



POLITECNICO DI TORINO

**Corso di Laurea Magistrale
in Ingegneria Gestionale**

A.A. 2024/2025

Digitalizzazione della Qualità

Sviluppo dell'I.A. come strumento a supporto del
miglioramento del processo di gestione dei reclami

Relatore:
Prof. Maurizio Galetto

Candidata:
Salvina Puliafito
S317982

Sommario

Sommario	2
Indice Figure.....	4
Indice Tabelle	4
Indice Figure Appendice	5
Indice Tabelle Appendice	5
Abstract	6
Introduzione.....	7
1. Panoramica sull'intelligenza artificiale.....	7
2. Profilo storico e operativo dell'azienda.....	9
Capitolo 1	11
Qualità e Performance Aziendale: Strategie e Misure.....	11
1. Concetto di qualità	11
1.1. Fasi e approccio alle azioni della qualità.....	13
2. KPIs e analisi dei dati.....	14
3. Miglioramento continuo e principali metodologie di problem solving	20
3.1. Metodo 8D	21
3.2. Analisi di Pareto	24
3.3. Ishikawa	26
3.4. 5 Why	28
3.5. 5W2H	29
Capitolo 2	31
Digitalizzazione del processo: l'intervento dell'Intelligenza Artificiale.....	31
1. Panoramica del problema e della soluzione	31
1.1. Raccolta dati tramite questionario	32
2. Use Case 1: Intelligent Database Enrichment.....	35
2.1. Obiettivo	35
2.2. Elementi del Database Intelligente	35
Capitolo 3	46
Verifica del Modello e implementazione della soluzione	46
1. Fase di verifica del modello	46
2. Fase post-processing.....	49
2.1. Assegnazione nomi ai cluster.....	49
2.2. Aggregazione cluster.....	49

3. Passato e Futuro: Punti di Forza e di Debolezza	50
3.1. Come si articola la nuova procedura	52
Capitolo 4	57
Applicazione su un caso reale: Pompa Urea	57
1. Obiettivo	57
2. Descrizione del componente	57
3. Valutazione del livello di affidabilità	58
4. Analisi economica	62
4.1. Definizione di “spending”	62
4.2. Recupero economico	64
Capitolo 5	66
Valutazione Economica del Progetto: Costi e Ricavi	66
1. Struttura dell’analisi	66
2. Costi e tempi	67
3. Recupero economico e benefici.....	69
Capitolo 6	71
Sviluppi Futuri - Use case 2: Self-service BI Platform.....	71
1. Obiettivo	71
2. Applicazione del modello – Qlik Sense	71
3. Come funziona la soluzione	72
Conclusioni.....	74
Appendice A.....	76
Questionario.....	76
Appendice B	78
Fase di controllo dei cluster	78
Appendice C	80
Grafici proposti.....	80
Appendice D.....	81
Calcolo livello di affidabilità per ogni informazione.....	81
Appendice E	84
Esempi analisi Qlik Sense	84
Bibliografia e Sitografia	87

Indice Figure

Figura 1 – Qualità durante il ciclo di vita del prodotto	14
Figura 2 – Grafico MERF	17
Figura 3 – Grafico CORF	19
Figura 4 – Step Metodo 8D	24
Figura 5 – Grafico di Pareto.....	25
Figura 6 – Diagramma di Ishikawa	27
Figura 7 – Distribuzione Informazioni estratte	34
Figura 8 - Distribuzione tipologie di difficoltà riscontrate	34
Figura 9 – Funzionamento algoritmi di Embeddings	40
Figura 10 – Rappresentazione Grafica PCA.....	42
Figura 11 – Esempio di Clustering con valore di Silhouette pari a 0.84.....	45
Figura 12 – Colonne aggiuntive.....	52
Figura 13 – Esempio grafici automatici	55
Figura 14 – Esempio riassunto automatico.....	56
Figura 15 – Configurazione impianto ATS	58
Figura 16 – Customer Complaint raggruppato per Issue Cause	60
Figura 17 – Performed Test raggruppato per Issue Cause	61

Indice Tabelle

Tabella 1 – Tier 4B Muffler	47
Tabella 2 – ERG Valve F5	47
Tabella 3 – Livello di affidabilità del modello.....	50
Tabella 4 – Tempi di lettura claim	50
Tabella 5 – Incidenza 'Others'	59
Tabella 6 – Livello di Affidabilità.....	60
Tabella 7 – Risultati economici.....	64
Tabella 8 – Tempi approccio on demand	67
Tabella 9 – Tempi approccio storico.....	67
Tabella 10 – Costi approccio on demand	68
Tabella 11 – Costi approccio storico	68
Tabella 12 – Totale costi	69
Tabella 13 – Recupero %	70

Indice Figure Appendice

Figura A. 1 – Distribuzione Customer Complaint.....	80
Figura A. 2 – Distribuzione Issue Cause	80
Figura A. 3 – Distribuzione Performed Test	80
Figura A. 4 – Combinazione Failure Code e Componenti	84
Figura A. 5 – Combinazione Worked Hours e Guasti	84
Figura A. 6 – Combinazione Worked Hours e Production Year.....	85
Figura A. 7 – Distribuzione Claims nel mondo	85
Figura A. 8 – Confronto tra Worked Hours e Labour Hour per Production Date.....	86
Figura A. 9 – Distribuzione Worked Hours per Production Year	86

Indice Tabelle Appendice

Tabella A. 1 – % Errore HDBScan (Tier 4B Muffler).....	78
Tabella A. 2 – % Errore Spectral (Tier 4B Muffler)	78
Tabella A. 3 – % Errore HDBScan (EGR Valve F5)	79
Tabella A. 4 – % Errore Spectral (EGR Valve F5)	79
Tabella A. 5 – Customer Complaint.....	81
Tabella A. 6 – Issue Cause	82
Tabella A. 7 – Performed Test	83

Abstract

L'Intelligenza Artificiale Generative (GenAI) offre nuove opportunità per ottimizzare l'analisi e la gestione dei dati e, nel contesto del Warranty Reporting System (WRS), può significativamente migliorare l'efficienza operativa e la qualità delle informazioni estratte. La proposta prevede lo sviluppo di una soluzione attraverso due use case principali, che si focalizzano sull'implementazione di soluzioni innovative per affrontare in modo più accurato sfide specifiche del sistema.

Il primo use case si concentra sull'analisi automatizzata dei commenti testuali inseriti dai concessionari nel database WRS. Attraverso modelli di IA generativa, il sistema sarà in grado di comprendere il contesto, estrarre parole chiave e classificare le informazioni in cluster predefiniti. Questo approccio consentirà di ottenere una categorizzazione più dettagliata delle richieste, velocizzando il processo di analisi rispetto alla procedura attuale. L'obiettivo è raggiungere un'accuratezza di almeno l'80%, con un modello adattabile ai cluster definiti tramite prompt engineering.

Il secondo use case mira a ridurre il tempo necessario per navigare tra i dati, generare grafici e produrre report basati su WRS e altri database correlati. L'utente potrà interagire con il sistema tramite linguaggio naturale, sfruttando una comprensione avanzata dei dati e della tecnologia per ottenere risposte mirate. Poiché gli strumenti di Business Intelligence attuali sono principalmente rivolti ai data analyst, il sistema dovrà garantire un controllo tecnico accurato sull'output generato, che sarà sottoposto a revisione sia tecnica che aziendale.

L'integrazione di questi due approcci consente di migliorare l'efficienza nell'analisi dei dati e nella gestione delle informazioni, offrendo un sistema più rapido e intuitivo per il supporto decisionale.

Introduzione

1. Panoramica sull'intelligenza artificiale

L'intelligenza artificiale (IA) è una disciplina moderna che ha subito una rapida evoluzione, trasformandosi da un ambito di ricerca specialistica a una tecnologia essenziale per la vita quotidiana e aziendale. Sebbene le sue radici affondino nel passato, la nascita ufficiale della disciplina risale al 1956, durante il seminario estivo di Dartmouth. In quell'occasione, furono raccolti i contributi scientifici sviluppati negli anni precedenti e si gettarono le basi per lo sviluppo futuro dell'IA.

Da allora, l'Intelligenza Artificiale ha subito l'influenza di diverse altre discipline, tra cui la filosofia, la matematica e l'economia, contribuendo significativamente all'evoluzione dell'informatica. Essa si occupa di studiare i fondamenti teorici e le tecniche per progettare sistemi hardware e software capaci di simulare prestazioni tipiche dell'intelligenza umana. Questi sistemi, quando osservati dall'esterno, sembrano in grado di prendere decisioni e risolvere problemi complessi, compiti che un tempo si ritenevano esclusivamente umani.

Negli ultimi dieci anni, l'evoluzione dell'intelligenza artificiale ha subito una notevole accelerazione. Oggi, non solo essa è diventata parte integrante della vita quotidiana, ma è anche un elemento essenziale per molteplici settori economici e ambiti di policy. La sua applicazione non è più limitata a contesti teorici o di nicchia: ha ormai un impatto significativo su vasta scala, sia nel mondo del lavoro che in quello personale, e le previsioni indicano che il suo ruolo continuerà a crescere.¹

Uno degli ambiti in cui l'intelligenza artificiale ha avuto un impatto particolarmente rilevante è quello aziendale. L'IA, infatti, non è vista solo come una tecnologia che sostituisce l'intervento umano, ma piuttosto come un supporto in grado di potenziare l'intelligenza e le capacità decisionali delle persone. In azienda, l'intelligenza artificiale sta contribuendo sempre di più a migliorare l'efficienza e la produttività, consentendo ai dipendenti di concentrarsi su attività che apportano un maggior valore aggiunto

¹ (Somalvico, 1987); (Willcocks, 2015)

all'azienda, mentre l'IA si occupa di compiti ripetitivi, complessi o che richiedono una rapida elaborazione dei dati.

Un esempio di questa sinergia è rappresentato dall'automazione dei processi aziendali (RPA - Robotic Process Automation), dove software basati su IA gestiscono attività standardizzate e di routine, come la gestione di documenti, l'inserimento dati o il controllo delle transazioni finanziarie. In questo modo, l'IA non sostituisce il lavoratore, ma lo libera da compiti meccanici, permettendogli di dedicarsi a mansioni più strategiche e creative, che richiedono un intervento umano critico. Oltre a ciò, l'intelligenza artificiale è utilizzata per analizzare grandi quantità di dati e fornire insight utili alle decisioni aziendali. I sistemi di IA sono in grado di elaborare dati complessi e individuare pattern o tendenze che possono sfuggire all'osservazione umana.

In definitiva, l'intelligenza artificiale in azienda non si limita a ridurre i costi operativi o a velocizzare i processi, ma diventa un vero e proprio partner nell'ottimizzazione delle performance e nella creazione di valore. L'obiettivo non è sostituire l'intelligenza umana, ma potenziarla: le persone e le macchine collaborano, creando un ecosistema in cui l'IA gestisce le attività che richiedono precisione, velocità e capacità di calcolo, mentre gli esseri umani si concentrano su aspetti creativi, strategici e relazionali.

Secondo la definizione fornita da Haenlein e Kaplan (2019), l'IA è la capacità di un sistema di interpretare correttamente i dati esterni, apprendere da essi e utilizzare tali conoscenze per raggiungere obiettivi specifici attraverso un adattamento flessibile, senza l'intervento diretto dell'essere umano. Questo rende l'intelligenza artificiale uno strumento formidabile per aumentare la capacità delle imprese di adattarsi alle mutevoli condizioni del mercato e di sviluppare strategie innovative.²

Accennare all'evoluzione dell'intelligenza artificiale è fondamentale per comprendere il ruolo che questa ha avuto nell'automazione di alcuni processi operativi all'interno di aziende aperte all'innovazione e al cambiamento e che sono state in grado di adattarsi rapidamente alle dinamiche di mercato. Tra queste rientra senz'altro FPT Industrial che ha scelto di impiegare l'intelligenza artificiale per ottimizzare il processo di gestione dei

² (Somalvico, 1987); (Willcocks, 2015)

reclami, dimostrando una spiccata capacità di adattamento alle trasformazioni del mercato. Con una visione orientata all'innovazione, l'azienda ha saputo affrontare le trasformazioni del mercato, utilizzando l'IA per rispondere a esigenze di modernizzazione e per risolvere problemi concreti. Questa strategia sottolinea l'impegno nel mantenere alta la competitività attraverso soluzioni tecnologiche avanzate.

2. Profilo storico e operativo dell'azienda

FPT Industrial nasce nel 2005 come divisione del gruppo Fiat, che già nel 1975 aveva fondato IVECO attraverso la fusione di tre marchi di camion all'interno del gruppo Fiat. FPT venne quindi istituita per consolidare tutte le attività di Fiat legate ai gruppi motopropulsori per applicazioni automobilistiche e industriali. Nel 2011, Fiat riorganizzò le proprie attività non automobilistiche, integrando IVECO, FPT Industrial e CNH Global N.V. Infine, nel 2022, è stato formalizzato il nuovo Iveco Group, a seguito della scissione da CNH Industrial.³

FPT Industrial, appartenente a Iveco Group N.V., è un'azienda di punta nella progettazione, produzione e vendita di soluzioni avanzate per la propulsione industriale, coprendo settori che spaziano dai veicoli on-road e off-road alle applicazioni marine e alla generazione di energia. Con una presenza consolidata in oltre cento paesi, FPT Industrial si distingue come leader globale nel settore, sostenuta da un solido impegno nell'innovazione e nella sostenibilità.

L'azienda, infatti, si è mossa con decisione verso una mobilità a zero emissioni di carbonio, guidando la transizione grazie alla divisione ePowertrain, dedicata esclusivamente allo sviluppo di sistemi di propulsione elettrica. Qui, vengono progettati e prodotti motori elettrici, pacchi batterie e avanzati sistemi di gestione delle batterie per rispondere alle esigenze del mercato in termini di efficienza e riduzione dell'impatto ambientale. Questo impegno verso la sostenibilità si affianca a una costante attenzione alla qualità e all'efficienza, caratteristiche che hanno permesso a FPT Industrial di distinguersi per affidabilità e innovazione nel settore delle soluzioni per la propulsione.

³ (https://www.ivecogroup.com/group/about_us/our_history, s.d.)

Grazie a un'intensa attività di ricerca e sviluppo, FPT Industrial è oggi uno dei quattro maggiori produttori di motori diesel al mondo, con una gamma di cilindrata che spazia da 2,3 a 20 litri e potenze da 42 a oltre 1.000 CV. I prodotti dell'azienda supportano anche carburanti alternativi come metano e biodiesel, anticipando e rispettando normative antinquinamento, tra cui Euro VI e Tier 4B. Le famiglie di motori comprendono soluzioni pensate per rispondere a esigenze diverse: il Vector, V8 compatto per applicazioni agricole; il Cursor, motore a 6 cilindri destinato a veicoli commerciali pesanti e disponibile anche in versione a gas naturale; il versatile NEF, utilizzato in settori che spaziano dall'agricoltura all'edilizia; il compatto F5, ideale per mezzi leggeri; e l'F28, riconosciuto per le sue dimensioni ridotte e l'innovazione tecnologica, premiato come Engine of the Year.

Oltre ai motori, FPT Industrial produce trasmissioni e sistemi di alimentazione avanzati, come il common rail e la sovralimentazione variabile, per offrire prestazioni elevate, efficienza nei consumi e riduzione delle emissioni. Questi elementi, associati a una gestione innovativa dei gas di scarico e all'utilizzo di tecnologie di trattamento avanzate, consentono di rispondere a criteri ambientali rigorosi, migliorando al contempo la durabilità e riducendo gli intervalli di manutenzione.

Nel corso degli anni, l'azienda ha costruito la propria reputazione grazie a una costante ricerca dell'eccellenza e ad un approccio che valorizza affidabilità, innovazione e qualità dei dettagli. L'attenzione alla qualità è un valore centrale per FPT Industrial, che ne assicura la conformità attraverso rigorosi controlli. Questi aspetti rappresentano un pilastro fondamentale per mantenere elevati gli standard aziendali e rispondere efficacemente alle aspettative del mercato.⁴

Per questi motivi, nel capitolo successivo, viene approfondito il concetto di qualità e le pratiche di controllo qualità adottate per garantire l'eccellenza del prodotto e il rispetto delle normative in continuo aggiornamento, oltre che diverse metodologie di problem solving e indicatori chiave per la misura delle prestazioni.

⁴ (https://www.fptindustrial.com/en/?sc_lang=it, s.d.)

Capitolo 1

Qualità e Performance Aziendale: Strategie e Misure

1. Concetto di qualità

La qualità rappresenta un pilastro fondamentale nella produzione di beni e servizi, con il cliente che assume un ruolo centrale nell'intero processo. L'obiettivo delle aziende non si limita a soddisfare le aspettative del consumatore, ma punta a superarle, offrendo prodotti e servizi che garantiscano affidabilità, efficienza e valore aggiunto. Un prodotto è considerato di qualità quando riesce a rispondere in modo ottimale alle esigenze del cliente, assicurando prestazioni elevate e continuità nel tempo.

Prima di arrivare al consumatore finale, un prodotto attraversa diverse fasi all'interno del ciclo produttivo, tra cui progettazione, realizzazione e distribuzione. Ogni fase ricopre un duplice ruolo: è sia fornitore della fase successiva sia cliente di quella precedente. Questa interdipendenza implica che ogni operatore aziendale debba considerarsi un cliente interno, responsabile della qualità del proprio contributo. Eventuali errori o imperfezioni in una singola fase del processo possono compromettere l'intero risultato finale, rendendo necessario un controllo costante e rigoroso in ogni passaggio.

Il concetto di cliente interno sottolinea l'importanza di ogni singola funzione e operatore nel determinare il valore complessivo di un prodotto o servizio. La qualità non è un elemento isolato, ma un principio che guida l'intera organizzazione, influenzando ogni aspetto della gestione aziendale. Con l'affermarsi della qualità come elemento strategico, si è sviluppato il concetto di Total Quality Management (TQM), un approccio che pone al centro le persone e punta al miglioramento continuo della soddisfazione del cliente, bilanciando al contempo i costi operativi.

Il TQM si configura come un sistema integrato che coinvolge orizzontalmente tutte le funzioni e i reparti aziendali, estendendosi anche alla rete di fornitori e clienti. Basandosi su principi di natura filosofica e su metodologie scientifiche, utilizza strumenti e tecniche di monitoraggio per garantire standard elevati. Tra i suoi elementi distintivi emergono il

valore del lavoro di squadra e il rispetto della dignità di ogni individuo all'interno dell'organizzazione.

L'importanza di un approccio globale alla qualità è stata evidenziata già negli anni '50 da Armand Feigenbaum, che ha introdotto il concetto di Total Quality Control (TQC). Secondo Feigenbaum, la qualità sarebbe diventata un fattore chiave nelle decisioni di acquisto, influenzando sia i consumatori privati sia le aziende e le istituzioni pubbliche. Egli definì la qualità come il risultato dell'integrazione di vari aspetti, tra cui marketing, progettazione, produzione e assistenza, tutti elementi necessari per rispondere pienamente alle esigenze del cliente. In questa visione, la gestione della qualità deve essere considerata un elemento centrale nella strategia aziendale, al pari del marketing e della gestione finanziaria. Il TQC, come concepito da Feigenbaum, rappresenta un sistema che coordina in modo efficace gli sforzi volti a sviluppare, mantenere e migliorare la qualità, con l'obiettivo di massimizzare la soddisfazione del cliente e ottimizzare l'efficienza economica dell'azienda.

In un contesto competitivo come quello attuale, le aziende non si limitano più a offrire prodotti tangibili, ma puntano a fornire un'esperienza complessiva che deriva dall'interazione tra processi, persone e tecnologie. Come affermava W. E. Deming, "la qualità è responsabilità di tutti", sottolineando che ogni singolo lavoratore, indipendentemente dal ruolo ricoperto, contribuisce al raggiungimento degli standard di eccellenza aziendale.

Il concetto di qualità, quindi, va ben oltre la semplice conformità a specifiche tecniche: rappresenta la capacità di un prodotto, un servizio o un processo di garantire risultati concreti in linea con requisiti espliciti e aspettative implicite. Questo approccio comprende sia obiettivi generali, come il miglioramento della reputazione aziendale, sia traguardi specifici, come la riduzione dei difetti e l'ottimizzazione dell'efficienza operativa.

In definitiva, la qualità si configura come un valore strategico che collega le esigenze dei clienti alla performance organizzativa, diventando un elemento chiave per la creazione di un vantaggio competitivo sostenibile nel lungo periodo.⁵

1.1. Fasi e approccio alle azioni della qualità

Durante il ciclo di vita del prodotto, il ruolo della squadra che si occupa della qualità si estende oltre il semplice monitoraggio delle prestazioni: essa agisce come un partner strategico sia per i clienti sia per i dipartimenti interni dell'azienda. L'obiettivo principale è garantire la conformità del prodotto agli standard di qualità e alle aspettative dei clienti, utilizzando un approccio che non si limiti a correggere i problemi, ma che li prevenga e li anticipi.

Per ottenere ciò, le azioni della qualità possono essere classificate in tre approcci distinti, ciascuno con finalità e metodi operativi specifici:

1. Approccio reattivo: questo approccio si concentra sulla gestione di problemi che si sono già presentati. L'obiettivo principale è comprendere le cause dei difetti, applicare azioni correttive efficaci e garantire che lo stesso problema non si ripresenti sul medesimo prodotto. È una strategia fondamentale per mantenere la fiducia dei clienti, ridurre i costi di non conformità e migliorare la gestione dei rischi associati ai prodotti già immessi sul mercato;
2. Approccio preventivo: l'approccio preventivo si sviluppa come una naturale estensione di quello reattivo, focalizzandosi sull'eliminazione della possibilità che difetti simili emergano su altri prodotti o processi aziendali. In questa fase, l'attenzione si sposta dalla risoluzione dei problemi già noti alla loro prevenzione in contesti differenti. Questa strategia si avvale dell'analisi delle cause profonde dei difetti (Root Cause Analysis) per progettare e implementare sistemi e processi più solidi, diminuendo così la probabilità di problematiche future;
3. Approccio proattivo: questo approccio, il più avanzato e strategico, mira a identificare potenziali difetti o problematiche teoriche prima ancora che si

⁵ (FPT Industrial, 2023)

verifichino, intervenendo già nelle fasi iniziali di progettazione, sviluppo o produzione. Le tecniche utilizzate includono strumenti di analisi predittiva, simulazioni e metodologie come il Failure Mode and Effect Analysis (FMEA). L'approccio proattivo consente di minimizzare i rischi, ottimizzare i costi di sviluppo e produzione e creare prodotti che rispondano meglio alle esigenze dei clienti, rafforzando il vantaggio competitivo dell'azienda.⁶

Il passaggio dall'approccio reattivo al proattivo riflette una crescente maturità nella gestione della qualità, dove l'obiettivo non è solo risolvere i problemi, ma anche prevenirli e anticiparli. Questa evoluzione consente all'azienda di migliorare continuamente i propri prodotti e processi, rafforzando il valore percepito dai clienti e garantendo una posizione di leadership sul mercato.

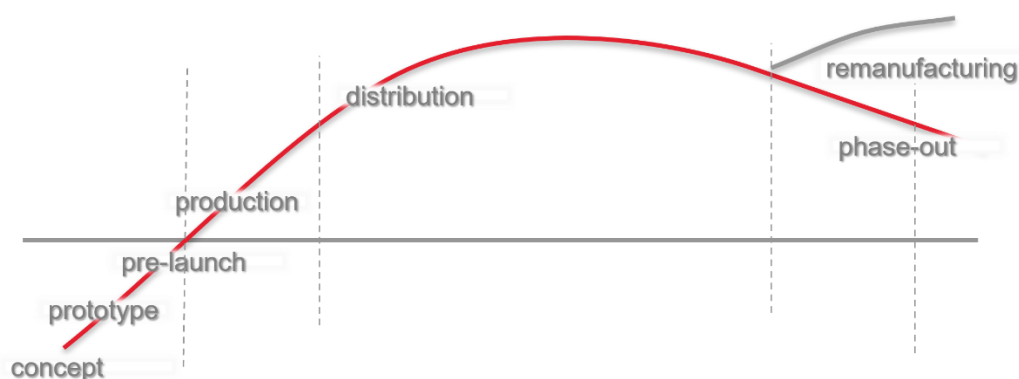


Figura 1 – Qualità durante il ciclo di vita del prodotto

2. KPIs e analisi dei dati

Nell'ambito della gestione della qualità, l'adozione di Indicatori Chiave di Prestazione (KPI) è essenziale per garantire che le operazioni aziendali siano monitorate con precisione e continuamente migliorate. I KPI forniscono informazioni che aiutano le organizzazioni a identificare rapidamente aree critiche, ottimizzare i processi e migliorare l'efficienza complessiva. Questi indicatori sono strumenti di misura che permettono di analizzare performance e risultati, con l'obiettivo di prendere decisioni mirate e

⁶ (FPT Industrial, 2023)

tempestive. Di seguito sono riportati alcuni dei KPI fondamentali utilizzati per monitorare la qualità e le prestazioni aziendali:

1. **Failures per 100 machines (F/100)** → Questo indicatore misura il numero di guasti o riparazioni ogni 100 macchine durante un determinato periodo di tempo. L'F/100 è utilizzato per valutare l'affidabilità del prodotto e l'esperienza del cliente, poiché un numero elevato di guasti può essere sintomo di problemi legati alla qualità. Inoltre, è uno strumento utile per identificare le aree critiche durante il periodo di garanzia, come la qualità iniziale del prodotto, le condizioni d'uso sul campo o l'approccio post-vendita. L'indicatore aiuta anche a individuare problemi ricorrenti e a prendere provvedimenti correttivi per evitare che si ripetano. Un esempio di calcolo: $8 \text{ guasti} / 75 \text{ macchine} \times 100 = 10,7 \text{ F/100}$;
2. **Average Cost per Unit (€) – ACPU** → L'ACPU misura la spesa media sostenuta per garanzia per ogni unità venduta in un periodo specifico. Questa metrica, che viene calcolata convertendo la valuta locale in euro tramite un tasso di cambio standard, è utile per monitorare quanto l'azienda sta spendendo per mantenere i prodotti entro gli standard di qualità durante il periodo di garanzia. L'ACPU consente di monitorare il trend dei costi di garanzia e di intervenire prontamente per ridurre le spese, migliorando la qualità del prodotto e ottimizzando i processi aziendali. Esempio di calcolo: $\text{€ } 277.560 \text{ (spesa totale)} / 401 \text{ macchine} = \text{€ } 692 \text{ ACPU}$;
3. **Time to Fix (TTF)** → Questo KPI misura il tempo necessario per risolvere un problema, calcolato dal momento in cui la squadra dedicata viene attivata per affrontare il problema fino a quando la soluzione è implementata sul campo. Il Time to Fix è fondamentale per monitorare l'efficienza del servizio post-vendita e l'agilità con cui l'azienda risponde alle problematiche dei clienti. Un tempo di risoluzione più breve non solo migliora la soddisfazione del cliente, ma riduce anche i costi complessivi e migliora la performance operativa;
4. **Accrual Rate (Tasso di Accantonamento)** → Il tasso di accantonamento rappresenta la parte delle entrate che viene destinata alla riserva per le garanzie. Questa riserva è utilizzata per coprire i costi relativi ai futuri interventi di riparazione o sostituzione durante il periodo di garanzia. Ogni volta che viene

venduto un prodotto, viene accantonato un determinato importo, che fornisce una stima dei costi di garanzia futuri. Il monitoraggio di questo indicatore consente all'azienda di pianificare finanziariamente i costi legati alla garanzia, assicurando una gestione più efficiente delle risorse;

5. **Claims Rate (Tasso di Reclami)** → Il tasso di reclami indica la spesa effettiva in denaro che l'azienda sostiene per rimborsare o coprire le riparazioni effettuate in garanzia durante un periodo specifico. Ogni volta che una riparazione viene eseguita, l'importo del rimborso viene prelevato dalla riserva per le garanzie, ed è importante monitorare il Claims Rate per valutare l'impatto finanziario delle garanzie sull'azienda. Un alto tasso di reclami può segnalare problemi di qualità a livello di prodotto, mentre un valore basso dell'indice può essere sintomo di una gestione efficace della qualità e di un buon servizio post-vendita.⁷

L'utilizzo di KPI permette alle aziende di avere una visione chiara e dettagliata delle proprie operazioni, facilitando l'individuazione delle aree di miglioramento, ottimizzando i costi e garantendo al cliente il soddisfacimento delle proprie esigenze. La capacità di misurare e analizzare costantemente questi indicatori è fondamentale per il miglioramento continuo dei processi aziendali e per mantenere elevati standard di qualità.⁸

Esistono due modalità di analisi differenti dei KPI precedentemente citati, che offrono due prospettive distinte ma complementari:

1. MERF (Manufacturing & Engineering Repair Frequency)

Questa modalità di analisi opera da un punto di vista interno all'azienda, mirando a monitorare e migliorare la qualità del prodotto attraverso una valutazione accurata delle riparazioni. Il MERF si basa sul raggruppamento dei reclami in garanzia secondo periodi di produzione specifici, fornendo così una visione chiara del rendimento dei prodotti in relazione al tempo di servizio e all'usura.

⁷ (FPT Industrial, 2023)

⁸ (<https://fractionalmanageritalia.it/operation/kpi-operations-sai-quali-dovresti-monitorare/>, 2024)

Il punto di partenza per garantire la qualità del prodotto è avere una misura chiara e trasparente delle performance, permettendo al team di monitorare costantemente i risultati e stabilire obiettivi misurabili da raggiungere. Il MERF viene calcolato in base a lotti di veicoli prodotti nello stesso periodo e con un grado di usura simile, così da ottenere dati omogenei e rappresentativi.

Questo indicatore non solo fornisce una valutazione sulla performance dei modelli attuali, ma consente anche di monitorare l'efficacia dei miglioramenti introdotti, come ad esempio quelli relativi al processo di produzione o alle modifiche progettuali di un nuovo modello. In altre parole, il MERF aiuta a comprendere se le modifiche o innovazioni implementate stanno effettivamente migliorando il prodotto o se è necessario intervenire ulteriormente.

Ogni report MERF include due grafici che mostrano l'andamento di F/100 (Failures per 100 machines) e ACPU (Average Cost per Unit), dove ogni linea del grafico rappresenta uno specifico MIS (Months in Service). Viene mostrata l'evoluzione dei parametri mese per mese. In particolare, in figura 2, i KPI sono monitorati ad intervalli di quattro mesi. Per ogni linea di MIS viene effettuata un'analisi di Pareto, che consente di individuare le aree più critiche e quelle che necessitano di un intervento prioritario, ottimizzando così le risorse e migliorando la gestione della qualità.

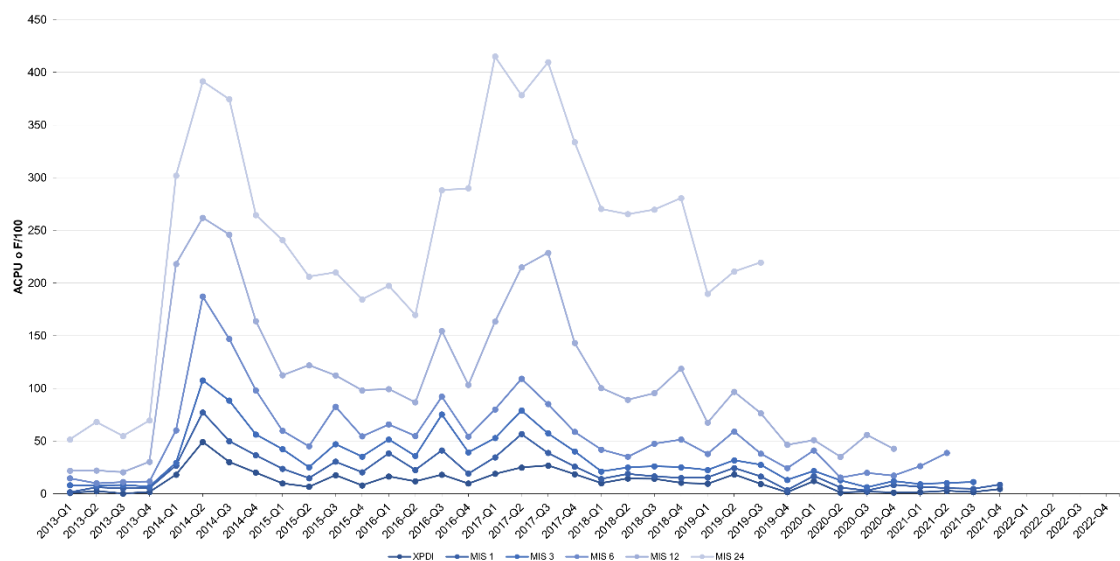


Figura 2 – Grafico MERF

2. CORF (Customer Observed Repair Frequency)

Il CORF rappresenta un'analisi focalizzata sull'esperienza del cliente, fornendo una visione diretta di come i prodotti si comportano nel tempo sul campo. In questo caso, l'analisi si concentra sui reclami dei clienti che si sono verificati negli ultimi 12 mesi, indipendentemente dal periodo di produzione del prodotto.

A differenza del MERF, che si concentra sull'analisi interna della produzione, il CORF prende in considerazione esclusivamente l'esperienza sul campo e i guasti segnalati dal cliente. L'analisi copre il periodo di garanzia standard, solitamente i primi 24 mesi di servizio del veicolo, durante i quali si osservano i reclami pagati in garanzia. Questo approccio consente di valutare come i clienti percepiscono la qualità e l'affidabilità del prodotto, oltre a individuare eventuali criticità che potrebbero influire sulla loro soddisfazione.

I tecnici specializzati esaminano l'andamento dei KPI, cercando di comprendere le ragioni alla base di eventuali miglioramenti o peggioramenti nelle prestazioni. Successivamente, queste informazioni vengono utilizzate per correlare l'andamento dei KPI con azioni specifiche intraprese o problemi legati al prodotto.

Il grafico di analisi del CORF (Figura 3) fornisce una visione chiara delle performance nel tempo. Le date, disposte da destra a sinistra, rappresentano il periodo di osservazione, suddiviso in intervalli di 12 mesi consecutivi di guasti. La linea rossa nel grafico mostra l'andamento dell'ACPU (Average Cost per Unit), che indica i costi medi di garanzia per unità, mentre la linea rossa tratteggiata rappresenta l'ACPU al netto degli effetti negativi sull'economia (come l'aumento dei costi delle parti e della manodopera). La linea grigia indica l'andamento del F/100 (Failures per 100 machines), mentre la linea grigia tratteggiata riflette l'F/100 al netto degli stessi effetti economici. Infine, le due colonne sulla sinistra del grafico rappresentano i target di performance stabiliti per l'anno corrente, fornendo un benchmark rispetto al quale vengono valutati i risultati ottenuti.

Questo tipo di analisi offre un importante strumento per monitorare la qualità percepita dal cliente e intervenire tempestivamente per risolvere problemi ricorrenti, migliorando l'affidabilità del prodotto e ottimizzando l'esperienza del cliente.⁹

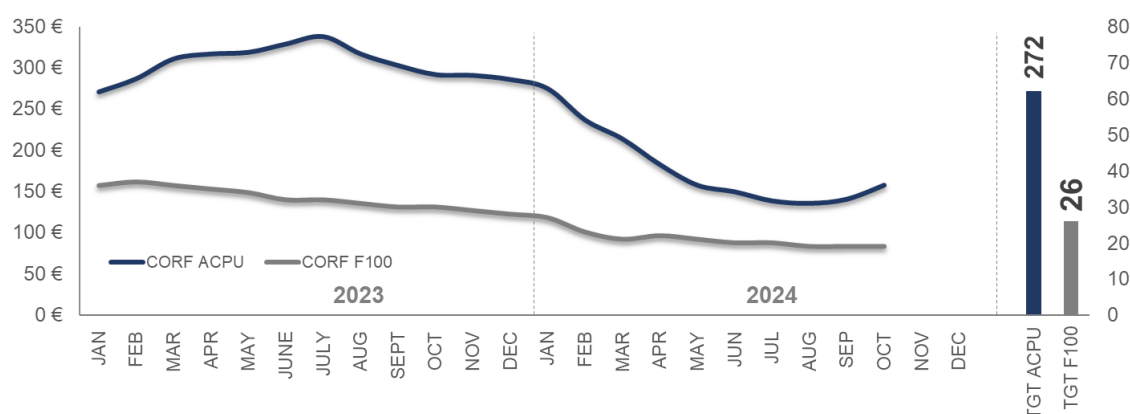


Figura 3 – Grafico CORF

L'analisi dei KPI attraverso il MERF e il CORF fornisce due prospettive complementari essenziali per un monitoraggio completo della qualità aziendale. Mentre il MERF offre una visione dettagliata delle performance interne, monitorando i risultati dei prodotti durante il loro ciclo di vita, il CORF si concentra sull'esperienza diretta del cliente, mettendo in evidenza la qualità percepita sul campo. L'integrazione di questi due approcci consente di ottenere una panoramica a 360 gradi sulla qualità del prodotto, supportando decisioni informate e azioni correttive mirate per ottimizzare sia i processi interni che l'affidabilità percepita dal cliente finale.

Un altro concetto che risulta fondamentale definire è la Tracking Family, che svolge un ruolo cruciale nella gestione delle informazioni, per garantire un'integrazione efficace dei dati. La Tracking Family rappresenta una combinazione unica di caratteristiche specifiche del veicolo e del prodotto. Definita dagli Utenti Chiave WRS (Warranty Reporting System) attraverso una Tabella di Gestione Dati, la Tracking Family costituisce un punto di riferimento centrale per l'organizzazione e il monitoraggio dei dati sulla qualità. La creazione, modifica e cancellazione di questa tabella sono prerogative esclusive degli utenti chiave, mentre gli altri utenti possono consultare le informazioni tramite la sezione "Master Data" della Landing Page, garantendo così il controllo e

⁹ (FPT Industrial, 2023)

l'integrità dei dati critici. Questo sistema permette di raccogliere e aggiornare le informazioni con precisione, facilitando un monitoraggio continuo delle performance e il miglioramento della qualità del prodotto.

Il Reporting Ufficiale WRS si basa interamente sulla Tracking Family, sfruttando la combinazione delle caratteristiche del veicolo e del prodotto per generare report accurati, come il MERF e il CORF. Questi report vengono elaborati per ciascuna Tracking Family in relazione a diverse dimensioni geografiche, come regione e mercato. In questo modo, l'azienda è in grado di monitorare la qualità e le prestazioni del prodotto in contesti specifici, adattandosi alle esigenze di ciascun mercato.

Inoltre, gli Utenti Chiave WRS possono associare a ciascuna Tracking Family un set di parametri personalizzati, affinando ulteriormente l'analisi e ottenendo report mirati che riflettono con precisione le variabili rilevanti per ciascuna combinazione di prodotto, mercato e regione. Grazie a questa funzionalità, è possibile ottenere una visione dettagliata delle performance in specifiche aree geografiche o segmenti di mercato, ottimizzando così la gestione dei dati e delle performance a livello globale.¹⁰

3. Miglioramento continuo e principali metodologie di problem solving

Il problem solving costituisce un pilastro essenziale per affrontare e risolvere situazioni critiche in modo strutturato ed efficace. Si tratta di un insieme di tecniche e metodologie progettate per analizzare problemi complessi, individuare le cause principali e implementare soluzioni mirate. In ambito aziendale, il problem solving è una competenza trasversale fondamentale, poiché garantisce continuità operativa, miglioramento della qualità e ottimizzazione dei processi.

Un problema si definisce come una deviazione da uno standard o da un obiettivo prestabilito, e il suo superamento richiede un'analisi approfondita e un approccio sistematico. Gli strumenti utilizzati per il problem solving combinano procedimenti mentali e tecniche di pensiero strutturato, pensati per generare soluzioni pratiche e

¹⁰ (FPT Industrial, 2023)

misurabili. Questi strumenti trovano applicazione in contesti diversi, dalla gestione della qualità alla produzione, fino allo sviluppo di nuovi prodotti.

Tra le metodologie più rinomate e utilizzate, in particolare in FPT Industrial, figurano approcci consolidati che offrono un framework chiaro per affrontare problemi di varia natura. Ogni metodologia è progettata per adattarsi a specifici contesti e obiettivi, ma tutte condividono un focus sulla raccolta e analisi dei dati, l'identificazione delle cause radice e l'implementazione di azioni correttive efficaci.¹¹

3.1. Metodo 8D

L'approccio 8D (Eight Disciplines) è una metodologia strutturata e collaborativa utilizzata per affrontare e risolvere problemi complessi, con un focus particolare sull'individuazione delle cause profonde e sull'implementazione di azioni correttive efficaci. Questo metodo si distingue per la sua capacità di analizzare le problematiche in modo sistematico, sviluppare soluzioni sostenibili e misurarne l'impatto sulle attività aziendali. Sebbene possa essere applicato a diverse tipologie di criticità, trova un utilizzo particolarmente efficace nella gestione delle non conformità e nel miglioramento della qualità. La sua applicazione è spesso richiesta in contesti specifici, quali:

- Reclami ricevuti dai clienti;
- Problemi legati alla sicurezza o alla conformità alle normative;
- Aumenti anomali di scarti, sprechi, difetti interni o fallimenti nei test;
- Tassi di guasto superiori a quelli previsti dai dati di garanzia.

Nonostante la sua flessibilità e le numerose applicazioni di successo, il metodo 8D presenta alcuni limiti. In particolare, può essere complesso e richiedere tempo per essere sviluppato e implementato in modo efficace. Inoltre, i membri del team devono essere adeguatamente formati per applicare correttamente la metodologia. È inoltre fondamentale mantenere una comunicazione costante tra i partecipanti e integrare il metodo in un programma di miglioramento continuo per garantirne l'efficacia a lungo termine.

¹¹ (Aldieri, 2024)

I passaggi in cui si articola la metodologia sono otto, come suggerisce il nome stesso, e sono riportati di seguito:

1. Costruire un Team di Lavoro (D1) → Un'adeguata pianificazione è essenziale per avviare correttamente il processo. Prima di formare il team che affronterà il problema, è importante raccogliere informazioni chiave, come il numero identificativo e la descrizione del componente segnalato, la data del guasto, i codici cliente e fornitore, e una breve analisi del problema. Inoltre, si utilizza una checklist per raccogliere i sintomi e si valuta la necessità di un'azione di risposta d'emergenza per proteggere i clienti da eventuali danni ulteriori. Il team deve essere multidisciplinare, comprendendo esperti provenienti da vari dipartimenti (come produzione, ingegneria e marketing), in modo da affrontare in maniera efficace i reclami dei clienti o le deviazioni dalla qualità;
2. Descrivere il Problema (D2) → In questa fase, il problema deve essere descritto con chiarezza, utilizzando l'approccio 5W+2H (Chi, Cosa, Quando, Dove, Perché, Come e Quanto) per analizzare e definire ogni aspetto del problema. Questa descrizione dettagliata assicura che tutti i membri del team abbiano una comprensione condivisa della situazione;
3. Sviluppare Azioni di Contenimento Temporaneo (D3) → I membri del team definiscono e implementano azioni correttive temporanee per limitare l'impatto del problema e proteggere il cliente fino all'applicazione di soluzioni permanenti. Le azioni devono essere in linea con i requisiti dei sistemi di gestione della qualità, e devono preservare le prove utili per la risoluzione definitiva del problema;
4. Identificare e Verificare le Cause Radice (D4) → In questa fase, si analizzano le cause profonde del problema utilizzando strumenti come il diagramma di Ishikawa e l'approccio 5W+2H. L'obiettivo è identificare con precisione le cause e verificare le ipotesi con dati concreti, evitando di fare supposizioni non confermate;
5. Sviluppare Azioni Correttive Permanenti (D5) → In base alle cause radice identificate, vengono sviluppate soluzioni correttive permanenti. Queste soluzioni devono essere validate attraverso programmi di preproduzione per assicurare una soluzione definitiva e sostenibile al problema;

6. Implementare e Validare le Azioni Correttive (D6) → Le azioni correttive selezionate vengono messe in atto e monitorate per verificare che il problema sia effettivamente risolto. È importante anche controllare che non si verifichino nuovi problemi o effetti indesiderati;
7. Prevenire le Recidive (D7) → Per evitare che il problema si ripresenti, è necessario modificare i sistemi di gestione, le pratiche operative e le procedure aziendali. Questi cambiamenti devono essere costantemente monitorati per garantirne l'efficacia e prevenire la comparsa di problemi simili in futuro;
8. Riconoscere e Premiare il Contributo del Team (D8) → Una volta risolto il problema, è importante condividere i risultati ottenuti e riconoscere l'impegno dei membri del team. Questa fase include anche la documentazione dei risultati, il feedback positivo e l'implementazione del ciclo PDCA (Plan-Do-Check-Act) per migliorare continuamente la soddisfazione del cliente.

Il metodo 8D è stato applicato con successo in numerosi contesti industriali, dimostrando la sua capacità di fornire soluzioni pratiche e di migliorare in modo significativo la gestione dei problemi di qualità.¹²

¹² (FPT Industrial, 2023); (Adaveesh, 2017)

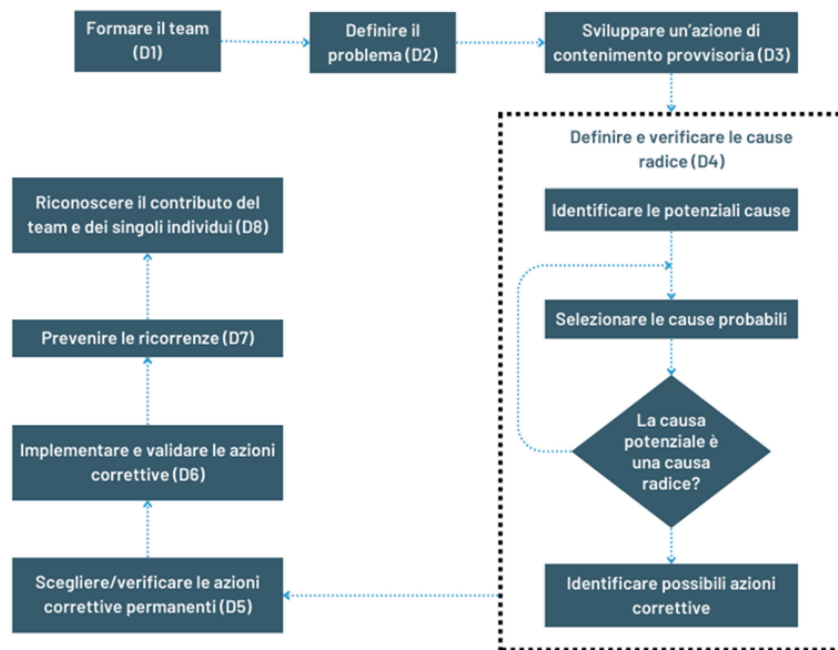


Figura 4 – Step Metodo 8D

Il metodo 8D si avvale di due strumenti supplementari particolarmente efficaci nell'approfondire l'analisi delle cause e nel facilitare la risoluzione dei problemi: il diagramma di Ishikawa e il grafico di Pareto. Questi strumenti contribuiscono a identificare e comprendere meglio le cause profonde dei problemi, favorendo un intervento mirato e sistematico.¹³

3.2. Analisi di Pareto

Tra le metodologie di problem solving, l'analisi di Pareto si distingue come uno strumento grafico efficace per rappresentare e interpretare i dati. Questo approccio consente di identificare i problemi più rilevanti in un contesto specifico e di stabilire priorità d'intervento, puntando a risolvere le cause principali che contribuiscono maggiormente all'effetto osservato. L'analisi si fonda sulla legge di Pareto, che afferma che circa l'80% degli effetti è generato dal 20% delle cause.

¹³ (Adaveesh, 2017)

Il processo di analisi di Pareto si articola in tre fasi principali:

1. **Classificazione dei dati:** in questa fase è fondamentale stabilire come suddividere i dati, ad esempio classificando i guasti in base alla loro modalità di occorrenza;
2. **Scelta del periodo di osservazione:** viene selezionato un intervallo temporale per raccogliere i dati. Ad esempio, si può analizzare la causa di un guasto monitorando il tempo che intercorre tra la produzione del pezzo e il guasto stesso, per ottenere una comprensione più approfondita dell'evento;
3. **Raccolta e ordinamento dei dati:** i dati vengono raccolti e ordinati in ordine decrescente, consentendo di individuare quali componenti si guastano più frequentemente e di determinare se il problema è stagionale o legato ad altri fattori.

Il grafico di Pareto, che rappresenta visivamente questi dati, è un diagramma a barre in cui ogni barra rappresenta una categoria o un fattore che contribuisce al problema. Le categorie sono disposte sull'asse orizzontale, mentre le frequenze sono riportate sull'asse verticale. Le barre sono ordinate in modo decrescente da sinistra a destra, mentre una linea sovrapposta rappresenta la percentuale cumulativa delle frequenze. Le categorie con le barre più alte sono quelle che contribuiscono maggiormente al problema.

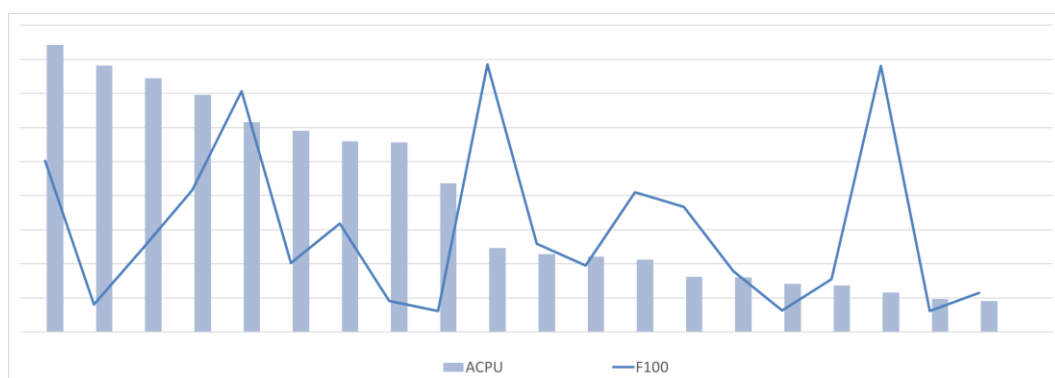


Figura 5 – Grafico di Pareto

Questo grafico permette non solo di identificare le categorie principali, ma anche di capire come diversi fattori possano influenzare il problema, aiutando così a focalizzare gli interventi sulle aree più critiche. Gli esperti suggeriscono di utilizzare il grafico di

Pareto principalmente per due scopi: decomporre un problema in fattori chiave e identificare le categorie più rilevanti per l'intervento.¹⁴

3.3. Ishikawa

Il diagramma di Ishikawa, noto anche come diagramma causa-effetto o a lisca di pesce, è uno strumento fondamentale per la comprensione e la risoluzione dei problemi, sviluppato da Kaoru Ishikawa negli anni '60. Questo approccio visivo aiuta a identificare le cause profonde di un problema e facilita la collaborazione all'interno del team, creando un ambiente di lavoro in cui ogni membro può contribuire attivamente all'analisi.

Il primo passo nell'applicazione del diagramma è la chiara definizione del problema. È essenziale formulare una descrizione precisa della questione da affrontare, identificando le aspettative e stabilendo obiettivi concreti e misurabili. Questo orienta le azioni future verso risultati tangibili e raggiungibili, favorendo una comprensione condivisa della situazione. Inoltre, è importante elencare i fattori noti e sconosciuti che possono influire sul problema, in modo da cogliere la complessità del contesto.

Il processo di brainstorming gioca un ruolo cruciale. Durante questa fase, vengono esplorate ipotesi, idee e suggerimenti che arricchiscono l'analisi e permettono di affrontare il problema da diverse angolazioni. La partecipazione attiva di tutti i membri del team è fondamentale per scoprire soluzioni innovative e creative, stimolando il pensiero divergente e migliorando la qualità dell'analisi.

Una volta raccolte le informazioni, è necessario validare i dati, i metodi e gli strumenti a disposizione. La validazione garantisce che l'analisi si basi su fondamenti solidi, migliorando l'affidabilità dei risultati e favorendo un approccio rigoroso e responsabile. Durante questa fase, il team è in grado di discernere le cause principali del problema, distinguendo gli elementi che hanno un impatto significativo da quelli marginali. Questo processo di convalida promuove una maggiore attenzione ai dettagli e facilita l'emergere di decisioni informate.

¹⁴ (Aldieri, 2024); (FPT Industrial, 2023)

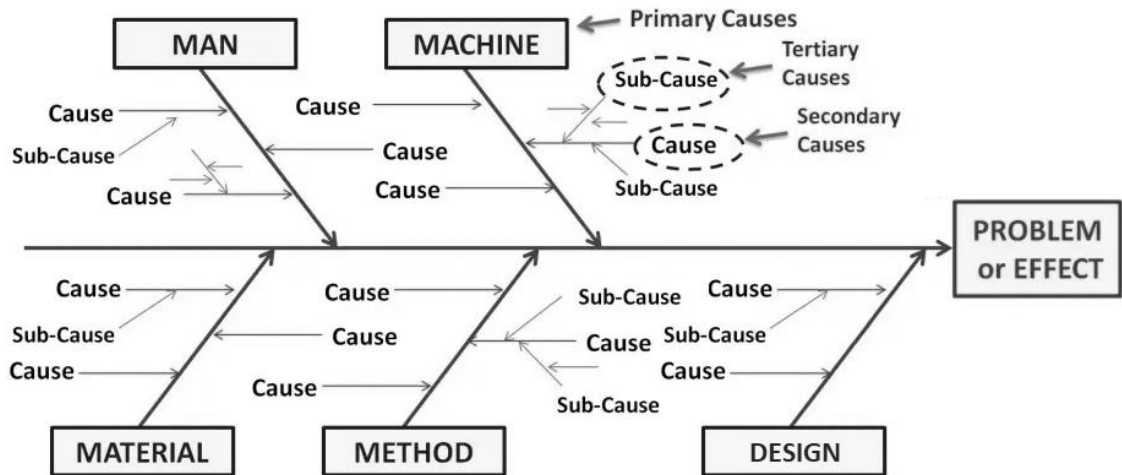


Figura 6 – Diagramma di Ishikawa

In sintesi, il diagramma di Ishikawa non è solo uno strumento utile per risolvere problemi, ma anche per promuovere un ambiente di apprendimento continuo all'interno dell'organizzazione. L'analisi accurata e la validazione delle informazioni consentono ai team non solo di risolvere i problemi attuali, ma anche di prevenire il ripetersi di situazioni simili, contribuendo ad una cultura orientata all'eccellenza.

Il diagramma di Ishikawa si distingue per la sua capacità di rappresentare visivamente le cause di un problema, rendendo evidente la loro interconnessione. Le frecce, che simboleggiano le "lische di pesce", si diramano dalla spina dorsale del diagramma, ognuna delle quali rappresenta una causa specifica, come materiali, metodi di lavoro, forza lavoro, misurazioni, macchinari, equipaggiamenti e ambiente. Questa struttura consente di classificare in modo chiaro tutte le cause legate al problema e di ridurre la complessità, offrendo una visione concisa delle relazioni causa-effetto.

Il diagramma di Ishikawa presenta numerosi vantaggi, tra cui:

- Classificare e visualizzare tutte le cause associate a un problema;
- Ridurre la complessità di un problema ampio;
- Stimolare la partecipazione attiva di tutti i membri del team nell'analisi;
- Aumentare il coinvolgimento del team nel processo di risoluzione dei problemi;
- Identificare le aree che necessitano di un approfondimento quando mancano informazioni;
- Fornire strumenti utili per sviluppare soluzioni adeguate;
- Offrire una visione chiara e strutturata delle cause e degli effetti.

Nel panorama degli strumenti di problem solving, le metodologie 5 Why e 5W2H si distinguono per la loro semplicità ed efficacia nell'analizzare problemi e identificare soluzioni mirate. Sebbene meno articolate rispetto ad approcci più complessi, come l'analisi di Pareto o il diagramma di Ishikawa, queste tecniche si rivelano estremamente utili in contesti in cui è necessario individuare rapidamente le cause profonde o pianificare interventi chiari e organizzati.¹⁵

3.4. 5 Why

La tecnica dei "5 Perché" è uno strumento fondamentale per individuare le cause profonde di un problema complesso o di un difetto. Questo metodo utilizza un approccio iterativo basato sulla semplice domanda "Perché?", esplorando le relazioni di causa-effetto fino a risalire all'origine della questione.

Il processo parte da una definizione chiara e precisa del problema, fondamentale per l'efficacia dell'analisi. Una corretta formulazione iniziale, spesso supportata da strumenti come il diagramma di Ishikawa o da una raccolta di dati accurata, permette di indirizzare correttamente l'indagine. Da qui si procede ponendo la domanda "Perché?", cercando di individuare la motivazione sottostante. La risposta può rappresentare la causa principale o un sintomo che conduce ad essa, e questa analisi viene ripetuta fino a

¹⁵ (Realyvásquez-Vargas, 2020); (FPT Industrial, 2023); (Aldieri, 2024)

identificare la radice del problema. L'indagine termina quando la catena delle domande non può più proseguire in modo significativo.

Un elemento cruciale è mantenere un approccio imparziale, evitando conclusioni affrettate influenzate da pregiudizi personali o dal bias di conferma. È essenziale considerare che non sempre il percorso delle domande porta a una singola causa radice: in molti casi emergono risposte multiple, indicando una rete di fattori interconnessi. Per risolvere completamente il problema, è necessario intervenire su tutte le cause identificate.

Tuttavia, eliminare direttamente la causa principale può non essere sempre possibile, soprattutto se comporta investimenti o risorse eccessive. In questi casi, si può agire su uno dei sintomi della causa radice, adottando soluzioni mirate e proporzionate.

Sebbene il metodo sia noto come "5 Perché", il numero di iterazioni necessarie non è fisso e dipende dalla complessità del problema. Ciò che conta è che ogni risposta derivi da un'analisi rigorosa e non da supposizioni. Questo approccio permette di superare le soluzioni superficiali, affrontando il problema alla radice e favorendo l'identificazione di strategie risolutive efficaci e sostenibili.¹⁶

3.5. 5W2H

Il metodo 5W2H è un efficace strumento strutturato per affrontare le problematiche; amplia il modello 5W1H con l'aggiunta dell'elemento *How Much* per analizzare costi e risorse. Si basa su sette domande fondamentali: *Who* (Chi), *What* (Cosa), *When* (Quando), *Where* (Dove), *Why* (Perché), *How* (Come) e *How Much* (Quanto), offrendo una visione completa di un problema o progetto.

Il processo inizia con il *What*, per identificare il problema da risolvere, analizzando ciò che va migliorato e le cause delle difficoltà. Con il *Why*, si esplorano le ragioni alla base della situazione e le logiche che hanno guidato le prassi esistenti. Il *Who* definisce i responsabili e le figure coinvolte, mentre il *Where* individua i luoghi critici e le aree d'intervento. Il *When* valuta le tempistiche e le condizioni che influenzano le azioni.

¹⁶ (MeeTheSkilled, 2018); (FPT Industrial, 2023)

L'*How* si concentra sulle modalità operative e strategie di risoluzione, esplorando anche approcci inediti, e l'*How Much* analizza risorse e costi necessari per una pianificazione sostenibile.

Sviluppato in Giappone, il 5W2H è apprezzato a livello globale per la sua semplicità e versatilità, adattandosi a contesti aziendali, educativi e personali. Questo approccio strutturato non solo favorisce una comprensione approfondita dei problemi, ma migliora anche la qualità delle decisioni, aiutando a organizzare informazioni e processi in modo chiaro ed efficace.¹⁷

¹⁷ (LBI, 2024); (FPT Industrial, 2023)

Capitolo 2

Digitalizzazione del processo: l'intervento dell'Intelligenza Artificiale

1. Panoramica del problema e della soluzione

FPT Industrial è una realtà specializzata nella progettazione, produzione e distribuzione di sistemi di propulsione e trasmissione destinati a camion, veicoli su strada e off-road, oltre che a settori come la power generation. La divisione Powertrain Quality & Product Behaviour svolge un ruolo cruciale nell'assicurare la qualità dei prodotti, concentrandosi sull'identificazione delle fasi del processo produttivo che necessitano di interventi migliorativi per risolvere i problemi di malfunzionamento riscontrati.

Operando su scala globale, FPT Industrial copre due macroregioni geografiche principali: AMECA, che include Europa, Africa, Medio Oriente e Americhe, e APAC, che comprende l'Asia e il Pacifico. L'obiettivo primario è individuare l'origine dei problemi segnalati dai clienti, analizzando il processo produttivo per identificare le cause principali. Questo avviene tramite un flusso di informazioni che parte dai *dealer*, ossia le officine che fungono da intermediari tra l'azienda e l'utente finale. I *dealer*, dopo aver riscontrato anomalie nei prodotti e dopo essere intervenuti per risolverle, inoltrano reclami all'azienda, che vengono utilizzati per determinare quali fasi del processo produttivo necessitano di attenzione.

La fase di analisi prevede l'utilizzo di due KPI fondamentali: *F/100* e *ACPU*, che consentono di monitorare l'andamento dei problemi. Un aumento significativo di uno di questi valori guida l'attenzione verso le componenti che causano tali variazioni. In questo contesto, si leggono e analizzano le *claim*, ovvero le descrizioni dei problemi inviate dai *dealer*. Tuttavia, quando il volume delle *claim* è elevato, l'analisi completa diventa difficoltosa e si tende a concentrarsi su un campione rappresentativo, rischiando di tralasciare dettagli importanti. Per facilitare questo processo, i problemi vengono raggruppati in *cluster*, rendendo più semplice l'identificazione delle macrocategorie di problematiche.

Nonostante questa metodologia offra una base strutturata, presenta limiti significativi: richiede tempi considerevoli e consente spesso di analizzare solo una quantità limitata di dati. Per superare queste inefficienze, è stata proposta l'introduzione dell'Intelligenza Artificiale (IA), come strumento a supporto del processo di gestione dei reclami, automatizzando la fase di aggregazione dei dati in cluster e migliorando sensibilmente l'efficienza operativa. Uno strumento così avanzato non è concepito per sostituire il lavoro umano, ma per affiancarlo, creando una sinergia che permetta di raggiungere livelli di efficienza superiori.

1.1. Raccolta dati tramite questionario

Prima di avviare l'analisi, si è ritenuto utile raccogliere informazioni su aspetti chiave del processo di gestione delle claim. Questo ha permesso di comprendere le attuali modalità operative e le aspettative degli operatori rispetto all'introduzione dell'intelligenza artificiale. L'indagine ha esplorato i seguenti temi:

1. La quantità media di claim analizzate in una settimana;
2. Il tempo necessario per leggere e analizzare una claim;
3. Le informazioni principali che gli operatori desiderano estrarre durante la lettura;
4. Il livello di utilità percepito riguardo l'introduzione dell'intelligenza artificiale in azienda.

Per raccogliere questi dati, è stato progettato un questionario tramite Microsoft Forms, rivolto agli operatori attualmente coinvolti o che in passato si sono occupati dell'analisi delle claim. Il questionario, composto da otto domande, è stato strutturato in modo da favorire risposte precise e mirate. Le domande sono per la maggior parte a risposta chiusa, con un'unica domanda aperta in cui i partecipanti hanno avuto la possibilità di suggerire ulteriori quesiti utili per l'analisi.

Nello specifico, il questionario include:

- Domande quantitative, come la stima del numero di claim analizzate settimanalmente e il tempo medio dedicato alla lettura di ciascuna claim, utili per comprendere i carichi di lavoro e le dinamiche operative;

- Domande qualitative, volte a individuare le informazioni principali ricercate durante l'analisi delle claim, fornendo così una panoramica sulle priorità e sugli obiettivi degli operatori;
- Valutazioni soggettive, con una scala da 1 a 5, per misurare il livello di utilità percepita rispetto all'introduzione dell'intelligenza artificiale come supporto al processo decisionale e analitico.

L'ultima domanda, aperta, ha permesso ai partecipanti di esprimere idee o proporre ulteriori spunti, arricchendo l'indagine con suggerimenti personalizzati. Questo approccio ha garantito una raccolta dati strutturata ma flessibile, utile per calibrare al meglio le soluzioni da implementare.

Grazie alla somministrazione online, il questionario è stato facilmente accessibile, favorendo una partecipazione ampia e rappresentativa. Le informazioni raccolte costituiranno una base fondamentale per analizzare in che modo l'introduzione dell'intelligenza artificiale possa migliorare i processi aziendali, ottimizzando i tempi e la qualità dell'analisi delle claim (v. [Appendice A](#)).

È stata ricavata un'informazione fondamentale riguardo ai tempi medi di lettura di una singola claim: dalle risposte degli operatori è emerso che il tempo medio impiegato si aggira intorno a 1-2 minuti per claim. Questa stima deriva dal fatto che la metà degli intervistati ha dichiarato di impiegare meno di un minuto, mentre l'altra metà ha indicato un tempo compreso tra uno e cinque minuti.

Inoltre, il sondaggio si è rivelato utile per individuare, in linea di massima, le informazioni più significative ricavate dalla lettura dei commenti dei dealer, oltre alle principali difficoltà che si possono riscontrare, dovute ad inesattezze ortografiche e grammaticali o legate alle traduzioni. I principali risultati ottenuti sono illustrati nelle figure 7 e 8.

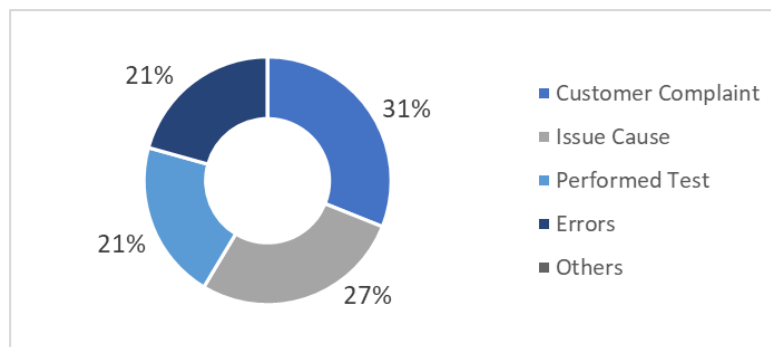


Figura 7 – Distribuzione Informazioni estratte

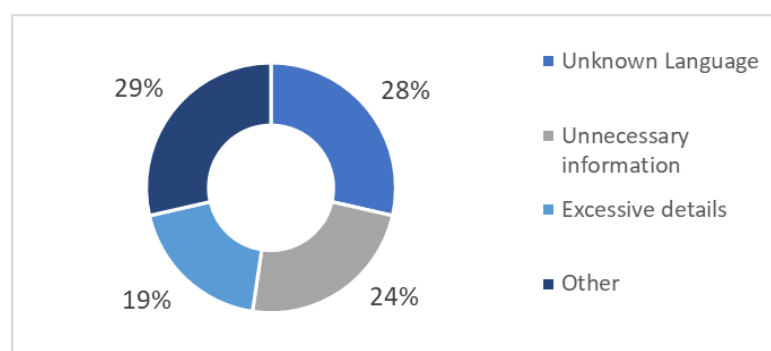


Figura 8 - Distribuzione tipologie di difficoltà riscontrate

Si è evidenziato che, nella fase di lettura e analisi dei reclami, l'operatore ha bisogno di estrapolare principalmente tre informazioni essenziali: customer complaint, issue cause e performed test.

Tuttavia, la lettura delle claim presenta alcune difficoltà, in particolare a causa della presenza di dettagli superflui, informazioni non necessarie e del linguaggio utilizzato. Spesso, infatti, i reclami provengono da diverse parti del mondo e vengono trascritti dai dealer nella loro lingua, che non sempre è l'inglese, rendendone complessa l'interpretazione. Questi dati di input costituiranno la base su cui costruire e implementare la soluzione.

2. Use Case 1: Intelligent Database Enrichment

2.1. Obiettivo

L'obiettivo è sfruttare i modelli di intelligenza artificiale generativa per analizzare in modo intelligente i commenti testuali dei concessionari presenti nel database WRS (Warranty Reporting System). I modelli dovranno essere in grado di comprendere il contesto, estrarre parole chiave e classificarle in cluster predefiniti. Il risultato di questo processo consentirà di ottenere una classificazione più dettagliata delle richieste, migliorando la velocità con cui gli utenti possono elaborare e generare l'output finale rispetto al metodo attuale.

L'obiettivo è raggiungere un'accuratezza di almeno l'80%. I cluster saranno predefiniti per consentire al modello di adattarsi a essi tramite il prompt engineering. La fase di industrializzazione mira ad applicare questi cluster su un ampio set di richieste, mentre eventuali modifiche future richiederanno lo sviluppo di prompt dedicati ad hoc.

2.2. Elementi del Database Intelligente

Gli strumenti utilizzati per ottenere una combinazione ottimale, in grado di estrarre e categorizzare le informazioni chiave dai reclami dei concessionari e di creare cluster di reclami simili, sono i seguenti, elencati nell'ordine di applicazione:

1. LLM (Large Language Model);
2. Algoritmi di Embeddings;
3. Algoritmi di Riduzione della Dimensionalità;
4. Algoritmi di Clustering.

Prima dell'elaborazione da parte degli algoritmi e dell'intelligenza artificiale, i dati sono stati sottoposti a un processo di pulizia, che ha incluso diverse operazioni. Sono stati rimossi i dati incoerenti, eliminati i duplicati derivanti da commenti ripetuti, corretti i caratteri speciali che potevano interferire con il codice, tradotti i commenti scritti in una lingua sconosciuta, e corrette eventuali sviste ortografiche e grammaticali. Inoltre, sono state rimosse informazioni irrilevanti e ridondanti.

2.2.1. LLM (*Large Language Models*)

Dopo aver completato la fase di "preparazione" e "pulizia" dei dati, si può procedere con lo sviluppo della soluzione, introducendo nel sistema i requisiti richiesti. Un elemento chiave in questo processo è rappresentato dai Large Language Models (LLM).

I Large Language Models sono sistemi di intelligenza artificiale avanzati progettati per comprendere ed elaborare il linguaggio naturale. Essi si basano su reti neurali profonde, caratterizzate da miliardi, e in alcuni casi migliaia di miliardi, di parametri. Tra le implementazioni più note di questi modelli rientra GPT (Generative Pre-trained Transformer), una famiglia di modelli dedicata all'elaborazione del linguaggio naturale. Il primo modello GPT è stato sviluppato da OpenAI nel 2018, segnando un netto miglioramento rispetto ai sistemi precedenti. Prima di questa innovazione, i chatbot disponibili sul mercato erano limitati: non riuscivano a cogliere il contesto delle parole, non ordinavano le informazioni in maniera efficace ed erano generalmente lenti nell'elaborazione del significato. Un esempio classico è rappresentato dai chatbot tradizionali, che fornivano risposte preimpostate basate sull'associazione di parole chiave a concetti già memorizzati.

L'introduzione di GPT ha permesso di superare queste limitazioni, grazie all'architettura Transformer, progettata per elaborare sequenze di dati in modo più efficiente rispetto ai modelli tradizionali. Questa struttura consente a GPT di rispondere con maggiore coerenza e fluidità, mostrando una comprensione del linguaggio molto simile a quella umana. La vasta memoria di questi modelli (pari a diversi petabyte di dati) consente di analizzare e interpretare il linguaggio naturale con un livello di profondità senza precedenti.

Negli anni, la tecnologia GPT si è evoluta attraverso diverse versioni. Il percorso ha avuto inizio nel 2018 con GPT-1, che era stato addestrato su libri e articoli per imparare a prevedere le parole successive in una sequenza. Nel 2019, è stato rilasciato GPT-2, un modello più avanzato capace di generare testi di senso compiuto; tuttavia, a causa di preoccupazioni sulla sua possibile utilizzazione impropria, OpenAI inizialmente ha scelto di non rendere pubblica la versione completa. Nel 2020, è stato introdotto GPT-3, con una quantità di parametri circa 200 volte superiore rispetto al suo predecessore,

rendendolo in grado di rispondere a domande, tradurre testi, riassumere contenuti e svolgere molte altre funzioni. Il 2022 ha visto la nascita di GPT-3.5, con cui è stata lanciata la prima versione di ChatGPT, un chatbot interattivo in grado di intrattenere conversazioni dinamiche con gli utenti. Nel 2023 è stato introdotto GPT-4, mentre nel maggio 2024 OpenAI ha rilasciato GPT-4o. Queste ultime versioni hanno introdotto la multimodalità, consentendo l'elaborazione di testo e immagini, e GPT-4o ha portato ulteriori miglioramenti, come il riconoscimento delle emozioni nella voce e la capacità di interagire in tempo reale.¹⁸

A dicembre 2024, OpenAI ha introdotto nuovi aggiornamenti rilevanti. Il 5 dicembre è stato lanciato ChatGPT Pro, un piano in abbonamento dal costo di 200 dollari al mese che garantisce accesso illimitato ai modelli più avanzati, tra cui OpenAI o1, o1-mini, GPT-4o e Advanced Voice. Inoltre, questo piano include l'o1 Pro Mode, una versione potenziata di o1 che utilizza una maggiore capacità di calcolo per fornire risposte più precise anche per problemi complessi. Il 9 dicembre, OpenAI ha presentato Sora, una tecnologia che consente la generazione di video a partire da testo, con output fino a 20 secondi in risoluzione 1080p. Il 12 dicembre, è stata introdotta la "Santa Voice" per ChatGPT, che permette agli utenti di interagire con un chatbot dotato di una voce a tema natalizio fino ai primi di gennaio. Infine, il 17 dicembre è stata rilasciata la versione completa del modello o1, con funzionalità avanzate come l'esecuzione di funzioni, l'analisi delle immagini e una modalità di "ragionamento avanzato" per migliorare la qualità delle risposte.

Questi aggiornamenti dimostrano il continuo sviluppo della tecnologia AI da parte di OpenAI, con l'obiettivo di rendere i modelli linguistici sempre più potenti, versatili e capaci di soddisfare le esigenze di utenti in ambiti sempre più diversificati.¹⁹

2.2.1.1. Utilizzo degli LLM

Lo sviluppo dello studio in questione ha previsto l'applicazione di questi modelli per la semplificazione dei commenti dei dealer e per l'estrazione di informazioni chiave dai commenti. Il prompt viene ottimizzato per essere più chiaro ed efficace, in modo da

¹⁸ (Morriello, 2023)

¹⁹ (Wong, 2024)

ottenere risposte più pertinenti e accurate dai modelli di linguaggio di grandi dimensioni (LLM). Vengono fornite al LLM informazioni chiave per migliorare la comprensione del contesto di lavoro; viene guidato il LLM verso una risposta standard per semplificare i passaggi successivi; si impone al LLM di rispondere in un formato standard.

Per ogni commento vengono estratte le seguenti sette informazioni:

- **Customer Complaint:** rappresenta il primo segnale di errore percepito dal cliente, ovvero come il problema si manifesta ai suoi occhi. Può trattarsi, ad esempio, di una spia accesa, un rumore anomalo o una riduzione delle prestazioni del componente interessato;
- **Issue Cause:** indica la causa radice del problema, cioè il fattore che lo ha generato e che ha portato alla sua manifestazione. È su questa causa che bisogna intervenire per risolvere definitivamente il problema;
- **Second Level Component:** identifica il sottocomponente specifico, situato all'interno del componente principale, che risulta danneggiato e richiede riparazione o sostituzione;
- **Errors Code:** si riferisce al codice di errore del componente danneggiato, espresso tramite codici alfanumerici in base 16 o numerici in base 10.
- **Performed Test:** comprende i test effettuati per valutare la situazione del componente e identificare eventuali anomalie;
- **Troubleshooting:** descrive le operazioni di risoluzione dei problemi eseguite prima della sostituzione del componente. Include l'elenco dettagliato delle azioni intraprese per cercare di risolvere il guasto;
- **THD:** acronimo di *Technical Help Desk Number*, è un numero di nove cifre riportato nei commenti. Rappresenta un identificativo univoco della richiesta di assistenza clienti.

I dati su cui si basa la categorizzazione sono i tre elementi più rilevanti tra quelli menzionati in precedenza: customer complaint, issue cause e performed test. Questi tre aspetti sono considerati i più significativi ai fini dell'analisi, sia per la loro utilità nel processo, sia per il fatto che rappresentano informazioni sempre disponibili, come è anche stato confermato dagli operatori che hanno risposto al questionario erogato.

Infatti, tali elementi possono essere estratti sistematicamente dai commenti forniti dai dealer, rendendoli una base solida e affidabile per l'elaborazione dei cluster.

Al contrario, le altre informazioni, pur essendo talvolta utili, possono risultare non sempre reperibili. Questa variabilità è dovuta alla natura non strutturata dei commenti, che non garantisce la presenza costante di tutti i dettagli desiderati. La scelta di focalizzarsi sui tre dati principali riflette quindi la necessità di lavorare su informazioni uniformemente presenti, massimizzando l'accuratezza e l'efficacia dell'analisi.

2.2.2. Algoritmi di Embeddings

Le tre categorie di informazioni che vengono considerate (Customer Complaint, Issue Cause e Performed Test) per effettuare l'analisi includono concetti linguistici, motivo per cui non è possibile utilizzare la matematica di base per capire se vi è similarità semantica tra i concetti, poiché i modelli matematici non performano bene su stringhe. Per analizzare, quindi, la vicinanza tra concetti linguistici vengono utilizzati gli algoritmi di Embeddings.

Gli algoritmi di embeddings sono una rappresentazione vettoriale numerica in uno spazio definito di una determinata porzione di testo. Viene data in pasto all'algoritmo una porzione di testo che restituisce un vettore di 1024 dimensioni e assume un valore compreso tra -1 e 1. L'insieme dei valori associati a questi concetti, all'interno di uno spazio predefinito, rappresenta il significato della parola in esame. Ogni valore esprime una determinata caratteristica del concetto. Questi valori indicano, dunque, la vicinanza tra concetti: se il valore è basso, significa che i concetti sono lontani dal punto di vista del significato; se il valore è alto, significa che sono vicini. Una volta che le parole vengono posizionate nello spazio, si può osservare quanto i punti siano vicini tra loro (i punti rappresentano le parole).

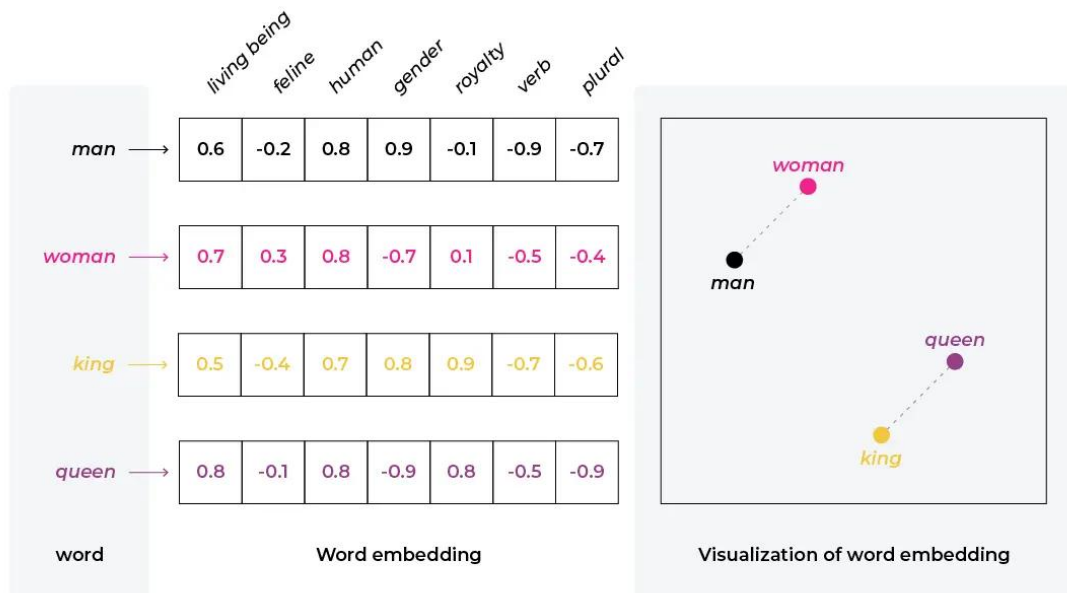


Figura 9 – Funzionamento algoritmi di Embeddings

L'immagine sopra (Figura 9) illustra il funzionamento degli embeddings attraverso un esempio significativo. Le parole "Re" e "Regina" hanno embeddings simili, poiché entrambe sono associate al concetto di "membri della famiglia reale". Al contrario, le parole "Regina" e "Uomo" presentano embeddings differenti, in quanto rappresentano concetti distinti.

Più in generale, il funzionamento prevede la creazione di reti neurali, addestrate per restituire degli embeddings che siano il più rappresentativi possibile della lingua, nel caso in esame la lingua utilizzata è l'inglese. Di seguito i due modelli utilizzati per generare embeddings:

- Word2Vec → si tratta di una tecnica che trasforma le parole in vettori, catturando le relazioni semantiche tra le parole;
- BERT → è una versione più smart del Word2Vec, poiché tiene conto anche della posizione delle parole nel testo.

Inizialmente l'analisi era stata effettuata utilizzando questi due algoritmi, in seguito sono state effettuate delle prove con altri due più efficienti che, infatti, hanno condotto a risultati ottimali. I due modelli utilizzati sono i seguenti:

- TITAN → Efficace per catturare le relazioni contestuali. Ideale per analizzare grandi set di dati di testo delle richieste all'interno dell'ecosistema di Amazon;
- Cohere → Potente per una comprensione semantica più sottile. Perfetto per generare embedding che migliorano i compiti di elaborazione del linguaggio naturale e ottimizzano le applicazioni basate su testo.

Questi due algoritmi sono più facili da integrare nelle applicazioni pratiche e sono più capaci di gestire il linguaggio naturale in modo più sofisticato, motivo per cui offrono prestazioni migliori rispetto ai primi due.²⁰

2.2.3. *Algoritmi di Riduzione della Dimensionalità*

Si è osservato che la dimensione dei vettori restituiti dall'algoritmo di Embeddings era eccessivamente elevata (1024 dimensioni), rendendo la gestione complessa e inefficiente. Con una dimensionalità così elevata, le prestazioni dell'algoritmo sono compromesse; pertanto, è necessario ridurre le dimensioni dei vettori, un processo che può essere ottimizzato mediante l'uso di algoritmi di riduzione della dimensionalità.

Esistono due tecniche volte alla riduzione della dimensionalità, si tratta delle seguenti:

1. PCA (Principal Component Analysis) → riduce la dimensionalità dei dati trasformandoli in un nuovo insieme di variabili (componenti principali) che catturano la massima varianza. Questa tecnica prova a cambiare lo spazio vettoriale in cui ci si muove; studia come sono distribuiti i punti nello spazio e cerca un sottospazio più piccolo che li descriva bene;
2. UMAP → tecnica di riduzione dimensionale non lineare che preserva la struttura globale dei dati mantenendo al contempo le relazioni sociali.²¹

²⁰ (Lovergine, 2022)

²¹ (Lovergine, 2022)

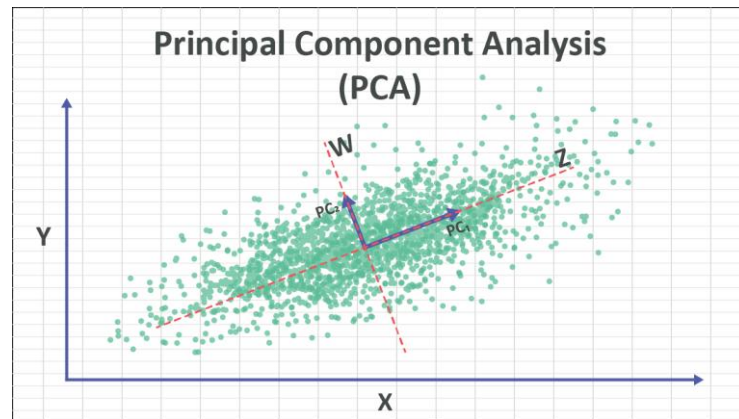


Figura 10 – Rappresentazione Grafica PCA

2.2.4. Algoritmi di Clustering

Una volta ridotta la dimensionalità dei vettori, si arriva alla fase di aggregazione dei dati vera e propria, tramite gli algoritmi di clustering.

Il clustering dei dati è una tecnica di apprendimento automatica che viene utilizzata per raggruppare tra loro dati simili in insiemi o in cluster, senza che necessariamente vengano assegnate delle etichette predefinite. L'obiettivo del clustering è identificare delle strutture o modelli nascosti all'interno dei dati, in modo che elementi dello stesso cluster siano più simili tra loro rispetto agli elementi di altri cluster. Ogni cluster rappresenta un gruppo di oggetti che condividono proprietà e/o elementi in comune.

Per effettuare l'aggregazione i dati esistono diversi algoritmi. Gli algoritmi di clustering sono delle procedure matematiche progettate per organizzare un insieme di dati in gruppi omogenei o "cluster", per l'appunto, in base alle loro caratteristiche in comune o simili.

Sono stati sviluppati numerosi algoritmi per la categorizzazione dei dati, di seguito sono riportati e descritti i tre che sono stati utilizzati per il caso in esame:

1. K-Means → Innanzitutto, si definisce il numero di cluster desiderato, k . Successivamente, i centroidi vengono posizionati in modo casuale nello spazio. Ogni punto viene assegnato al cluster il cui centroide è il più vicino in termini di distanza. Una volta completata questa fase, si aggiorna la posizione dei centroidi calcolando la media delle coordinate dei punti assegnati a ciascun cluster. Questo

processo di riassegnazione dei punti e di aggiornamento dei centroidi viene ripetuto iterativamente fino a quando i centroidi non di muoversi. Quando ciò avviene, si è raggiunto un minimo locale, che rappresenta una soluzione stabile, ma non necessariamente quella ottimale, poiché il risultato può dipendere dalle posizioni iniziali dei centroidi. Per questo motivo, è consuetudine eseguire più prove per determinare il numero ottimale di centroidi.

2. HDBSCAN (Hierarchical Density-Based Spatial Clustering of Applications with Noise) → è un'estensione di un altro algoritmo di clustering, il DBSCAN (Density-Based Spatial Clustering of Application with Noise). L'HDBSCAN è un algoritmo che individua gruppi di punti in uno spazio multidimensionale basandosi sulla densità locale dei punti. Quando all'interno di un'area si trovano tanti punti vicini, si ha un cluster. Con questo algoritmo è facile trovare forme di cluster irregolari e si possono individuare punti di dati isolati che vengono identificati come rumore (come si vedrà più avanti i cluster rumore sono i cluster identificati come 'Others').
3. Spectral clustering → è un algoritmo che ha un approccio basato sulla teoria dei grafi e utilizza una matrice di similarità per rappresentare le relazioni tra i punti dati. Questo metodo è utile per individuare cluster con forme irregolari e funziona bene su dati ad alta dimensionalità. Il punto centrale dello spectral clustering è la trasformazione dei dati in uno spazio alternativo, dove le strutture dei cluster risultano più chiare.²²

²² (Lovergine, 2022)

2.2.4.1. Silhouette

Per valutare la qualità di un clustering si utilizza spesso il coefficiente di Silhouette, una metrica ampiamente riconosciuta. Questo coefficiente permette di determinare se i punti dati sono stati assegnati correttamente ai rispettivi cluster, valutando contemporaneamente due aspetti fondamentali: la coesione interna di ciascun cluster e la separazione tra cluster differenti.

Di seguito è riportata la formula per calcolare il valore della silhouette:

$$s(x) = \frac{b(x) - a(x)}{\max\{a(x), b(x)\}}$$

Dove:

- $a(x)$ è la distanza media di x da tutti gli altri punti del suo stesso cluster.
- $b(x)$ è la distanza media di x da tutti i punti del cluster più vicino, ovvero il cluster con la distanza media più bassa da x .

Il valore della silhouette può variare tra -1 e 1, e ogni valore ha un significato specifico: quanto più si avvicina a 1, tanto più il clustering è accurato. Di seguito sono elencati i valori di silhouette e i significati associati a ciascun range di valori:

- Tra 0.75 e 1 → clustering eccellente; i punti sono ben separati dagli altri cluster e coesi all'interno del proprio cluster;
- Tra 0.5 e 0.75 → buona qualità del clustering;
- Tra 0 e 0.5 → clustering subottimale; i punti potrebbero essere vicini ai bordi dei cluster;
- Vicino a 0 → i punti sono equidistanti da due cluster, indicando probabilmente una sovrapposizione tra i cluster;
- Sotto lo 0 (vicino a -1) → clustering scadente; i punti sono probabilmente assegnati ai cluster sbagliati.²³

²³ (Lovergine, 2022)

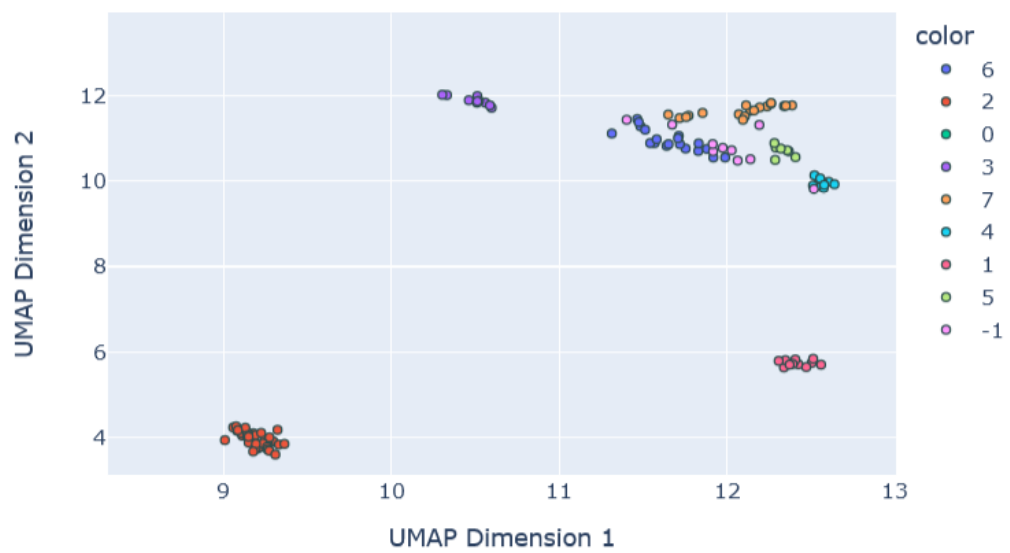


Figura 11 – Esempio di Clustering con valore di Silhouette pari a 0.84

Capitolo 3

Verifica del Modello e implementazione della soluzione

1. Fase di verifica del modello

Come già evidenziato, sono stati esaminati diversi tipi di algoritmi di embeddings, riduzione della dimensionalità e clustering. È stato quindi fondamentale analizzare quale tra queste tecniche offrisse i risultati migliori, individuando la combinazione più efficace per ridurre al minimo il margine di errore e garantire la massima affidabilità. A tal fine, sono state condotte prove su vari gruppi di claim per determinare la miglior combinazione dei tre algoritmi. La fase di controllo si è rivelata cruciale per fornire al sistema indicazioni più precise su come eseguire correttamente il clustering. Inizialmente, il sistema non riconosceva alcuni termini simili nel significato. Per questo motivo, si è reso necessario un intervento che migliorasse la qualità delle indicazioni fornite, consentendo così una aggregazione dei dati più precisa. Lo scopo del controllo manuale è stato, quindi, quello di perfezionare ulteriormente le istruzioni per l'algoritmo, permettendogli di interpretare i dati in maniera più scrupolosa.

La procedura di controllo manuale si articola in due fasi:

1. **Primo step:** fornire per ogni cluster una piccola descrizione delle caratteristiche comuni dei dati raggruppati in cluster dall'algoritmo. Sono stati individuati dei cluster con all'interno dati non omogenei tra loro, questi sono stati definiti dal sistema stesso "Others", sono cioè i cosiddetti cluster 'rumore' che contengono dati tra loro molto diversi e che non possono essere contenuti all'interno di nessun altro cluster, perché sono dei casi rari e isolati. È proprio dei processi di clustering trovare dei gruppi di dati che racchiudano al loro interno degli elementi non categorizzabili in altre macro-individuazioni. Fornire una piccola descrizione per ogni cluster è un passaggio utile per individuare quali sono i dati che sono stati assegnati in maniera errata ai cluster, in modo da poter indirizzare meglio l'algoritmo, fornendogli delle indicazioni più accurate. In questo modo, si è altresì in grado di stabilire quanto è stato accurato l'algoritmo nella creazione dei cluster.

2. **Secondo step:** dopo aver assegnato a ciascun cluster una breve descrizione, è possibile identificare gli errori commessi dall'algoritmo. Da qui, si possono calcolare percentuali utili a valutare la qualità dell'output. Per ogni cluster, sono stati individuati gli elementi che avrebbero dovuto appartenere ad altri cluster; questi vengono considerati errori che, rapportati al numero totale di elementi nel cluster, forniscono la percentuale di errore associata a quel cluster. Sommando tutti gli errori sul totale dei dati, si ottiene la percentuale di errore complessiva relativa all'algoritmo utilizzato.

$$\% \text{ errore} = \frac{n. \text{errori}}{\text{totale}} \times 100$$

Al termine di questa procedura si è in grado di determinare quale combinazione di algoritmi performa meglio in termini di percentuale di errore.

Il modello è stato inizialmente testato su due gruppi di claim:

- *Tier 4B Muffler* → silenziatore per motori diesel conformi agli standard Tier 4B, progettato per ridurre rumore ed emissioni;
- *Valvola EGR del motore F5* → è un componente che circola i gas di scarico per abbassare la temperatura di combustione e ridurre le emissioni di Nox nei motori diesel.

Tier 4B Muffler						
	Customer Complaint		Issue Cause		Performed Test	
	N. Cluster	% Errore	N. Cluster	% Errore	N. Cluster	% Errore
HDBScan	12	1%	20	6%	19	3%
Spectral	7	0%	12	11%	8	5%

Tabella 1 – Tier 4B Muffler

EGR Valve F5						
	Customer Complaint		Issue Cause		Performed Test	
	N. Cluster	% Errore	N. Cluster	% Errore	N. Cluster	% Errore
HDBScan	11	3%	8	4%	11	4%
Spectral	2	0%	10	8%	2	0%

Tabella 2 – EGR Valve F5

Nelle Tabelle 1 e 2 sono riportati i risultati dei test condotti per valutare l'accuratezza del modello su entrambi i gruppi di claim citati. Dall'analisi dei dati emerge che

l'algoritmo di clustering più performante è HDBSCAN. Sebbene in alcuni casi lo Spectral mostri percentuali di errore inferiori, la valutazione deve considerare anche il numero di cluster generati: un numero troppo basso potrebbe indicare una scarsa capacità di distinguere correttamente le diverse tipologie di errore riscontrabili su un singolo componente. Le celle colorate nelle tabelle evidenziano i valori chiave sulla base dei quali è stata selezionata la combinazione ottimale di algoritmi.

Il confronto è stato effettuato testando due diverse combinazioni di algoritmi: TITAN + UMAP + HDBSCAN e TITAN + UMAP + Spectral (rispettivamente algoritmo di embeddings, tecnica di riduzione della dimensionalità e algoritmo di clustering). I risultati confermano che la configurazione più efficace è TITAN + UMAP + HDBSCAN.

Le tecniche di embeddings e riduzione della dimensionalità erano già state valutate dal team di sviluppo, il quale aveva individuato TITAN come la soluzione più performante per la rappresentazione vettoriale dei dati e UMAP come il metodo più efficace per la riduzione della dimensionalità. Per quanto riguarda gli algoritmi di clustering, invece, è stato necessario condurre test specifici direttamente sulle claim, verificandone i risultati attraverso controlli manuali basati sulla conoscenza del prodotto.

Di conseguenza, la soluzione definitiva, che verrà implementata in azienda come strumento di supporto per l'analisi dei reclami, sarà sviluppata sulla base di questa combinazione ottimale. I dettagli sui passaggi analitici che hanno portato ai valori riportati nelle tabelle sono disponibili in [Appendice B.](#)

2. Fase post-processing

2.1. Assegnazione nomi ai cluster

A seguito dell'individuazione della combinazione più performante, sono stati effettuati dei raffinamenti, tra cui l'assegnazione di nomi ai cluster. Inizialmente, l'identificazione dei cluster avveniva tramite numeri univoci, seguendo un ordine crescente (da -1 a N). Questa modifica non introduce errori, ma consiste semplicemente nell'assegnare un nome testuale ai cluster già creati dalla rete. Per eseguire questa operazione, è stata utilizzata la rete, in particolare la chatbot GPT, che ha suggerito nomi coerenti con le informazioni aggregate da ciascun cluster. Inizialmente, l'associazione dei nomi ai cluster è stata realizzata manualmente, come descritto nel primo step della metodologia di verifica del clustering, che prevedeva l'intervento di un operatore umano. Successivamente, fornendo questa prima associazione come input al sistema, è stato possibile ottenere un'assegnazione più precisa e accurata dei nomi ai cluster.

2.2. Aggregazione cluster

Alcuni cluster risultavano essere molto simili tra loro, un problema emerso sia nella fase di verifica che successivamente, durante la fase di assegnazione dei nomi. Per ovviare a questa problematica ed evitare ridondanze, nonché ridurre la necessità di ulteriori controlli nella fase di analisi dei reclami, è stata implementata una soluzione che accorpasse i cluster simili tra loro.

Tuttavia, se da un lato questa operazione migliora l'efficienza e facilita il lavoro, dall'altro introduce un margine di errore del 5%, che rimane comunque inferiore alla soglia di tolleranza stabilita per l'industrializzazione della soluzione, fissata al 20% (Vedi tabella sotto). Nonostante l'incremento di errore, la fase di aggregazione è ritenuta essenziale per garantire un buon livello di efficienza, anche a costo di sacrificare qualche punto percentuale in termini di affidabilità. Questo passaggio finale riduce significativamente il numero di cluster creati in precedenza dal sistema.

Con l'ultima fase di affinamento del modello, è stata ottenuta una soluzione in grado di semplificare il processo di lettura e analisi dei reclami, rendendolo meno complesso e ripetitivo. La combinazione individuata presenta una percentuale di errore che

rappresenta il prodotto tra quella ottenuta dal controllo manuale effettuato prima dell'aggregazione e quella registrata dopo il processo di fusione dei cluster. È stato calcolato un valore percentuale di successo del modello per ciascuna delle tre informazioni considerate fondamentali per l'analisi:

	Pre-Agg.	Post-Agg.	Total
Customer Complaint	99%	98%	97%
Issue Cause	94%	96%	90%
Performed Test	97%	90%	87%

Tabella 3 – Livello di affidabilità del modello

Quindi, il valore medio è pari al 91,3%, questo significa che la probabilità che il sistema commetta degli errori nella fase di aggregazione delle informazioni in cluster è inferiore al 10%.

3. Passato e Futuro: Punti di Forza e di Debolezza

Prima dell'adozione dell'intelligenza artificiale, l'analisi delle claim era un processo lento e dispendioso in termini di tempo. In base ai dati raccolti tramite questionario, la lettura di una singola claim richiedeva in media tra uno e due minuti. La tabella sottostante riporta i tempi necessari per esaminare diverse quantità di reclami, espressi sia in minuti che in ore.

N. claim	Time	
	Mins	Hours
1	1,5	0,025
300	450	7,5
500	750	12,5
1000	1500	25

Tabella 4 – Tempi di lettura claim

È evidente come tale approccio comporti un considerevole spreco di risorse, dal momento che il personale risulta coinvolto nello svolgimento di attività ripetitive e gravose in termini di tempo ed energie. Dati questi risultati, si è ritenuto utile implementare una soluzione che prevedesse l'adozione di una tecnologia volta ad

ottimizzare l'analisi, consentendo di elaborare un volume di dati notevolmente superiore nel medesimo arco temporale. Se in precedenza, in una giornata, si poteva esaminare solo un numero ridotto di claim, oggi è possibile processarne molte di più. Il vantaggio guadagnato in termini di tempo permette di potersi dedicare ad altre attività che richiedono un'attenzione maggiore, aumentando il livello di efficienza generale del ramo della qualità in azienda. Si assiste, dunque, ad una riorganizzazione delle tempistiche dell'analisi, che si sposta dalla fase di lettura e categorizzazione manuale delle claim alla combinazione delle informazioni per configurare i problemi.

Questo miglioramento non si limita semplicemente ad un incremento del numero di dati analizzati, ma garantisce anche una maggiore completezza, con un conseguente aumento del tasso di completamento dell'analisi: grazie all'automazione, tutte le claim possono essere analizzate, senza omissioni o ipotesi basate su somiglianze presunte con reclami già esaminati. In passato, infatti, i vincoli di tempo obbligavano a selezionare solo una parte dei dati, compromettendo una visione complessiva e dettagliata delle problematiche riscontrate.

Se con la metodologia precedente le valutazioni che venivano effettuate potevano variare sulla base del livello di esperienza dell'operatore, con il nuovo approccio l'interpretazione delle informazioni è univoca ed inequivocabile, non soggetta ad interpretazioni.

Inoltre, in precedenza l'analisi era focalizzata principalmente sui cluster predominanti e le informazioni secondarie venivano trascurate, in base all'ambito di studio. È anche richiesto un numero di iterazioni minore, dovuto ad esempio al fatto che non è più necessario effettuare la traduzione da lingue sconosciute all'inglese, visto che il sistema ha tradotto tutti i commenti dei dealer.

In passato, il processo di analisi delle claim seguiva una sequenza manuale, articolata come segue:

1. Tramite un portale, si richiedevano le claim da analizzare;
2. Si estraevano manualmente le informazioni principali per ciascun tema rilevante;

3. Si categorizzavano i problemi e si procedeva con un'aggregazione manuale dei dati in cluster.

Con l'introduzione dell'IA, il flusso operativo è stato rivoluzionato:

1. Si richiedono le claim da analizzare tramite il medesimo portale;
2. Il file scaricato contiene già i cluster predefiniti e un riepilogo automatico dell'analisi, con relativi grafici;
3. Si richiede una fase di verifica finale da parte dell'operatore che assicura l'affidabilità del clustering eseguito dall'IA.

In [Appendice C](#), sono riportati, a scopo esemplificativo, tre grafici che mostrano la distribuzione dei cluster per ogni informazione estratta (customer complaint, issue cause, performed test). Sono dei grafici proposti al fine dell'analisi, in quanto risultano utili per comprendere quali cluster sono più ricorrenti e quali meno.

3.1. Come si articola la nuova procedura

Step 1

Si richiede al portale il file di claim di cui si ha bisogno.

Step 2

Con la nuova soluzione implementata, il file scaricato dagli utenti includerà **9 colonne aggiuntive** rispetto al formato originale. Queste colonne sono state progettate per facilitare l'analisi delle informazioni relative ai reclami e organizzare i dati in modo più efficace, come mostrato in figura 12.



...	Dealer Comments	Customer Complaint	Issue Cause	Performed Test	...
-----	-----------------	--------------------	-------------	----------------	-----

Figura 12 – Colonne aggiuntive

Le 9 colonne aggiuntive contengono le seguenti informazioni:

1. Colonna 1 (Customer Complaint): Riporta in modo chiaro il contenuto principale del reclamo espresso dal cliente;

2. Colonna 2 (Issue Cause): Identifica e descrive la causa del problema segnalato;
3. Colonna 3 (Performed Test): Documenta i test eseguiti per analizzare o risolvere il problema.

Queste tre colonne rappresentano l'individuazione delle informazioni fondamentali, senza alcuna elaborazione tramite clustering.

4. Colonna 4 (Customer Complaint_Clusterized): Classifica i reclami in cluster specifici e ben definiti;
5. Colonna 5 (Issue Cause_Clusterized): Organizza le cause dei problemi in base a categorie specifiche;
6. Colonna 6 (Performed Test_Clusterized): Raccoglie i diversi test che sono stati effettuati sui componenti in cluster.

Queste colonne rappresentano una prima clusterizzazione delle informazioni fondamentali. I cluster sono già ben definiti ma mantengono una categorizzazione dettagliata, con nomi che possono risultare simili.

7. Colonna 7 (Customer Complaint_Macro-Aggregated): Consolida i cluster delle lamentele in categorie più ampie e univoche;
8. Colonna 8 (Issue Cause_Macro-Aggregated): Raccoglie le cause dei problemi in cluster aggregati e privi di duplicati;
9. Colonna 9 (Performed Test_Macro-Aggregated): Fornisce una visione sintetica dei test eseguiti, con una classificazione più chiara.

Queste ultime tre colonne costituiscono la divisione in cluster definitiva, in cui i cluster iniziali, che presentavano ridondanze, sono stati macro-aggregati per garantire maggiore chiarezza e leggibilità.

La suddivisione in tre gruppi di informazioni permette all'operatore di accedere al livello di dettaglio necessario per l'analisi da svolgere. In caso di problematiche complesse da interpretare o analizzare, l'operatore può concentrarsi sul gruppo di colonne più utile, sfruttando la granularità dei dati a disposizione.

Questa organizzazione consente agli utenti di:

- Accedere rapidamente alle informazioni fondamentali attraverso le prime tre colonne.
- Analizzare i dati in dettaglio con i cluster specifici.
- Ottenere una visione d'insieme semplificata tramite i cluster macro-aggregati.

Step 3

Poiché il sistema presenta una probabilità di errore leggermente inferiore al 10% nella fase di creazione dei cluster, è richiesto all'operatore di eseguire un controllo per valutare, in modo generale, l'affidabilità dei risultati prodotti. Questo step consente di identificare eventuali anomalie e garantire la qualità dell'output prima di procedere con le fasi successive.

Questa fase di controllo è stata introdotta perché, sebbene il sistema abbia un margine di errore inferiore al 10% – una percentuale considerata accettabile ma comunque presente – si è ritenuto necessario un rapido check per garantire la massima attendibilità dell'analisi.

Inoltre, il nuovo file contiene un foglio Excel aggiuntivo che riporta due elementi:

1. **Grafici Automatici:** si tratta di due grafici che mostrano le combinazioni tra le informazioni, si hanno i customer complaints e i test raggruppati per issue cause. Quindi i grafici mostrano, per ogni causa radice, quante volte si presenta un determinato customer complaint, stessa cosa per i test.

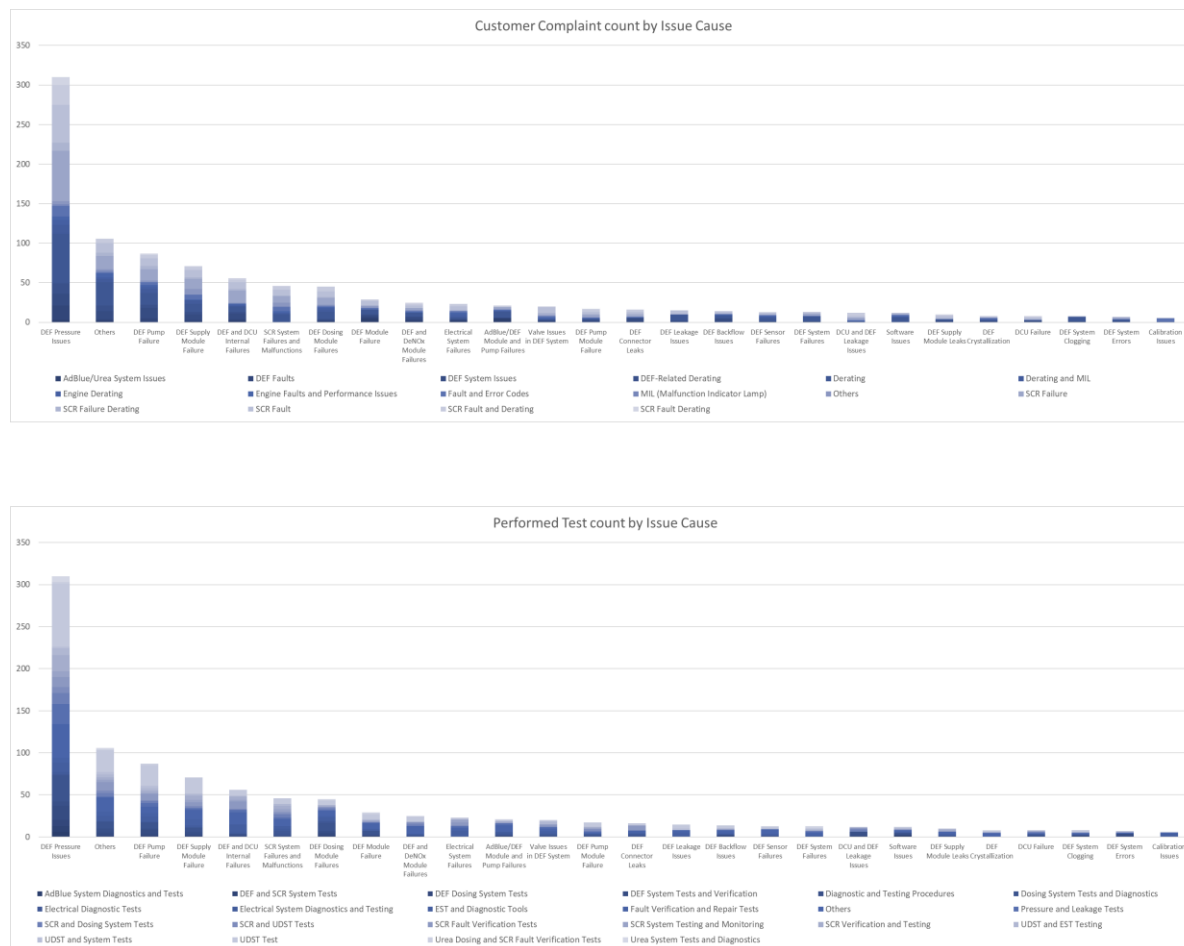


Figura 13 – Esempio grafici automatici

2. **Riassunto Automatico:** si tratta di un riassunto di qualche riga, in cui vengono riportate le varie occorrenze per ogni tipologia di informazione e le correlazioni più ricorrenti.

1. Principali Cause dei Problemi

Sulla base del numero totale di occorrenze, le principali cause dei problemi sono:

- a. Accumulo di idrocarburi (69 occorrenze)
- b. Ostruzione di SCR e catalizzatore (57 occorrenze)
- c. Bassa efficienza del catalizzatore (39 occorrenze)
- d. Problemi al sistema SCR (32 occorrenze)
- e. Altri (30 occorrenze)

2. Principali correlazioni tra le cause dei problemi e i reclami dei clienti

- a. Derating è il reclamo più comune tra tutti i problemi, con un totale di 154 occorrenze. Questo suggerisce che molte delle problematiche riscontrate stanno limitando le prestazioni dei veicoli.
- b. Guasti e malfunzionamenti del sistema SCR rappresentano il secondo reclamo più frequente (54 occorrenze), indicando una diffusa incidenza di problemi legati al sistema di Riduzione Catalitica Selettiva (SCR).
- c. Codici di errore e guasto sono il terzo reclamo più ricorrente (41 occorrenze), suggerendo che molte problematiche stanno attivando codici di errore diagnostico nel sistema di controllo del veicolo.
- d. L'accumulo di idrocarburi, la principale causa dei problemi, sembra essere correlato a diversi reclami, tra cui Derating, Codici di errore e guasto e Guasti e malfunzionamenti del sistema SCR.
- e. L'ostruzione di SCR e catalizzatore, seconda causa più comune, mostra una forte correlazione con i reclami di Derating.
- f. La bassa efficienza del catalizzatore, terza causa principale, è anch'essa strettamente correlata ai reclami di Derating.
- g. I problemi al sistema DEF/Urea e le attivazioni della spia MIL (Malfunction Indicator Light) sono meno frequenti ma comunque rilevanti, con rispettivamente 19 e 13 occorrenze. Questi problemi sono probabilmente legati al funzionamento dei sistemi di controllo delle emissioni.

Figura 14 – Esempio riassunto automatico

Capitolo 4

Applicazione su un caso reale: Pompa Urea

1. Obiettivo

Lo studio analizza un caso reale affrontato in azienda, focalizzandosi sulla gestione di un gruppo di claim relativi alla pompa urea, un componente dell'ATS. L'obiettivo è valutare l'impatto dell'introduzione dell'intelligenza artificiale nell'analisi dei reclami, confrontando l'approccio tradizionale con una nuova soluzione automatizzata.

L'analisi si basa sul confronto tra le due metodologie, evidenziando le differenze nell'identificazione dei cluster e verificando l'affidabilità della nuova soluzione attraverso un controllo successivo alla generazione del file dal portale. Oltre agli aspetti operativi, lo studio approfondisce anche i benefici economici derivanti dall'uso dell'intelligenza artificiale, con particolare attenzione all'aumento dell'efficienza e alla riduzione dei costi

2. Descrizione del componente

La pompa urea di un sistema di trattamento post-combustione, noto come After-Treatment System (ATS), rappresenta un elemento fondamentale per la gestione e la distribuzione dei fluidi necessari al funzionamento del sistema stesso. Associato spesso alla tecnologia di Riduzione Catalitica Selettiva (SCR), svolge un ruolo cruciale nella riduzione delle emissioni di ossidi di azoto (NOx) nei motori diesel, assicurando che il Diesel Exhaust Fluid (DEF), comunemente noto come AdBlue, venga dosato con precisione. Questo fluido, una soluzione di urea, viene iniettato nel sistema di scarico, dove reagisce con i gas di scarico, trasformando i NOx in azoto e acqua.

La pompa urea è progettata per garantire una distribuzione uniforme del fluido attraverso l'utilizzo di sensori e attuatori, migliorando così l'efficienza del catalizzatore SCR. Inoltre, opera in stretta sinergia con altre componenti del sistema ATS, come il filtro antiparticolato diesel (DPF), il catalizzatore di ossidazione diesel (DOC) e la camera di

decomposizione dell'urea, per ottimizzare la pulizia delle emissioni e la durata del sistema.²⁴

In Figura 15 è cerchiato il componente a cui ci si sta riferendo.

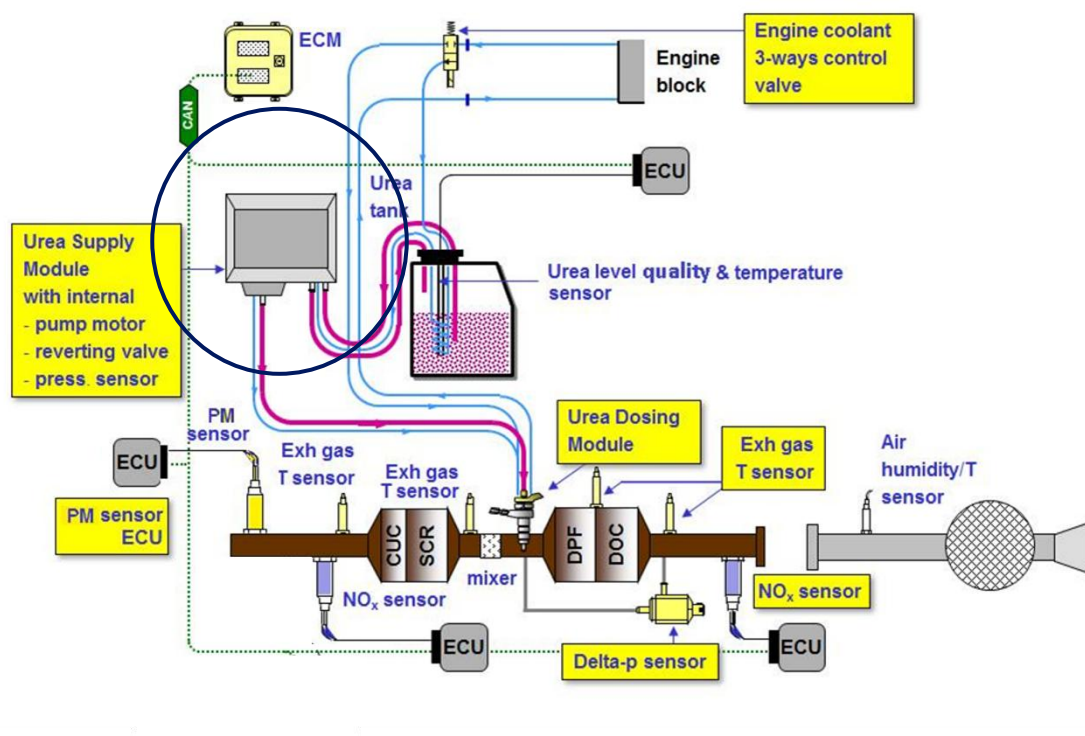


Figura 15 – Configurazione impianto ATS

3. Valutazione del livello di affidabilità

Come già riportato nella descrizione della procedura da seguire, il flusso dell'analisi prevede allo step 3 una fase di verifica, effettuata dall'operatore, del livello di affidabilità del raggruppamento dei dati in cluster effettuato dall'IA. In questa fase, viene mostrato il livello di plausibilità ottenuto, espresso in termini di percentuale di successo.

Dopo che l'operatore ha richiesto al sistema il file di interesse e ha ottenuto in output il file delle claim arricchito con le colonne relative ai cluster assegnati, viene verificata l'accuratezza della classificazione. Questa operazione si basa sull'analisi degli errori presenti in ciascun cluster, ovvero delle informazioni erroneamente assegnate a un

²⁴ (<https://highwayandheavyparts.com/what-are-supply-and-dosing-modules-a-guide-to-bosch-denoxtronic-systems/>, s.d.); (<https://www.cummins.com/components/aftertreatment/modular-aftertreatment-system>, s.d.)

gruppo anziché a quello corretto. In [Appendice D](#) sono riportate le percentuali di errore per ogni cluster, suddivise in base alle tre informazioni estratte: *Customer Complaint*, *Issue Cause* e *Performed Test*.

La fase di controllo è stata condotta a livello macro-aggregato, esaminando tutti i cluster generati dal sistema. Tuttavia, l'operatore ha la possibilità di effettuare verifiche mirate in base alle proprie esigenze. Ad esempio, se necessita di informazioni relative a un gruppo specifico di claim, come quelli associati a un determinato componente, può concentrare il controllo esclusivamente sulle categorizzazioni riferite a quel componente.

Come si può osservare dai risultati ricavati e mostrati in Appendice D, il cluster 'Others' presenta percentuali di errore più elevate per tutte e tre le informazioni considerate. Le ragioni per cui si presenta un tale livello di errore possono essere attribuite a diversi fattori:

- il cluster 'Others' è soggetto ad un'ampia variabilità interna, la mancanza di coerenza può aumentare le probabilità di errori;
- il cluster 'Others' contiene dati che presentano anomalie, proprio perché sono più difficili da classificare correttamente;
- alcuni dati appartenenti al cluster 'Others' potrebbero avere caratteristiche che li rendono simili a più cluster contemporaneamente.

Va, inoltre, specificato che sono stati considerati "errori" i dati che sarebbero dovuti appartenere ad un cluster e che, invece, il sistema ha assegnato ad altri cluster.

La Tabella 3 mostra il livello di incidenza del cluster 'Others' che ha generato l'Intelligenza Artificiale e la somma dei cluster vuoti e non chiari che l'operatore non è riuscito a identificare o a identificare parzialmente.

No IA		IA
Vuoto	Not Clear	Others
47%	22%	
69%		11%

Tabella 5 – Incidenza 'Others'

Questi risultati mostrano come il sistema dotato di IA sia in grado, in maniera automatica, di raggruppare informazioni che, per l'operatore, erano poco chiare o di difficile interpretazione.

La tabella 4 riporta le percentuali di successo per ogni informazione e la media totale, il risultato finale è un livello di affidabilità medio del 92%, quindi si rientra perfettamente negli obiettivi fissati ex-ante (percentuale di successo dell'80%).

Informazione	Livello di affidabilità
Customer Complaint	96%
Issue Cause	89%
Performed Test	90%
Valore medio	92%

Tabella 6 – Livello di Affidabilità

Si è visto dalla guida presentata nel capitolo precedente che il file di claim che viene richiesto al portale, oltre a riportare le colonne con i cluster già creati, genera un foglio con due grafici utili all'analisi e una panoramica generale delle occorrenze che si verificano nel caso analizzato. Per l'esempio esaminato, queste informazioni aggiuntive sono riportate di seguito: le figure 16 e 17 riportano i due grafici automatici generati.

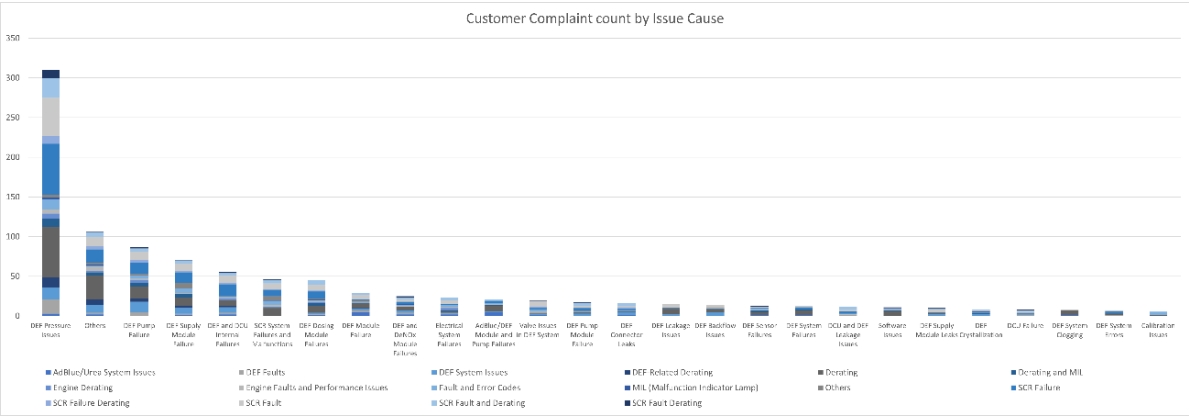


Figura 16 – Customer Complaint raggruppato per Issue Cause

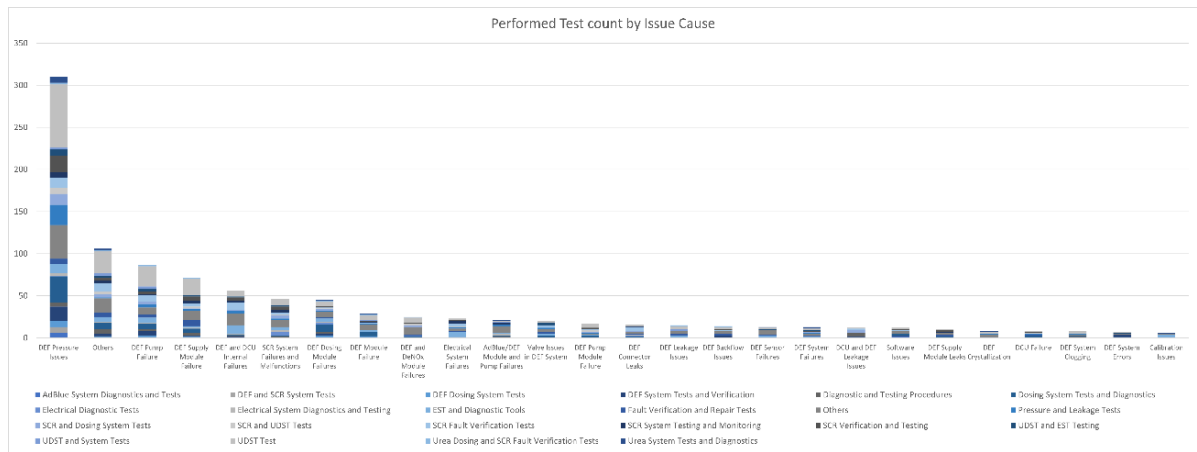


Figura 17 – Performed Test raggruppato per Issue Cause

Mentre, di seguito è riportato il riassunto delle occorrenze:

“1. Principali cause del problema: sulla base dei dati disponibili, le principali cause del problema sono: 1. DEF Pressure Issues 2. DEF Pump Failure 3. DEF Supply Module Failure 4. DEF and DCU Internal Failure 5. SCR System Failures and Malfunctions. È importante notare che questa analisi si basa su un campione limitato e che un’analisi dell’intero dataset potrebbe rivelare frequenze differenti. 2. Correlazioni principali tra le cause dei problemi e i reclami dei clienti: DEF Pressure Issues sono fortemente correlati con SCR Failure (64 occorrenze) e con Derating (63 occorrenze), mentre sono moderatamente correlati con SCR Fault (48 occorrenze). DEF Pump Failure è fortemente correlato con Derating (15 occorrenze), con SCR Failure (14 occorrenze) e con DEF System Issues (13 occorrenze), mentre è moderatamente correlato con SCR Fault (10 occorrenze). DEF Supply Module Failure è fortemente correlato con SCR Failure (13 occorrenze) e con Derating (10 occorrenze), mentre è moderatamente correlato con SCR Fault (9 occorrenze). DEF and DCU Internal Failures sono fortemente correlati con SCR Failure (15 occorrenze), moderatamente correlati con SCR Fault (9 occorrenze), con Derating (7 occorrenze) e con DEF System Issues (7 occorrenze). SCR System Failures and Malfunctions sono fortemente correlati con Derating (10 occorrenze), moderatamente correlati con SCR Failure (8 occorrenze) e con SCR Fault (7 occorrenze). 3. Principali intuizioni: Derating è il Customer Complaint più comune e si verifica in relazione a più cause, suggerendo che rappresenta un sintomo critico per diversi problemi. SCR Fault e SCR Failure sono anche reclami che si presentano più frequentemente, indicando che il

sistema SCR è spesso affetto da questi problemi. I problemi correlati al DEF (Diesel Exhaust Fluid) sembrano essere una causa significativa, apparendo in quattro delle cinque principali cause. I problemi relativi al DCU (Dosing Control Unit) sono anche prominenti, suggerendo che questo componente potrebbe essere un punto comune di guasto nel sistema. Queste intuizioni possono essere utili per dare priorità alle aree di miglioramento della qualità e per guidare ulteriori indagini sulle cause radice di questi problemi.”

4. Analisi economica

L’analisi di questo caso reale ha l’obiettivo di dimostrare come la nuova soluzione rappresenti un vantaggio economico concreto, quantificando il recupero ottenibile attraverso l’adozione di questa metodologia. L’implementazione del sistema ha portato a un miglioramento significativo nella gestione e nell’analisi dei reclami dei clienti, automatizzando alcune fasi del processo e riducendo il margine di errore. Questo ha avuto un impatto diretto sul recupero economico, permettendo all’azienda di ottimizzare tempi e risorse. Per valutare in modo accurato l’effettiva entità di questo beneficio, è necessario definire lo “spending”, ovvero il costo associato a ciascuna claim, elemento chiave per misurare il ricavo ottenuto.

4.1. Definizione di “spending”

Noto come “spending”, il costo che deve sostenere l’azienda per ogni claim ricevuta, viene riportato testualmente come “Total Amount with Dealer Net CCR Dollars”. Si tratta, infatti, del costo associato ad una claim, che è il risultato della somma di altri tre costi:

- Materiale: il costo dei componenti utilizzati per l’intervento;
- Manodopera: il costo delle ore lavorative impiegate per eseguire la riparazione o la sostituzione;
- Riparazioni extra: eventuali spese aggiuntive dovute a interventi straordinari o particolarmente complessi.

La somma di questi tre costi costituisce, per l’appunto, lo spending totale della claim, ovvero il valore complessivo che l’azienda deve sostenere per soddisfare un reclamo. Il costo da sostenere dipende dell’intervento effettuato dal dealer: ad esempio, la

riparazione di un pezzo, la pulizia, la sostituzione di un componente o, nei casi più estremi, dell'intero prodotto. Ogni tipologia di intervento ha un costo specifico, che varia in funzione della complessità e della natura del problema riscontrato.

Prima dell'introduzione dell'intelligenza artificiale, l'analisi e la gestione delle claim erano interamente affidate agli operatori, con tutte le limitazioni che questo comportava. Il principale rischio era che l'operatore potesse tralasciare qualche reclamo o interpretarlo in maniera errata, soprattutto nei casi di volume elevato o di descrizioni ambigue e incomplete. Inoltre, un'analisi non sempre accurata poteva portare a un'attribuzione errata delle responsabilità, con il risultato che i costi finivano spesso per ricadere sull'azienda, anche in situazioni in cui il fornitore avrebbe dovuto farsi carico della spesa. Questo generava perdite economiche evitabili, oltre a rallentare l'intero processo di gestione delle claim.

L'introduzione dell'intelligenza artificiale ha permesso di superare queste criticità, rendendo il processo più preciso ed efficiente. Grazie all'elaborazione avanzata dei dati, la soluzione implementata è in grado di individuare tutte le claim associate a un problema specifico, analizzando i commenti dei dealer, anche quelli più ambigui e di difficile interpretazione, e riconoscendo automaticamente pattern e parole chiave che potrebbero sfuggire a un operatore umano. Questo consente di attribuire correttamente le responsabilità, distinguendo i casi in cui il costo deve essere sostenuto dall'azienda da quelli in cui è imputabile a un fornitore. La riduzione degli errori nell'assegnazione dei costi si traduce in un recupero economico significativo, che prima dell'adozione della soluzione non era possibile ottenere.

Un ulteriore vantaggio è la capacità del sistema di fornire insight utili al miglioramento continuo del processo. Attraverso l'analisi automatizzata, è possibile individuare tendenze ricorrenti e intervenire in modo proattivo per prevenire problemi futuri. L'ottimizzazione della gestione delle claim porta quindi benefici concreti, sia in termini di riduzione delle perdite che di efficienza operativa.

4.2. Recupero economico

A conferma dell'efficacia della soluzione adottata, è stata condotta un'analisi per quantificare il potenziale recupero economico ottenibile grazie all'introduzione dell'intelligenza artificiale nel processo di gestione delle claim. Ogni claim correttamente individuata e analizzata rappresenta un valore che, in passato, rischiava di non essere riconosciuto e che ora potrebbe essere recuperato e riscattato nei confronti del fornitore. È importante sottolineare che la percentuale di recupero ottenuta non corrisponde a un ricavo netto immediato per l'azienda, ma esprime il valore economico teorico che può essere oggetto di trattativa e recupero.

L'analisi è stata condotta su un campione di 1000 claim, con un focus su due problematiche specifiche riguardanti la pompa urea dell'ATS. Per determinare il numero effettivo di claim recuperabili, è stato seguito un processo articolato in più fasi. Inizialmente, attraverso l'intelligenza artificiale, sono stati individuati i cluster costruiti sui problemi di interesse, aggregando automaticamente i dati e assegnando ciascuna claim alla categoria corrispondente. Successivamente, è stato effettuato un controllo di validazione, sfruttando i codici errore, per verificare se alcune claim fossero state erroneamente escluse dal sistema. Questa verifica ha permesso di affinare ulteriormente l'analisi, portando all'identificazione di un numero di claim superiore rispetto a quello individuato manualmente dall'operatore.

I risultati ottenuti sono riportati nella Tabella 5 e mostrano che il sistema AI ha rilevato:

	Problema 1		Problema 2	
	Con IA	Senza IA	Con IA	Senza IA
N. claim	162	148	170	144
Spending	271.724,00 €	248.242,11 €	229.447,55 €	194.355,57 €
Δ	23.481,89 €		35.091,98 €	
%	9,5%		18,1%	
Recupero medio	13,8%			

Tabella 7 – Risultati economici

- 162 claim relative al problema 1, rispetto alle 148 individuate manualmente;
- 170 claim relative al problema 2, rispetto alle 144 senza AI.

L'incremento del numero di claim individuate evidenzia la maggiore accuratezza del sistema nell'intercettare le anomalie, riducendo il rischio di errori o omissioni nell'analisi. Per quantificare il potenziale recupero economico, è stata calcolata la spesa media associata a ciascuna claim e applicata ai dati ottenuti. L'analisi ha evidenziato un incremento medio del 14% rispetto al metodo tradizionale, indicando il valore potenziale che l'azienda potrebbe recuperare grazie a una gestione più efficiente delle claim.

$$\% \text{ Recupero} = \frac{\textit{Spending con IA} - \textit{Spending senza IA}}{\textit{Spending senza IA}} \times 100$$

Questo risultato non solo conferma il ruolo dell'intelligenza artificiale nell'ottimizzazione del processo di analisi, ma ne sottolinea anche l'importanza strategica. La nuova soluzione, infatti, consente di ottenere una classificazione più accurata e affidabile, offrendo una base solida per eventuali trattative con i fornitori finalizzate al recupero economico delle anomalie individuate.

Capitolo 5

Valutazione Economica del Progetto: Costi e Ricavi

1. Struttura dell'analisi

Nel capitolo 3 è stata condotta un'analisi economica su un caso specifico, esaminando un business case specifico e valutando il recupero economico ottenuto grazie all'implementazione della nuova soluzione. Tuttavia, per avere una panoramica più completa del progetto e delle implicazioni in termini economici, è necessario considerare l'investimento iniziale, i costi di realizzazione e il recupero economico derivante da più casi analizzati.

Tuttavia, risulta complesso realizzare un'analisi quantitativa basata su valori esatti dei ricavi ottenuti, poiché il recupero economico varia in base a diversi fattori e non è possibile determinare con precisione la distribuzione delle diverse casistiche. Per questo motivo, si è scelto di adottare un approccio qualitativo, basato sull'analisi delle percentuali di recupero dei costi delle claim, come già evidenziato nel caso reale analizzato nel capitolo precedente. Questa scelta è motivata da diversi elementi:

- Elevata variabilità dei dati: il recupero economico varia notevolmente in base a fattori come la tipologia di claim, il fornitore coinvolto e le condizioni contrattuali; quindi, non è possibile determinare con precisione la distribuzione delle diverse casistiche;
- Impossibilità di definire una probabilità esatta: in alcuni casi, il costo viene interamente recuperato dal fornitore, in altri solo in parte, e talvolta rimane a carico dell'azienda. Tuttavia, non esiste un modello predittivo che permetta di assegnare una probabilità precisa a ciascuno di questi scenari;
- Focus sull'impatto strategico e operativo: l'obiettivo principale dell'analisi non è solo quantificare il beneficio economico in termini assoluti, ma dimostrare come la soluzione migliori la gestione delle claim, riducendo il rischio di errori e aumentando la capacità di recupero.

Infine, è importante sottolineare che l'investimento iniziale e i costi operativi risultano marginali rispetto ai benefici ottenuti, rafforzando ulteriormente la validità dell'approccio adottato.

2. Costi e tempi

Nell'analisi dei costi sostenuti, oltre al valore dell'investimento iniziale, sono stati presi in considerazione due approcci differenti per stimare costi operativi e tempi di esecuzione: approccio storico e approccio on-demand.

- **Approccio storico:** si tratta dell'estrazione dei casi che sono già stati analizzati; quindi, viene sostenuto un costo per estrarre claim che il sistema ha già in memoria, è stato considerato un arco temporale di cinque anni (dal 2020 al 2024);
- **Approccio on demand:** il sistema analizza nuove claim, mai elaborate prima. In questo caso il processo richiede il download e la gestione di nuovi dati. Con un impatto diverso in termini di tempi e costi.

Quando viene avviata un'analisi, il sistema verifica automaticamente se le claim sono già state trattate:

- Se sì, utilizza i dati esistenti (approccio storico);
- Se no, avvia il processo di acquisizione e analisi (approccio on-demand).

I tempi di esecuzione del processo variano tra i due approcci, nelle tabelle di seguito sono riportati i tempi necessari dall'avvio della richiesta all'ottenimento del risultato.

Approccio on demand			
Tempo Processo senza IA (minuti)	Totale Processo con IA (minuti)	Totale Tempo (minuti)	Totale Tempo(ore)
40	20	60	1

Tabella 8 – Tempi approccio on demand

Approccio storico			
Tempo Processo senza IA (minuti)	Totale Processo con IA (minuti)	Totale Tempo (minuti)	Totale Tempo (ore)
40	1560	1600	26,7

Tabella 9 – Tempi approccio storico

È stato valutato il tempo di elaborazione con la procedura tradizionale, ovvero senza l'utilizzo dei cluster automatici basati sull'IA. A questo valore si somma il tempo necessario affinché il sistema generi il file con le colonne aggiuntive.

La differenza nei tempi di elaborazione tra i due approcci dipende dal volume di dati trattato:

- Con l'approccio on-demand, il sistema genera solo i file relativi alle claim specificamente richieste dall'operatore;
- Con l'approccio storico, invece, vengono elaborate tutte le claim relative al periodo di analisi considerato (cinque anni), aumentando così il tempo di elaborazione.

Anche i costi cambiano in base al tipo di approccio, come si può vedere dalle tabelle 6 e 7.

Approccio on demand		
Costo per claim	Numero claim	Costo totale
0,005 €	1.000	4,720 €

Tabella 10 – Costi approccio on demand

Approccio storico			
	Costo per claim	Numero claim	Costo totale
Analisi reclami (una tantum)	0,005 €	3.807.000	17.969,040 €
Analisi claim (al mese)	0,005 €	104.000	490,880 €

Tabella 11 – Costi approccio storico

Nell'approccio on-demand, il costo è proporzionale al numero di claim da analizzare: ad esempio, per una richiesta di mille claim, si sostiene esattamente il costo corrispondente a quelle mille claim.

Per l'approccio storico, si sostiene un costo fisso iniziale relativo al totale dei reclami analizzati nei cinque anni considerati e un ulteriore costo mensile, avendo stimato che il numero di claim da analizzare sia in media 104.000.

La differenza tra i due approcci risiede nel costo fisso richiesto dall'approccio storico. In questo caso, il sistema analizza in un'unica elaborazione tutti i reclami degli ultimi cinque anni, memorizzandone i dati. Questo consente, in caso di un nuovo reclamo simile, di effettuare un'analisi più rapida, attingendo direttamente alle informazioni già archiviate.

Oltre ai costi sostenuti per lo scarico delle claim, va considerato anche il costo dello sviluppo degli algoritmi utilizzati per implementare la soluzione. Questa spesa corrisponde al compenso versato all'ente terzo incaricato di sviluppare gli algoritmi, che hanno reso possibile la creazione della soluzione, includendo grafici e cluster. Il valore ammonta a 84.960,00 €. Dunque, considerando i costi sostenuti ex-ante, ovvero i costi fissi precedenti all'introduzione della soluzione in azienda e prima della richiesta di nuove claim da analizzare, il totale delle spese da sostenere è riportato di seguito in figura 18.

COSTI EX-ANTE	
Sviluppo algoritmi (una tantum)	84.960,00 €
Analisi reclami (una tantum)	17.969,040 €
TOTALE	102.929,040 €

Tabella 12 – Totale costi

3. Recupero economico e benefici

Il ricavo economico ottenuto non può essere quantificato con un valore assoluto, ma va considerato in termini di delta valore delle claim. A questo scopo è stata condotta un'analisi su un altro caso studio, relativo a un problema riguardante un componente dell'ATS. Si è ritenuto opportuno studiare un altro problema relativo ad un altro caso rispetto a quello già studiato per la pompa urea, per confermare e avvalorare i dati già ottenuti, in questo modo si ha una visione più completa dei risultati che si possono ottenere con la nuova soluzione. Quindi, si confrontano le percentuali di recupero per ciascun problema dei tre problemi analizzati e dalla media dei tre valori considerati è emerso un valore di recupero pari al 13,7%.

Il sistema è in grado di individuare un numero di claim superiore rispetto a quelli rilevati dall'operatore per uno specifico problema. Sulla base dei costi associati a tali reclami, si calcola il valore aggiunto ottenuto in termini percentuali, definito come valore di recupero. Le percentuali corrispondenti sono riportate nella tabella seguente.

	Problema 1	Problema 2	Problema 3
Spending	23.481,89 €	35.091,98 €	8.992,02 €
% recupero	9,5%	18,1%	13,6%
Valore medio	13,7%		

Tabella 13 – Recupero %

I dati presentati si riferiscono a tre problemi appartenenti a gruppi di claim distinti. Tuttavia, il numero di claim ricevuti settimanalmente è molto elevato e ogni gruppo può includere numerose problematiche segnalate dai dealer. Considerando che potrebbero verificarsi molteplici problemi su altrettanti casi di claim, i valori di spending riportati in tabella 11 dovrebbero essere moltiplicati per tutti i casi settimanali. Questo porterebbe a cifre nettamente superiori rispetto ai costi di realizzazione del progetto, rendendo un'analisi economica dettagliata poco significativa, dato il divario tra costi e ricavi.

Infine, le percentuali di recupero non rappresentano ricavi diretti per l'azienda, ma piuttosto il potenziale recupero di valore dal fornitore. I valori riportati in tabella indicano il massimo teoricamente recuperabile, ma il valore effettivo può variare tra 0% e 13,7%, a seconda dei singoli casi.

Capitolo 6

Sviluppi Futuri - Use case 2: Self-service BI Platform

1. Obiettivo

L'obiettivo di questo secondo use case è implementare una soluzione che riduca il tempo necessario per navigare tra i dati, generare grafici e produrre report basati sul database WRS e su altri database correlati (elaborati nello Use Case 1). L'utente si aspetta di interagire con il sistema attraverso un linguaggio naturale, con un livello di consapevolezza sui domini di dati e sulla tecnologia, e di poter gestire in modo critico l'output generato.

Gli strumenti di Business Intelligence attualmente disponibili sono progettati principalmente per i data analyst. Pertanto, è fondamentale mantenere un controllo tecnico sul processo e sull'output del modello per garantire l'accuratezza delle informazioni prodotte. L'output generato su richiesta dell'utente deve essere revisionato sia dal punto di vista tecnico che aziendale.

2. Applicazione del modello – Qlik Sense

Qlik Sense è una piattaforma di Business Intelligence (BI) sviluppata dall'azienda svedese Qlik, fondata nel 1993. Rilasciata nel 2014, rappresenta l'evoluzione di QlikView, con l'obiettivo di offrire uno strumento più moderno, intuitivo e accessibile anche agli utenti senza particolari competenze tecniche. Si tratta di una soluzione di BI self-service, che consente di analizzare e visualizzare dati provenienti da diverse fonti in modo interattivo, attraverso dashboard e report personalizzabili.

L'utilizzo di Qlik Sense è estremamente intuitivo grazie alla sua interfaccia drag-and-drop, che permette agli utenti di creare report e analisi senza necessità di scrivere codici complessi. Il processo inizia con il caricamento dei dati, che possono provenire da file di diverso formato (come Excel, CSV, XML e JSON), database SQL e NoSQL, oppure servizi cloud. Una volta importati i dati, la piattaforma consente di esplorarli dinamicamente, creare grafici e visualizzazioni interattive e applicare filtri per ottenere analisi più dettagliate. Inoltre, grazie alla sua architettura associativa, Qlik Sense permette di

scoprire correlazioni e tendenze che potrebbero non emergere con strumenti di analisi tradizionali.

Uno dei principali vantaggi di Qlik Sense è la sua capacità di integrare ed elaborare grandi volumi di dati, offrendo agli utenti una visione chiara e immediata delle informazioni. Inoltre, la piattaforma fornisce strumenti avanzati come la modellizzazione predittiva, l'analisi delle serie storiche e la possibilità di definire regole di business, migliorando ulteriormente l'efficacia delle analisi.

Un altro punto di forza di Qlik Sense è la collaborazione. Gli utenti possono condividere report e dashboard con colleghi e team di lavoro, favorendo un processo decisionale più rapido e basato su dati concreti. Inoltre, la possibilità di personalizzare le visualizzazioni e di utilizzare estensioni esterne consente di adattare lo strumento alle esigenze specifiche di ogni organizzazione.

Infine, Qlik Sense è pensato per essere utilizzato in diversi contesti aziendali, sia in ambienti desktop che in cloud, garantendo flessibilità e scalabilità. La sua capacità di adattarsi a diverse esigenze e di semplificare l'analisi dei dati lo rende uno strumento potente per tutte le realtà che vogliono migliorare il processo decisionale basandosi su informazioni affidabili e facilmente accessibili.²⁵

3. Come funziona la soluzione

Questo secondo use case mira alla realizzazione di un sistema avanzato in grado di combinare e analizzare i dati ricevuti in input in modo dinamico e interattivo. L'obiettivo è fornire agli operatori uno strumento che consenta di sfruttare al meglio i risultati dell'intelligenza artificiale per ottenere insights strutturati e immediatamente utilizzabili.

In particolare, il sistema elabora il file generato dall'intelligenza artificiale, che contiene i dati di base arricchiti dalle colonne aggiuntive relative ai cluster già realizzati (frutto dello studio dello Use Case 1). Una volta caricato questo file nella piattaforma digitale, il sistema organizza e combina le informazioni per restituire tabelle strutturate in base alle richieste specifiche dell'operatore. Quest'ultimo può interagire con la

²⁵ (Melia, 2023); (<https://www.qlik.com/it-it/products/qlik-sense>, s.d.)

piattaforma ponendo domande in linguaggio naturale, specificando le informazioni che intende correlare e il tipo di analisi di cui ha bisogno. Il sistema interpreta la richiesta e genera in output sia le tabelle con i dati combinati sia rappresentazioni grafiche che ne evidenziano l'andamento nel tempo o altre correlazioni rilevanti.

Un ulteriore sviluppo futuro della piattaforma potrebbe essere l'ottimizzazione dell'usabilità del sistema e la riduzione del tempo necessario per ottenere risposte. È stata inoltre prevista la creazione di dashboard preimpostate, progettate sulla base delle analisi più frequenti e rilevanti per gli operatori. Queste dashboard consentono agli utenti di accedere rapidamente alle informazioni di cui hanno bisogno, senza dover formulare ogni volta nuove richieste. L'operatore può così selezionare la dashboard più adatta al proprio scopo e inserire i dati disponibili per visualizzare immediatamente i risultati.

Attualmente, la soluzione è ancora in fase di test, con l'obiettivo di valutare il funzionamento del sistema e affinare la sua capacità di interpretare correttamente le richieste degli operatori. Si stanno esaminando diversi aspetti, tra cui la modalità con cui vengono poste le domande, la correttezza delle risposte fornite dal sistema e la coerenza con cui i dati vengono combinati in base ai criteri richiesti. L'analisi di questi test consentirà di migliorare ulteriormente la piattaforma, garantendo un'interazione sempre più fluida e precisa tra l'operatore e il sistema di analisi dati.

In [Appendice E](#) sono riportati gli output delle analisi che l'operatore può richiedere alla piattaforma, inclusi sia tabelle con combinazioni di informazioni che grafici utili per individuare le possibili cause dei problemi di qualità.

Inoltre, dalla proposta di unificare alcune di queste analisi, potrebbero nascere delle dashboard preimpostate, particolarmente utili per effettuare analisi standard, che si presentano più frequentemente rispetto ad altre.

Conclusioni

L'implementazione del sistema descritto in questa tesi ha portato a un'evoluzione significativa nel processo di gestione delle claim, migliorandone l'efficienza, la precisione e l'affidabilità. Il progetto è ormai operativo in azienda, con gli operatori che hanno accesso al portale e possono beneficiare delle funzionalità avanzate offerte dal sistema. Per ogni team del dipartimento qualità è stato individuato un referente incaricato di consultare la piattaforma, arricchita dal commento generato dall'intelligenza artificiale. Questa innovazione non solo ha semplificato il lavoro degli operatori, ma ha anche introdotto un approccio più strutturato all'analisi dei reclami, con vantaggi evidenti sia in termini di riduzione dello sforzo richiesto sia di ottimizzazione del recupero economico.

In passato, l'analisi delle claim era un'attività estremamente onerosa, che prevedeva diversi passaggi manuali: la lettura e l'interpretazione di ogni singolo reclamo, la traduzione quando necessario, l'estrazione delle informazioni chiave, la categorizzazione dei problemi e l'aggregazione manuale in cluster. Solo dopo aver completato queste fasi, l'operatore poteva avviare un'analisi approfondita. L'introduzione del nuovo sistema ha rivoluzionato questo processo, automatizzando la fase di clustering e consentendo agli operatori di concentrarsi direttamente sulla verifica dei risultati e sull'analisi delle problematiche.

Uno degli aspetti più rilevanti è la riduzione della soggettività nell'interpretazione delle claim. L'elaborazione automatizzata su tutti i commenti disponibili elimina il rischio che alcune informazioni vengano trascurate, un limite tipico dell'approccio tradizionale basato sull'esperienza dell'operatore. In questo modo, il sistema permette di individuare un numero maggiore di reclami riconducibili alla stessa problematica, aumentando la qualità e l'affidabilità dell'analisi. Inoltre, il margine di errore inferiore al 20% assicura che la categorizzazione automatizzata sia sufficientemente precisa da poter essere validata con un intervento minimo da parte degli operatori.

I principali vantaggi ottenuti con questa soluzione includono:

- Maggiore accuratezza nell'identificazione delle problematiche, con un incremento medio di circa venti claim per categoria rispetto al metodo tradizionale, determinando un impatto positivo sul recupero economico;
- Maggiore efficienza, grazie alla riduzione del tempo dedicato alla fase di aggregazione dei dati, permettendo agli operatori di concentrarsi direttamente sull'analisi e sulla risoluzione dei problemi;
- Eliminazione della soggettività nell'interpretazione delle claim, garantendo un'analisi più uniforme e affidabile attraverso l'elaborazione automatizzata.

L'implementazione di questa soluzione rappresenta un significativo passo avanti nell'ottimizzazione del processo di gestione dei reclami, rendendolo più rapido, preciso e strutturato. Questo caso dimostra come l'intelligenza artificiale non si limiti a migliorare l'efficienza operativa, ma possa anche trasformare il ruolo degli operatori, consentendo loro di dedicarsi a compiti a più alto valore aggiunto. La capacità di analizzare rapidamente grandi volumi di dati e di individuare correlazioni in modo sistematico permette di affrontare le problematiche con un approccio più strategico, migliorando la qualità complessiva delle decisioni aziendali.

L'integrazione di soluzioni basate sull'intelligenza artificiale nei processi aziendali è destinata a diventare sempre più centrale. Questo progetto dimostra come l'adozione di strumenti innovativi possa non solo aumentare l'efficienza, ma anche ridurre la variabilità delle analisi e migliorare la capacità di risposta dell'azienda alle problematiche emergenti. Il percorso di innovazione intrapreso in questo caso potrebbe essere ulteriormente sviluppato con l'introduzione di nuove funzionalità, come l'integrazione con modelli predittivi o l'automazione di ulteriori fasi dell'analisi. In definitiva, l'adozione dell'intelligenza artificiale non rappresenta solo un miglioramento del processo, ma un'opportunità concreta per ripensare il modo in cui i dati vengono utilizzati per supportare le decisioni aziendali.

Appendice A

Questionario

Link:

<https://forms.office.com/Pages/ResponsePage.aspx?id=BblMYpEg5EGQuedozyKFGkGMidgCYltHncZLFFuBvFFURVRKUEo4V0NCMktZMkRGMjg3MURUSDIBTC4u>

Testo:

Introduzione IA

Benefici che si ottengono dall'introduzione dell'IA per la lettura dei reclami

1. Da quanto tempo lavori per FPT Industrial?
 - ☐ Meno di 1 anno
 - ☐ Tra 1 e 5 anni
 - ☐ Più di 5 anni

2. Quante claim analizzi, in media, alla settimana?
 - ☐ Meno di 100
 - ☐ 100-200
 - ☐ 200-300
 - ☐ 300-400
 - ☐ 400-500
 - ☐ 500+

3. Quanto tempo impieghi a leggere una claim/commento in media?
 - ☐ Meno di un minuto
 - ☐ Tra 1 e 5 minuti
 - ☐ Più di 5 minuti

4. A quale scopo leggi i reclami?
- Configurazione del problema
 - Identificazione della causa radice
 - Impatto sul cliente
5. Quali informazioni vuoi estrarre di solito?
- Customer complaint (il problema che il cliente nota)
 - Issue cause (il problema che ha generato il reclamo del cliente)
 - Performed test/Troubleshooting (cosa è stato fatto per risolvere il problema)
 - Errori
6. Quali difficoltà incontri maggiormente nella lettura dei reclami?
- Lingua conosciuta
 - Informazioni superflue sul viaggio (indirizzi, costi)
 - Dettagli eccessivi
7. Quanto pensi che sia utile l'introduzione dell'Intelligenza Artificiale per la lettura delle claim?
1. - Inutile
 - 2.
 - 3.
 - 4.
 5. – Estremamente utile
8. Se hai qualche suggerimento riguardo le domande del sondaggio, che ritieni possa essere utile per condurre l'analisi, scrivilo di seguito:

Appendice B

Fase di controllo dei cluster

HDBscan - Tier 4B Muffler												
	Customer complaint				Issue cause				Performed test			
	N. Cluster	Errori	Totale	Valore %	N. Cluster	Errori	Totale	Valore %	N. Cluster	Errori	Totale	Valore%
	-1	6	12	50%	-1	11	30	37%	-1	9	30	30%
	0	1	29	3%	0	0	12	0%	0	0	18	0%
	1	0	25	0%	1	0	39	0%	1	0	32	0%
	2	0	122	0%	2	2	24	8%	2	0	6	0%
	3	0	13	0%	3	0	7	0%	3	0	25	0%
	4	0	18	0%	4	0	7	0%	4	0	16	0%
	5	0	14	0%	5	0	17	0%	5	0	12	0%
	6	0	20	0%	6	0	19	0%	6	0	16	0%
	7	0	19	0%	7	1	7	14%	7	2	11	18%
	8	0	8	0%	8	0	13	0%	8	0	11	0%
	9	1	7	14%	9	0	12	0%	9	2	19	11%
	10	0	13	0%	10	4	25	16%	10	0	14	0%
					11	2	8	25%	11	0	9	0%
				12	1	14	7%	12	0	14	0%	
				13	2	8	25%	13	1	17	6%	
				14	0	9	0%	14	4	15	27%	
				15	0	13	0%	15	0	20	0%	
				16	0	6	0%	16	2	10	20%	
				17	0	8	0%	17	0	5	0%	
				18	3	22	14%					
Totale	12	2	288	1%	20	15	261	6%	19	7	255	3%

Tabella A. 1 – % Errore HDBScan (Tier 4B Muffler)

Spectral - Tier 4B Muffler												
	Customer complaints				Issue cause				Performed test			
	N. Cluster	Errori	Totale	Valore %	N. Cluster	Errori	Totale	Valore %	N. Cluster	Errori	Totale	Valore %
	0	21	80	26%	0	7	15	47%	0	0	17	0%
	1	0	26	0%	1	6	101	6%	1	0	126	0%
	2	0	118	0%	2	10	21	48%	2	2	41	5%
	3	0	19	0%	3	0	11	0%	3	0	40	0%
	4	0	23	0%	4	1	18	6%	4	2	30	7%
	5	0	13	0%	5	0	12	0%	5	5	16	31%
	6	0	21	0%	6	6	21	29%	6	0	12	0%
					7	3	20	15%	7	0	18	0%
					8	0	28	0%				
					9	0	12	0%				
					10	1	22	5%				
				11	0	19	0%					
Totale	7	0	220	0%	12	34	300	11%	8	9	174	5%

Tabella A. 2 – % Errore Spectral (Tier 4B Muffler)

HDBScan - EGR Valve F5												
	Customer complaint				Issue cause				Performed test			
	N.cluster	Errori	Totale	Valore %	N. cluster	Errori	Totale	Valore %	N. cluster	Errori	Totale	Valore %
	-1	8	16	50%	0	1	24	4%	-1	0	13	0%
	0	0	75	0%	1	0	26	0%	0	0	42	0%
	1	0	10	0%	2	1	13	8%	1	0	18	0%
	2	0	7	0%	3	1	25	4%	2	1	14	7%
	3	0	15	0%	4	0	20	0%	3	0	6	0%
	4	1	10	10%	5	7	27	26%	4	3	10	30%
	5	1	8	13%	6	0	7	0%	5	2	21	10%
	6	0	7	0%	7	2	26	8%	6	0	16	0%
	7	3	8	38%					7	1	6	17%
	8	0	5	0%					8	0	10	0%
	9	0	7	0%					9	1	12	8%
Totale	11	5	152	3%	8	5	141	4%	11	5	135	4%

Tabella A. 3 – % Errore HDBScan (EGR Valve F5)

Spectral - EGR Valve F5												
	Customer complaints				Issue cause				Performed test			
	N. cluster	Errori	Totale	Valore %	N. cluster	Errori	Totale	Valore %	N. cluster	Errori	Totale	Valore %
	0	9	93	10%	0	0	20	0%	0	0	126	0%
	1	0	75	0%	1	2	23	9%	1	0	42	0%
					2	1	25	4%				
					3	0	16	0%				
					4	1	13	8%				
					5	1	24	4%				
					6	3	13	23%				
					7	2	10	20%				
					8	4	14	29%				
					9	0	10	0%				
Totale	2	0	75	0%	10	12	158	8%	2	0	42	0%

Tabella A. 4 – % Errore Spectral (EGR Valve F5)

N.B. Le celle evidenziate rappresentano i cluster etichettati dal sistema come “Others”.
Le percentuali di errore sono state calcolate escludendo questi cluster.

Appendice C

Grafici proposti

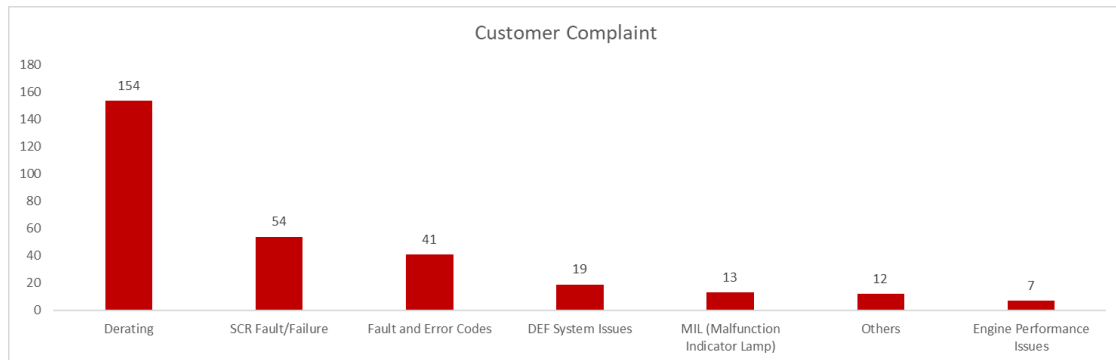


Figura A. 1 – Distribuzione Customer Complaint

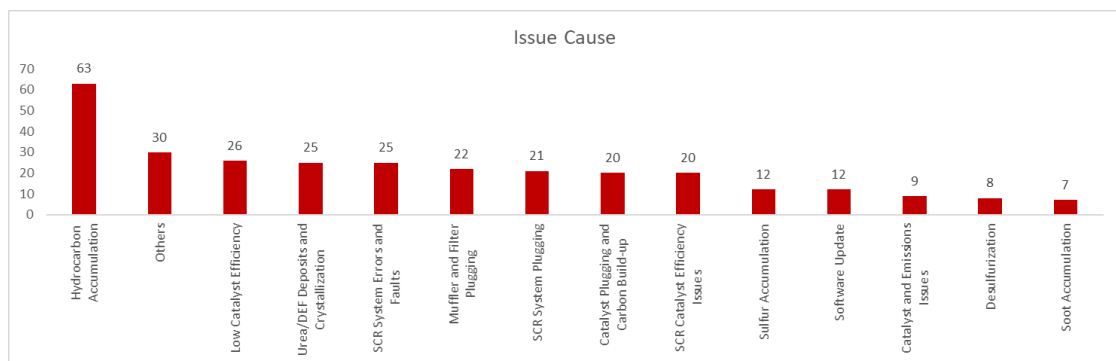


Figura A. 2 – Distribuzione Issue Cause

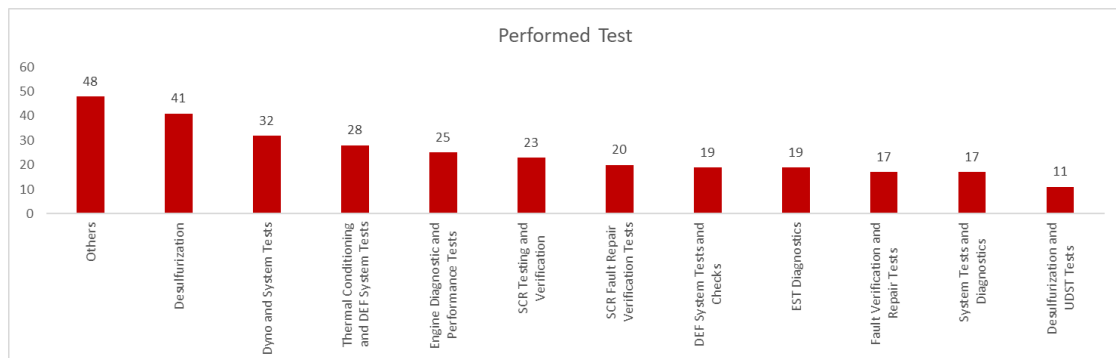


Figura A. 3 – Distribuzione Performed Test

Appendice D

Calcolo livello di affidabilità per ogni informazione

Customer Complaint			
Cluster	Errori	Totale	Valore %
AdBlue/Urea System Issues	0	28	0%
DEF Faults	0	38	0%
DEF System Issues	12	82	15%
DEF-Related Derating	1	30	3%
Derating	0	205	0%
Derating and MIL	0	41	0%
Engine Derating	0	19	0%
Engine Faults and Performance Issues	6	27	22%
Fault and Error Codes	4	49	8%
MIL (Malfunction Indicator Lamp)	0	9	0%
Others	15	26	58%
SCR Failure	0	166	0%
SCR Failure Derating	0	28	0%
SCR Fault	0	152	0%
SCR Fault and Derating	2	74	3%
SCR Fault Derating	0	24	0%
TOTALE	40	998	4%

Tabella A. 5 – Customer Complaint

Issue Cause			
Cluster	Errori	Totale	Valore %
AdBlue/DEF Module and Pump Failures	0	21	0%
Calibration Issues	1	6	17%
DCU and DEF Leakage Issues	0	12	0%
DCU Failure	0	8	0%
DEF and DCU Internal Failures	0	56	0%
DEF and DeNOx Module Failures	0	25	0%
DEF Backflow Issues	0	14	0%
DEF Connector Leaks	0	16	0%
DEF Crystallization	0	8	0%
DEF Dosing Module Failures	0	45	0%
DEF Leakage Issues	0	15	0%
DEF Module Failure	0	29	0%
DEF Pressure Issues	5	310	2%
DEF Pump Failure	0	87	0%
DEF Pump Module Failure	0	17	0%
DEF Sensor Failures	1	13	8%
DEF Supply Module Failure	0	71	0%
DEF Supply Module Leaks	0	10	0%
DEF System Clogging	1	8	13%
DEF System Errors	0	7	0%
DEF System Failures	1	13	8%
Electrical System Failures	3	23	13%
Others	85	106	80%
SCR System Failures and Malfunctions	1	46	2%
Software Issues	0	12	0%
Valve Issues in DEF System	7	20	35%
TOTALE	105	998	11%

Tabella A. 6 –Issue Cause

Performed Test			
Cluster	Errori	Totale	Valore %
AdBlue System Diagnostics and Tests	0	13	0%
DEF and SCR System Tests	0	11	0%
DEF Dosing System Tests	1	15	7%
DEF System Tests and Verification	8	39	21%
Diagnostic and Testing Procedures	0	26	0%
Dosing System Tests and Diagnostics	14	85	16%
Electrical Diagnostic Tests	0	9	0%
Electrical System Diagnostics and Testing	1	15	7%
EST and Diagnostic Tools	7	65	11%
Fault Verification and Repair Tests	0	39	0%
Others	60	161	37%
Pressure and Leakage Tests	4	49	8%
SCR and Dosing System Tests	0	36	0%
SCR and UDST Tests	0	20	0%
SCR Fault Verification Tests	0	63	0%
SCR System Testing and Monitoring	0	27	0%
SCR Verification and Testing	0	51	0%
UDST and EST Testing	0	25	0%
UDST and System Tests	3	12	25%
UDST Test	0	209	0%
Urea Dosing and SCR Fault Verification Tests	0	11	0%
Urea System Tests and Diagnostics	0	17	0%
TOTALE	98	998	10%

Tabella A. 7 – Performed Test

Appendice E

Esempi analisi Qlik Sense

Failure Code for Components			
Failure Index Code	Q	Component Code	Q
220		10500AR	
220		10500AO	
220		55000AA	
220		10501AE	
220		10500AU	
220		10001AD	
220		10254AI	
220			1050099
220			5598899
220		10AAAAA	

Figura A. 4 – Combinazione Failure Code e Componenti

Frequency of Breakdowns by Worked Hours			
WorkedHoursBucket	Q	#Faults	
Totals		300	
0		59	
200		69	
400		60	
600		29	
800		18	
1000		9	
1200		8	
1400		14	
1600		8	
1800		5	

Figura A. 5 – Combinazione Worked Hours e Guasti

Average worked hours per production year		
	Production Year 	Avg([Worked Hours\ Mileage km])
Totals		710.88
	2014	2046.6667
	2015	881.5
	2016	909.90476
	2017	770.4717
	2018	682.72308
	2019	618.7069
	2020	426.57895
	2021	633.94118
	2022	341.11765

Figura A. 6 – Combinazione Worked Hours e Production Year



Figura A. 7 – Distribuzione Claims nel mondo

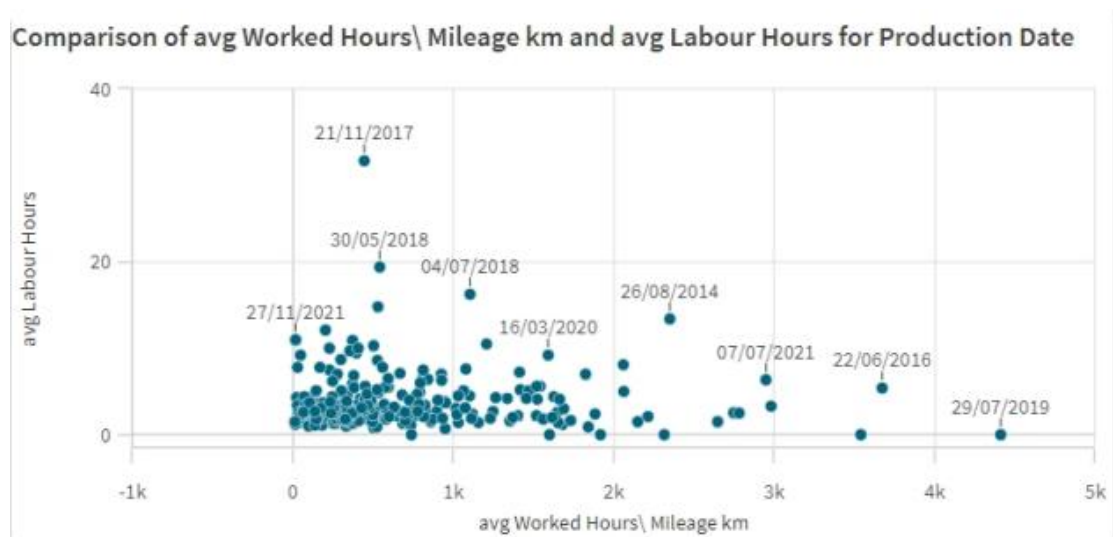


Figura A. 8 – Confronto tra Worked Hours e Labour Hour per Production Date

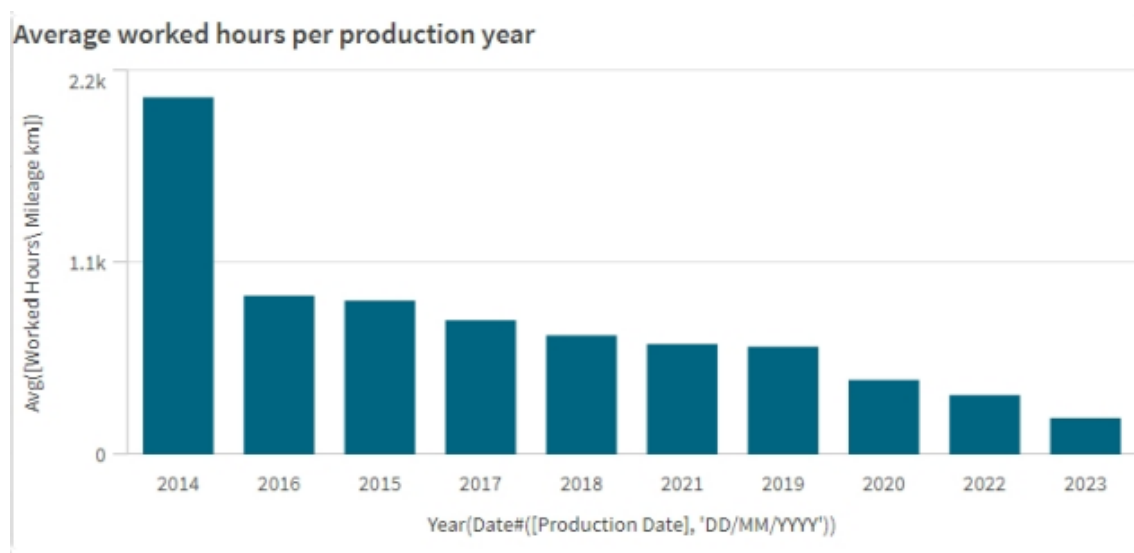


Figura A. 9 – Distribuzione Worked Hours per Production Year

Bibliografia e Sitografia

(s.d.). Tratto da https://www.ivecogroup.com/group/about_us/our_history.

(s.d.). Tratto da https://www.fptindustrial.com/en/?sc_lang=it.

Adaveesh, T. S. (2017). Application of “8D Methodology” for the Root Cause Analysis.

Aldieri, I. (2024). Problem Solving: Significato, Strumenti e Metodologie. *Carriera Vincente*.

FPT Industrial. (2023). Intro to Quality. Torino.

<https://fractionalmanageritalia.it/operation/kpi-operations-sai-quali-dovresti-monitorare/>. (2024). Tratto da <https://fractionalmanageritalia.it/operation/kpi-operations-sai-quali-dovresti-monitorare/>

<https://highwayandheavyparts.com/what-are-supply-and-dosing-modules-a-guide-to-bosch-denoxtronic-systems/>. (s.d.).

<https://www.cummins.com/components/aftertreatment/modular-aftertreatment-system>. (s.d.).

<https://www.qlik.com/it-it/products/qlik-sense>. (s.d.).

LBi. (2024). 5W2H: A Simple Tool with Limitless Possibilities for Problem-Solving.

Lovergine, S. (2022). *Breve disamina degli algoritmi di intelligenza artificiale*. Roma.

MeeTheSkilled. (2018). Problem solving: utilizzo della tecnica dei 5 perchè.

Melia, C. (2023). Introduzione a Qlik Sense, piattaforma di Business Intelligence avanzata. *Orbyta Technologies*.

Morriello, R. (2023). *OpenAI e ChatGPT: funzionalità, evoluzione e questioni aperte*.

Realyvásquez-Vargas, A. (2020). *Improving a Manufacturing Process Using the 8Ds Method. A Case Study in a Manufacturing Company*.

Silvestri, C. (2010). *Le determinanti del rapporto tra parametri di qualità e customer satisfaction*.

Somalvico, M. (1987). *Progetto di Intelligenza Artificiale*. Milano.

Willcocks, L. P. (2015). *Robotic Process Automation: The Next Transformation Lever for Shared Services*.

Wong, M. (2024). The GPT Era is already ending. *The Journal*.