms,» *ENGINEERING APPLICATIONS OF COMPUTATIONAL FLUID MECHANICS*, vol. 16, n. 1, pp. 2002-2034, 03 Oct 2022.
- [101] F. Harrou e F. K. a. Y. Sun, «Forecasting of Photovoltaic Solar Power Production Using LSTM Approach,» in *Advanced Statistical Modeling, Forecasting, and Fault Detection in Renewable Energy Systems*, Fouzi Harrou and Ying Sun, 17/02/2023.
- [102] R. Siddiqui, H. Anwar, F. Ullah, R. Ullah, M. A. Rehman, N. Jan e a. F. Zaman, «Power

- Prediction of Combined Cycle Power Plant (CCPP) Using Machine Learning Algorithm-Based Paradigm,» *Wireless Communications and Mobile Computing*, vol. 2021, n. 9966395, 23 December 2021.
- [103] H. Sharma , L. Marinovici, V. Adetola e H. T. Schaef, «Data-driven modeling of power generation for a coal power plant under cycling,» *Energy and AI* 11, n. 100214, 13 November 2022.
- [104] . Y. Rodríguez, A. A. Gamboa, E. C. Rodríguez e A. F. d. Silva, «Comparison of Adaptive Neuro-Fuzzy Inference System (ANFIS) and Machine Learning Algorithms for Electricity Production Forecasting,» *IEEE Latin America Transactions*, vol. 20, n. 10, pp. 2288 - 2294, 09 September 2022.
- [105] E. A. Elfaki e A. H. Ahmed, «Prediction of Electrical Output Power of Combined Cycle Power Plant Using Regression ANN Model,» *Journal of Power and Energy Engineering*, vol. 6, pp. 17-38, December 18, 2018.
- [106] R. B. Chokshi, N. K. Chavda e D. A. D. Patel, «Prediction of Performance of Coal-Based KWU Designed Thermal Power Plants using an Artificial Neural Network,» *International Journal of Applied Engineering Research*, vol. 13, n. 5, pp. 3093-3110, 2018.
- [107] L. Guo, J. Chen e F. W. a. M. Wang, «An electric power generation forecasting method using support vector machine,» *Systems Science & Control Engineering*, vol. 6, n. 3, pp. 191-199, 16 Nov 2018.
- [108] W. M. Ashraf, G. M. Uddin , M. F. 2, F. Riaz , H. A. Ahmad, A. H. Kamal, S. Anwar , A. M. El-Sherbeeney , M. H. Khan , N. Hafeez , A. Ali , A. Samee e M. A. N. al., «Construction of Operational Data-Driven Power Curve of a Generator by Industry 4.0 Data Analytics,» *Energies* 2021, vol. 14, n. 1227, 24 February 2021.
- [109] Z. Liu e I. A. Karimi, «Gas turbine performance prediction via machine learning,» *Energy*, vol. 192, n. 116627, 1 February 2020.
- [110] A. Jędrzejewski, J. Lago, G. Marcjasz e R. Weron, «Electricity Price Forecasting: The Dawn of Machine Learning,» *IEEE Power & Energy Magazine*, May/June 2022.
- [111] V. Sridharan, M. Tuo e X. Li, «Wholesale Electricity Price Forecasting Using Integrated,» *Energies MDPI*, 14 October 2022.
- [112] R. Das, R. Bo, W. U. Rehman, H. Chen e D. Wunsch, «Cross-Market Price Difference Forecast Using Deep Learning for Electricity Markets,» The Hague, Netherlands, 26-28 October 2020.
- [113] M. M. Alshater, I. Kampouris, H. Marashdeh e O. F. Atayah, «Early warning system to predict energy prices: the role of artificial intelligence and machine learning,» *Annals of Operations Research*, 2 August 2022.
- [114] R. Syah, M. Rezaei, M. Elveny e M. M. Nezhad, «Day-ahead electricity price Day-ahead electricity price LSSVM-based self adaptive fuzzy kernel and modified HBMO algorithm,» *Nature scientific reports*, 2021.
- [115] S. S. & A. Wagner, «Electricity Price Forecasting with Neural Networks on EPEX Order Books,» *Applied Mathematical Finance*, 24 Aug 2020.
- [116] A. Arifoğlu e T. Kandemir, «Electricity price forecasting in turkish day-ahead market via deep learning techniques,» *Mehmet Akif Ersoy University - Journal of Economics and*

Administrative Sciences Faculty, vol. 8, n. 3, July 2022.

- [117] A. Marszałek e T. Burczyński, «Forecasting day-ahead spot electricity prices using deep neural networks with attention mechanism,» *Journal of Smart Environments and Green Computing*, vol. 1, pp. 21-31, 30 Mar 2021.
- [118] M. Castelli, A. Groznik e A. Popovič, «Forecasting Electricity Prices: a Machine Learning Approach,» *Algorithms MDPI*, 8 May 2020.
- [119] A. Lucas, K. Pegios, E. Kotsakis e D. Clarke, «Price Forecasting for the Balancing Energy Market Using Machine-Learning Regression,» *Energies MDPI*, 16 October 2020.
- [120] D. Zaroni, A. Piazzzi, T. Tettamanti e Á. Sleisz, «Investigation of Day-ahead Price Forecasting Models in the Finnish Electricity Market,» in *ICAART 2020 - 12th International Conference on Agents and Artificial Intelligence*, January 2020.