

A Traceable Engineering-Evidence Framework for Maintenance Prioritization of Public Direct Current Fast Chargers

Herlangga Arisanto

Department of Electrical Engineering, Faculty of Applied Science and Technology, Institut Sains dan Teknologi Nasional, Jakarta, Indonesia

Masbah R. T. Siregar

Department of Electrical Engineering, Faculty of Applied Science and Technology, Institut Sains dan Teknologi Nasional, Jakarta, Indonesia

Baskoro Abie Pandowo

Department of Electrical Engineering, Faculty of Applied Science and Technology, Institut Sains dan Teknologi Nasional, Jakarta, Indonesia

Abstract: Public direct current fast chargers generate heterogeneous operational snapshots, service records, and platform-specific diagnostics that cannot be treated directly as field-supported failures. Existing monitoring identifies abnormal states, while conventional Failure Mode and Effects Analysis (FMEA) ranks an already defined failure taxonomy; neither explicitly operationalizes a qualification step between raw diagnostics and maintenance mechanisms. This study develops a pre-FMEA engineering-evidence framework integrating data normalization, subsystem classification, evidence registration, service-report correlation, manufacturer-document verification, and structured score review. The bounded case comprised 30 ABB Terra T54HV diagnostic datasets collected across multiple public sites during a January–December 2025 collection window (19,866 records and 877 unique keys) and 30 associated field service reports. The evidence set retained a common cross-source schema and linked diagnostic context to documented interventions; it is not claimed to represent all charger populations or a complete failure census. Of the diagnostic keys, 546 formed a common cross-source core and 217 varied across sites. Three primary service categories—communication and authorization, global interlock, and voltage-related cases—accounted for 25 of 30 documented interventions (83.3%). Their internally reviewed screening RPNs were 105, 90, and 72; no score was independently validated. The framework supports inspectable maintenance prioritization by preserving a source-to-mechanism evidence path, excluding

Correspondents Author:

Herlangga Arisanto, Department of Electrical Engineering, Faculty of Applied Science and Technology, Institut Sains dan Teknologi Nasional, Jakarta, Indonesia

Email: herlanggkr6io6@gmail.com

Received July 18, 2026; Revised August 22, 2026; Accepted August 22, 2026; Published September 1, 2026.

snapshot repetition from occurrence, disclosing overlapping mechanism links, and retaining unresolved scores as provisional.

Keywords: direct current fast charger, engineering evidence, failure mode and effects analysis, maintenance prioritization, service report.

Introduction

Public DC fast-charger maintenance is practically significant because a port reported as available may still fail to initiate or sustain charging. In an independent field test of 655 public CCS ports, only 73.3% completed the defined functionality test; 23.5% were nonfunctional because of screen/error, payment, initiation, network, or connector problems ([Rempel et al., 2024](#)). This result does not estimate the performance of the present fleet, but it demonstrates why maintenance decisions require verifiable port-level evidence rather than nominal status alone. A charger is also a cyber-physical system: power conversion, vehicle communication, charging-interface control, interlocks, isolation monitoring, thermal management, authorization, and backend connectivity can each interrupt service ([Arena et al., 2024](#); [Garofalaki et al., 2022](#); [Johnson et al., 2022](#)). Three terms are distinguished before the framework is introduced. Diagnostic variation means that a normalized key has more than one value across the 30 snapshots; it may reflect configuration, capability, software version, instantaneous state, measurement variation, or degradation. An abnormal condition is a value or state outside its expected functional context but is not, by itself, evidence of a field intervention. A failure mechanism is accepted for risk analysis only when a retained candidate is correlated with a documented service intervention and its engineering interpretation is supported by platform documentation. Monitoring and anomaly detection identify departures from expected behaviour ([Chen et al., 2024](#); [Fang et al., 2024](#); [Sakwa et al., 2024](#)), whereas conventional or data-driven FMEA ranks an already defined failure taxonomy ([AIAG & VDA, 2019](#); [Ervural & Ayaz, 2023](#)). Across these streams, the neglected link is not detection or scoring, but the explicit qualification of what a raw platform-specific observation is allowed to mean before it becomes an occurrence input. The methodological innovation is the integration of three evidence sources—normalized diagnostics, documented service interventions, and manufacturer documentation—with two analytical controls—subsystem interpretation and structured score review—as a conservative pre-FMEA engineering-evidence qualification layer. The study is bounded to 30 ABB Terra T54HV diagnostic datasets collected at multiple public sites during January–December 2025 and 30 associated service reports; it does not

infer fleet-wide failure rates or multi-vendor generalizability. Its measurable outputs are a unique-key dictionary, a 30-source comparison matrix, an evidence register with mechanism-qualification status, mutually exclusive primary service-category counts and Pareto shares, and an eight-row FMEA screening table with traceable S–O–D rationales and separate score-review status. The contribution is a specified decision chain that separates observation, abnormal condition, documented intervention, bounded mechanism, and risk.

Related Work and Research Gap

Monitoring, Anomaly Detection, and Fault Prediction

Charging-infrastructure monitoring studies increasingly use data-driven techniques to detect abnormal states. Power-module input-current characteristics have been used for open-circuit diagnosis, while fleet and equipment monitoring systems combine sensor data, statistical indicators, and communication status to support maintenance decisions ([Chen et al., 2024](#); [Fang et al., 2024](#); [Martins & Rodrigues., 2025](#)). Autoencoder-based approaches identify deviations from learned normal behavior, and more recent work applies contrastive learning, hierarchical attention, recurrent optimization, and unsupervised anomaly detection to improve diagnosis under sparse or imbalanced fault data ([García-Garvı et al., 2026](#); [Huang et al., 2025](#); [Lei et al., 2026](#); [Murshed et al., 2026](#); [Sakwa et al., 2024](#); [Wu et al., 2026](#)).

These methods can process large operational datasets and reveal patterns that are difficult to detect manually. For the present study, however, the unresolved issue occurs earlier in the analytical chain: an algorithm can classify a pattern only after the meaning of the feature and the validity of the label have been established. A heterogeneous diagnostic key may describe a configuration value, software capability, instantaneous state, protective status, or an actual failure indication. Without engineering qualification, repeated snapshot values can become misleading features or labels.

Reliability, Maintenance, and Service Evidence

Field reliability research demonstrates that charger accessibility and successful transaction completion should be evaluated at connector level rather than inferred only from reported uptime ([Rempel et al., 2024](#)). Reliability and economic studies show that redundancy, configuration, and support arrangements influence service continuity ([Vaishali & Prabha, 2024](#)). A systematic review of after-sales and maintenance services further identifies service capability as a hidden pillar of electric-vehicle deployment ([Panciu et al., 2026](#)). Service reports provide stronger maintenance evidence than snapshots because they record a

technician observation, corrective action, and affected equipment. They nevertheless remain incomplete: service coding may be inconsistent, one physical mechanism may produce several error strings, and one service action may address more than one diagnostic symptom. Service records therefore require normalization and engineering interpretation before they can define an occurrence basis.

Engineering Risk Analysis and FMEA

Failure Mode and Effects Analysis is widely used to structure functions, failure modes, effects, causes, controls, and actions. Conventional FMEA depends substantially on expert judgment, which can produce inconsistency if the evidence basis and scoring rationale are not documented. Data-driven and ontology-supported FMEA studies seek to improve objectivity and updateability by incorporating empirical records or formal knowledge structures ([Ervural & Ayaz, 2023](#); [Komarski et al., 2026](#); [Younus et al., 2026](#)). These developments improve risk assessment after a failure taxonomy has been established. The unresolved issue in charger diagnostics is taxonomy formation itself. Raw keys and snapshot values cannot be transferred directly into occurrence because their repetition may reflect file structure rather than repeated failure. The present framework therefore places evidence qualification before FMEA and treats risk scoring as the final stage rather than the starting point.

Gap-to-Contribution Positioning

Table 1 synthesizes the conceptual gap across the principal literature streams and identifies the specific response provided by the proposed framework. Its purpose is to distinguish a pre-FMEA evidence-qualification contribution from fault detection, service-record description, manufacturer guidance, and conventional risk ranking.

Table 1 Study Gap and Engineering Contribution

Literature stream	Typical limitation	Response in this paper
Monitoring and fault prediction	Features and labels are often assumed to be analytically valid	Creates an auditable key dictionary, subsystem context, and bounded engineering labels
Field reliability and service studies	Document availability or interventions but do not reconstruct diagnostic meaning	Correlates diagnostic candidates with documented field actions
Conventional FMEA	Starts from an established failure taxonomy and relies heavily on expert judgment	Qualifies failure mechanisms before scoring and records mechanism-qualification and score-review status

Literature stream	Typical limitation	Response in this paper
Manufacturer documentation	Explains platform functions but cannot prove field occurrence	Uses documents only to bound interpretation after field correlation
This study	Multi-source evidence is fragmented across separate records	Connects snapshots, service reports, technical documents, and structured review in one traceable workflow

Table 1 shows that previous streams each contribute necessary but incomplete evidence. Monitoring provides scalable detection but normally presupposes meaningful features and labels; service studies provide field relevance but not always mechanism reconstruction; FMEA provides structured prioritization but begins with a predefined taxonomy; and manufacturer documents bound functional interpretation without proving occurrence. The proposed framework connects these sources in sequence and makes evidence qualification—not a new prediction algorithm or a modified RPN equation—the central novelty. Where Table 1 presents the conceptual synthesis, Table 2 benchmarks the proposed framework against representative study-level approaches. The comparison focuses on the evidence each study uses, its primary contribution, and the unresolved boundary addressed by the present work.

Table 2 Compact Benchmark of Related Approaches

Study	Primary evidence or method	Contribution	Boundary addressed here
Rempel et al. (2024)	Field connector testing	Shows the gap between reported status and user-accessible charging	Adds traceability from diagnostic evidence to maintenance action
Chen et al. (2024)	Power-module current signatures	Diagnoses open-circuit faults	Extends interpretation beyond one component and one signal family
Sakwa et al. (2024)	Autoencoder-based monitoring	Detects early abnormal behaviour	Provides engineering qualification for features and labels
Martins and Rodrigues (2025)	Intelligent charging monitoring	Integrates operational monitoring functions	Adds field-service and technical-document correlation
Ervural and Ayaz (2023)	Data-driven FMEA	Reduces purely subjective risk assessment	Adds a pre-FMEA evidence-qualification stage

Study	Primary evidence or method	Contribution	Boundary addressed here
This study	Snapshots, service reports, documents, and structured review	Produces auditable maintenance priorities	Does not claim predictive accuracy or chronological reliability

The benchmark confirms that the proposed framework occupies an upstream methodological position. It extends field testing by reconstructing the diagnostic-to-maintenance evidence path, extends component-specific diagnosis to system-level evidence qualification, and complements anomaly detection by establishing defensible labels. Compared with data-driven FMEA, its distinctive contribution is to validate the taxonomy and occurrence basis before risk scoring. This advantage is accompanied by a deliberate boundary: the framework improves traceability and decision defensibility, not predictive accuracy or population-level reliability estimation.

Research Method

Research Design Rationale

A bounded, multi-source engineering case-study design was selected because the research question asks how heterogeneous industrial evidence can be transformed into a traceable maintenance decision within its real operational context. Case-study methodology is appropriate for a contemporary process whose evidence and context cannot be separated cleanly and supports triangulation across records, documents, and observations (Yin, 2018). A purely statistical design was unsuitable because the operational files are snapshots rather than event histories; an experimental design could not reproduce field interventions; and an expert-only design would weaken source traceability. The case design therefore supports analytical generalization of the evidence process rather than statistical generalization to all charger populations. The unit of analysis was an engineering failure mechanism supported by an evidence chain, not an individual snapshot row. Five sources were integrated sequentially: fleet diagnostics identified common and varying keys; subsystem interpretation assigned functional context; service reports established documented intervention evidence and the occurrence basis; manufacturer documentation constrained the proposed mechanism to platform functions; and structured review checked the consistency of the three highest-ranked severity–occurrence–detection results. No source was accepted as complete proof by itself.

The evidence set was purposive rather than random. A diagnostic file was eligible when it originated from the investigated Terra T54HV platform within the January–December 2025 collection window, was readable and nonempty, and retained the Key–Update Time–Value schema and a source identifier. A service report was eligible when it referred to the investigated platform and collection window and retained an identifiable charger, technical observation or error string, and documented intervention. The delivered inventory comprised 30 diagnostic files and 30 associated service reports. No exclusion count or statistical sampling frame is inferred beyond that inventory. The set is sufficient to demonstrate the reported transformation stages on a common schema, but it is not a complete failure census and does not support population prevalence or future failure probability. Procedural auditability was supported through the fixed ten-step workflow, preservation of original source–key–time–value fields, a source-by-key matrix, explicit subsystem and evidence fields, primary service-category reconciliation, the stated occurrence rule, written S–O–D rationales, and separate mechanism-qualification and score-review status. The reported arithmetic is independently checkable from the published counts, while record-level computational reproduction remains constrained by proprietary-data access. Conflicts were retained rather than forced to agreement: variation without a service intervention remained a review trigger; a service intervention without adequate technical support remained provisional; manufacturer documentation bounded function but did not prove occurrence; and no score was described as independently validated.

Data Sources and Evidence Boundaries

Table 3 defines the role and evidential boundary of each source used in the framework. Making these boundaries explicit prevents a source that is informative for one analytical step from being treated as proof for another.

Table 3 Evidence Sources, Roles, and Boundaries

Evidence source	Role in the framework	Boundary
Operational diagnostic snapshots	Key inventory, source-by-key comparison, variation screening	Not a chronological failure history
Field service reports	Documented interventions and occurrence basis	Limited to 30 recorded cases
Manufacturer documentation	Subsystem function and diagnostic interpretation	Platform-specific and not proof of field occurrence
Structured expert review	Consistency check and confirmation of the three highest-ranked S–O–D results	Single reviewer who was also the study author

Operational snapshots provide cross-source breadth and reproducible key comparison but cannot establish chronological occurrence. Service reports document interventions but may combine symptoms. Manufacturer documents define platform functions but cannot prove that a function failed in the field. Subsystem interpretation and structured score review are analytical controls rather than independent data sources. A candidate enters mechanism screening only when its diagnostic context, intervention record, and documented function converge; contradictions remain provisional or rejected. Mechanism qualification and score review are recorded separately, and no independent audit or inter-rater validation was completed.

Data Collection and Preprocessing

Preprocessing preserved the original source, key, update-time, and value fields and created separate comparison fields. Blank keys were excluded; text encoding, surrounding whitespace, and capitalization variants were normalized where required for comparison, while original wording and values remained retained. The recorded update-time field was preserved as source metadata and was not interpreted as a chronological failure-occurrence timestamp. Let $N=19,866$ denote retained nonblank rows and $K = \{k: k \text{ appears in at least one source}\}$; set union produced $|K|=877$ unique normalized keys. For key k and source s , $A_{ks}=1$ when k occurred in s . A common key satisfied $\sum_s A_{ks}=30$, producing 546 keys ($546/877=62.3\%$). A varying key had more than one normalized comparison value, producing 217 keys ($217/877=24.7\%$). Common and varying are independent attributes and may overlap; $546+217$ is therefore not a partition of 877. Service wording was normalized only for primary-category grouping, with the original report text retained for mechanism interpretation.

Table 4 Data-selection and preprocessing rules

Evidence type	Inclusion rule	Exclusion rule	Trace retained
Diagnostic snapshot	T54HV; 2025 collection window; readable; nonempty; Key–Time–Value schema; source ID	File outside the defined platform/schema/window	Source ID, original row, normalized comparison fields
Service report	T54HV; collection window; charger ID; technical observation/error; documented intervention	Record without the minimum technical evidence needed for correlation	Report ID, original wording, action, normalized primary category

Engineering Classification and Evidence Registration

Diagnostic keys were mapped to power conversion, charging interface, communication, safety, thermal management, control, monitoring, user interaction, or auxiliary functions using the key name, recorded value, associated service text, and manufacturer-defined function. The mapping rule was: assign the narrowest documented subsystem consistent with the key and intervention; retain a cross-subsystem flag when more than one mechanism remained plausible; and keep the mechanism provisional when the ambiguity could not be resolved. For example, a communication-status key linked to a modem/SIM/certificate reconnection action was mapped to communication, whereas a global-interlock key linked to a door switch or residual-current condition was mapped to safety. For key i , F_i denotes explicit failure semantics, G_i an engineering-critical function, and V_i cross-site variation. Candidate retention is $C_i=1$ when $F_i \vee G_i \vee V_i$ is true; otherwise $C_i=0$.

$$C_i = 1 \text{ if } (F_i \vee G_i \vee V_i); \text{ otherwise } C_i = 0 \quad (1)$$

Candidate status does not imply failure. The evidence register stores the original key/value/source, subsystem and mapping rationale, variation status, service-report link, technical-document basis, proposed mechanism, operational effect, and decision status. Let $R_i=1$ only when a qualifying service report documents a relevant intervention and $M_i=1$ only when a manufacturer source supports the proposed functional interpretation. A mechanism is eligible for FMEA when $E_i=C_i R_i M_i=1$. If $C_i=1$ but $R_i=0$, the item remains a diagnostic review trigger; if $R_i=1$ but $M_i=0$ or the mapping is ambiguous, it remains provisional; if the technical evidence contradicts the interpretation, it is rejected. These gates prevent variation, documentation, or expert opinion alone from creating occurrence.

$$E_i = 1 \text{ if } (C_i = 1 \wedge R_i = 1 \wedge M_i = 1); \text{ otherwise } E_i = 0 \quad (2)$$

A binary qualification rule was selected instead of an arbitrary weighted evidence score. The rule is deliberately conservative: variation alone cannot confirm a failure, and a documented platform function alone cannot confirm field occurrence. A candidate enters risk analysis only after field correlation and technical verification are both present.

Service-Report Correlation and Occurrence Conversion

The 30 reports were assigned exactly once to seven mutually exclusive primary service categories: communication (10), global interlock (9), voltage drop (6), maximum voltage (2), fan failure (1), general charger error (1), and vehicle-protocol error (1); their sum is 30. A

report could also support a secondary engineering-mechanism link, but that link did not add a Pareto case. For mechanism j , n_j denotes distinct reports linked to that mechanism and may overlap across mechanisms; it is therefore not additive across FMEA rows. Occurrence was assigned by $O(n)=1$ for $n=0$, 2 for $1 \leq n \leq 2$, 3 for $3 \leq n \leq 5$, 4 for $6 \leq n \leq 8$, and 5 for $n \geq 9$. This deterministic five-level rule preserves a separate baseline for zero observed reports and progressively orders positive intervention counts within the bounded 30-report set. The thresholds are operational screening bands, not calibrated probabilities, failure rates, or AIAG–VDA frequency criteria.

$$O(n) = 1, n = 0; 2, 1 \leq n \leq 2; 3, 3 \leq n \leq 5; 4, 6 \leq n \leq 8; 5, n \geq 9 \quad (3)$$

The conversion creates a transparent relationship between documented cases and the occurrence score. It does not estimate probability or future failure rate. It was used only to support internal screening within the available evidence set.

Evidence-Driven FMEA and Review Status

Severity S represents the highest consequence supported by the documented platform function and bounded mechanism; Occurrence O is generated only by the stated report-count conversion; and Detection D represents the capability of existing diagnostics or inspection controls to indicate the condition or affected function. A lower D score denotes stronger detection capability. The risk priority number is $RPN_j = S_j \times O_j \times D_j$. Internal screening bands (high ≥ 80 , medium 40–79, low < 40) organize follow-up only; they are not acceptance criteria and do not replace the AIAG–VDA Action Priority method. One author-engineer internally reviewed the three highest S – O – D combinations. No row was independently validated, and the remaining score combinations are provisional.

Table 5 Study-Specific Severity and Detection Anchors

Factor	Study-specific anchor used in this case
$S=10$	Potential safety-critical loss of required protective integrity
$S=9$	Complete charging inhibition or controlled protective shutdown involving safety or DC-output control
$S=8$	Essential charging interface, protocol, or thermal function unavailable
$S=7$	Service unavailable through communication/authorization without the maximum electrical-safety consequence
$D=2$	Direct onboard indication, protection status, or targeted inspection identifies the condition/affected function

D=3	Logs indicate the condition, but root-cause isolation requires additional cross-system checks
-----	---

$$RPN_j = S_j \times O_j \times D_j \quad (4)$$

The score anchors and internal bands were used only for bounded follow-up ordering. The structured review was conducted by the study author in his professional role as a certified electric-vehicle-charger service engineer with more than eight years of relevant experience. This single author-engineer review assessed the consistency of the three highest S–O–D combinations; it did not constitute independent confirmation, inter-rater validation, or an AIAG–VDA Action Priority determination.

Pareto Concentration and Algorithmic Procedure

For the Pareto analysis, the cumulative contribution of the first k ranked service categories was calculated as:

$$P_k = \frac{\sum_{i=1}^k n_i}{\sum_{i=1}^m n_i} \times 100\% \quad (5)$$

Algorithm 1 Traceable engineering-evidence workflow

Input: operational snapshots D , field service reports R , and manufacturer documents M .

Output: evidence-linked screening rows with mechanism-qualification and internally reviewed/provisional score status.

1. Parse D and normalize key, time, and value fields.
2. Build the unique-key dictionary and source-by-key matrix.
3. Mark common and varying keys and retain candidates using C_i .
4. Assign each candidate to an engineering subsystem and create the evidence register.
5. Normalize R and correlate candidate mechanisms with documented interventions.
6. Verify subsystem functions and mechanisms against M and determine E_i .
7. Aggregate service-report counts n_j and calculate O_j .
8. Assign S_j and D_j from engineering effects and existing controls.
9. Calculate RPN_j , internal priority, and cumulative Pareto contribution P_k .
10. Record mechanism qualification separately from internally reviewed or provisional score status.

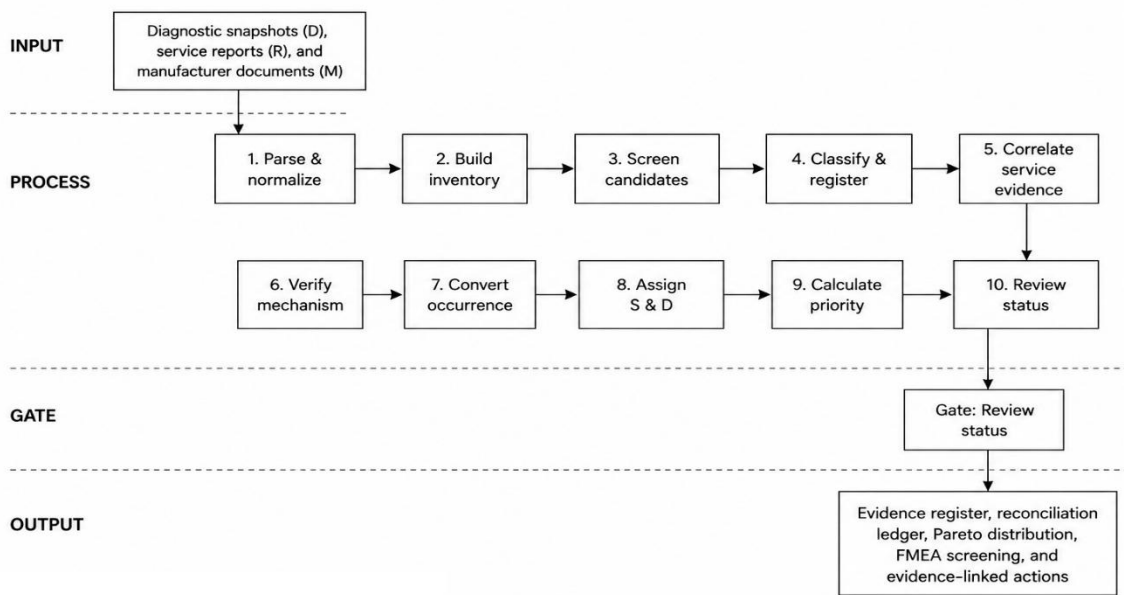


Figure 1 Evidence-to-priority workflow

Figure 1 summarizes the logical sequence from raw diagnostics to ranked maintenance priorities. Preprocessing establishes a normalized evidence base; engineering classification and service correlation then qualify candidate mechanisms; technical verification and structured review are completed before risk ranking. The workflow therefore prevents a direct and unsupported conversion of record counts into RPN values.

Result and Discussion

Stage-by-Stage Framework Outputs

Application of the fixed workflow produced auditable intermediate outputs rather than one final score. The parser retained $N=19,866$ nonblank source rows. Set union across normalized keys produced 877 unique keys; the 30-column presence matrix identified 546 common keys, and normalized value-set cardinality identified 217 varying keys. Candidate screening then assigned subsystem context and a decision status without treating either subset as failure counts. Separately, 30 qualifying service reports were normalized into seven mutually exclusive primary categories ($\Sigma n=30$), linked to eight mechanism-level FMEA rows through primary and non-counting secondary links, and converted to occurrence scores. Thus, the diagnostic branch supplies context and candidate evidence, while the service branch supplies the bounded occurrence denominator.

Table 6 Stage-by-stage reconciliation and validation checks

Stage	Operation/check	Output	Validation identity
Row retention	Count nonblank normalized records	19,866 rows	$N=19,866$
Dictionary	Union of normalized keys	877 unique keys	$ K =877$
Presence matrix	$\Sigma A_{ks}=30$	546 common keys	$546/877=62.3\%$
Value variation	Normalized value-set size >1	217 varying keys	$217/877=24.7\%$
Service normalization	One primary category/report	30 reports in 7 categories	$10+9+6+2+1+1+1=30$
Mechanism mapping	Primary + non-counting secondary links	8 FMEA rows	Overlap flagged; no duplicate Pareto count
Risk calculation	$S \times O \times D$	8 ranked RPNs	Row-wise arithmetic check

Operational Evidence Base

Data preprocessing reduced 19,866 raw records to 877 unique diagnostic keys without discarding source traceability. Table 4 shows that 546 keys, or 62.3% of the dictionary, were present in every source file. This common core indicates a stable platform-wide diagnostic schema that supports a standardized analysis pipeline. Independently, 217 keys, or 24.7%, exhibited more than one value across the fleet. Because the common and varying sets overlap, the percentages are not additive. Variation may reflect configuration, enabled options, software version, instantaneous state, measurement uncertainty, or degradation; it was therefore used to trigger engineering review rather than being counted as failure.

Table 7 Characteristics of the Operational Evidence Base

Item	Verified value
Public DC fast-charger datasets	30
Observation period	January-December 2025
Operational records	19,866
Unique nonblank diagnostic keys	877
Keys present in all source files	546
Keys with multiple recorded values	217
Field service reports	30

The analytical significance of Table 7 is twofold. The 546-key common core shows that cross-site comparison is technically feasible on a shared schema, while the 217-key varying subset narrows the engineering review to parameters that may differentiate site configuration or condition. This separation prevents the 19,866 snapshot rows from being misrepresented as 19,866 failure observations and supplies the first control against artificial occurrence inflation. In contrast with monitoring studies that begin with preselected features, this stage establishes the auditable dictionary from which defensible features and labels can later be derived.

Service-Report Distribution and Pareto Concentration

Pareto analysis was used to identify the small number of mutually exclusive service-report categories accounting for most documented interventions and to direct subsequent mechanism qualification and FMEA effort. It is a prioritization screen rather than a causal model or forecast. Communication and authorization were the largest category with 10 cases, followed by global interlock with 9 cases and charger voltage drop with 6 cases; together they contributed 25 of 30 interventions (83.3%). This concentration identifies where standardized inspection, and recovery procedures can provide the broadest coverage in the investigated evidence set before lower-frequency categories are considered.

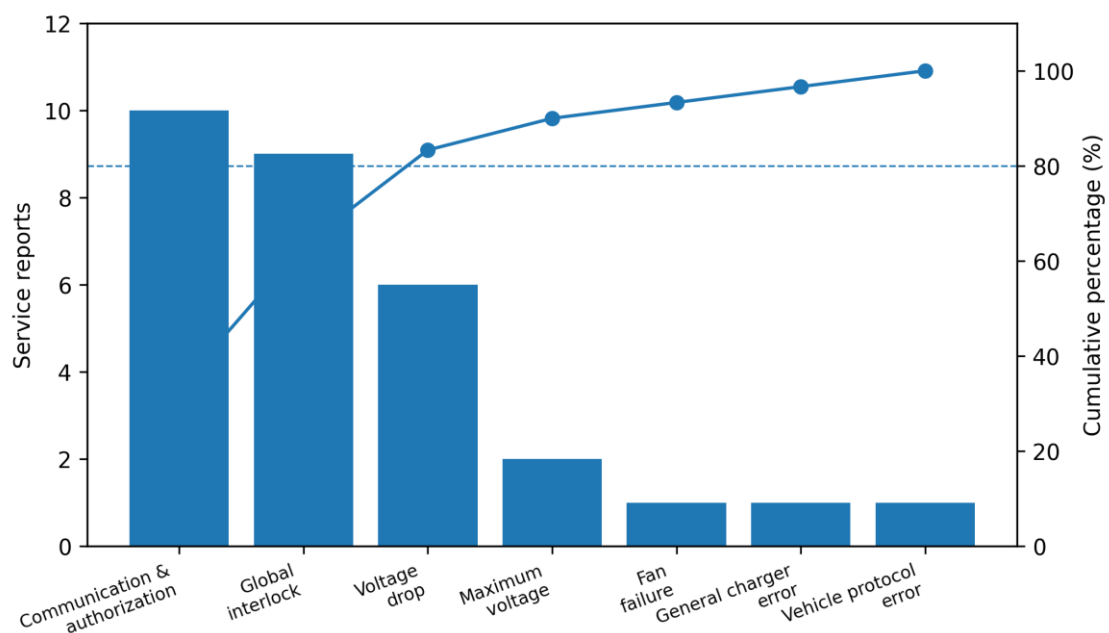


Figure 2 Pareto distribution of mutually exclusive service-report error categories. Communication and authorization, global interlock, and voltage-drop cases account for 83.3% of the reports

Table 8 Distribution of Field Service Reports

Service-report category	Cases	Percentage
Communication and authorization	10	33.3%
Global interlock	9	30.0%
Charger voltage drops	6	20.0%
Charger maximum voltage	2	6.7%
Power-port fan failure	1	3.3%
General charger error	1	3.3%
Vehicle-protocol error	1	3.3%
Total	30	100.0%

Figure 2 and Table 8 reveal a mixed cyber-physical service burden rather than a purely power-electronic failure pattern. Communication and authorization cases imply that endpoint configuration, certificates, modem or subscriber-identity-module connectivity, backend availability, and post-outage recovery must be diagnosed before hardware replacement. Global-interlock cases emphasize the availability consequence of safety-chain conditions even when the power stage is functional. Voltage-drop cases require component-level discrimination among power modules, auxiliary supplies, and associated assemblies. This distribution is qualitatively consistent with literature that treats charger communication, cybersecurity, availability, and power conversion as interacting reliability concerns ([Arenas et al., 2024](#); [Garofalaki et al., 2022](#); [Johnson et al., 2022](#); [Rempel et al., 2024](#)), although direct percentage comparison is inappropriate because those studies use different units, platforms, and observation designs. The Pareto output therefore determines analytical attention but not the final risk order. Its three dominant categories were reconstructed into bounded engineering mechanisms, transferred to the evidence-driven FMEA, and evaluated together with severity and detectability. This linkage explains why communication and authorization, global interlock, and power-conversion voltage instability became the three highest internally reviewed screening priorities, while recommendations were translated into communication-recovery checks, interlock inspection, and component-level voltage diagnosis. Pareto concentration supplies breadth of service demand; FMEA adds consequence and control awareness.

Engineering Interpretation and Traceability

The primary service distribution and mechanism taxonomy are related but not identical. Every report is counted once in the Pareto denominator of 30, whereas a secondary link may support provisional mechanism screening without creating another Pareto case. Communication/authorization reports support recovery of the external backend path; global-interlock reports share the operational effect of charging inhibition but retain heterogeneous causes; and the six voltage-related reports require component-level separation rather than being assumed to represent one physical failure. The generic-error report remains a primary service category but is not itself a mechanism; it contributes only a provisional CPI/interface link when the supporting context permits. Isolation monitoring is retained solely as a document-supported preventive reference with no observed service-report occurrence. Consequently, mechanism-linked counts sum to 31 and are explicitly non-additive across FMEA rows.

Table 9 Service-report-to-mechanism mapping and overlap control

Primary service category	Distinct reports	Mechanism disposition	Secondary link	Counting rule
Communication/authorization	10	External OCPP communication/reconnection	None used for occurrence	10 once
Global interlock	9	Safety interlock/global interlock	Cause may be door/RCD/connectivity stop condition	9 once
Voltage drops	6	Voltage-related category; component-bounded interpretation required	Power conversion/auxiliary/display context retained	6 primaries; mechanism status bounded
Maximum voltage	2	DC output-voltage control	None	2 once
Fan failure	1	Power-module cooling	None	1 once
General charger error	1	Insufficiently specific as a mechanism	CCS/CPI interface (provisional)	1 primary; no extra Pareto case
Vehicle-protocol error	1	Vehicle protocol/CCB	CCS/CPI interface (secondary)	1 primary; no extra Pareto case
Isolation fault	0	Document-supported preventive reference	No field-occurrence link	0 observed reports

Communication-related cases also required system-level interpretation. ABB's dual-uplink concept permits direct OCPP integration with the operator backend while retaining a manufacturer-platform connection for service support (ABB., 2019). In the investigated operator configuration, the normal path used an external modem and operator-specific SIM arrangement. After power restoration, several chargers required reconnection through the

manufacturer communication path and certificate updating before normal external-backend communication was restored. The observed occurrence therefore reflects the deployed recovery architecture and maintenance procedure, not a universal defect of OCPP or the charger platform. Global-interlock cases combined several field causes, including door-switch contamination, residual-current protection events, and connectivity-related stop conditions. A high severity score was justified because operation is intentionally inhibited when the safety chain cannot be confirmed, whereas favorable detection reflected direct interlock monitoring. These examples demonstrate why frequency must be interpreted together with function, effect, and existing controls.

Evidence-Driven FMEA and Maintenance Priority

Table 10 contains seven field-linked screening rows and one document-supported preventive reference. Communication/reconnection was assigned S=7, O=5 from n=10, and D=3 because logs indicate the condition but root-cause isolation spans the charger-backend path; RPN=105. Global interlock was assigned S=9, O=5 from n=9, and D=2 because interlock status and targeted inspection directly indicate the affected safety function; RPN=90. The voltage-related screening row was assigned S=9, O=4 from six linked reports, and D=2; RPN=72, while its component-level mechanism remains bounded by the report-specific evidence. These three score combinations were internally reviewed by one author-engineer, not independently validated. The remaining score combinations are provisional.

Table 10 Evidence-Linked FMEA Screening and Preventive Reference Mode

Engineering function / failure mode	Linked reports (n)	S	O	D	RPN	Internal priority	Score rationale (S / O / D)
External OCPP communication and authorization / failure to reconnect after power restoration	10	7	5	3	105	High (internally reviewed)	Service unavailable; 10 reports→O5; logs/connectivity checks detect but root cause spans systems.
Safety interlock / global interlock failure	9	9	5	2	90	High (internally reviewed)	Safety-chain inhibition; 9 reports→O5; direct interlock indication gives strong detection.
Power conversion / voltage instability or module failure	6	9	4	2	72	Medium (internally reviewed)	Controlled output lost/safe shutdown; 6 reports→O4; voltage/module diagnostics are direct.

Engineering function / failure mode	Linked reports (n _i)	S	O	D	RPN	Internal priority	Score rationale (S / O / D)
DC output-voltage control / maximum-voltage fault	2	9	2	3	54	Medium (provisional)	Potential output-control interruption; 2 reports→O2; logs detect but mechanism remains provisional.
CCS/CPI charging interface / control-pilot or interface failure	2*	8	2	3	48	Medium (provisional)	Charge-interface interruption; 2 linked reports→O2; interface codes aid detection; secondary links not double-counted.
Vehicle protocol and charge control / CCB communication failure	1	8	2	3	48	Medium (provisional)	Protocol negotiation prevents charging; 1 report→O2; protocol logs available; provisional.
Power-module cooling / fan or cooling failure	1	8	2	2	32	Low (provisional)	Thermal protection/derating risk; 1 report→O2; fan monitoring is direct; provisional.
Isolation monitoring / isolation fault	0	10	1	2	20	Preventive reference	Highest supported safety consequence; 0 observed reports→O1; document-supported preventive reference only.

*The CPI/interface value represents two overlapping report links used for provisional mechanism screening. The mechanism-linked count is non-additive across FMEA rows and sums to 31, while the mutually exclusive primary-category denominator remains 30. Score-review status does not imply independent validation or a complete failure census.

The arithmetic is verifiable from the reported counts: communication/authorization gives $7 \times 5 \times 3 = 105$; global interlock gives $9 \times 5 \times 2 = 90$; and the voltage-related screening row gives $9 \times 4 \times 2 = 72$. For a qualification example, a varying communication-status key gives $C_i = 1$ when $V_i = 1$; a documented reconnection intervention gives $R_i = 1$; and manufacturer communication documentation supports the functional interpretation, giving $M_i = 1$ and $E_i = 1$. Without the service link, $E_i = 0$ and the key remains a review trigger. Mechanism qualification and score review are recorded separately: field-supported evidence does not constitute independent S–O–D validation. The isolation reference has $n = 0$, $O = 1$, $S = 10$, $D = 2$, and $RPN = 20$; it remains a mandatory preventive-control reference rather than an observed field-occurrence mechanism.

Comparison with Previous Approaches

Table 11 compares the proposed framework with common analytical approaches to clarify its novelty, contribution, and limits. The comparison asks what each approach produces, what

evidence supports that output, and what remains unresolved when it is used for maintenance prioritization.

Table 11 Comparison of the Proposed Framework with Common Analytical Approaches

Approach	Primary output	Evidence strength	Limitation for maintenance prioritization	Role of the proposed framework
Snapshot counting	Frequency of recorded values	Broad fleet coverage	Can confuse file repetition with failure occurrence	Rejects snapshot rows as chronological event counts
Anomaly detection	Abnormality score or label	Scalable pattern detection	Depends on defensible features and labels	Creates traceable engineering features and labels
Service-report analysis	Intervention categories	Direct field relevance	Coding inconsistency and incomplete mechanism detail	Normalizes categories and reconstructs mechanism context
Conventional FMEA	Risk ranking	Structured consequence-control logic	Often starts from expert-defined taxonomy	Qualifies the taxonomy before S-O-D scoring
Proposed framework	Evidence-linked maintenance priority	Multi-source traceability	Platform and evidence-set specific	Provides bounded, auditable decisions and review status

The application reconciled 19,866 rows into 877 unique keys, identified 546 common and 217 varying keys without treating either as occurrences, assigned 30/30 reports to one primary Pareto category, and recorded arithmetic, rationale, and status for 8/8 FMEA rows. Three of eight score combinations were internally reviewed and none was independently validated. These are reporting-coverage and traceability outputs, not reliability improvements. Direct performance comparison with anomaly-detection accuracy or population reliability is invalid because prior studies use different outcomes. (Rempel et al., 2024), for example, measured whether 655 ports completed a two-minute charge and found 73.3% functional; this case study does not retest port functionality. The contribution is therefore procedural and inspectable: snapshot repetition is excluded from occurrence, primary service counts remain reconciled to 30, overlapping mechanism links are disclosed, and unresolved scores remain provisional. No claim is made that the framework improved availability, downtime, diagnostic accuracy, or causal inference.

Practical Maintenance Implications

For maintenance managers, the results support three immediate decisions. First, post-outage recovery procedures should include endpoint, access-point name, certificate, modem, SIM, and uplink check because communication recovery was the largest documented service category. The procedure should distinguish charger hardware failure from an unavailable authorization or backend path before component replacement is initiated. Second, routine inspection should explicitly cover door switches, emergency inputs, residual-current-device status, and interlock wiring because interruption of the safety chain can make a charger unavailable even when the power stage remains functional. The inspection record should preserve the diagnostic code, physical cause, corrective action, and post-repair verification so that future scoring is based on mechanism-level evidence rather than error-string frequency. Third, voltage instability should be diagnosed at component level so that power-module cases are separated from auxiliary-supply and display faults. This distinction improves spare-parts planning and prevents one generic voltage category from masking different corrective actions. These recommendations translate the Pareto concentration and risk ranking into inspectable work instructions. They remain evidence-based priorities rather than demonstrated improvements in availability or downtime.

Limitations and Future Work

The principal limitations are the purposive, platform-specific multi-site case-study design, the 30 available service reports, the snapshot nature of the operational data, and the use of a single author-engineer reviewer who was also the study author. The study cannot estimate chronological failure rate, downtime, mean time between failures, causal probability, or achieved reliability improvement, and no independent inter-rater agreement was measured. The communication-architecture findings apply to the investigated operator configurations and should not be generalized to every charging network without validation. Future research should extend the evidence register with longitudinal event histories, explicit start and recovery timestamps, repair outcomes, and repeated observations across multiple charger vendors. Independent expert panels should be used to test inter-rater agreement for subsystem mapping and S–O–D scoring. A prospective field study should finally test whether the prioritized procedures reduce diagnosis time, repeat intervention, and charger unavailability.

Conclusions

The bounded case produced four principal findings. The 19,866 retained snapshot rows formed 877 unique keys, including a 546-key common cross-source core and 217 varying keys; neither row repetition nor variation was treated as occurrence. All 30 service reports were assigned to seven mutually exclusive primary categories, and the three leading categories accounted for 25 of 30 documented interventions (83.3%) within this report set. The three highest internally reviewed screening results were 105, 90, and 72; the remaining score combinations were provisional and the isolation row was retained only as a preventive reference. The methodological contribution is a pre-FMEA engineering-evidence qualification layer that separates diagnostic observation, documented intervention, manufacturer-bounded interpretation, mechanism qualification, and risk scoring. It operates upstream of predictive monitoring by establishing bounded engineering features and labels, and it complements conventional FMEA by qualifying the mechanism set and occurrence basis before S–O–D scoring. For the investigated maintenance context, the findings support prioritizing validation of communication-recovery checks, cause-specific safety-interlock inspection, and component-bounded voltage diagnosis. These are evidence-linked maintenance priorities, not measured post-intervention improvements in availability, downtime, diagnosis time, or repeat intervention. The numerical distributions, subsystem mappings, and RPN results are limited to a purposive, platform-specific multi-site case study, 30 available service reports, snapshot data, and a single author-engineer reviewer. The reports are not asserted to be a complete census of all faults, unsuccessful charging attempts, or downtime episodes during 2025. The workflow is proposed for procedural transfer only and requires platform-specific revalidation. Longitudinal multi-vendor data and independent reviewers with reported inter-rater agreement are required before numerical generalizability or operational effectiveness is claimed.

Acknowledgements

The authors acknowledge the industrial data owner and the engineering personnel who supported charger maintenance, operational-data collection, technical-document review, and preparation of the engineering evidence base. The study received no external funding. The first author is professionally affiliated with the data-owning organization; this relationship is disclosed, and the authors report no other financial or personal competing interests. The operational and service datasets contain proprietary industrial information. Aggregated data

may be shared through the corresponding author when the request is scientifically justified, and the data owner grants permission.

AI Use Declaration

OpenAI ChatGPT was used to assist with English-language editing, structural refinement, citation-consistency checking, and document formatting. The tool was not treated as an author and did not replace scientific judgment. All data, engineering classifications, service-report interpretations, FMEA scores, and conclusions were independently checked by the corresponding author and are presented under the authors' responsibility.

References

- ABB. (2019). Electric vehicle infrastructure: Global product portfolio-Smarter mobility (Document No. 4EVC901502-BREN). ABB EV Infrastructure.
- ABB. (2022). Terra 54 and Terra 54HV installation manual (Document No. 001, 22 November 2022). ABB E-mobility.
- AIAG & VDA. (2019). Failure mode and effects analysis: FMEA handbook (1st ed.). Automotive Industry Action Group.
- Arena, G., Chub, A., Lukianov, M., Strzelecki, R., Vinnikov, D., & De Carne, G. (2024). A comprehensive review on DC fast charging stations for electric vehicles: Standards, power conversion technologies, architectures, energy management, and cybersecurity. *IEEE Open Journal of Power Electronics*, 5, 1573-1611.
<https://doi.org/10.1109/OJPEL.2024.3466936>
- Chen, Y., Tang, Z., Weng, X., He, M., Zhou, S., Liu, Z., & Jin, T. (2024). A diagnostic method for open-circuit faults in DC charging piles based on power-module input-current characteristics. *Energies*, 17(2), 404. <https://doi.org/10.3390/en17020404>
- Ervural, B., & Ayaz, H. I. (2023). A fully data-driven FMEA framework for risk assessment on manufacturing processes using a hybrid approach. *Engineering Failure Analysis*, 152, 107525. <https://doi.org/10.1016/j.engfailanal.2023.107525>
- Fang, H., Liao, J., Huang, S., & Zhang, M. (2024). Research on status assessment and operation and maintenance of electric vehicle DC charging stations based on XGBoost. *Smart Cities*, 7(6), 3055-3070. <https://doi.org/10.3390/smartcities7060119>
- García-Garvía, A., Arroyo-Torres, B., & Tormo-Domènech, C. (2026). Predictive maintenance of DC fast-charging stations using unsupervised anomaly detection. *Applied Sciences*, 16(14), 7052. <https://doi.org/10.3390/app16147052>

- Garofalaki, Z., Kosmanos, D., Moschoyiannis, S., Kallergis, D., & Douligeris, C. (2022). Electric vehicle charging: A survey on the security issues and challenges of the Open Charge Point Protocol (OCPP). *IEEE Communications Surveys & Tutorials*, 24(3), 1504-1533. <https://doi.org/10.1109/COMST.2022.3184448>
- Huang, Y., Ngaopitakkul, A., & Yoomak, S. (2025). Charging pile fault prediction method combining whale optimization algorithm and long short-term memory network. *Energy Informatics*, 8, 70. <https://doi.org/10.1186/s42162-025-00530-8>
- Johnson, J., Berg, T., Anderson, B., & Wright, B. (2022). Review of electric vehicle charger cybersecurity vulnerabilities, potential impacts, and defenses. *Energies*, 15(11), 3931. <https://doi.org/10.3390/en15113931>
- Komarski, D., Vassilev, V., Nikolov, S., Dimitrova, R., & Dimitrov, S. (2026). Data-driven process FMEA for flexible manufacturing systems: Framework and industrial case study. *Applied Sciences*, 16(8), 3760. <https://doi.org/10.3390/app16083760>
- Lei, Z., Xing, B., Liu, J., Yang, Y., Miao, T., & Lu, Y. (2026). Sustainable and reliable operation of EV charging infrastructure: A lightweight prototype-driven contrastive learning framework for fault diagnosis under class-imbalanced conditions. *Sustainability*, 18(11), 5783. <https://doi.org/10.3390/su18115783>
- Martins, J. A., & Rodrigues, J. M. F. (2025). Intelligent monitoring systems for electric vehicle charging. *Applied Sciences*, 15(5), 2741. <https://doi.org/10.3390/app15052741>
- Murshed, I., Islam, F., Joarder, M. S. A., Islam, M. M., Ali, M. Y., & Ahshan, R. (2026). Operational anomaly detection in EV charging systems using multi-parameter data analysis. *IEEE Access*, 14, 2552-2565. <https://doi.org/10.1109/ACCESS.2025.3649308>
- Nasri, S., Mansouri, N., Mnassri, A., Lashab, A., Vasquez, J., & Rezk, H. (2025). Global analysis of electric vehicle charging infrastructure and sustainable energy sources solutions. *World Electric Vehicle Journal*, 16(4), 194. <https://doi.org/10.3390/wevj16040194>
- Panciu, A., Kifor, C.-V., Ință, M., Lobont, L., & Zerbes, M. V. (2026). After-sales and maintenance services: The hidden pillar behind a successful electric vehicle deployment—A systematic literature review. *Systems*, 14(6), 642. <https://doi.org/10.3390/systems14060642>
- Rempel, D., Cullen, C., Bryan, M. M., & Cezar, G. V. (2024). Reliability of open public electric vehicle direct current fast chargers. *Human Factors*, 66(11), 2528-2538. <https://doi.org/10.1177/00187208231215242>
- Sadeghian, O., Oshnoei, A., Mohammadi-Ivatloo, B., Vahidinasab, V., & Anvari-Moghaddam, A. (2022). A comprehensive review on electric vehicles smart charging: Solutions,

- strategies, technologies, and challenges. *Journal of Energy Storage*, 54, 105241.
<https://doi.org/10.1016/j.est.2022.105241>
- Sakwa, M., Nespoli, A., Matrone, S., Leva, S., Guerini, A., Demartini, A., & Ogliari, E. (2024). Electric vehicle supply equipment monitoring and early fault detection through autoencoders. *Sustainable Energy, Grids and Networks*, 40, 101497.
<https://doi.org/10.1016/j.segan.2024.101497>
- Sarda, J., Patel, N., Patel, H., Vaghela, R., Brahma, B., Bhoi, A. K., & Barsocchi, P. (2024). A review of the electric vehicle charging technology, impact on grid integration, policy consequences, challenges and future trends. *Energy Reports*, 12, 5671-5692.
<https://doi.org/10.1016/j.egyr.2024.11.047>
- Shern, S. J., Sarker, M. T., Ramasamy, G., Thiagarajah, S. P., Al Farid, F., & Suganthi, S. T. (2024). Artificial intelligence-based electric vehicle smart charging system in Malaysia. *World Electric Vehicle Journal*, 15(10), 440. <https://doi.org/10.3390/wevj15100440>
- Vaishali, K., & Prabha, D. R. (2024). The reliability and economic evaluation approach for various configurations of EV charging stations. *IEEE Access*, 12, 26267-26280.
<https://doi.org/10.1109/ACCESS.2024.3367133>
- Wu, Q., Li, D., Li, J., & Wu, F. (2026). Fault prediction method of electric vehicle charging pile based on hierarchical attention mechanism. *Peer-to-Peer Networking and Applications*, 19, 109. <https://doi.org/10.1007/s12083-026-02269-9>
- Younus, H., Kabir, S., Campean, F., Bonnaud, P., & Delaux, D. (2026). AI- and ontology-based enhancements to FMEA for advanced systems engineering: Current developments and future directions. *Applied Sciences*, 16(5), 2464.
<https://doi.org/10.3390/app16052464>
- Zahedieh, S., Oulis Rousis, A., & Bourazeri, A. (2025). Toward a unified framework for electric vehicle infrastructure: Insights into charging management, communication and emerging technologies. *Energy Reports*, 14, 1834-1854.
<https://doi.org/10.1016/j.egyr.2025.08.011>
- Yin, R. K. (2018). *Case study research and applications: Design and methods* (6th ed.). SAGE Publications.