

Vorhersage der Einspeiseleistung lokaler Photovoltaikanlagen mit Methoden des Maschinellen Lernens für das Lademanagement von Elektrofahrzeugen

von der

Fakultät für Elektrotechnik und Informationstechnik
der Technischen Universität Dortmund

genehmigte

DISSERTATION

zur Erlangung des akademischen Grades

Doktor der Ingenieurwissenschaften (Dr.-Ing.)

vorgelegt von

Katrin Schulte, M. Eng.

Referent: Prof. Dr.-Ing. habil. Christian Rehtanz, Technische Universität Dortmund

Korreferent: Prof. Dr.-Ing. Jens Haubrock, Hochschule Bielefeld

Tag der mündlichen Prüfung: 11.05.2026

Danksagung

Diese Dissertation ist im Rahmen meiner Tätigkeit als wissenschaftliche Mitarbeiterin an der Hochschule Bielefeld in der Arbeitsgruppe Netze und Energiesysteme (AGNES) in einer kooperativen Promotion mit der Technischen Universität Dortmund am Institut für Energiesysteme, Energieeffizienz und Energiewirtschaft entstanden. In dieser Zeit hatte ich die Möglichkeit, mich sowohl fachlich als auch persönlich weiterzuentwickeln. Dies verdanke ich zahlreichen Menschen, denen ich an dieser Stelle meinen herzlichen Dank aussprechen möchte.

Mein besonderer Dank gilt Herrn Prof. Dr.-Ing. Jens Haubrock, der mich während meiner Promotion an der Hochschule Bielefeld fachlich und persönlich begleitet hat. Seine wertvollen Anregungen und seine kontinuierliche Unterstützung bereits während des Studiums haben maßgeblich zur Entstehung dieser Arbeit beigetragen. Ebenfalls danken möchte ich Herrn Prof. Dr.-Ing. Christian Rehtanz für die Betreuung seitens der Technischen Universität Dortmund. Durch seine Unterstützung konnte ich meine kooperative Promotion in einem spannenden wissenschaftlichen Umfeld durchführen.

Ein großes Dankeschön geht an meine Kolleginnen und Kollegen aus der Arbeitsgruppe AGNES, die mit ihrer offenen und motivierenden Art ein großartiges Arbeitsumfeld geschaffen haben. Besonders die anregenden Diskussionen, die gemeinsame Teilnahme an Konferenzen und unsere Unternehmungen abseits der Arbeit haben diese Zeit besonders wertvoll gemacht. Zudem danke ich allen Studierenden, die mich in meiner Forschung unterstützt haben. Deren Mitarbeit hat einen wichtigen Beitrag zu dieser Arbeit geleistet.

Schließlich möchte ich mich bei meiner Familie und meinem Freund bedanken, die mir immer den Rücken freigehalten, mich unterstützt und für die nötige Ablenkung gesorgt haben.

Kurzfassung

Der Ausbau dezentraler erneuerbarer Energien wie Photovoltaik (PV)-Anlagen sowie die zunehmende Elektrifizierung des Wärme- und Verkehrssektors bedingen einen Ausbau der Verteilnetze im Rahmen der deutschen Energiewende. Da der Netzausbau nicht mit der erforderlichen Geschwindigkeit voranschreitet, steigt der Bedarf nach einem intelligent koordinierten Energiesystem. Eine wesentliche Grundlage dafür bildet die vorausschauende Planung mittels lokaler PV-Erzeugungsprognosen. Allerdings sind in Deutschland derzeit nicht ausreichend PV-Leistungsmessdaten für derartige Prognosen verfügbar, da der flächendeckende Einsatz von Smart Metern sich weiter verzögert und der Anteil installierter Smart Meter gering ist. Das Ziel dieser Arbeit ist die Entwicklung und Untersuchung eines zweistufigen PV-Prognosemodells, bestehend aus einer PV-Leistungsprognose für einen Prognosehorizont von 24 Stunden und einem Korrekturmodell zur Reduzierung von Prognosefehlern. Die PV-Prognose kombiniert eine Vorhersage der Bestrahlungsstärke auf Basis einer Wettervorhersage mithilfe eines Künstlichen Neuronalen Netzes und ein physikalisches PV-Modell, das die PV-Leistung anhand der Bestrahlungsstärke anlagenspezifisch berechnet. Die PV-Prognose ist dadurch im ersten Schritt für verschiedene PV-Anlagen anwendbar und benötigt keine PV-Leistungsmessdaten. Im zweiten Schritt reduziert ein Stacking-Ensemble als Korrekturmodell Prognosefehler, indem es Leistungsmessdaten berücksichtigt. Die Nutzung ausschließlich frei verfügbarer Wetterdaten gewährleistet die Praktikabilität der Methode. Die Korrektur ist durch die geringe Anzahl an PV-Leistungsmessdaten, die direkt in der Anwendung gesammelt werden können, möglich. Die Validierung des zweistufigen PV-Prognosemodells erfolgt anhand einer realen PV-Anlage. Bei einer Day-Ahead-Prognose der PV-Leistung beträgt der Prognosefehler 9,33 % für einen Testzeitraum von 629 Tagen. Eine stündliche Aktualisierung der Prognose führt zu einer leichten Reduzierung des Fehlers. Das Korrekturmodell reduziert den Prognosefehler auf 7,67 % im Vergleich zu 8,76 % bei einer PV-Prognose ohne Korrektur für einen Testzeitraum von 231 Tagen. Die Wirksamkeit des Korrekturmodells zeigt sich insbesondere in einer Winterwoche, in der die PV-Anlage teilweise verschattet ist. Das Modell identifiziert den wiederkehrenden Fehler durch Verschattung und reduziert ihn. Prognosefehler, die durch die Wettervorhersage bedingt sind, sind aufgrund ihres zufälligen Charakters schwerer durch das Korrekturmodell zu identifizieren. Bereits etwas mehr als sechs Monate chronologisch gesammelter PV-Leistungsmessdaten für das Training des Korrekturmodells reichen aus, um die Prognose für verschiedene Wetterbedingungen zu optimieren. Der Nutzen der PV-Prognose wird durch ihre Anwendung im Lademanagement von Elektrofahrzeugen demonstriert, wobei eine Optimierung auf Grundlage der PV-Prognose erfolgt.

Abstract

The expansion of decentralized renewable energies such as photovoltaic (PV) systems and the increasing use of electrical consumers such as electric vehicles as part of the German energy transition require an expansion of the distribution grids. As grid expansion is not progressing at the required speed, the need for an intelligently coordinated energy system is increasing. An essential basis for this is predictive planning using local PV generation forecasts. However, there is currently insufficient PV power measurement data available in Germany for such forecasts, as the widespread use of smart meters continues to be delayed and the share of installed smart meters is low. The aim of this work is to develop and investigate a two-stage PV power forecast model consisting of a day-ahead PV power forecast model and a correction model. The PV forecast combines a prediction of the solar irradiance based on a weather forecast using an artificial neural network and a physical PV model that calculates the PV power based on the solar irradiance for a specific PV system. In the first step, the PV forecast can therefore be used for different PV systems and does not require any PV power measurement data. In the second step, a stacking ensemble as a correction model reduces forecast errors by taking power measurement data into account. The use of only publicly available weather data ensures the practicability of the method. The correction is enabled by a small amount of PV power measurement data that can be collected directly in the application. The two-stage PV forecasting model is validated using a real PV system. The day-ahead PV power forecast has a forecast error of 9.33% over a test period of 629 days. An hourly update of the forecast leads to a slight reduction in the error. The correction model reduces the forecast error to 7.67% compared to 8.76% for a PV forecast without correction for a test period of 231 days. The effectiveness of the correction model is particularly evident in a winter week in which the PV system is partially shadowed. The model identifies the recurring error due to shadowing and reduces it. Forecast errors caused by the uncertainties in the weather forecast are more difficult to identify using the correction model due to their random nature. A little more than six months of chronologically collected PV power measurement data for training the correction model is sufficient to optimize the forecast for different weather conditions. The benefits of the PV forecast are demonstrated by its application in the charging management of electric vehicles, where optimization is based on the PV forecast.

Inhaltsverzeichnis

Danksagung	III
Kurzfassung.....	V
Abstract.....	VII
1 Einleitung	1
1.1 Hintergrund	1
1.2 Ziele der Arbeit und Forschungsthemen	4
1.3 Struktur der Arbeit	7
2 Stand der Wissenschaft und Technik.....	9
2.1 Methoden des Maschinellen Lernens.....	9
2.1.1 Künstliche Neuronale Netze.....	10
2.1.2 Lineare Regression	18
2.1.3 K-Nächste-Nachbarn	21
2.1.4 Entscheidungsbaum.....	22
2.1.5 Clusteranalyse	24
2.1.6 Ensemble-Lernen.....	25
2.2 Methoden der PV-Prognose in der Wissenschaft	29
2.3 Zusammenfassung.....	41
3 PV-Prognose	43
3.1 Methode der Prognose der Bestrahlungsstärke	44
3.1.1 Datenaufbereitung und -aufteilung.....	45
3.1.2 Training und Validierung eines ANN	49
3.2 Physikalisches PV-Modell	57
3.2.1 Berechnung der Sonnenposition.....	59
3.2.2 Berechnung der PV-Leistung	61
3.2.3 Darstellung der PV-Anlage und Validierung des PV-Modells	65
3.3 Ergebnisse der Day-Ahead-Prognose einer PV-Anlage	69
	IX

3.4	Zusammenfassung.....	74
4	Korrekturmodell der PV-Prognose.....	77
4.1	Methode des Stacking-Ensembles.....	78
4.1.1	Datenaufbereitung und -aufteilung.....	79
4.1.2	Auswahl und Training der ML-Modelle	85
4.2	Ergebnisse der korrigierten Day-Ahead-Prognose einer PV-Anlage	92
4.3	Zusammenfassung.....	98
5	Anwendung der PV-Prognose für das Lademanagement von Elektrofahrzeugen.....	101
5.1	Definition einer Optimierungsaufgabe mittels linearer Optimierung.....	102
5.2	Ergebnisse der Elektrofahrzeug-Ladesteuerung auf Basis der PV-Prognose	107
5.3	Zusammenfassung.....	113
6	Zusammenfassung und Ausblick.....	115
6.1	Zusammenfassung.....	115
6.2	Bezugnahme zu den Forschungsthemen.....	117
6.3	Ausblick	121
	Literaturverzeichnis	123
	Abkürzungs- und Formelzeichenverzeichnis	133
	Abbildungs- und Tabellenverzeichnis.....	139
A.	Anhang PV-Prognose.....	145
B.	Anhang Korrekturmodell.....	157
C.	Anhang Wissenschaftlicher Tätigkeitsnachweis	165
D.	Hinweis zur Nutzung generativer KI	171

1 Einleitung

Zu Beginn des Kapitels wird der Hintergrund dargestellt, um den Ausgangspunkt der vorliegenden Arbeit zu erläutern. Hierbei wird die Notwendigkeit lokaler Prognosen der erzeugten Leistung durch Photovoltaikanlagen einschließlich der damit verbundenen Herausforderungen hervorgehoben. Aufbauend auf dieser Einführung wird das Ziel der Arbeit definiert, um die Neuheiten und den wissenschaftlichen Beitrag zu verdeutlichen. Anschließend werden die zentralen Forschungsthemen formuliert, die im weiteren Verlauf der Arbeit untersucht und überprüft werden. Abschließend bietet das Kapitel einen strukturierten Überblick über den Aufbau der Arbeit.

1.1 Hintergrund

Im Zuge der Energiewende in Deutschland und Europa sowie dem damit verbundenen Ausbau erneuerbarer Energien gewinnt die Photovoltaik (PV) zunehmend an Bedeutung. Mit einem Anteil von 14 % des Bruttostromverbrauchs im Jahr 2024 gehört die PV bereits zu den wichtigsten Stromerzeugungsarten in Deutschland [1]. PV-Anlagen entwickeln sich zu einem essenziellen Bestandteil der Energieversorgung. Während zwischen 2013 und 2018 im Durchschnitt jährlich lediglich 1,9 GW PV-Leistung neu installiert wurde, ist der Zubau für 2024 bereits auf 16,2 GW gestiegen [1, 2]. Gemäß den im Erneuerbare-Energien-Gesetz (EEG) 2023 festgelegten Ausbauzielen ist eine jährliche Steigerung der installierten PV-Leistung auf 22 GW vorgesehen, mit dem Ziel, bis 2030 eine installierte Leistung von 215 GW zu erreichen [3]. Diese Zahlen verdeutlichen den schnellen Anstieg des PV-Zubaus in Deutschland, insbesondere in den letzten Jahren, sowie die ambitionierten Ausbaupläne für die Zukunft. Etwa zwei Drittel der aktuell installierten PV-Leistung in Deutschland entfallen auf PV-Dachanlagen [4]. Ein signifikanter Anteil der installierten Leistung ist somit auf kleinere Anlagen in der Niederspannung zurückzuführen, die dezentral auf Gebäuden installiert sind.

In Niederspannungsnetzen kann es in Regionen mit einer hohen Konzentration von PV-Anlagen zu besonderen Herausforderungen kommen. Vor allem zur Mittagszeit an sonnigen Tagen speisen viele Anlagen bedingt durch den hohen Gleichzeitigkeitsfaktor große Mengen Strom ins Netz ein. Gleichzeitig ist der lokale Stromverbrauch in dieser Zeit häufig niedrig, was dazu führt, dass die Erzeugung lokal höher ist als der Verbrauch. Dadurch fließt der überschüssige Strom zurück in das Mittelspannungsnetz, was zu einer Überlastung des Netzes sowie der Transformatoren führen kann. Solche Herausforderungen treten vor allem in ländlichen Regionen auf, in denen entweder große PV-Anlagen betrieben werden oder viele kleinere Anlagen auf engem Raum installiert sind [1].

Der zunehmende Ausbau dezentraler erneuerbarer Energien wie PV-Anlagen sowie die zunehmende Elektrifizierung des Wärme- und Verkehrssektors durch Elektrofahrzeuge und Wärmepumpen erfordern sowohl den Ausbau der Verteilnetze als auch deren Digitalisierung [5, 6]. Dieser Prozess kann jedoch nicht sofort umgesetzt werden und bringt zahlreiche Herausforderungen mit sich. Dazu zählen die unterschiedlichen Strukturen im Verteilnetz, die sich durch vielfältige Erzeuger und Verbraucher auszeichnen, sowie die hohen Kosten des Netzausbaus. Zusätzlich stehen die Verteilnetzbetreiber vor Problemen wie Personalmangel, umfangreichen Genehmigungsverfahren und einer begrenzten Akzeptanz der Bevölkerung für den Netzausbau [7]. Aus diesen Gründen gewinnen alternative Ansätze im Rahmen einer Ausgestaltung eines intelligent koordinierten Netzes an Bedeutung, um den Netzausbau zeitlich zu verschieben und Netzengpässe zu vermeiden.

Ein intelligentes Energiesystem zeichnet sich durch eine hohe Komplexität bedingt durch die Vielzahl dezentraler Erzeuger, neuer Verbraucher, Speicher sowie der Sektorenkopplung aus. Dadurch steigt der Bedarf an einer koordinierten Steuerung aller Komponenten. Eine wesentliche Grundlage dafür bildet die vorausschauende Planung mittels verlässlicher Prognosen [8]. Insbesondere bei der Integration von PV-Anlagen ergeben sich aus der Volatilität der Stromerzeugung Herausforderungen, die zu Netzengpässen und Überlastungen führen können. Durch präzise Prognosen können Flexibilitäten durch Lastverschiebung genutzt werden, indem der Verbrauch planbar an die volatile Einspeisung angepasst werden kann oder Batteriespeicher rechtzeitig geladen oder entladen werden [1]. Daher gewinnen lokale Erzeugungsprognosen zunehmend an Bedeutung für die Integration erneuerbarer Energien in das Energiesystem.

Regionale Prognosen zur Einspeisung aus PV-Anlagen werden bereits seit mehreren Jahren im Übertragungsnetz in der jeweiligen Regelzone für das Fahrplanmanagement verwendet [9]. Prognosen für die regionale PV-Erzeugung sind dabei weniger fehleranfällig als lokale Prognosen oder Vorhersagen für einzelne Anlagen, da sich durch die dezentrale Verteilung der PV-Anlagen lokale Bewölkungsveränderungen gegenseitig ausgleichen. Dadurch schwankt die gesamte Stromproduktion aus PV vergleichsweise geringfügig [1]. Mit dem weiteren Ausbau erneuerbarer Energien und dem zunehmenden Anschluss leistungsstarker Verbraucher wächst jedoch auch auf Verteilnetzebene der Bedarf an präzisen Engpassvorhersagen. Um potenzielle Netzüberlastungen frühzeitig zu erkennen und deren Auftreten zu vermeiden, gewinnen lokale Einspeiseprognosen auf Verteilnetzebene zunehmend an Bedeutung [9]. Aufgrund ihrer räumlichen Begrenzung und der hohen Abhängigkeit von lokalen Wetterveränderungen sind lokale Prognosen bedeutend komplexer. Der Anwendungsfall lokaler PV-Leistungsprognosen erfordert unterschiedliche Prognosehorizonte, die jeweils mit einer bestimmten Prognosegenauigkeit einhergehen. Die Prognose der Leistungserzeugung einzelner PV-Anlagen über den Tag hinweg ermöglicht eine gezielte Verschiebung von Lasten wie die

Ladung von Elektrofahrzeugen oder die Ladung und Entladung von Batteriespeichern zu netzdienlichen Zwecken im Rahmen von Energiemanagementsystemen im Haushalt [10–12].

Angeichts dieser Komplexität stoßen klassische Prognosemethoden zunehmend an ihre Grenzen. Im Zuge der Energiewende kommt dem Maschinellen Lernen (*engl. Machine Learning, ML*) eine entscheidende Bedeutung zu. ML ist dazu fähig, große Datenmengen zu analysieren und komplexe, nichtlineare Zusammenhänge zu erkennen und zu erlernen. Insbesondere für die Prognose der lokalen PV-Erzeugung eignet sich dieser datengetriebene Ansatz, da er in der Lage ist, verschiedene Informationen, wie beispielsweise Wetterdaten bezüglich der Bestrahlungsstärke oder der Temperatur sowie historische PV-Leistungsmessdaten zu integrieren. Die zunehmende Verfügbarkeit hochaufgelöster Messdaten durch die Digitalisierung schafft dabei die notwendige Datenbasis für den Einsatz von ML im Energiesystem [8].

In Deutschland mangelt es jedoch an Messdaten der erzeugten Leistung einzelner PV-Anlagen, die über einen längeren Beobachtungszeitraum vorliegen. Dies ist darauf zurückzuführen, dass der Anteil verbauter Smart Meter in Deutschland zum Stichtag am 30. September 2024 bei lediglich 1,9 % lag [13]. In der Hälfte der EU-Mitgliedsstaaten sind es 2023 bereits über 80 % [14]. Smart Meter messen sowohl den Stromverbrauch als auch den erzeugten Strom in Haushalten mit Erzeugungsanlagen wie PV-Anlagen in kurzen Zeitintervallen, üblicherweise alle 15 Minuten. Die erfassten Daten werden in einem bestimmten Rhythmus zunächst an den Messstellenbetreiber übermittelt, der sie anschließend je nach Bedarf an den Netzbetreiber, den Energieversorger oder weitere berechnigte Akteure weiterleitet [9, 15]. Durch das am 27. Mai 2023 in Kraft getretene Gesetz zum Neustart der Digitalisierung der Energiewende ist das Messstellenbetriebsgesetz (MsbG) novelliert worden, um den Rollout von Smart Metern in Deutschland zu beschleunigen. Ein Einbau eines Smart Meters ist verpflichtend für Haushalte mit einem Jahresverbrauch ab 6.000 kWh, Haushalte mit steuerbaren Verbrauchseinrichtungen gemäß §14a Energiewirtschaftsgesetz (EnWG), wie Wärmepumpen oder Wallboxen zur Ladung von Elektrofahrzeugen und Haushalte mit Erzeugungsanlagen wie PV-Anlagen ab 7 kW installierter Leistung [15]. Die Umsetzung soll bis 2032 abgeschlossen sein. Die Dynamik des Rollouts bleibt allerdings vorerst unverändert und eine Beschleunigung ist aktuell nicht in Sicht. Zu den Hauptursachen für das anhaltend langsame Tempo beim Ausbau von Smart Metern zählen die hohen Anforderungen an die Smart Meter, die hohen Sicherheitsstandards für Lieferketten sowie die geringe Wirtschaftlichkeit des Einbaus [5]. Derzeit ist daher nur eine geringe Menge hochaufgelöster Messdaten zur erzeugten PV-Leistung einzelner PV-Anlagen verfügbar. Zudem ist unklar, wie viele PV-Anlagen über private Energiemanagementsysteme verfügen, die langfristig Leistungsdaten speichern. Besonders in älteren Bestandsanlagen ist die Wahrscheinlichkeit gering, dass solche Systeme vorhanden sind. Selbst wenn Messdaten erfasst werden,

können diese zwar im privaten Bereich wie für Energiemanagementanwendungen im Haushalt genutzt werden, jedoch nicht für netzdienliche Zwecke durch den Netzbetreiber. Ein flächendeckender Einsatz von Smart Metern in Deutschland wird sich voraussichtlich weiterhin verzögern [5].

1.2 Ziele der Arbeit und Forschungsthemen

Ausgehend von den beschriebenen Herausforderungen bei lokalen Prognosen der PV-Leistung verfolgt diese Arbeit das Ziel, ein neuartiges zweistufiges PV-Prognosemodell zu entwickeln und zu untersuchen. Das Prognosemodell besteht aus einer PV-Leistungsprognose, die im weiteren Verlauf als PV-Prognose bezeichnet wird, und einem Korrekturmodell. Die Funktionsweise des zweistufigen PV-Prognosemodells im Betrieb ist in Abbildung 1.1 dargestellt. Die PV-Prognose besteht aus einer Prognose der Bestrahlungsstärke und einem physikalischen PV-Modell, das die prognostizierte Bestrahlungsstärke in die erzeugte Leistung umrechnet. Die Prognose der Bestrahlungsstärke erfolgt auf Basis von Wettervorhersagedaten, die als Eingangsdaten in ein mit historischen Wetterdaten trainiertes Künstliches Neuronales Netz (*engl. Artificial Neural Network, ANN*) integriert werden. Der Prognosehorizont ist dabei abhängig vom Prognosehorizont der numerischen Wettervorhersage (*engl. Numerical Weather Prediction, NWP*) und beträgt in dieser Arbeit 24 Stunden, was einer Prognose von einem Tag im Voraus entspricht. Eine Prognose, die sich auf einen Zeitraum von 24 Stunden im Voraus bezieht, wird auch als Day-Ahead-Prognose bezeichnet. Die für den Standort der PV-Anlage prognostizierte Bestrahlungsstärke wird im physikalischen PV-Modell in die von der PV-Anlage erzeugte Leistung umgerechnet. Da die Bestrahlungsstärke für einen ausgewählten Standort prognostiziert wird, kann die PV-Prognose auf verschiedene PV-Anlagen angewendet werden und ist nicht auf eine einzelne Anlage beschränkt. Das Korrekturmodell basiert auf einem Stacking-Ensemble verschiedener ML-Modelle, die mit einem geringen Datensatz von PV-Leistungsmessdaten trainiert sind. Die prognostizierte Leistung eines Tages einer PV-Anlage sowie weitere zeitliche Angaben werden als Eingabe für das Korrekturmodell verwendet, um saisonale und tägliche Trends zu berücksichtigen. Die Ausgabe des Korrekturmodells ist eine korrigierte prognostizierte PV-Leistung. Das Ziel besteht in der Reduzierung wiederkehrender Prognosefehler, die beispielsweise durch Verschattung oder Alterung der PV-Anlage bedingt sind und allein durch die Prognose auf Basis von Wetterdaten nicht erfassbar sind. Dafür werden zusätzlich Leistungsmessdaten der PV-Anlage zum Training eines Stacking-Ensembles genutzt. Die Anwendung der Day-Ahead-PV-Prognose für das Lademanagement von Elektrofahrzeugen wird untersucht, um deren Nutzen für die Lastverschiebung zu verdeutlichen.

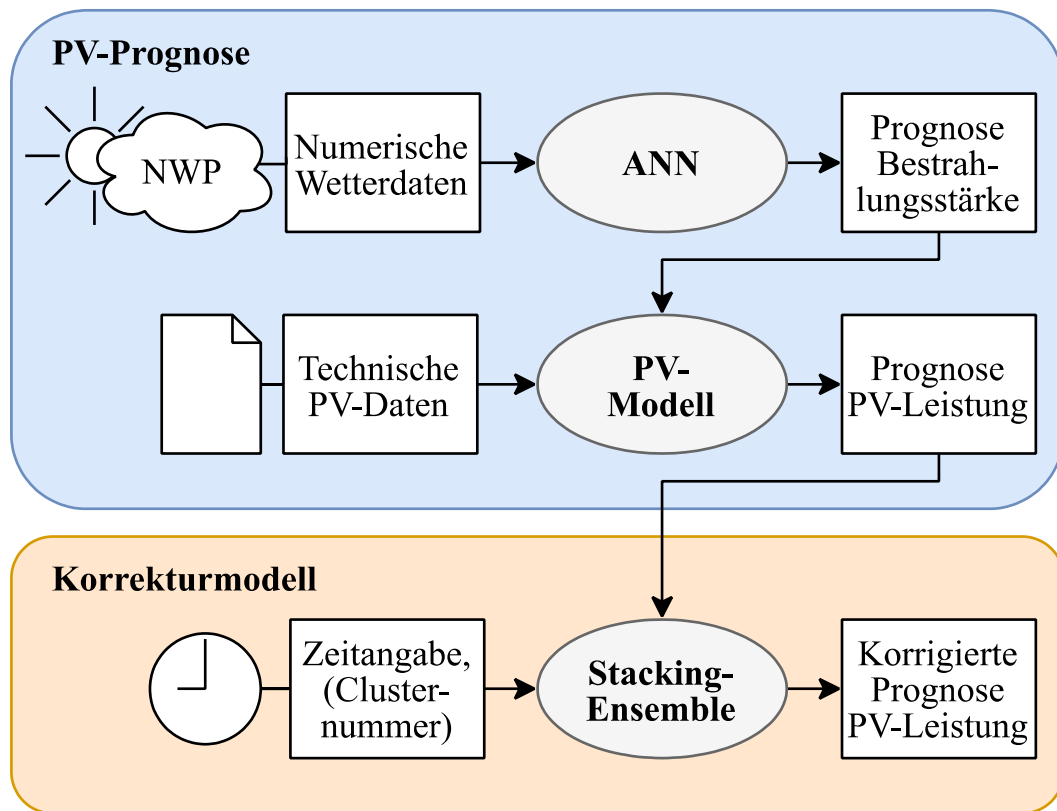


Abbildung 1.1: Übersicht über die Funktionsweise des zweistufigen PV-Prognosemodells im Betrieb

Im Rahmen dieser Arbeit werden die nachfolgenden Forschungsthese formuliert und überprüft. In Deutschland existieren bislang nur wenige gespeicherte Messdaten zur erzeugten Leistung einzelner PV-Anlagen. Daher fehlen geeignete Leistungswerte für das Training eines ML-Modells und selbst ein trainiertes Modell wäre nur auf die jeweilige PV-Anlage anwendbar. Aus diesem Grund wird ein Ansatz betrachtet, der ausschließlich auf historischen Wetterdaten basiert. Das Modell nutzt Wettervorhersagen, um die erzeugte Leistung für einen bestimmten Zeitraum vorherzusagen. Da Angaben zur Bestrahlungsstärke in klassischen Wettervorhersagen in der Regel nicht kostenfrei verfügbar sind, ergibt sich daraus die erste Forschungsthese:

Die lokal erzeugte PV-Leistung einer einzelnen PV-Anlage kann mit Hilfe von ML-Methoden unter der Verwendung von frei verfügbaren Wetterdaten prognostiziert werden.

Prognosefehler können auf unterschiedliche Faktoren, wie beispielsweise die verwendete Wettervorhersage oder anlagenspezifische Gegebenheiten durch eine Verschattung oder eine Alterung der PV-Anlage, zurückzuführen sein. Solche Fehler können durch eine PV-Prognose alleine auf Basis von Wetterdaten nicht erkannt oder durch ein PV-Modell nur schlecht abgebildet werden. Wenn die PV-Prognose auf Basis von Wetterdaten bereits im Haushalt im Einsatz ist, können parallel Messdaten der erzeugten PV-

Leistung gesammelt werden. Diese Daten können dazu genutzt werden, die PV-Prognose schrittweise zu optimieren und Prognosefehler zu reduzieren. Die Reduzierung dieser Prognosefehler sollte insbesondere bereits bei einer geringen Anzahl an Messdaten erfolgen, da ansonsten eine unmittelbare Prognose auf Basis der Leistungsdaten möglich wäre. Die zweite Forschungsthese lautet daher wie folgt:

Der Fehler der PV-Prognose kann durch die Verwendung weniger Messdaten der erzeugten Leistung der PV-Anlage reduziert werden.

Die Notwendigkeit lokaler PV-Prognosen wird am Anwendungsfall der Steuerung von Elektrofahrzeugen auf Basis einer PV-Prognose demonstriert. Durch die Integration solcher Prognosen in Energiemanagementsysteme kann der Stromverbrauch gezielt auf Basis der erwarteten PV-Erzeugung geplant und die volatile Einspeisung direkt vor Ort genutzt werden. Im Rahmen der fortschreitenden Elektrifizierung des Mobilitätssektors ermöglicht dies eine intelligente Verteilung der Ladevorgänge von Elektrofahrzeugen über den Tag. So kann der Anteil von PV-Strom beim Laden deutlich erhöht, die Last durch Elektrofahrzeuge flexibel verschoben und die verfügbare Flexibilität optimal genutzt werden. Zudem kann die direkte Nutzung des lokal erzeugten PV-Stroms dazu beitragen, Netzüberlastungen zu reduzieren und potenziell zu vermeiden. Die dritte Forschungsthese lautet daher wie folgt:

Die Anwendung der PV-Prognose für das Lademanagement von Elektrofahrzeugen ermöglicht eine Maximierung des Anteils von PV-Leistung in der Ladeleistung.

In der vorliegenden Arbeit erfolgt erstmals auf Basis vorhandener Datenlage die Entwicklung und Untersuchung eines zweistufigen PV-Prognosemodells, das sich durch Praktikabilität auszeichnet und mit einem umfangreichen Datensatz validiert wird. Der neuartige Ansatz besteht darin, dass ausschließlich frei verfügbare Daten wie historische Wetterdaten und kostenfreie Wettervorhersagen für das Training und die Prognose genutzt werden. Diese Vorgehensweise gewährleistet nicht nur die Praktikabilität der Methodik, sondern ermöglicht auch eine unmittelbare Anwendung der PV-Prognose. Der erste Teil des zweistufigen PV-Prognosemodells kann direkt zu Beginn eingesetzt werden, ohne dass Messdaten der erzeugten PV-Leistung vorab gesammelt werden müssen. Im Rahmen der Anwendung der PV-Prognose besteht die Möglichkeit, Leistungsmessdaten schrittweise zu sammeln, um eine Reduzierung von Prognosefehlern zu erzielen. Für die Validierung des Prognosemodells anhand einer realen PV-Anlage wird eine umfassende Datenbasis herangezogen, sodass ein großer Zeitraum betrachtet wird, in dem alle Wetterbedingungen abgebildet werden.

1.3 Struktur der Arbeit

Die vorliegende Dissertation ist in sechs Kapitel unterteilt. In Kapitel 1 wird eine Einführung in die Thematik gegeben, die den Hintergrund und die Notwendigkeit von PV-Prognosen beleuchtet. Zudem werden die Forschungsthemen präsentiert, die im weiteren Verlauf der Arbeit untersucht werden. Die drei formulierten Forschungsthemen werden in den Kapiteln 3, 4 und 5 jeweils einzeln untersucht, sodass jeder These ein eigenes Kapitel gewidmet ist. In Kapitel 2 wird der Stand der Wissenschaft und Technik erörtert und der theoretische Hintergrund für die Methodik dieser Arbeit dargelegt. Dazu wird der Begriff ML definiert und die für diese Arbeit relevanten ML-Methoden ANN, Lineare Regression, K-Nächste-Nachbarn, Entscheidungsbaum, Clusteranalyse und Ensemble-Lernen erläutert. Darüber hinaus wird ein Überblick über den aktuellen Stand der Wissenschaft und damit über verschiedene wissenschaftliche Arbeiten zum Thema Methoden der PV-Prognose gegeben, um die Neuheiten dieser Dissertation zu belegen. Da das zweitstufige PV-Prognosemodell in zwei Teile unterteilt ist, widmen sich die Kapitel 3 und 4 jeweils der Methodik und den Ergebnissen der PV-Prognose und des Korrekturmodells. In Kapitel 3 werden die Funktionsweise und die Elemente der PV-Prognose beschrieben, die aus der Prognose der Bestrahlungsstärke mittels eines ANN und einem physikalischen Modell zur Berechnung der PV-Leistung besteht. Im Rahmen dessen erfolgt eine Erläuterung der Methode zur Prognose der Bestrahlungsstärke. Zunächst werden die verwendeten Daten dargestellt und das Training sowie die Validierung eines ANN erläutert. Im Anschluss wird die Methode für das physikalische PV-Modell präsentiert. Die Ergebnisse der PV-Prognose werden anhand einer realen PV-Anlage präsentiert und diskutiert. In Kapitel 4 wird anschließend das Korrekturmodell präsentiert, das auf einem Stacking-Ensemble aus ML-Methoden basiert und dazu dient, wiederkehrende Prognosefehler zu reduzieren. Diese können beispielsweise durch Verschattung oder Alterung von PV-Anlagen entstehen. Es wird die Methodik des Korrekturmodells, einschließlich der Datenaufbereitung, Auswahl der Basismodelle sowie der Hyperparameter beschrieben. Im Anschluss erfolgt die Darstellung und Diskussion der Ergebnisse. Die Prognosen mit und ohne Korrektur werden für die bereits in Kapitel 3 betrachtete PV-Anlage miteinander verglichen, um zu analysieren, ob die Korrektur eine Reduktion der Prognosefehler bewirkt. Kapitel 5 widmet sich der Demonstration der Anwendung der PV-Prognose für das Lademanagement von Elektrofahrzeugen. In diesem Kapitel wird zunächst die Methodik der linearen Optimierung auf Basis der PV-Prognose erläutert und die Zielfunktion sowie die Nebenbedingungen für die Optimierungsaufgabe definiert. Im Anschluss werden die Ergebnisse der Validierung der Ladesteuerung auf Basis der PV-Prognose präsentiert. Das untersuchte Szenario umfasst die Ladung privater Elektrofahrzeuge von Mitarbeitern eines Unternehmens auf dem Un-

ternehmensgelände während der Arbeitszeit. Abschließend werden in Kapitel 6 die Ergebnisse zusammengefasst, ein Bezug zu den anfangs definierten Forschungsthemen hergestellt und ein Ausblick auf mögliche zukünftige Forschungsansätze gegeben.

2 Stand der Wissenschaft und Technik

Der Wandel des Energiesystems sowie die zunehmende Komplexität durch die Integration erneuerbarer Energien und die Elektrifizierung des Wärme- und Mobilitätssektors führen zu einer steigenden Bedeutung der Künstlichen Intelligenz (KI) im Kontext der Energiewende. Die Anwendung von ML-Methoden bietet aufgrund der zunehmenden Datenmenge und steigenden Rechenleistung ein großes Potenzial, den Wandel zu unterstützen [16]. In diesem Kapitel erfolgt daher zunächst eine Definition des Begriffs ML und eine Aufteilung in verschiedene Lernverfahren. Im weiteren Verlauf werden die für diese Arbeit relevanten ML-Methoden ANN, Lineare Regression, K-Nächste-Nachbarn, Entscheidungsbaum, Clusteranalyse und Ensemble-Lernen näher erläutert. Im Anschluss daran folgt ein Überblick über den aktuellen Stand der Wissenschaft und damit über verschiedene wissenschaftliche Arbeiten zu Methoden der PV-Prognose. Die wissenschaftliche Literatur zu Prognosemethoden für PV-Anlagen ist sehr umfangreich. Zur Strukturierung werden die Arbeiten nach in der wissenschaftlichen Literatur gängigen Merkmalen zur Gruppierung von PV-Prognosemethoden gruppiert.

2.1 Methoden des Maschinellen Lernens

ML ist ein Teilbereich von KI, dessen Ziel es ist, Algorithmen zu entwickeln, die in der Lage sind, aus vorhandenen Daten zu lernen. Dabei sollen auf Basis dieses Wissens Muster erkannt oder Vorhersagen für getroffen werden. Diese Methode findet insbesondere in Bereichen Anwendung, in denen die Prozesse komplex sind und herkömmliche Analysemethoden an ihre Grenzen kommen [17]. ML lässt sich hinsichtlich des Lernverfahrens in die drei Hauptkategorien überwachtes Lernen, unüberwachtes Lernen und bestärkendes Lernen unterteilen [18, 19]. Teilweise wird auch die Kategorie des teilüberwachten Lernens unterschieden [17]. In Abbildung 2.1 ist eine Übersicht über die Lernverfahren des ML sowie eine Einteilung ausgewählter Methoden dargestellt. Unter anderem werden darin die Methoden aufgezeigt, die in dieser Arbeit verwendet werden.

Beim überwachten Lernen werden Daten verwendet, bei denen sowohl die Eingabewerte als auch die Ausgabewerte bekannt sind. Das überwachte Lernen zielt darauf ab, Zusammenhänge zwischen diesen Daten zu erkennen, um auf dieser Grundlage Vorhersagen für neue Eingangsdaten treffen zu können. In diesem Kontext ist zwischen den zwei Varianten des überwachten Lernens der Klassifikation und der Regression zu unterscheiden. Bei der Klassifikation trifft das Modell die Entscheidung, welcher Klasse ein Datenpunkt zugeordnet wird. Im Falle der Regression werden wie auch in dieser Arbeit kontinuierliche Werte, wie beispielsweise Temperaturen, vorhergesagt [17, 19].

In Abbildung 2.1 sind ausgewählte ML-Methoden innerhalb dieses Lernverfahrens dargestellt. Es existieren Methoden, die sowohl für eine Klassifikation als auch für eine Regression Anwendung finden, wie beispielsweise die K-Nächste-Nachbarn-Methode oder ANNs. Es existieren Methoden, die ausschließlich für die Klassifikation oder ausschließlich für die Regression eingesetzt werden können. Diesbezüglich ist die lineare Regression für die Regression und die logistische Regression für die Klassifikation zu nennen.

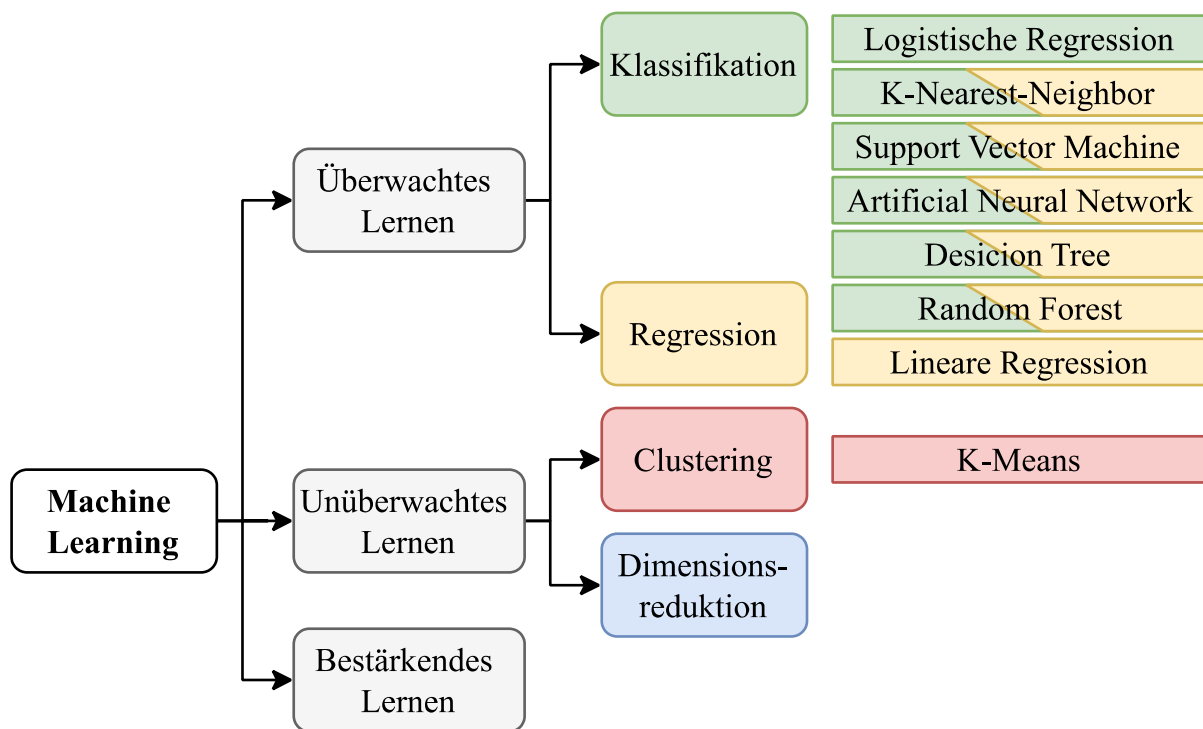


Abbildung 2.1: Übersicht über Lernverfahren und Methoden des ML nach [17, 20]

Im Gegensatz zum überwachten Lernen sind die Ausgabewerte beim unüberwachten Lernen dem System nicht bekannt. Es arbeitet lediglich mit den Eingabedaten und versucht, eigenständig Muster oder Regelmäßigkeiten zu identifizieren [17, 19]. Weit verbreitete Ansätze des unüberwachten Lernens sind die Clusteranalyse und die Dimensionsreduktion [21]. Im Rahmen des bestärkenden Lernens agiert ein Agent in einer Umgebung, führt verschiedene Aktionen aus und erhält auf Basis dieser Aktionen positive oder auch negative Belohnungen. Das übergeordnete Ziel besteht in der Identifikation einer optimalen Strategie, welche dem Agenten dabei hilft, in der Umgebung die Summe an Belohnungen zu maximieren. Das System lernt folglich aus seinen Erfahrungen und passt seine Aktionen entsprechend an [17, 19].

2.1.1 Künstliche Neuronale Netze

ANNs orientieren sich am Aufbau des menschlichen Gehirns und verarbeiten Informationen ähnlich dem Nervensystem. Dabei erfolgt die Informationsverarbeitung parallel

durch eine Vielzahl miteinander verbundener Neuronen, die Informationen über Aktivierungssignale weiterleiten [22]. Im Unterschied zum menschlichen Gehirn bearbeiten Computer Aufgaben meist sequenziell. Trotz Fortschritten in der Parallelverarbeitung bleibt die Leistung und Art der Datenverarbeitung von Computern weit hinter dem menschlichen Gehirn zurück. Dadurch können komplexe Aufgaben wie Muster- und Spracherkennung oder Vorhersagen nur schwer von klassischen Programmen bewältigt werden, während das menschliche Gehirn solche Aufgaben dank seiner neuronalen Vernetzung und Lernfähigkeit meistert [23].

ANNs bestehen aus einer Vielzahl von Neuronen, die die Grundlage eines solchen Netzes bilden. Ein Neuron weist die beiden Zustände aktiv oder inaktiv auf, was dem Ausgang eines Neurons entspricht. Die Eingangssignale in ein Neuron resultieren entweder aus anderen Neuronen oder sind direkt die Eingangssignale des gesamten Systems [23]. In Abbildung 2.2 ist der Aufbau eines einzelnen Neurons mit seinen Eingangssignalen x_j und dem Ausgangssignal y dargestellt.

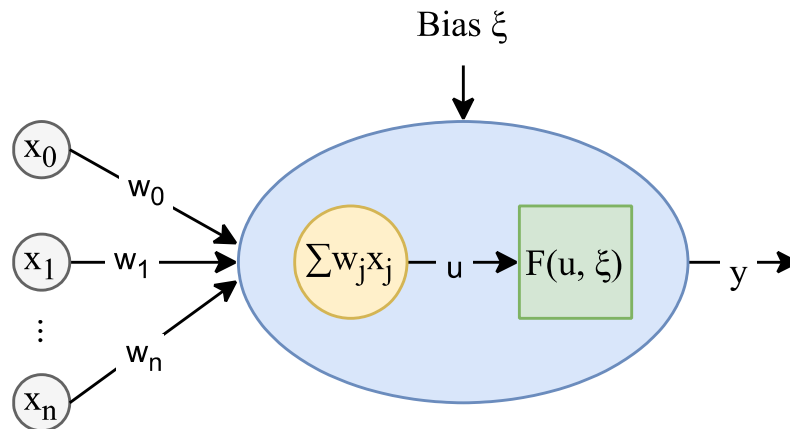


Abbildung 2.2: Aufbau eines Neurons in einem ANN nach [18, 23]

Die Eingangssignale mit der Anzahl n , die in ein Neuron fließen, werden gemäß der eingestellten Gewichte w_j gewichtet und innerhalb des Neurons gemäß Formel (2.1) aufsummiert [17, 23]. Eine Aktivierungsfunktion bzw. Übertragungsfunktion F rechnet diese Summe u entsprechend der Aktivierungsfunktion im Neuron um und aktiviert das Neuron, wenn ein Schwellwert innerhalb der Aktivierungsfunktion überschritten wird. Die Aktivierungsfunktion gibt an, wie stark das Neuron aktiviert wird [20]. In dieser Funktion findet ebenfalls die Verrechnung des sogenannten Bias-Wertes ξ , des Verschiebungswertes der Aktivierungsfunktion, statt. Bei Nichtüberschreiten des Schwellwertes befindet sich das Neuron im Zustand inaktiv. Dieser Zusammenhang ist in Formel (2.2) definiert [23]. Die Übertragungsfunktion F weist der gewichteten Summe hinsichtlich der gewählten Funktion einen Ausgangswert y zu.

$$u = \sum_{j=0}^n w_j \cdot x_j \quad (2.1)$$

$$y = F(u, \xi) \quad (2.2)$$

Das Verhalten eines ANN kann durch verschiedene anpassbare Parameter gesteuert werden. Wesentliche Parameter dabei sind die Gewichte der Eingangssignale w_j , der Bias-Wert ξ und die Aktivierungsfunktion F . Im Verlauf des Trainingsprozesses werden die Gewichte und Bias-Werte der Neuronen anhand eines bestimmten Algorithmus schrittweise so angepasst, bis das Netzwerk die gewünschten Ergebnisse erreicht. Damit speichert das ANN sein erlerntes Wissen in den Gewichten und Bias-Werten der Neuronen [23].

Die Auswahl der Aktivierungsfunktion erfolgt in der Regel in Abhängigkeit des zu lösenden Problems vor Beginn des Trainings [23]. In einigen Fällen ist die Ausgabe von Werten, die den Zahlen Eins oder Null entsprechen, von entscheidender Bedeutung. In anderen Fällen hingegen ist die Stetigkeit der Funktion von essenzieller Relevanz, da kontinuierliche Werte als Ausgabe erforderlich sind [20, 22, 23]. Eine Übersicht ausgewählter Aktivierungsfunktionen mit Bezeichnung, Verlauf und Formel ist in Abbildung 2.3 dargestellt.

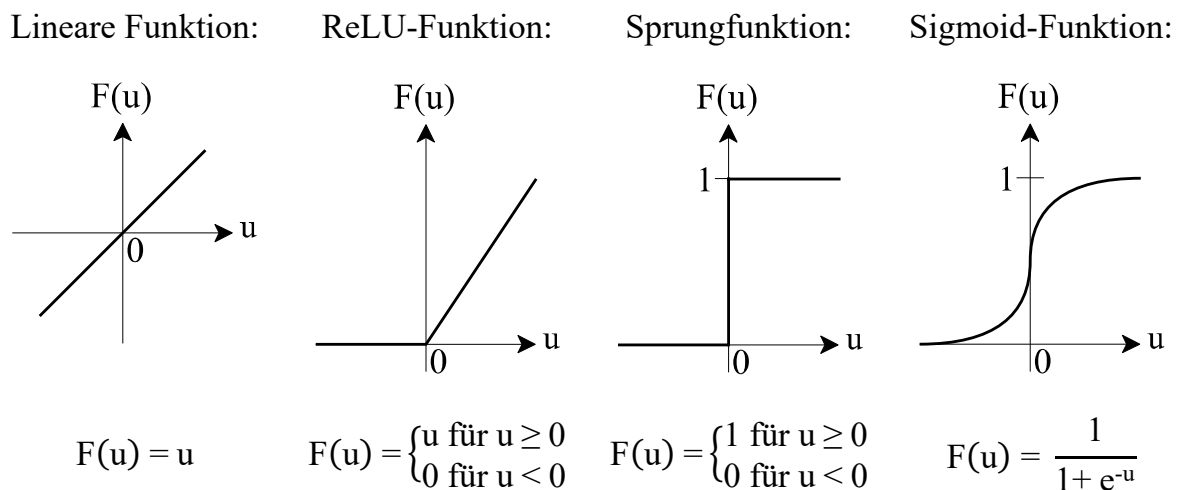


Abbildung 2.3: Übersicht über häufige Aktivierungsfunktionen nach [20, 22, 23]

Die Verwendung der Sigmoid-Funktion ermöglicht beispielsweise die Generierung von Wahrscheinlichkeitswerten durch die Ausgangsneuronen eines ANN [22]. Die Aktivierungsfunktion begrenzt den Bereich der Ausgangswerte auf Werte zwischen null und eins. Des Weiteren ist bei dieser Funktion die Stetigkeit von essenzieller Bedeutung, da in einigen Lernalgorithmen aufgrund des Einsatzes partieller Ableitung stetige Funktionen erforderlich sind [23]. Im Falle der ReLU-Funktion (Rectified Linear Units bzw.

rektifizierte Lineareinheiten) erfolgt eine Setzung negativer Werte auf null, wodurch dem ANN lediglich die Ausgabe positiver Werte ermöglicht wird [20]. Die Nichtlinearität von Aktivierungsfunktionen spielt eine zentrale Rolle, da sie ANNs erlaubt, komplexe Beziehungen zwischen Eingabe- und Ausgabedaten abzubilden [17]. Der Bias-Wert ξ in der Aktivierungsfunktion bewirkt eine Verschiebung der Funktion auf der horizontalen Achse, was eine Modifikation des Schwellenwerts zur Aktivierung des Neurons zur Folge hat [23].

Die Fähigkeit von ANNs, komplexe Probleme zu lösen, basiert in der Regel auf der Verwendung mehrerer Neuronen und Schichten, um komplexe Zusammenhänge zwischen Eingabe- und Ausgabedaten zu erkennen [17, 23]. In Bezug auf die Struktur sowie die Richtung des Informationsflusses lassen sich verschiedene Netzarten unterscheiden. Im Allgemeinen kann eine Unterteilung von ANN in Netze ohne Rückkopplung und Netze mit Rückkopplung vorgenommen werden [23]. Der Informationsfluss erfolgt bei den sogenannten Feed-Forward-Netzen, den Netzen ohne Rückkopplung, von den Eingangsneuronen zu den Ausgangsneuronen. Bei ANNs mit Rückkopplung, auch als rekurrente Netze bezeichnet, besteht die Möglichkeit, dass die Information auch zu den Neuronen zurückgeführt wird [22, 23]. Des Weiteren existieren kombinierte Netztypen, welche beide Netzformen miteinander vereinen.

In Bezug auf die Feed-Forward-Netze kann eine Unterteilung in einschichtige Netze sowie mehrschichtige Netze vorgenommen werden [23]. Einstufige Netze umfassen lediglich eine Schicht von Neuronen, wie dies beim Perzeptron oder dem Adaline-Modell der Fall ist. Die Schicht besteht entweder aus einem Neuron oder mehreren Neuronen. Beide Modelle unterscheiden sich hauptsächlich im Trainingsalgorithmus. Während beim Perzeptron die Anpassung der Gewichte nach dem Ausgang nach der Aktivierungsfunktion erfolgt, erfolgt die Gewichtsbestimmung bei Adaline vor der Aktivierungsfunktion [23]. Mehrstufige Netze bestehen aus mehreren, hintereinanderliegenden Schichten. Diese lassen sich in eine Eingabeschicht (*engl. Input-Layer*), eine versteckte Schicht (*engl. Hidden-Layer*) und eine Ausgabeschicht (*engl. Output-Layer*) unterteilen. Während die Eingabe- und die Ausgabeschicht Eingabewerte aus der Umgebung erhält bzw. Ausgabewerte an die Umgebung abgibt, ist die versteckte Schicht vor der Umgebung verborgen [22]. Ein Netz, das mehr als zwei Hidden-Layer umfasst, wird als tiefes Netz (*engl. Deep Network*) bezeichnet [17]. In Abbildung 2.4 ist ein Beispiel für ein Feed-Forward-Netz mit einer Eingangsschicht, zwei versteckten Schicht und einer Ausgangsschicht dargestellt. Die verborgene Schicht kann aus einer beliebigen Anzahl an Schichten mit beliebig vielen Neuronen bestehen, was die Repräsentation komplexer Problemstellungen ermöglicht [17, 23]. Mehrstufige Netze wie das Multi-Layer-Perzeptron und das Madaline-Modell bestehen aus mehreren Schichten, wobei die Eingangssignale durch die versteckte Schicht weitergeleitet und verarbeitet wird bis sie die Ausgangsschicht erreichen. Dadurch wird eine Abbildung der Eingangsdaten auf die

Zieldaten erreicht. Das Madaline-Modell besteht aus mehreren Schichten von Adalinen [23].

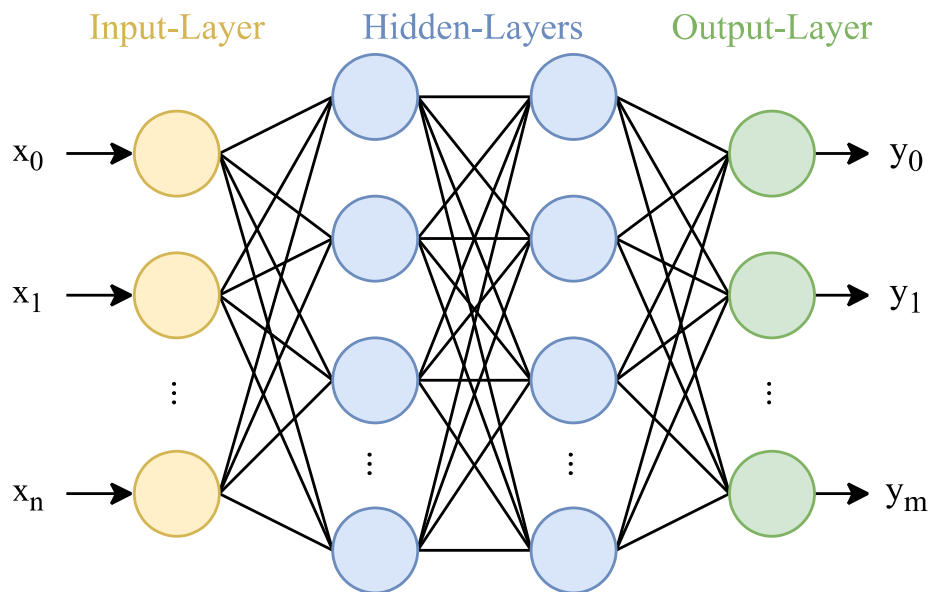


Abbildung 2.4: Aufbau eines Feed-Forward-Netzes mit zwei Hidden-Layers nach [17, 23]

Rekurrente Netze sind Netze, bei denen Neuronen Signale an Neuronen derselben oder einer vorherigen Schicht zurückleiten. Dadurch können zeitlich verschlüsselte Informationen in den Daten erkannt werden [20]. Rekurrente Netze lassen sich in Netze ohne und mit Selbstorganisation unterteilen. Zu den Netzen ohne Selbstorganisation sind die Hopfield-Netze und die Boltzmann-Maschine zu nennen. Kohonen-Netze stellen eine Form selbstorganisierender ANNs dar [22, 23]. Kombinierte Netze, wie beispielsweise das Counterpropagation-ANN und das Neocognitron, vereinen die Eigenschaften von Feed-Forward-Netzen und rekurrenten Netzen. Das Counterpropagation-Netz basiert auf einem Kohonen-Netz sowie einem Feed-Forward-Netz, um eine optimale Einstellung der Gewichte zu erreichen. Das Neocognitron wird primär zur Mustererkennung eingesetzt und besteht aus hierarchisch organisierten Schichten, welche einfache und komplexe Neuronen umfassen [23].

Damit ANNs Probleme lösen können, müssen diese erst mit Beispieldaten trainiert werden. Das Lernen beinhaltet dabei die Anpassung der Verbindungsgewichte und der Schwellenwerte zur Aktivierung der Neuronen [22, 23]. Eine Veränderung der Aktivierungsfunktion während der Trainingsphase ist eher unüblich und zudem nicht einfach umsetzbar [23]. Das Training von ANNs orientiert sich an der Hebb'schen Lernregel. Laut dieser Regel soll die Verknüpfung von zwei Neuronen stärker sein, wenn diese Neuronen gleichzeitig aktiv sind [20, 23]. Perzeptrons können auf Basis dieser Lernregel trainiert werden, indem Verbindungen gestärkt werden, die den Ausgabefehler reduzieren [20]. Dieser Zusammenhang ist in Formel (2.3) dargestellt. Die Veränderung des Gewichtes $\Delta w_{\lambda v}$ zwischen den Neuronen λ und v ergibt sich aus der Lernrate η , dem

Ausgangssignal y_λ des Neurons λ und der Aktivierung a_v des Neurons v , welche gleichzeitig Eingangswert für das Folgeneuron ist [23].

$$\Delta w_{\lambda v} = \eta \cdot y_\lambda \cdot a_v \quad (2.3)$$

Trainingsmethoden auf Basis der Hebb'schen Regel lassen sich in überwachtes und unüberwachtes Lernen unterteilen. Beim überwachten Lernen, das auf dem Fehler zwischen dem Ausgangssignal des ANN und dem erwarteten Ausgangssignal basiert, werden Gewichte und Bias-Werte angepasst, um den Fehler zu minimieren. Damit Eingangsdaten und Zieldaten bekannt sind, wird ein Trainingsdatensatz benötigt. Das Konzept des überwachten Lernens ist in Formel (2.4) dargestellt, wobei $w_{\lambda v}(\tau)$ das Gewicht der Verbindung zwischen dem Neuron λ und dem Ausgabeneuron v zum τ -ten Schritt des Trainingsalgorithmus ist [20, 23]. Die Änderung des Gewichtes $\Delta w_{\lambda v}(\tau)$ ergibt sich dabei gemäß Formel (2.5) entsprechend der Hebb'schen Regel aus der Lernrate η , dem Fehler $e(\tau)$ und dem Eingangssignal $x_\lambda(\tau)$ zum τ -ten Trainingsschritt [20, 23]. Der Fehler ist dabei definiert als die Differenz zwischen dem Ausgangssignal $y_v(\tau)$ des Neurons v und dem erwarteten Ausgangssignal $\hat{y}_v(\tau)$ des Neurons v .

$$w_{\lambda v}(\tau+1) = w_{\lambda v}(\tau) + \Delta w_{\lambda v}(\tau) \quad (2.4)$$

$$\Delta w_{\lambda v}(\tau) = \eta \cdot (y_v(\tau) - \hat{y}_v(\tau)) \cdot x_\lambda(\tau) = \eta \cdot e(\tau) \cdot x_\lambda(\tau) \quad (2.5)$$

In jedem Trainingsschritt wird das Gewicht in ein Neuron um den Wert $\Delta w_{\lambda v}(\tau)$ geändert. Infolgedessen entsteht für jeden weiteren Schritt $\tau+1$ ein neues Gewicht $w_{\lambda v}(\tau+1)$. Die Lernrate gibt dabei an, wie schnell ein Neuron lernt [23]. Eine große Lernrate bedingt eine stärkere Änderung der Gewichte. Dies kann in manchen Fällen dazu führen, dass die optimalen Gewichte nicht gefunden werden, da diese übersprungen werden. Eine kleine Lernrate bewirkt dagegen ein sehr langsames Training [22, 23]. Daher ist die Auswahl einer optimalen Lernrate erforderlich.

Beim Training ist die Qualität und die Menge der Trainingsdatensätze von großer Bedeutung [17, 23]. Wenn der Trainingsdatensatz einen Sonderfall darstellt oder zu komplex ist, kann ein ANN diesen auswendig lernen und liefert bei leicht abgewandelten Eingangsdaten fehlerhafte Werte. Dieses Phänomen wird als Überanpassung (*engl. Overfitting*) bezeichnet [17, 22, 23]. Unteranpassung (*engl. Underfitting*) tritt dagegen auf, wenn ein ANN ein Problem aufgrund von zum Beispiel einer zu geringen Anzahl versteckter Neuronen nicht optimal darstellen kann. Das Modell ist entsprechend des Problems nicht komplex genug [17, 22]. Eine Aufteilung des Datensatzes in mehrere Teile zum Trainieren und Validieren, gewährleistet, dass das ANN auch auf unbekannte Daten verlässlich reagiert [23].

Backpropagation stellt aktuell die am häufigsten eingesetzte Methode zum Training von ANNs dar [20]. Backpropagation stellt eine Verallgemeinerung der Delta-Regel dar. Im Folgenden wird diese Methode näher erläutert. Die Delta-Regel wird allerdings nur bei einschichtigen Perzeptrons eingesetzt [23]. Für das Training eines mehrschichtigen ANNs ist das Backpropagation-Verfahren erforderlich, da eine Anpassung der Gewichte der versteckten Neuronen mit der Delta-Regel nicht möglich ist. Im Rahmen der Delta-Regel erfolgt lediglich eine Betrachtung der Ausgangsneuronen, um die Veränderung der Gewichte zu berechnen. In einem mehrschichtigen ANN müssen jedoch auch die versteckten Neuronen in die Betrachtung mit einbezogen werden [20, 23]. Das Ziel des Backpropagation-Algorithmus besteht in der Minimierung des Fehlers zwischen der Ausgabe des ANN und der erwarteten Ausgabe [23]. Dies erfolgt mittels des Verfahrens des Gradientenabstiegs [22]. Zu diesem Zweck wird der Gradient einer definierten Fehlerfunktion bestimmt. Der Gradient stellt das Steigungsverhalten einer Funktion dar und zeigt dementsprechend in die Richtung des steilsten Anstiegs. Da jedoch das Minimum der Fehlerfunktion gesucht ist, muss eine Bewegung in die Gegenrichtung erfolgen. Daraus lässt sich ableiten, in welche Richtung eine Veränderung der Gewichte und Bias-Werte zu erfolgen hat, um den Fehler zu minimieren [17, 22]. Nach Anpassung der Gewichte und Bias-Werte wird der Gradient erneut berechnet und die Netzparameter dementsprechend angepasst. Dieser Ablauf wird wiederholt, bis das Minimum der Fehlerfunktion erreicht ist.

Zu Beginn des Trainings erhalten die Gewichte aller Neuronen zufällige Werte, sodass die Ausgabe des ANN mit den zufälligen Gewichten mit der erwarteten Ausgabe verglichen werden kann [20, 23]. Das Ziel besteht in der Minimierung des Fehlervektors E als Differenz zwischen dem Ausgabevektor y des ANN und dem erwarteten Ausgangsvektor $\hat{y}$ gemäß Formel (2.6) [23]. Die einzelnen Komponenten werden durch Vektoren dargestellt, da ein gesamtes Netz mit mehr als einem Neuron betrachtet wird.

$$\min E = \min (y - \hat{y}) \rightarrow 0 \quad (2.6)$$

Das Grundkonzept der Delta-Regel ist die Minimierung des Gesamtfehlers. Dieser besteht aus den einzelnen Ausgangsfehlern. Der Gesamtfehler wird in Abhängigkeit der vorliegenden Gewichte betrachtet. Die Veränderung der gesamten Gewichte ΔW in einem ANN hängt vom Gesamtfehler $E(W)$ ab. Der Einfluss der einzelnen Gewichte auf die Minimierung des Gesamtfehlers ist durch die partielle Ableitung des Gesamtfehlers gegeben. Die Veränderung der Gewichte ΔW ergibt sich dementsprechend aus dem Gradienten des Gesamtfehlers E entlang des Gewichtsvektors W und der Lernrate gemäß Formel (2.7) [23]. Dabei beschreibt $w_{\lambda v}$ ein Gewicht zwischen einem Neuron λ und einem Neuron v .

$$\Delta W = \eta \cdot \nabla E(W) = \eta \cdot \frac{\partial E(W)}{\partial w_{\lambda v}} \quad (2.7)$$

Um die Veränderung der Gewichte bestimmen zu können, muss zunächst die Fehlerfunktion definiert werden. In der Regel wird dafür die quadratische Abweichung zwischen dem tatsächlichen Ausgangssignal des ANN und dem erwarteten Ausgangssignal gebildet [22, 23]. Der Vorteil dieser Fehlerfunktion begründet sich zum einen darin, dass sich negative und positive Abweichungen nicht gegenseitig aufheben und zum anderen große Abweichungen durch das Quadrieren stärker ins Gewicht fallen [22]. Die Fehlerfunktion als Summe der quadratischen Abweichungen ist in Formel (2.8) definiert, wobei N die Anzahl der Ausgangsneuronen bezeichnet [23].

$$E(W) = \frac{1}{2} \sum_{m=1}^N (y_m - \hat{y}_m)^2 \quad (2.8)$$

In Bezug auf die Lösung der partiellen Ableitung in Formel (2.7) ist eine Unterscheidung zwischen zwei Fällen erforderlich. Hierbei geht es darum, ob v die Ausgangsneuronen darstellt oder die Neuronen in den versteckten Schichten. Nach einigen Umformungen und Vereinfachungen lässt sich die Änderung der einzelnen Gewichte gemäß der Formel (2.9) darstellen, wobei η die Lernrate und y_λ die Ausgabe aus dem Vorgängerneuron λ ist [23]. Diese Definition zeigt, wie die Gewichte der Verbindungen zwischen Neuronen angepasst werden müssen.

$$\Delta w_{\lambda v} = \eta \cdot \delta_v \cdot y_\lambda \quad (2.9)$$

Der Wert für δ_v wird nach der Formel (2.10) berechnet [23]. In diesem Zusammenhang ist eine Differenzierung erforderlich, ob der δ_v -Wert für die Neuronen der Ausgangsschicht oder für die Neuronen in den versteckten Schichten benötigt wird. Dabei entspricht F der Aktivierungsfunktion. Die Differenzierbarkeit der Aktivierungsfunktion ist eine notwendige Voraussetzung für die Berechnung ihres Gradienten. Das ist nur der Fall, wenn die Aktivierungsfunktion ableitbar ist. Als Aktivierungsfunktion kann beispielsweise die Sigmoid-Funktion verwendet werden, die einen kontinuierlichen Verlauf aufweist [23]. Folglich kann dieser Trainingsalgorithmus lediglich für differenzierbare Aktivierungsfunktionen verwendet werden.

$$\delta_v = \begin{cases} \frac{\partial F(u_v)}{\partial u_v} \cdot (y_v - \hat{y}_v) & \text{für Ausgabeneuronen} \\ \frac{\partial F(u_v)}{\partial u_v} \cdot \sum_{\lambda=1}^M (\delta_\lambda \cdot w_{\lambda v}) & \text{für versteckte Neuronen} \end{cases} \quad (2.10)$$

Im Folgenden wird der Ausdruck für die gewichtete Summe u_v nach Formel (2.11) definiert. Diese führt von den Neuronen der Vorgängerschicht mit der Neuronenanzahl M in ein nachfolgendes Neuron v [23]. Dabei bezeichnet y_λ die Ausgabe aus dem Neuron λ und $w_{\lambda v}$ das Gewicht zwischen den Neuronen λ und v .

$$u_v = \sum_{\lambda=1}^M y_\lambda \cdot w_{\lambda v} \quad (2.11)$$

Die Formeln (2.9) und (2.10) machen deutlich, dass zur Berechnung der Gewichtsänderung bei der Betrachtung der Ausgabeneuronen die Ausgabeneuronen selbst und die erwartete Ausgabe, die sich aus dem Trainingsdatensatz ergibt, benötigt werden. Damit die Gewichte auch in den versteckten Schichten angepasst werden können, wird die Gewichtsänderung durch den δ -Wert und die Gewichte zu den nachfolgenden Neuronen bestimmt. Der δ -Wert wird durch die vorherige Berechnung für die nachfolgende Schicht ermittelt. In einer anschaulichen Betrachtung wird der Fehler von der Ausgabenschicht über alle versteckten Schichten in Richtung Eingangsschicht weitergereicht [22, 23]. Durch die Bestimmung der δ -Werte in jeder Schicht werden alle Fehler zwischen der Ausgabenschicht und den versteckten Schichten nacheinander durch Veränderung der Gewichte minimiert [20]. Von dieser Fehlerrückführung stammt auch der Name Backpropagation [23]. Das Backpropagation-Verfahren findet Anwendung für mehrschichtige Perzeptrone (*engl. Multi-Layer Perceptron*) als auch für einschichtige Perzeptrone (*engl. Single-Layer Perceptron*), wobei die allgemeine Form wieder zur Delta-Regel zurückgeführt wird.

2.1.2 Lineare Regression

Die Lineare Regression (LR) ist ein parametrisches Lernverfahren, bei dem die Anzahl der Parameter sowie die grundsätzliche Struktur vor dem Training festgelegt wird [17]. Die LR-Methode stellt eine einfache ML-Methode dar, da eine Interpretierbarkeit gegeben ist [24]. Das Ziel besteht in der Identifikation eines linearen Zusammenhangs zwischen einer abhängigen Variable y und einer oder mehreren unabhängigen Variablen x , um den Wert der abhängigen Variable auf Basis der unabhängigen Variable vorherzusagen zu können. Das Ziel ist es, eine Ausgleichsgerade durch die Punktwolke der Datenpunkte wie in Abbildung 2.5 zu legen, um den Zusammenhang zwischen x und y möglichst genau abzubilden [17, 18, 25]. Bei mehreren unabhängigen Variablen, auch als Merkmale bezeichnet, tritt an die Stelle der Geraden eine Hyperebene, die den Zusammenhang darstellt.

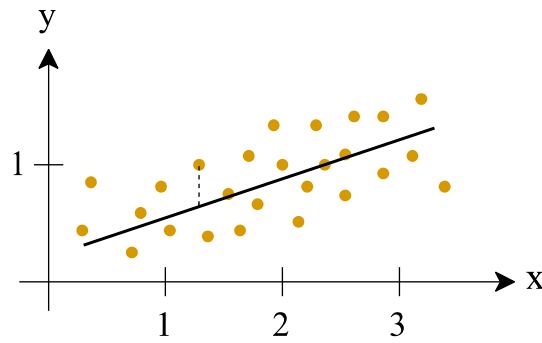


Abbildung 2.5: Beispiel für die LR-Methode nach [18, 26]

Die LR-Methode ist nach Formel (2.12) definiert [25]. Dabei sind die Datenpaare (x_i, y_i) die unabhängigen Variablen sowie die abhängigen Variablen bzw. Zielvariablen. Die Koeffizienten α und β stellen den zu bestimmenden Achsenabschnitt sowie die Steigung der Geraden dar. Der zufällige Fehlerterm ε_i wird eingeführt, da die Beziehung zwischen den unabhängigen Variablen und der Zielvariablen nicht exakt ist. Dieser Fehler soll möglichst klein gehalten werden, um den Großteil der Variabilität der Daten durch die Gerade abbilden zu können. Der Fehlerterm selbst kann nicht direkt beeinflusst oder minimiert werden und ist eine zufällige Größe. Bei mehreren unabhängigen Variablen handelt es sich um eine multiple Regression und ist nach Formel (2.13) definiert [17, 20]. Dabei stellt m die Anzahl der unabhängigen Variablen dar.

$$y_i = \alpha + x_i \cdot \beta + \varepsilon_i \quad \forall i \in [1, n] \quad (2.12)$$

$$y_i = \alpha + x_{i,1} \cdot \beta_1 + x_{i,2} \cdot \beta_2 + \dots + x_{i,m} \cdot \beta_m + \varepsilon_i \quad \forall i \in [1, n] \quad (2.13)$$

Bei der LR-Methode erfolgt das Training des Modells anhand der Methode der kleinsten Quadrate durch die Minimierung der Summe der quadratischen Fehler zwischen den tatsächlichen Werten y_i und den vorhergesagten Zielwerten $\hat{y}_i$. Diese wird auch als Residuenquadratsumme bezeichnet. Das Quadrieren der Fehler verhindert, dass sich positive und negative Abweichungen ausgleichen, und führt dazu, dass größere Fehler stärker ins Gewicht fallen [18, 26]. In der Abbildung 2.5 zeigt der gestrichelte Abstand zwischen einem Punkt und der Geraden den Fehler, der reduziert werden soll. Die Bestimmung der Parameter α und β erfolgt mit dem Ziel, dass die Datenpunkte möglichst nahe an der Geraden liegen. Die Parameter α und β müssen folglich so gewählt werden, dass die Kostenfunktion $J(\alpha, \beta)$ nach Formel (2.14) minimiert wird [18, 24, 25].

$$J(\alpha, \beta) = \frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2 = \frac{1}{n} \sum_{i=1}^n (y_i - (\alpha + x_i \cdot \beta))^2 \quad (2.14)$$

Um eine Überanpassung bei der LR-Methode zu vermeiden, wird häufig ein Regularisierungsterm zur Kostenfunktion hinzugefügt. Die sogenannte Ridge-Regression erreicht dies durch einen Strafterm, der die Größe der Regressionskoeffizienten β_j einschränkt und somit eine zu starke Anpassung des Modells an spezifische Merkmale verhindert [17, 20, 24]. Die Ridge-Regression ist entsprechend Formel (2.15) definiert [20]. Die Stärke der Regularisierung wird durch den Hyperparameter λ gesteuert. Bei einem Wert von $\lambda = 0$ entspricht die Ridge-Regression der klassischen LR-Methode ohne Regularisierung. Ein sehr hoher λ -Wert führt dazu, dass die Regressionskoeffizienten gegen null tendieren, wodurch das Modell stark vereinfacht wird. Der Achsenabschnitt bzw. Bias-Term α wird nicht regularisiert. Vor dem Training ist eine entsprechende Anpassung des Regularisierungsparameters λ erforderlich. Eine weitere gängige Methode zur Regularisierung ist die Lasso-Regression. Der wesentliche Unterschied der beiden Regularisierungsmethoden ist der Strafterm. In [24] wird ausgeführt, dass mit der Ridge-Regression in der Regel bessere Vorhersagen erzielt werden.

$$J_{\text{Ridge}}(\alpha, \beta) = J(\alpha, \beta) + \frac{\lambda}{n} \sum_{j=1}^m \beta_j^2 \quad (2.15)$$

Im Rahmen der LR-Methode existieren verschiedene Lernverfahren, welche die Regressionskoeffizienten auf unterschiedliche Art und Weise bestimmen. Dies erfolgt entweder auf Basis der Normalengleichung oder des Gradientenabstiegsverfahrens [17, 20]. Zum einen kann die Ermittlung einer analytischen Lösung unter Zuhilfenahme einer mathematischen Gleichung, welche die Normalengleichung darstellt, erfolgen. Es handelt sich hierbei im Gegensatz zum Gradientenabstieg nicht um ein iteratives Optimierungsverfahren. Allerdings erweist sich dieses Verfahren als sehr aufwendig bei der Anwendung auf umfangreiche Datensätze mit einer Vielzahl an Merkmalen. Demgegenüber stellt das Gradientenabstiegsverfahren eine iterative Methode dar, bei der die Regressionskoeffizienten schrittweise angepasst werden, um eine Minimierung der Kostenfunktion zu erreichen. Die Auswahl des Lernverfahrens ist beispielsweise von der Struktur und Größe des Datensatzes abhängig. Dabei spielen Faktoren wie die vorhandene Menge an Trainingsdaten oder die Anzahl an Merkmalen eine entscheidende Rolle. Die Normalengleichung ermöglicht beispielsweise die Berechnung schneller Lösungen für kleinere Probleme, während das Gradientenabstiegsverfahren eine effektivere Methode zur Handhabung großer oder komplexer Datensätze darstellt. In Bezug auf die Vorhersage zeigen die verschiedenen Algorithmen keine großen Unterschiede,

da alle Algorithmen zu sehr ähnlichen Modellen führen und ähnliche Vorhersagen aufweisen [20].

2.1.3 K-Nächste-Nachbarn

Die Methode der K-Nächsten-Nachbarn (*engl. K-Nearest Neighbors, KNN*) stellt eine nichtparametrische Lernmethode dar. Diese Methode bestimmt die Anzahl der Parameter erst während des Trainings und macht keine Annahmen über die Verteilung der Daten. Daher sind nichtparametrische Lernverfahren flexibler und eignen sich gut, wenn die Datenstruktur unbekannt oder komplex ist [17]. Die KNN-Methode ist ein instanzbasiertes Verfahren, das davon ausgeht, dass ähnliche Datenpunkte nahe beieinanderliegen. Die Vorhersage für eine Eingabe erfolgt basierend auf den k nächstgelegenen Punkten, auch als Nachbarn bezeichnet. Die KNN-Methode zeichnet sich durch ihre Einfachheit aus und findet Anwendung in der Klassifikation sowie der Regression. Bei der KNN-Regression wird der Wert eines Eingabepunkts durch den Durchschnitt der Werte seiner k nächsten Nachbarn oder der gewichteten Summe der Werte bestimmt [17, 18, 26]. Der KNN-Algorithmus basiert auf dem Konzept des trägen Lernens (*engl. Lazy Learning*). Beim Lazy Learning wird kein Modell vorab erstellt, sondern alle Trainingsdaten werden gespeichert. Die Verarbeitung erfolgt erst bei einer Anfrage. Im Gegensatz dazu erzeugt das eifrige Lernen (*engl. Eager Learning*) bereits während des Trainings ein Modell, das anschließend zur Vorhersage neuer Daten genutzt wird, wie es beispielsweise bei ANNs der Fall ist, ohne erneut auf die Trainingsdaten zugreifen zu müssen [18, 26].

Für die Vorhersage wird zunächst die Distanz zwischen den Trainingsdaten und einem neuen Datenpunkt berechnet. Auf Basis von diesem Abstand werden anschließend die k nächsten Nachbarn bestimmt. In Abbildung 2.6 ist ein Beispiel für einen zweidimensionalen Raum mit Datenpunkten dargestellt. Für einen neuen Datenpunkt, in der Abbildung als blauer Datenpunkt dargestellt, werden die $k = 3$ nächsten Nachbarn über den Abstand bestimmt.

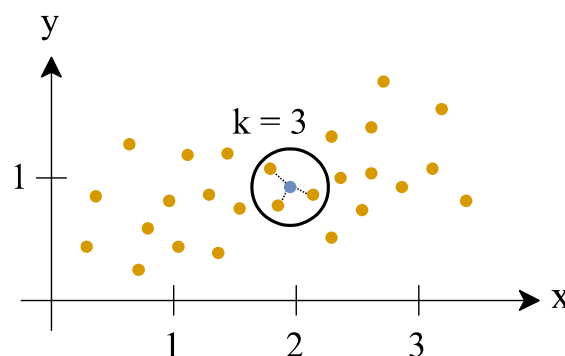


Abbildung 2.6: Beispiel der KNN-Methode für $k = 3$ Nachbarn nach [26, 27]

Bei der Distanzberechnung werden lediglich die x -Werte bzw. die Merkmale berücksichtigt [27]. Aus den der k Nachbarn zugehörigen y -Werten bzw. Zielgrößen, wird ein (gewichteter) Mittelwert bestimmt, welcher als Prognose für den neuen blauen Datenpunkt dient. Für die Berechnung der Distanzen zwischen den Trainingsdaten und den neuen Datenpunkten stehen verschiedene Abstandsmetriken zur Verfügung. Eine häufig verwendete Metrik ist der euklidische Abstand, dessen Definition in Formel (2.16) dargestellt ist [17, 28]. Eine weitere Metrik ist der in Formel (2.17) dargestellte Manhattan-Abstand [24]. Der Abstand d zwischen den Punkten x und y wird in einem n -dimensionalen Raum berechnet. Dabei stellt x den Eingabevektor und y einen Trainingsvektor aus dem Trainingsdatensatz dar. Die Dimensionen bzw. Merkmale werden durch n beschrieben.

$$d_{\text{euklidisch}}(x,y) = \sqrt{\sum_{i=1}^n (x_i - y_i)^2} \quad (2.16)$$

$$d_{\text{Manhattan}}(x,y) = \sum_{i=1}^n |x_i - y_i| \quad (2.17)$$

Wichtig ist, dass die Daten vorab skaliert werden. Eine Skalierung ist erforderlich, um eine Verzerrung der Abstandsberechnung durch Merkmale mit größeren Werten zu vermeiden. Die Skalierung gewährleistet, dass alle Merkmale gleich gewichtet werden [24, 28]. Die KNN-Methode gilt als einfach, da nur die zwei Parameter Anzahl der Nachbarn k und die Abstandsmetrik erforderlich sind. Zudem ist kein Training notwendig. Allerdings ist zu berücksichtigen, dass dies zu einem hohen Speicherbedarf führt, da alle Trainingsdaten zur Vorhersage gespeichert werden müssen.

2.1.4 Entscheidungsbaum

Entscheidungsbäume (*engl. Decision Tree, DT*) zählen zu den Methoden des Eager Learning und sind nichtparametrische Verfahren. Sie finden sowohl in der Klassifikation als auch in der Regression Anwendung [17, 20, 28]. Ein wesentlicher Vorteil von DTs besteht in der Möglichkeit, den zugrundeliegenden Entscheidungsprozess in einer grafischen Darstellung zu visualisieren. Die grafische Darstellung ermöglicht eine leichte Interpretation und Nachvollziehbarkeit der Vorhersagen. Aufgrund der transparenten Darstellung der Entscheidungen werden DT oft als White-Box-Modelle bezeichnet, im Gegensatz zu den Black-Box-Modellen wie ANNs [20, 26, 28].

Ein DT besteht wie in Abbildung 2.7 dargestellt aus einer Wurzel, inneren Knoten und Blättern. Die Wurzel bildet den Startpunkt, an dem die Verarbeitung beginnt. Sowohl

die Wurzel als auch die inneren Knoten enthalten Entscheidungsregeln, die sich auf die Merkmale eines Problems beziehen. Die Verbindungen zwischen den Knoten stellen die Überprüfung von Bedingungen dar. Die Blätter fungieren als Endpunkte des Baumes, an denen entweder eine Klassifikation durchgeführt oder bei der Regression ein Wert ausgegeben wird. Nach dem Training wird ein DT verwendet, indem neue Daten von der Wurzel durch die Knoten bis zu einem Blatt geleitet werden. An jedem Knoten wird eine Entscheidung getroffen, bis die endgültige Vorhersage eines Zielwertes am Blatt erreicht ist [17, 18].

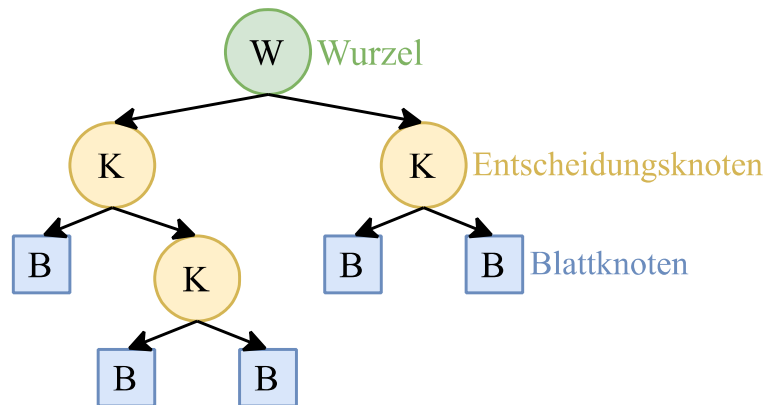


Abbildung 2.7: Beispiel für die Struktur eines DTs nach [18]

Der Lernprozess eines DT ist rekursiv. Für jeden Knoten wird aus allen verfügbaren Merkmalen jenes ausgewählt, das den höchsten Informationsgewinn liefert. Das Ziel besteht in der Aufteilung der Daten in Teilmengen mit dem Zweck der möglichst genauen Vorhersage der Zielwerte. Der Prozess wird für jeden Knoten fortgesetzt, bis entweder eine sinnvolle Klassifikation erreicht oder ein Abbruchkriterium erfüllt ist [28]. Der Mittelwert der Zielwerte in einem Blatt stellt die Vorhersage dar. Bei Regressionsaufgaben erfolgt eine Teilung des Trainingsdatensatz mit dem Ziel, den mittleren quadratischen Fehler (*engl. Mean Squared Error, MSE*) gemäß Formel (2.18) zu minimieren [20, 24]. Dabei ist y_i der wahre Zielwert und $\hat{y}_i$ der Vorhersagewert. Der MSE findet beim Training von DTs, wie bei zahlreichen weiteren ML-Methoden, häufig als Fehlermetrik Anwendung. Eine Normalisierung der Daten ist beim Training von DTs in der Regel nicht erforderlich, da sie keinen Vorteil für den Lernprozess bietet. Dadurch bleibt die Interpretierbarkeit des DT erhalten [28].

$$\text{MSE} = \frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2 \quad (2.18)$$

DTs neigen zu Overfitting, da sie dazu tendieren, die Details der Trainingsdaten zu stark zu lernen. Um dies zu vermeiden, ist eine Regularisierung notwendig, bei der die Freiheitsgrade des DTs während des Trainings begrenzt werden. Dies erfolgt durch die Anpassung bestimmter Hyperparameter [17, 20]. Die spezifischen Hyperparameter, die zur

Regularisierung eines DTs zum Einsatz kommen, sind abhängig vom jeweils gewählten Algorithmus. In jedem Fall besteht die Möglichkeit, die maximale Tiefe des Baums anzupassen, um auf diese Weise dessen Komplexität zu begrenzen. Weitere Hyperparameter, die zur Regularisierung eingesetzt werden können, umfassen:

- die maximale Anzahl an Blättern,
- die minimale Anzahl an Datenpunkten pro Blatt
- die minimale Anzahl an Datenpunkten pro Knoten, damit eine Aufteilung erfolgen kann sowie
- die maximale Anzahl an Merkmalen, die beim Aufteilen eines Knoten berücksichtigt werden.

Die Varianz von DTs ist hoch, da bei geringen Veränderungen der Hyperparameter oder der Daten unterschiedliche Modelle entstehen. Die Kombination der Vorhersagen mehrerer DTs ermöglicht die Lösung des Problems der Varianz und verhindert Overfitting [20, 26]. Ein solches Ensemble wird als Zufallswald (*engl. Random Forest, RF*) bezeichnet. Die Methode des Ensemble-Lernens wird in Kapitel 2.1.6 näher erläutert.

2.1.5 Clusteranalyse

Die Clusteranalyse (*engl. Clustering*) stellt ein Verfahren des unüberwachten Lernens dar, welches ähnliche Datenpunkte anhand ihrer Merkmale in Gruppen bzw. Cluster einteilt, ohne dass die Daten mit Zieldaten versehen sind [17, 18]. Ein häufig verwendetes Verfahren des Clustering ist das K-Means-Verfahren, welches Datenpunkte in K Cluster einteilt. Das Ziel besteht in der Maximierung der Gemeinsamkeiten innerhalb eines Clusters sowie der Unterschiede zwischen den Clustern. Innerhalb von jedem Cluster wird ein Mittelpunkt der Datenpunkte definiert, der als Centroid bezeichnet wird. Die Gruppierung erfolgt beispielsweise durch die Minimierung der Summe der quadratischen Abstände zwischen den Datenpunkten und dem Centroid eines Clusters gemäß Formel (2.19) [17]. Dabei sind x_i und x_j zwei Punkte, zwischen denen der Abstand d berechnet wird und Ω ist die Anzahl der Merkmale der Datenpunkte. Es existieren aber auch weitere Abstandsmaße wie der Euklidische Abstand oder der Manhattan-Abstand [28].

$$d(x_i, x_j)^2 = \sum_{n=1}^{\Omega} (x_{i,n} - x_{j,n})^2 \quad (2.19)$$

Der K-Means-Algorithmus basiert auf der Wahl der Centroide der Cluster mit dem Ziel, die Kostenfunktion J zu minimieren. Die Kostenfunktion J ist gemäß Formel (2.20) definiert [17, 18]. Dabei ist K die Anzahl der Cluster, C_k ein Cluster mit der Nummer k und μ_k der Centroid des Clusters C_k .

$$J = \sum_{k=1}^K \sum_{x_i \in C_k} d(x_i, \mu_k)^2 \quad (2.20)$$

Im Rahmen des K-Means-Clusterings ist es erforderlich, den Wert von K vorab zu definieren. Das Vorgehen des K-Means-Algorithmus umfasst zu Beginn die Auswahl von K Centroids zufällig aus dem Datensatz wie in Abbildung 2.8 dargestellt. Im Anschluss erfolgt die Berechnung der Distanz zwischen jedem Datenpunkt und den Centroiden. Anschließend erfolgt die Zuordnung jedes Datenpunkts zu dem Cluster, dessen Centroid die geringste Distanz aufweist. In der Folge wird für jedes Cluster ein neuer Centroid berechnet, wobei der Mittelwert aller zugehörigen Datenpunkte ermittelt wird [17]. Der zuvor beschriebene Prozess wird wiederholt, wobei eine vorgegebene Anzahl an Iterationen durchlaufen wird oder bis sich die Cluster-Zuweisungen der Datenpunkte nicht mehr ändern und wie in Abbildung 2.8 Cluster entstanden sind. Sobald dies der Fall ist, gilt der Algorithmus als konvergiert und wird beendet [24]. Ein Nachteil des Algorithmus besteht darin, dass eine zufällige Auswahl der Centroiden dazu führen kann, dass bei jeder erneuten Ausführung unterschiedliche Cluster identifiziert werden. Es besteht die Möglichkeit, dass lediglich ein lokales Minimum ermittelt wird [20]. Diesem Nachteil kann durch eine mehrmalige Ausführung des Algorithmus begegnet werden [17].

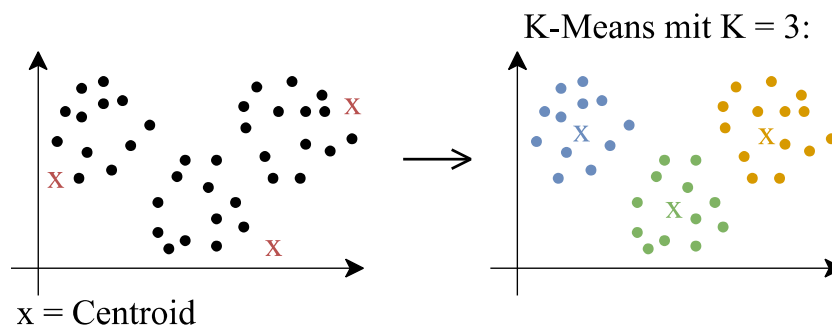


Abbildung 2.8: Beispiel für den K-Means-Algorithmus nach [20]

2.1.6 Ensemble-Lernen

Ensemble-Lernen (*engl. Ensemble Learning*) bezeichnet ein Verfahren im Bereich des ML, bei dem mehrere ML-Modelle kombiniert werden, um die Genauigkeit zu erhöhen. Der Grundgedanke des Ensemble Learning ist, dass die Kombination der Vorhersagen einer Gruppe von ML-Modellen oft bessere Ergebnisse liefert als das beste Einzelmodell [20]. Die Methode basiert auf der Verknüpfung der Ergebnisse der einzelnen Modelle. Dies erfolgt entweder durch einen Mehrheitsentscheid bei der Klassifikation oder durch beispielsweise eine Mittelwertbildung bei Regressionsaufgaben. Ensemble-Methoden nutzen eine Vielzahl von Modellen, die sich entweder anhand ihres Verfahrens oder der Trainingsdaten unterscheiden [18, 20, 26]. Diese Vielfalt trägt dazu bei, dass

die Prognosevarianz des Ensembles im Vergleich zu einem einzelnen Modell deutlich reduziert wird [19]. Das Ensemble Learning kann dabei für eine Vielzahl unterschiedlicher ML-Methoden angewendet werden. Die besten Ergebnisse werden mit Ensemble-Methoden erzielt, wenn die im Ensemble verwendeten ML-Modelle möglichst unterschiedlich sind. Der Einsatz verschiedener ML-Algorithmen, die unterschiedliche Fehler machen, führt zu einer Verbesserung der Genauigkeit des Ensembles [20]. Beliebte Methoden des Ensemble Learning sind das Bagging, Boosting und Stacking, die im weiteren Verlauf dieses Kapitels näher erläutert werden.

Beim Bagging, kurz für Bootstrap Aggregating, werden mehrere Modelle desselben Lernalgorithmus mit unterschiedlichen, zufällig ausgewählten Teilmengen der Trainingsdaten trainiert, um auf diese Weise unterschiedliche Modelle zu generieren [19, 20]. Die Teilmengen, auch als Bootstrap-Samples bezeichnet, werden nach dem Zufallsprinzip entsprechend des Bootstrapping-Verfahrens mit der Methode des Ziehens mit Zurücklegen erstellt. Dies impliziert, dass einige Datenpunkte mehrfach in einer Teilmenge auftreten, während andere fehlen [18, 26]. Die Auswahl erfolgt mit gleicher Wahrscheinlichkeit für jede Teilmenge. Erfolgt die Erstellung der Teilmengen ohne Zurücklegen, wird dies als Pasting bezeichnet [20]. Nach dem Training der einzelnen Modelle werden die Vorhersagen der Modelle bei neuen Daten kombiniert, um eine finale Prognose zu erzeugen. Bei Klassifikationsaufgaben erfolgt dies in der Regel durch einen Mehrheitsentscheid, bei Regressionsaufgaben durch eine Mittelwertbildung. In Abbildung 2.9 ist die Methode des Bagging für mehrere DTs in einem Ensemble, die mit unterschiedlichen Teilmengen aus dem Trainingsdatensatz trainiert werden, dargestellt.

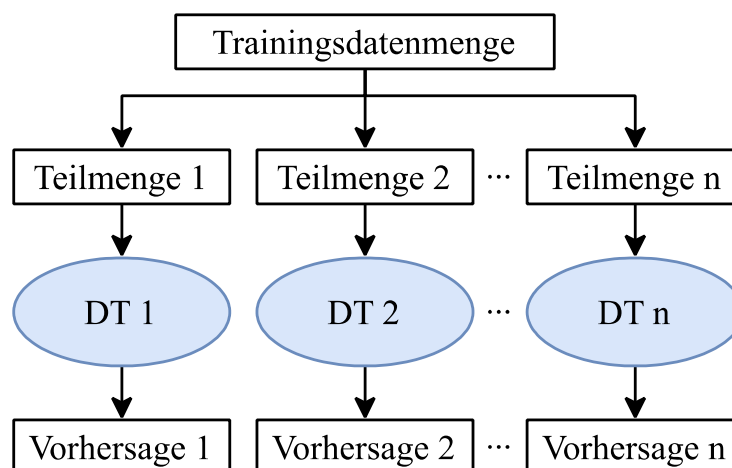


Abbildung 2.9: Ablauf des Trainings beim Bagging nach [20, 26]

Ein Ensemble von DTs, das mithilfe des Bagging-Verfahrens trainiert wird, ist ein RF. Die Pasting-Methode findet beim Training eines RF nur selten Anwendung [20]. Beim RF erfolgt nicht nur eine zufällige Auswahl der Trainingsdaten für jeden DT, sondern ebenfalls eine zufällige Auswahl der Teilmenge der Merkmale, welche der Bildung der Äste an einem Knoten eines DT dient. Folglich stehen nicht sämtliche Merkmale für

eine Aufteilung kontinuierlich zur Verfügung [18, 24]. Die Vorteile eines RF bestehen in der Parallelisierbarkeit während des Trainings mehrerer DTs, seiner Einfachheit und der Vermeidung von Overfitting durch die Kombination mehrerer DTs. Obgleich die Methodik eine hohe Einfachheit aufweist, zählt RF zu den weit verbreiteten und leistungsstarken Verfahren des ML [18, 20].

Die Erstellung von Teilmengen aus einem Trainingsdatensatz birgt das Risiko, dass relevante Datenpunkte unberücksichtigt bleiben. Dies kann zu einer Verzerrung beim Bagging führen, auch wenn eine Verringerung der Varianz erzielt wird [17]. Dies stellt einen Nachteil des Bagging dar. Um diesem entgegenzuwirken, kann das Boosting-Verfahren eingesetzt werden.

Die Boosting-Methode basiert auf der Kombination mehrerer schwacher Lerner. Jedes Modell wird nacheinander trainiert, sodass jedes Modell die Fehler des vorherigen Modells korrigiert [18, 20]. Im Gegensatz zum parallel arbeitenden Bagging erfolgt die Verarbeitung beim Boosting sequenziell. Zu den populärsten Verfahren zählen AdaBoost (Adaptive Boosting) und Gradient Boosting. Die Verfahren unterscheiden sich in Bezug auf die Durchführung der Korrektur. In den meisten Fällen findet beim Boosting der Einsatz von DTs statt, jedoch ist auch die Verwendung alternativer Modelle möglich [18]. Ein Nachteil der sequenziellen Lernmethode des Boosting besteht darin, dass das Training nicht parallelisiert werden kann, wie dies beim Bagging möglich ist und daher längere Trainingszeiten erfordert.

Das AdaBoost-Verfahren basiert auf einer Korrektur des vorherigen Modells, indem die Gewichte der falsch vorhergesagten Datenpunkte erhöht werden. Im Falle korrekter Vorhersagen erfolgt eine Reduktion des Gewichts der jeweiligen Datenpunkte. Bei jedem neuen Modell erfolgt das Training mit den entsprechend gewichteten Trainingsdaten, wodurch sich das Modell zunehmend auf die Daten fokussiert, die falsch vorhergesagt wurden [20, 26]. Der beschriebene Ablauf des Trainings bei AdaBoost ist in Abbildung 2.10 dargestellt. Dieses Verfahren weist eine Ähnlichkeit zum Gradientenverfahren auf, da ebenfalls iterativ Optimierungen erfolgen, jedoch durch das Hinzufügen neuer Modelle anstelle der Anpassung einzelner Parameter mit dem Ziel der Minimierung der Kostenfunktion [20]. Wenn das Training abgeschlossen ist, können die Vorhersagen für neue Daten für eine finale Vorhersage kombiniert werden. Es stehen verschiedene Aggregationsmethoden zur Verfügung. Beim Weighted Voting werden die Vorhersagen der Modelle gewichtet zusammengeführt [17]. Modelle mit höherer Genauigkeit werden dabei mit einer höheren Gewichtung versehen, sodass sie in stärkerem Maße zur finalen Vorhersage beitragen. Beim Stacking erfolgt die finale Vorhersage über einen weiteren DT, der die einzelnen Vorhersagen der DTs als Eingabe erhält. Die Methode des Stacking wird im weiteren Verlauf dieser Arbeit näher erläutert.

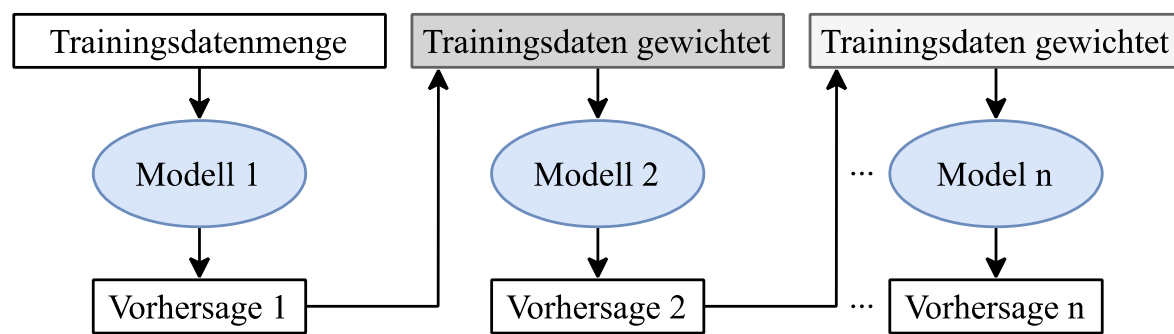


Abbildung 2.10: Ablauf des Trainings bei AdaBoost nach [20]

Gradient Boosting ist ein weit verbreiteter Algorithmus beim Boosting. Analog zu AdaBoost erfolgt das Training der Modelle nacheinander. Im Gegensatz zu AdaBoost minimiert der Algorithmus eine Fehlerfunktion unter Verwendung des Gradientenabstiegs, um die Restfehler, welche die Differenz zwischen den tatsächlichen Werten und den Vorhersagen darstellen, schrittweise zu reduzieren [18]. In der Regel wird bei Regressionsaufgaben der MSE verwendet. Das Training beginnt mit einer einfachen Initialisierung, bei der die anfängliche Vorhersage als Durchschnitt der Zielvariablen definiert wird. Anschließend werden die Fehler, auch als Residuals bezeichnet, berechnet, welche die Fehler der aktuellen Vorhersage darstellen. Die Residuals dienen als Grundlage für das Training des nächsten Modells, dessen Ziel es ist, die genannten Fehler vorherzusagen. Im Rahmen jeder Iteration erfolgt eine Anpassung der Residuals unter Verwendung eines neuen Modells. Die finale Vorhersage erfolgt durch die Kombination der Vorhersagen der einzelnen Modelle. Dazu werden die Vorhersagen mit einer Gewichtung, auch Lernrate genannt, zum ersten Modell addiert. Die Lernrate definiert, in welchem Umfang die einzelnen Vorhersagen bei der Gesamtvorhersage Berücksichtigung finden. Ein bekanntes Problem von Gradient Boosting ist seine Neigung zum Overfitting. Um dies zu vermeiden, ist es wichtig, die Hyperparameter gezielt anzupassen [20]. Dies umfasst bei der Verwendung von DTs beim Gradient Boosting beispielsweise die Anzahl der DTs, deren maximale Tiefe sowie die Lernrate. XGBoost (eXtreme Gradient Boosting) stellt eine schnellere und optimierte Version des klassischen Gradient Boosting Algorithmus dar. Andere Optimierungsverfahren werden genutzt, um die Suche nach dem Minimum der Fehlerfunktion zu beschleunigen [18]. CatBoost (Categorical Boosting) stellt einen weiteren Algorithmus des Gradient Boosting dar, welcher sich dadurch auszeichnet, dass er kategoriale Merkmale automatisch verarbeiten kann, ohne dass vorher eine manuelle Kodierung erfolgen muss [29]. Obwohl sich diese Arbeit nicht auf kategoriale Daten bezieht, können die Vorteile von CatBoost von Nutzen sein.

Eine weitere weit verbreitete Methode des Ensemble Learning ist das Stacking. Beim Stacking wird ein zusätzliches Modell eingesetzt, um die Vorhersagen der einzelnen Modelle optimal zu kombinieren. Stacking funktioniert wie das Aufeinanderstapeln von Modellen. Eine Schicht von Vorhersagen wird als Grundlage für die nächste Schicht

genutzt [20]. Der Ablauf des Trainings ist in Abbildung 2.11 dargestellt. Im ersten Schritt werden mehrere Modelle, die sogenannten Basismodelle, parallel trainiert, um individuelle Vorhersagen zu erstellen. Es empfiehlt sich, dabei auf verschiedene Modelle zurückzugreifen wie beispielsweise DTs, ANNs oder die KNN-Methode. Die Basismodelle weisen jeweils unterschiedliche Stärken und Schwächen bezüglich eines bestimmten Datensatzes auf. In der Regel werden auf diese Weise bessere Prognoseergebnisse erzielt. Auf einer zweiten Ebene wird ein Metamodell integriert, welches die Ergebnisse der Basismodelle als Trainingsdaten nutzt, um die finale Vorhersage zu treffen [30].

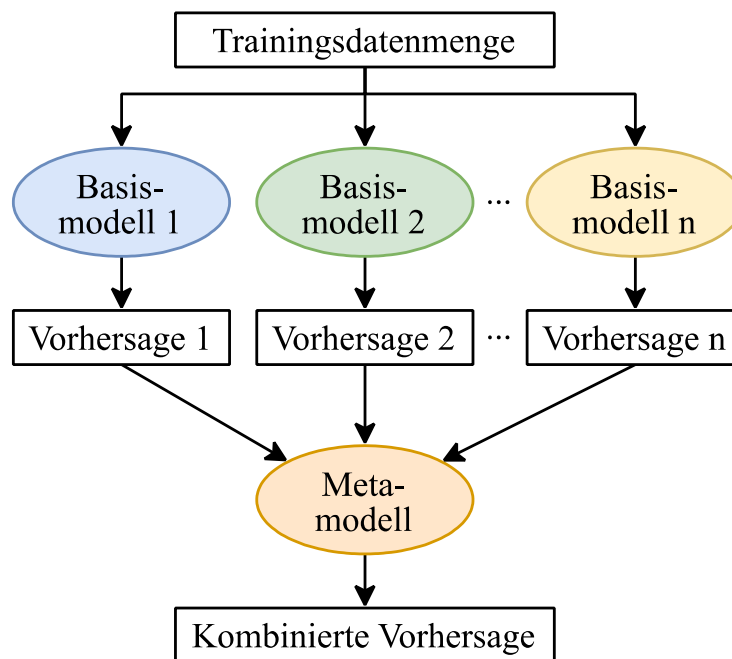


Abbildung 2.11: Ablauf des Trainings beim Stacking nach [20, 30]

2.2 Methoden der PV-Prognose in der Wissenschaft

Für die Prognose der erzeugten PV-Leistung wird in der Wissenschaft eine Vielzahl von Modellen angewendet, welche sich in Abhängigkeit von den verfügbaren Daten, dem betrachteten Zeithorizont sowie dem Anwendungskontext der Vorhersage unterscheiden. Die verschiedenen Modelle zur Vorhersage der erzeugten PV-Leistung lassen sich in unterschiedliche Gruppen unterteilen, um eine sinnvolle Einordnung der Modelle in den Stand der Wissenschaft zu erhalten. In der Literatur werden verschiedene Gruppen zur Unterteilung genutzt, die sich von Arbeit zu Arbeit unterscheiden. Es gibt folglich keine einheitliche Unterteilung, jedoch gruppieren viele Arbeiten die Prognosemodelle nach ähnlichen Merkmalen. Mögliche Cluster wie in [11, 12] sind zum Beispiel die angewendete Vorhersagemethodik, der zeitliche Prognosehorizont, die räumliche Größen-

ordnung und die Ausgabeart des Prognosemodells. In Abbildung 2.12 sind die möglichen Gruppen zur Unterteilung dargestellt. Die einzelnen Gruppen werden im weiteren Verlauf näher erläutert.

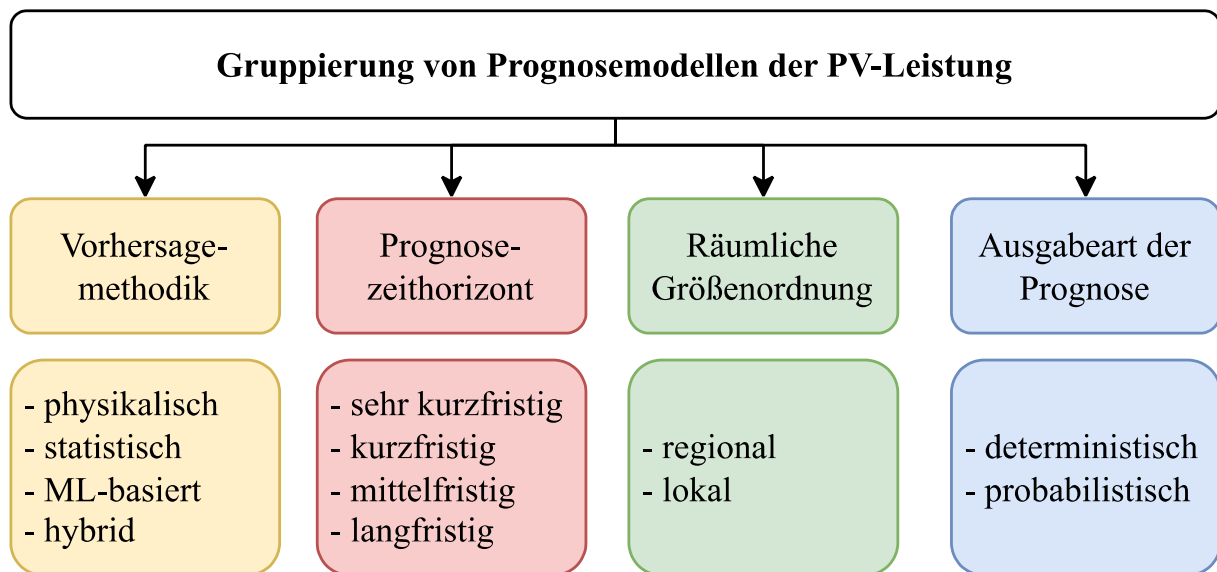


Abbildung 2.12: Übersicht über die Gruppierung von Prognosemodellen der PV-Leistung

Die Vorhersagemodelle lassen sich anhand der verwendeten Vorhersagemethodik in physikalische, statistische und hybride Modelle einteilen (siehe Abbildung 2.12, gelber Bereich), wobei sich die Einteilung auf Basis der verwendeten Methodik in der Forschung in verschiedenen Arbeiten auch geringfügig unterscheiden kann [11, 12].

Physikalische Modelle simulieren PV-Anlagen. Solche Modelle beruhen auf spezifischen Anlagenparametern, wie der Ausrichtung der Anlage, dem Neigungswinkel, dem Wirkungsgrad und dem Standort. Diese Parameter werden in mathematischen Zusammenhängen integriert, um die erzeugte Leistung der PV-Anlage zu berechnen [11, 12]. Die Berechnungen basieren auf der Vorhersage von meteorologischen Daten wie der Bestrahlungsstärke und der Temperatur aus einer NWP [12]. Die Genauigkeit und Auflösung der Vorhersage der erzeugten PV-Leistung ist abhängig von der NWP [11].

Zu den statistischen Modellen werden in der Literatur sowohl die klassischen Zeitreihenvorhersagen als auch die Vorhersagen basierend auf ML-Methoden gezählt [11, 31]. Die Basis der statistischen Methoden sind historische Daten, die es ermöglichen, Muster innerhalb des Datensatzes zu erlernen. Dabei werden sowohl Wetterdaten als auch Messdaten der erzeugten Leistung einer PV-Anlage herangezogen. Für eine hohe Genauigkeit der Vorhersage für die entsprechende Anwendung ist die Größe und Qualität des Datensatzes von entscheidender Bedeutung [12].

Hybride Modelle sind eine Kombination aus mehreren Modellen. Hierbei sind sowohl die Kombination von verschiedenen Vorhersagemethoden wie physikalischen Modellen

und statistischen Modellen als auch mehrere Methoden eines Modelles als Ensemble möglich [12]. Aufgrund der Komplexität der PV-Leistungsprognose können hybride Modelle hilfreich sein, indem sie die Vorteile der jeweiligen Methode nutzen [11, 12]. Allerdings sind hybride Modelle komplizierter und benötigen mehr Rechenleistung sowie einen höheren Zeitaufwand. In der Wissenschaft wurde in verschiedenen Arbeiten wie in [32, 33] gezeigt, dass der Vorhersagefehler durch die Nutzung hybrider Modelle geringer ist und dadurch eine genauere Prognose erstellt werden kann.

Der zeitliche Prognosehorizont von PV-Prognosemodellen lässt sich, wie in Abbildung 2.12 im roten Bereich dargestellt, in sehr kurzfristige, kurzfristige, mittelfristige und langfristige Prognosen unterteilen (siehe Tabelle 2.1) [11, 12]. Die Nutzung der verschiedenen Prognosemodelle hinsichtlich Zeithorizont hängt von den für die Prognose zu Verfügung stehenden und verwendeten Daten ab. Dabei spielt zum Beispiel die Abhängigkeit von einer Wettervorhersage oder die Nutzung von Himmelsbildern eine Rolle für die Länge des Vorhersagehorizonts. Darüber hinaus hängt der Zeithorizont von der spezifischen Anwendung im Energiesystem ab. Kurzfristige Prognosen werden eher für betriebliche Anwendungen eingesetzt während langfristige Prognosen in der langfristigen Planung von Stromerzeugung, Übertragung und Verteilung Anwendung finden. Die Länge des Prognosehorizonts hat zudem einen Einfluss auf die Genauigkeit der Prognose. Kurzfristige Prognosen sind in der Regel exakter als längerfristige [12, 34]. Die Definitionen bezüglich des Zeithorizontes und der Anwendung unterscheiden sich in den Arbeiten leicht [11, 12, 35–37].

Tabelle 2.1: Einteilung des zeitlichen Prognosehorizonts nach Zeithorizont und Anwendung nach [11, 12]

Gruppe	Zeithorizont	Anwendung
Sehr Kurzfristig	1 Sekunde bis 1 Stunde	Echtzeitüberwachung, Echtzeit-Stromeinsatz, Leistungsglättung, Strompreisbildung
Kurzfristig	1 Stunde bis 1 Woche	Sicherung des Netzbetriebs, wirtschaftliche Lastverteilung, Lastausgleich, Anlagenmanage- ment und Energiemanagementsysteme
Mittelfristig	1 Woche bis 1 Monat	Planung des zukünftigen Stromeinsatzes, War- tung
Langfristig	1 Monat bis 1 Jahr	Langfristige Planung von Erzeugung, Übertra- gung und Verteilung

Prognosemodelle können anhand der räumlichen Größenordnung der Vorhersage in regionale und lokale Vorhersagemodelle (siehe Abbildung 2.12, grüner Bereich) basierend auf der Anzahl der betrachteten PV-Anlagen und der betrachteten Fläche unterteilt

werden. Regionale Vorhersagen betrachten mehrere PV-Anlagen in einem Gebiet während lokale Vorhersagen eine einzelne PV-Anlage betrachten [11]. Die Vorhersage der Leistungserzeugung einer einzelnen Anlage ist komplexer, da lokale Wetteränderungen schwer vorhersagbar sind. Bei der regionalen Vorhersage werden mehrere PV-Anlagen aggregiert betrachtet.

Zwei verschiedene Ausgabearten der Prognose können unterschieden werden. Dabei kann die Ausgabe der Prognose entweder deterministisch oder probabilistisch sein (siehe Abbildung 2.12, blauer Bereich) [11]. Während bei einer deterministischen Prognose genaue Werte für jeden Zeitschritt vorhergesagt werden, werden bei einer probabilistischen Prognose zusätzlich Vorhersageintervalle, Quantile oder Dichtefunktionen vorhergesagt, um Unsicherheiten anzugeben [11, 38]. In der Literatur werden probabilistische Prognosen häufig als sehr hilfreich für Entscheidungsträger wie Netzbetreiber und Energiehandelsunternehmen bewertet, um auf obere und untere Grenzen der Prognose eingehen zu können [11, 31].

Im weiteren Verlauf erfolgt eine Einordnung von aktuellen wissenschaftlichen Arbeiten anhand der vorgestellten Einteilung. Die Arbeiten werden anhand der verwendeten Methoden gruppiert und vorgestellt. Diese unterscheiden sich anhand der verwendeten Daten, des Zeithorizonts der Prognose, der Ausgabeart und der räumlichen Größenordnung.

In [39] wird ein physikalisches Modell vorgestellt, welches mittels Wettervorhersagedaten aus einer NWP, in diesem Fall das Europäische Zentrum für mittelfristige Wettervorhersage (*engl. European Centre for Medium-Range Weather Forecast, ECMWF*), die Erzeugung einer PV-Anlage mittels technischen Merkmalen der Anlage und Solartheorie berechnet. Die Autoren kommen zu dem Schluss, dass die Wettervorhersage der verwendeten NWP zuverlässig genutzt werden kann für Prognosen der PV-Leistung. Allerdings stellen NWPs meistens nicht alle notwendigen Wettervorhersagedaten frei zur Verfügung. Daten wie die Bestrahlungsstärke müssen für einen solchen Ansatz gegeben sein. Außerdem ist dadurch eine Abhängigkeit von der Genauigkeit und der räumlichen Auflösung der NWP gegeben. Daher finden physikalische Modelle meist in hybriden Prognosemodellen Anwendung. Dabei werden die physikalischen Modelle zur Umrechnung der durch andere Modelle prognostizierten Bestrahlungsstärke in die erzeugte PV-Leistung wie in [10, 40, 41] genutzt, um eine Vorhersage der PV-Leistung zu erhalten.

Klassische statistische Modelle werden in [34, 42–45] für die PV-Erzeugungsprognose angewendet. In [34] wird mittels eines LR-Modells ein PV-Vorhersagemodell mit Messungen der Bestrahlungsstärke erstellt. Häufig werden in den Arbeiten der klassischen statistischen Modelle Methoden der Zeitreihenvorhersage wie in [43, 44] verwendet.

Nach [11] sind ARIMA (Auto Regressive Integrated Moving Average) und SARIMA (Seasonal ARIMA) bekannte Methoden zur Zeitreihenvorhersage.

ML-basierte Methoden haben in den letzten Jahren an großer Bedeutung für die PV-Erzeugungsprognose gewonnen, wodurch eine Vielzahl wissenschaftlicher Arbeiten sich mit dieser Methodik beschäftigen. Aufgrund der Fähigkeit, mit nichtlinearen Zusammenhängen in meteorologischen Daten und deren Komplexität umzugehen, sind ML-Methoden für die Prognose geeignet.

In [32, 46–48] werden klassische ANNs genutzt, um die erzeugte PV-Leistung mit Wetterdaten und vorhandenen PV-Leistungsmessdaten vorherzusagen. In [46] wird eine regionale Vorhersage der PV-Leistung untersucht. Es wird ein gemeinsamer Datensatz mit Messdaten von verschiedenen sich in der Nähe befindenden PV-Anlagen erstellt, wodurch eine regionale Prognose erreicht wird. In [48] werden mehrere ANNs für eine Day-Ahead-PV-Prognose trainiert und für die Prognose der Mittelwert mehrerer Modelle gebildet. Mittels eines ANN und historischen Wetterdaten wird in [49] die Globalstrahlung für einen Tag im Voraus prognostiziert. Das ANN sagt die Globalstrahlung mit den Wetterdaten der vorherigen 24 Stunden voraus, wodurch keine NWP notwendig ist. Allerdings müssen verschiedene Wetterdaten wie Temperatur, Luftfeuchtigkeit und Druck an der PV-Anlage bekannt sein.

In [31, 35, 38, 50, 51] werden Long Short-Term Memory (LSTM) Netze verwendet, um mit PV-Leistungsmessdaten sowie Wetterdaten eine PV-Leistungsprognose zu erstellen. Die wissenschaftlichen Arbeiten [50, 51] beschäftigen sich mit verschiedenen Typen von Tagen, die sich in den Wetterbedingungen unterscheiden. In [51] wird eine Zeitreihenvorhersage mittels LSTM für eine Stunde im Voraus für die PV-Stromproduktion erstellt. Dabei werden verschiedene Modelle für überwiegend bewölkte Tage, teilweise bewölkte Tage und klare Tage trainiert. Eine Vorhersage des Klarheitsindex wird vorausgesetzt. In [50] wird eine synthetische Wettervorhersage mit einer Vorhersage des Himmeltyps und historischen Bestrahlungsstärkedaten erstellt, um für verschiedene Himmelstypen mittels LSTM eine Vorhersage der PV-Leistung zu treffen. In [31, 38] werden mittels LSTM probabilistische Prognosen der PV-Stromerzeugung untersucht. Die Autoren in [31] nutzen ausschließlich historische PV-Leistungsmessdaten, um ein robustes Modell zu trainieren, das mit fehlenden Daten arbeiten kann. In [38] wird eine probabilistische Day-Ahead-Prognose mit vorhandenen PV-Leistungsmessdaten und Wettervorhersagedaten wie Bestrahlungsstärke und Temperatur mit einem Fehler als der normalisierte mittlere absolute Fehler (*engl. normalized Mean Absolute Error, nMAE*) von 7,7 % erreicht. Auch in [35] werden neben Wettervorhersagedaten solare Bestrahlungsstärkevorhersagen sowie eine ausreichende Menge historischer PV-Leistungsmessdaten einer einzelnen Anlage vorausgesetzt. Mit einem NARX-LSTM, welches eine Kombination aus einem Nonlinear Auto-Regressive Neural Network with exo-

genous input und einem LSTM ist, wird ein Prognosefehler als die normalisierte Quadratwurzel des mittleren quadratischen Fehlers (*engl. normalized Root Mean Square Error, nRMSE*) von 1,33 % für eine regionale Vorhersage mit mehreren PV-Anlagen erreicht.

In [37, 52, 53] werden Convolutional Neural Networks (CNN) nicht wie gewöhnlich für die Bilderkennung sondern für die Zeitreihenvorhersage der erzeugten PV-Leistung verwendet. Alle Arbeiten gehen von einer umfassenden Datengrundlage mit historischen PV-Messdaten und vorhandenen Wettervorhersagedaten mit insbesondere der Bestrahlungsstärke aus.

Die Bewölkung über einer PV-Anlage hat einen großen Einfluss auf die PV-Leistungserzeugung. Wenn CNNs klassisch für die Bildverarbeitung verwendet werden, gibt es zwei Arten von Bildern, die in der Wissenschaft verwendet werden, um die Wolken über einer PV-Anlage darzustellen. Es werden entweder Satellitenbilder oder Bilder einer Weitwinkelkamera, die den Himmel über einer PV-Anlage aufnimmt, verwendet. Während Satellitenbilder eher größere Gebiete abbilden, werden mit Bildern aus Weitwinkelkameras lokale Bereiche des Himmels betrachtet und damit sehr kurzfristige Vorhersagen realisiert. Dabei haben Himmelsbilder eine höhere Auflösung als Satellitenbilder [54]. In [55] werden CNNs zur Verarbeitung von Satellitenbildern verwendet. Diese werden mit meteorologischen Daten kombiniert, um eine Vorhersage der Bestrahlungsstärke für ein bis vier Stunden im Voraus zu erhalten. Dabei werden Bewölkungsfaktoren aus Satellitenbildern bestimmt und diese zusammen mit Wetterdaten in einem Vorhersagemodell für die Bestrahlungsstärke genutzt. In [56, 57] werden Himmelsbilder einer Weitwinkelkamera mittels CNNs verarbeitet und die Optical-Flow-Methode für die Bestimmung von Wolkenbewegungsvektoren zwischen zwei Himmelsbildern genutzt. In [57] werden Himmelsbilder verwendet, um ein Warnsystem ein bis zwei Minuten vor der Verdeckung der Sonne auszulösen. Wetterdaten oder die Bestrahlungsstärke werden nicht berücksichtigt. Ebenso werden in [58] Himmelsbilder zur Darstellung der Wolkenverteilung auf der Basis von Convolutional Autoencoder mit einem Prognosehorizont von bis zu vier Minuten ohne Berücksichtigung von Wetterdaten oder PV-Leistungsmessdaten vorhergesagt. In [56] hingegen werden zusätzlich Messungen der Bestrahlungsstärke am Ort der Aufnahme der Himmelsbilder einbezogen, um eine Vorhersage der Bestrahlungsstärke zehn Minuten im Voraus zu erhalten. Für einen einmonatigen Testzeitraum ergibt sich für sonnige Tage ein Fehler, ausgedrückt als mittlerer absoluter prozentualer Fehler (*engl. Mean Absolute Percentage Error, MAPE*), von 2,91 %, für teilweise bewölkte Tage ein Fehler von 4,84 % und für bedeckte Tage ein Fehler von 6,6 %. In [54] wird eine probabilistische Vorhersage der Bestrahlungsstärke für vier Stunden im Voraus auf der Basis von bodengestützten Wolkenbildern mittels DenseNet (Densely Connected Convolutional Networks) erstellt.

Neben CNNs werden auch andere Methoden eingesetzt, um aus Satellitenbildern und Himmelsaufnahmen Vorhersagen über Bewölkung und Sonneneinstrahlung zu treffen. In [59] wird eine Vorhersage der Wolkenbewegung mit Hilfe von Satellitenbildern durchgeführt, die für die PV-Leistungsvorhersage verwendet werden kann. Aus aufeinanderfolgenden Bildern werden mit der Optical-Flow-Methode Bewegungsvektoren bestimmt und mit einem ANN die Vektoren für die Wolkenbewegung vorhergesagt. Wetter- oder PV-Leistungsdaten werden nicht berücksichtigt. Für eine lokale PV-Leistungsvorhersage ist die Qualität und Auflösung der Satellitenbilder entscheidend. Wolkenbewegungen werden auch in [60] mit dem der von der Natur inspirierten Methode Flocking Behavior Algorithm (oder Boids Algorithm) vorhergesagt. Der Vorhersagehorizont liegt zwischen 84 und 168 Stunden. Damit wird ein eher großes Gebiet für die Wolkenbewegungen betrachtet. In [43, 61] werden Himmelsbilder verarbeitet. In [43] erfolgt eine Vorhersage der Bestrahlungsstärke mit einem Prognosehorizont von bis zu fünf Minuten mit einem Fehler von 18,9 % für einen ausgewählten Tag ausgedrückt als Quadratwurzel des mittleren quadratischen Fehlers (*engl. Root Mean Square Error, RMSE*). Wolkenbewegungsvektoren werden durch Kreuzkorrelation aus Himmelsbildern ermittelt und mit einer Zeitreihenvorhersage der erzeugten PV-Leistung kombiniert. In [61] erfolgt eine Vorhersage der Bestrahlungsstärke mittels Wettermustererkennung auf Basis von der Support Vector Machine (SVM) Methode aus Himmelsbildern für einen Prognosehorizont von einer Stunde mit einem Fehler als nMAE von 6,7 %.

Weitere ML-Methoden wie die Support Vector Regression (SVR) Methode in [62, 63] und die Gaussian Process Regression Methode in [64] werden für die Prognose der erzeugten PV-Leistung genutzt. In [62] wird die Prognose der erzeugten PV-Leistung auf Basis von Big Data, bestehend aus PV-Messdaten und Wetterdaten, durchgeführt. Die Daten werden mittels der Minimum Redundancy Maximum Relevance Merkmalsreduktionstechnik optimiert und reduziert. Es werden vier SVR-Modelle für die vier Jahreszeiten genutzt. Der Fehler der Prognose beträgt 9,48 % bis 12 % für klares Wetter, 12,58 % bis 18,08 % für teilweise bewölktes Wetter, 10,67 % bis 15,59 % für bedecktes Wetter und 13 % bis 19,7 % für regnerisches oder schneereiches Wetter. In [63] wird ein auf Particle Swarm Optimization basierendes SVR-Modell verwendet, um eine Day-Ahead-PV-Prognose zu entwickeln, die historische Daten der erzeugten PV-Leistung sowie Wetterdaten wie die Bestrahlungsstärke, Temperatur und Windgeschwindigkeit nutzt. Die Bestrahlungsstärke wird nicht, wie in vielen anderen Arbeiten, durch eine Wettervorhersage als gegeben angenommen, sondern vorab durch einen Datenaufbereitungsalgorithmus basierend auf Wetterdaten und der Wolkenbedeckung aus einer Online-Wettervorhersage bestimmt. Der Fehler der Prognose als nRMSE liegt bei 2,841 % und als MAPE bei 10,776 %. In [64] werden Datensätze bestehend aus meteorologischen Daten wie Temperatur und Bestrahlungsstärke sowie PV-Leistungsmessdaten mittels der K-Means Clusteringmethode in Cluster gruppiert. Für jedes Cluster wird ein

Matérn 5/2 Gaussian Process Regression Modell trainiert. Mit vier Clustern entsteht ein Prognosefehler als nMAE von 1,85 %, womit ein ausgewogenes Verhältnis zwischen der Komplexität des Modells und der Genauigkeit erreicht wird.

Eine weitere Möglichkeit, die PV-Leistung zu prognostizieren, ist, die Nutzung der räumlichen und zeitlichen Abhängigkeit von PV-Anlagen, was in [42, 65, 66] untersucht wird. In [65] werden zwei Graph Neural Networks, Graph-Convolutional LSTM und Graph-Convolutional Transformer für die deterministische PV-Leistungsprognose mit einem Prognosehorizont von sechs Stunden verwendet. PV-Anlagen werden als ein enges Netz virtueller Wetterstationen betrachtet. Die Prognose nutzt PV-Leistungsmessdaten und erfolgt unter der Berücksichtigung der räumlichen und zeitlichen Abhängigkeit von anderen PV-Anlagen. Der Fehler der Prognose als nRMSE beträgt 8,3 % bis 13,6 %. In [66] wird ein räumlich-zeitliches probabilistisches Prognosemodell basierend auf einem Monotone Broad Learning System and Copulatheorie untersucht. In dieser Arbeit wird ein Prognosehorizont von einer Stunde betrachtet. In [42] wird ein statistisches Verfahren verwendet, um die PV-Leistung bis zu sechs Stunden im Voraus zu prognostizieren, indem Leistungsmessungen anderer PV-Anlagen in der Nähe und die räumlich-zeitliche Abhängigkeit berücksichtigt werden. Es ist ebenfalls möglich, die Bestrahlungsstärke über die räumlich-zeitliche Abhängigkeit vorherzusagen. In [40] wird eine Zeitreihenvorhersage der Bestrahlungsstärke auf Basis der räumlichen und zeitlichen Abhängigkeit für einen Prognosehorizont von drei Stunden untersucht. Die Prognose verwendet Wetterdaten wie die Bestrahlungsstärke, Temperatur und Windgeschwindigkeit der Zielwetterstation und umliegender Wetterstationen. Eine Vorhersage der Wolkenbewegung ist daher nicht notwendig, da dies durch die räumlichen Daten erfolgt.

In [12, 67–71] werden verschiedene ML-Methoden für die PV-Prognose miteinander verglichen. In [67] werden die vier Methoden Multiple lineare Regression, Multiple Polynomielle Regression, KNN und SVR für die Prognose der erzeugten PV-Leistung für zehn Tage im Voraus miteinander verglichen. Der geringste Prognosefehler wird mit der KNN-Methode für einen Testdatensatz von zehn Tagen mit einem nMAE von unter 7,5 % erreicht. Als Datenbasis dient ein kleiner Datensatz von einem Monat, der Messdaten der PV-Leistung und Wetterdaten wie die Globalstrahlung, die Windgeschwindigkeit und die Temperatur enthält. In [68] werden vier Modelle von ANNs, das Non-Linear Autoregressive Network, das Feedforward Neural Network (FNN), das LSTM-Netz und das Echo State Network, für die Zeitreihenvorhersage der Globalstrahlung für einen Prognosehorizont von bis zu 120 Minuten untersucht, wobei das Echo State Network die besten Ergebnisse liefert. Dabei werden nur historische Messwerte der Bestrahlungsstärke und keine Wetterdaten verwendet. In [69] erweist sich die LSTM-Methode als die beste für die Day-Ahead-PV-Leistungsprognose unter den vier getesteten Methoden: LR, FNN, CNN und LSTM. In dieser Arbeit wird die Methode des Transfer

Learning verwendet, bei der ein Prognosemodell mit einem umfangreichen historischen Messdatensatz einer PV-Anlage trainiert wird. Dieses Modell wird auf eine neu errichtete PV-Anlage in derselben Region übertragen, wo noch kein oder nur ein geringer Datensatz vorhanden ist. Der Ansatz stellt eine praktikable Lösung dar und ist direkt anwendbar, da in den meisten Fällen keine oder nicht genügend historische PV-Leistungsmessdaten von einzelnen Anlagen vorliegen. In [70] wurden die zum Vergleich ausgewählten Methoden nach dem Prinzip der Einfachheit ausgewählt, da sie im Vergleich zu komplexen Modellen mindestens gleich gute Vorhersagen liefern. Die Methoden RBFNN (radial basis function neural networks), LS-SVM (least square support vector machines), KNN und Wk NN (weighted KNN) werden für eine Day-Ahead-Prognose der PV-Leistung verglichen. Die Methoden KNN und Wk NN liefern die besten Ergebnisse. Mit der KNN-Methode wird ein Fehler als nMAE von 6,38 % für viele aggregierte PV-Anlagen und mit der Wk NN Methode ein Fehler als nMAE von 7,74 % für eine einzelne PV-Anlage erreicht. Als Daten werden sowohl PV-Leistungsmessdaten als auch verschiedene historische Wetterdaten und Wettervorhersagedaten einschließlich der Bestrahlungsstärke verwendet. In [71] werden verschiedene Sequenz-zu-Sequenz-Encoder-Decoder-Modelle wie die LSTM-Methode und rekurrente neuronale Netze mit den Methoden Gradient Boosted Regression Trees und FNN für die Vorhersage der Bestrahlungsstärke für einen Prognosehorizont von einer Stunde bis 24 Stunden verglichen. Die LSTM-Methode zeigt die besten Ergebnisse mit einem mittlere absoluten Fehler (*engl. Mean Absolute Error, MAE*) von 29,9 W/m² und einem RMSE von 81,1 W/m² für eine 24 Stunden Prognose unter der Berücksichtigung von Daten mehrerer Standorte in der Umgebung. Insbesondere bei einem Vorhersagehorizont von einer Stunde haben räumlich, zeitliche Eigenschaften wie Windrichtung, Windgeschwindigkeit und Bestrahlungsstärke des benachbarten Standortes einen positiven Einfluss auf die Vorhersage. In [12] werden drei ML-Methoden für die Vorhersage der erzeugten PV-Leistung mit einem Prognosehorizont von einer Stunde miteinander verglichen. Dabei werden ein RF, ein ANN und ein LSTM mit historischen PV-Leistungsmessdaten und Wetterdaten wie Temperatur, Druck, Feuchte, Bewölkung und Windgeschwindigkeit trainiert. Das ANN und der RF zeigen die besten Ergebnisse für die Prognose an sonnigen Tagen. An regnerischen oder bewölkten Tagen ist die Vorhersage schwieriger. Insgesamt ist die Methodik leicht anzuwenden, da allgemein zugängliche Wetterdaten verwendet werden. Die Bestrahlungsstärke, die in der Regel nicht kostenfrei über Wettervorhersagen erhältlich ist, wird hier nicht für die Prognose genutzt. In vielen Arbeiten wird die Bestrahlungsstärke jedoch als Eingangsparameter für die Prognose verwendet. Dennoch wird in dieser Arbeit eine große Menge an PV-Leistungsmessdaten verwendet. Die Autoren kommen zu dem Schluss, dass die Genauigkeit der PV-Leistungsprognose sowohl von der verwendeten Methodik, den verfügbaren Daten als auch von der geografischen Lage der PV-Anlage abhängt. Die Arbeiten zeigen, dass keine Methode die

beste ist, sondern dass es auf den jeweiligen Anwendungsfall und die verwendeten Daten ankommt.

In der Wissenschaft werden daher auch Hybridmodelle untersucht. Dabei können sowohl verschiedene Vorhersagemethoden wie physikalische Modelle mit zum Beispiel ML-Modellen als auch mehrere Methoden eines Modells als Ensemble kombiniert werden. Ein Ensemble kann eine Kombination mehrerer Methoden der gleichen Art oder eine Kombination verschiedener Methoden sein. Neben Hybridmodellen, bei denen das Prognosemodell aus einem Zeitreihenmodell wie in [40] oder einem ML-Modell wie in [10, 41] zur Vorhersage der Bestrahlungsstärke und einem physikalischen PV-Modell zur Umrechnung der Bestrahlungsstärke in die erzeugte PV-Leistung besteht, gibt es weitere Hybridmodelle, die verschiedene Vorhersagemethoden im Kombination nutzen. In [36] wird ein hybrides Modell aus einem ANN und einer Methode, die auf ähnlichen Stunden basiert, zur Vorhersage der Bestrahlungsstärke für 24 Stunden im Voraus vorgestellt. Durch die Kombination der beiden Methoden kann eine optimale Vorhersage in Abhängigkeit von der Größe des Datensatzes erreicht werden, da die ANN-Methode alleine die besten Vorhersagen mit einem größeren Datensatz liefert und die Methode der ähnlichen Stunden bessere Vorhersagen mit einem kleineren Datensatz liefert. Das Hybridmodell erreicht einen Fehler ausgedrückt als MAPE von 21,37 % und als nRMSE von 30,99 %. In [72] wird ein hybrider Ansatz aus zwei ML-Modellen untersucht. Die CNN-Methode berücksichtigt Informationen aus den Leistungsmessdaten vorheriger Tage, die gleichzeitig an verschiedenen Tagen auftreten. Mit der LSTM-Methode wird die erzeugte Leistung für den nächsten Zeitschritt aus den vorherigen PV-Leistungsmessdaten des gleichen Tages vorhergesagt. Es werden nur Leistungsmessdaten einer einzelnen sehr großen PV-Anlage und keine Wetterdaten für einen Prognosehorizont von 15 Minuten bis 180 Minuten verwendet.

In [10, 48, 73, 74] werden mehrere gleiche ML-Methoden zu einem Ensemble für die Prognose verbunden. In [73] wird ein Ensemble aus mehreren Modellen der SVR-Methode gebildet. Die Modelle unterscheiden sich in der Parametrisierung und der Menge und Kombination der Wetterdaten. Die SVR-Modelle prognostizieren die erzeugten PV-Leistung für 24 Stunden im Voraus und die RF-Methode kombiniert die SVR-Modelle, um die Prognose zu optimieren und die beste Kombination der Modelle zu erhalten. Als Daten werden meteorologische Daten, unter anderem die Bestrahlungsstärke und Messdaten der PV-Leistung verwendet. Nach Ansicht der Autoren sind die Prognosefehler hauptsächlich auf die NWP und die technische Verschlechterung der Anlagen wie Wirkungsgrad und Ausrichtung zurückzuführen. Daher werden in dieser Arbeit Prognosen aus der Vergangenheit hinzugefügt, was die Genauigkeit der kombinierten Prognosen erhöht. In [10, 48] werden mehrere ANNs für eine Day-Ahead-Prognose der Bestrahlungsstärke in [10] sowie der erzeugten PV-Leistung in [48] trainiert und der Mittelwert des Ensembles aus ANNs für die Prognose gebildet, da sich bei jedem Training eines

ANN andere Gewichte einstellen. In [10] erreicht die Prognose der Globalstrahlung für einen Zeithorizont von 24 Stunden einen MAPE von 3,4 % an einem sonnigen Tag und 23 % an einem bewölkten Tag. Die Prognose basiert jedoch nicht auf einer realen Wettervorhersage und wird nicht auf eine reale PV-Anlage angewendet. In [74] wird eine probabilistische Ensemble-Methode untersucht, die für alle ML-Methoden anwendbar ist, da die stündlichen Prognosewerte einer Day-Ahead-Prognose verwendet werden. Aufgrund der allgemein stochastischen Eigenschaften von ML-Methoden ist es sinnvoll, mehrere Modelle in einem Ensemble zu kombinieren. Die Methode schließt die weniger wahrscheinlichen Modelle in einem Ensemble aus und erreicht eine Verbesserung des nRMSE von bis zu 4,79 % für vollständig bewölkte Tage im Vergleich zu einem Ensemble, das auf dem Mittelwert basiert. Als Datengrundlage dienen sowohl PV-Leistungsmessdaten als auch Wetterdaten mit der Bestrahlungsstärke.

In [33, 75–77] werden verschiedene ML-Methoden zu einem Ensemble für die Prognose verbunden. In [33] werden die Methoden Grey-Box Model, ANN, KNN, Quantile Random Forest und SVR sowohl einzeln als auch kombiniert in einem Ensemble über den Mittelwert für die Day-Ahead-PV-Leistungsprognose betrachtet. Es werden verschiedene PV-Anlagen von niedriger bis großer Leistung betrachtet, für die sowohl Leistungsmessdaten als auch Wetterdaten wie die Bestrahlungsstärke vorliegen. Die Ensemble-Methode erreicht insgesamt die höchste Prognosegenauigkeit und der Fehler ausgedrückt als nMAE beträgt im Mittel für alle PV-Anlagen 3,07 %. Bei unterschiedlichen Wetterbedingungen und Jahreszeiten sind einzelne Methoden besser. In [76] wird ein Ensemble basierend auf einem autoregressives Modell, einer radialen Basisfunktion und einem FNN-Modell für die PV-Leistungsprognose gebildet. Die Aggregation der Methoden in einem Ensemble wird durch Gleichgewichtung vollzogen. Der Prognosefehler ausgedrückt als nMAE beträgt 7,72 % für eine Day-Ahead-Prognose. In [75] werden die fünf verschiedenen ML-Methoden SVR, DT, Ensemble-Bagging, Ensemble-Boost und das auf dem Grey Wolf Optimizer basierende Prognosemodell für eine Day-Ahead-PV-Leistungsprognose untersucht, wobei die Ensemble-Bagging Methode die höchste Prognosegenauigkeit ausweist. In beiden Arbeiten dienen Wetterdaten, wie unter anderem die Bestrahlungsstärke und PV-Leistungsmessdaten, als Datengrundlage. In [77] wird eine Vorhersage der PV-Leistung mit den drei verschiedenen probabilistischen Methoden Bayesian Model, Markov Chain Model und Quantile Regression Model in einem Ensemble für einen Prognosehorizont von bis zu 24 Stunden erstellt. Die Gewichte der Modelle für die Prognose werden durch eine Multi-Objective Optimierung bestimmt.

Prognosefehler entstehen vor allem durch Fehler in den verwendeten Wettervorhersagen sowie durch die technische Degradation der PV-Systeme aufgrund von Alterung und äußeren Randbedingungen und Ereignissen wie zum Beispiel Verschattung [73]. Daher

ist eine Fehlerkorrektur wie in [45, 78, 79] sinnvoll. In [45] wird eine Regressionsanalyse verwendet, um die Vorhersage der Bestrahlungsstärke aus einer numerischen Wettervorhersage unter Berücksichtigung der Wettervorhersage aus der Vergangenheit für einen Prognosehorizont von bis zu sieben Tagen zu korrigieren. In [78] wird eine regionale PV-Leistungsprognose für eine Aggregation mehrerer PV-Anlagen mit einem Prognosehorizont von bis zu zwei Tage im Voraus untersucht. Als Daten für ein ML-Modell zur Leistungsprognose werden numerische Wettervorhersagen verwendet, die Wetterdaten wie Bestrahlungsstärke, Temperatur, Luftdruck und Windgeschwindigkeit liefern. Die Fehlerkorrektur für die prognostizierte Leistung wird mit der LR-Methode durchgeführt, sodass eine Korrektur auf Basis verschiedenen NWP Modelle erfolgt. Der nRMSE beträgt 5,28 % für eine Day-Ahead-Prognose. In [79] werden die 24 Stunden Prognosen der erzeugten Leistung, die mit einem ANN-Modell erstellt werden, nach der Prognose mit der K-Means-Clusteringmethode gruppiert und mit der LR-Methode korrigiert. Die Ergebnisse mit historischen Datensätzen zeigen einen MAPE von 4,7 %. Als Daten für die Prognose werden alle Wetterdaten wie die Bestrahlungsstärke und die gemessene erzeugte Leistung verwendet.

Obwohl die wissenschaftliche Literatur zu Prognosemethoden für die Erzeugung von PV-Anlagen sehr umfangreich ist, fehlt es meist an Anwendbarkeit und Praktikabilität. In vielen wissenschaftlichen Arbeiten wird von einer optimalen Datenlage ausgegangen. In der Realität sieht dies jedoch oft anders aus, insbesondere in Deutschland. Historische Leistungsmessdaten von einzelnen PV-Anlagen sind nur sehr begrenzt verfügbar, ebenso wie kostenlose Wettervorhersagedaten zur Bestrahlungsstärke. Dies führt dazu, dass die Methoden zur Prognose in der Praxis kaum anwendbar und die angewandten Verfahren wenig praktikabel sind. Darüber hinaus verwenden viele Arbeiten kleine Datensätze und validieren ihre Methode an wenigen ausgewählten Tagen. Dadurch bleiben viele Wetterbedingungen unberücksichtigt, was die Aussagekraft der Ergebnisse einschränkt. Einige wissenschaftliche Publikationen betrachten nur mehrere aggregierte oder sehr große PV-Anlagen, die im Gegensatz zu kleinen, lokal verteilten PV-Anlagen einfacher zu prognostizieren sind, da ein größeres Gebiet betrachtet wird. In vielen wissenschaftlichen Arbeiten, die sich mit der Vorhersage der Bestrahlungsstärke beschäftigen, werden die Methoden nicht anhand realer PV-Anlagen getestet. Darüber hinaus vernachlässigen viele Arbeiten die Verwendung von realen Wettervorhersagen und setzen stattdessen ausschließlich auf historische Wetterdaten. Da jedoch auch reale Wettervorhersagen Fehler enthalten können, ist es sinnvoll, die Vorhersagegenauigkeit anhand realer Wettervorhersagedaten zu überprüfen. Faktoren wie die Verwendung von Wettervorhersagen, die Alterung der PV-Anlage oder Verschattungen können zu Ungenauigkeiten in der Prognose führen. Daher ist es sinnvoll, Methoden zur Fehlerminimierung in die Analyse einzubeziehen, was aktuell in der Forschung kaum Anwendung fin-

det. In dieser Arbeit wird ein PV-Leistungsprognosemodell untersucht, das die in Kapitel 1.2 aufgeführten Neuheiten enthält und über den aktuellen Stand der Wissenschaft hinausgeht.

Die Definition eines dem aktuellen Stand der Wissenschaften entsprechenden Fehlerwertes für die PV-Prognose aus den zuvor dargestellten wissenschaftlichen Arbeiten gestaltet sich als schwierig. Ein zentraler Grund für diese Schwierigkeit liegt darin, dass in den wissenschaftlichen Arbeiten unterschiedliche Zeitbereiche und Testtage untersucht werden. Für diese Zeiträume werden spezifische Fehler angegeben, die sich je nach Untersuchungszeitraum stark unterscheiden können. Beispielsweise variieren die Testdaten von einzelnen Tagen bis hin zu mehreren Wochen oder Monaten. Einzelne Tage decken jedoch nicht zwangsläufig alle möglichen Wetterbedingungen ab, was die Vergleichbarkeit einschränkt. In einigen Arbeiten werden die Nachtstunden für die Fehlerberechnungen nicht entfernt, was zu einer Verzerrung der Daten führt. Da der Fehler während der Nachtstunden gleich null ist, ist der mittlere Fehler über den gesamten Tag hinweg geringer, als bei einer Fehlerberechnung, bei der die Nachtstunden nicht berücksichtigt werden. Ein weiteres Problem besteht darin, dass unterschiedliche Fehlermetriken verwendet werden. Dies erschwert den direkten Vergleich zwischen den Arbeiten. Zudem gibt es Unterschiede in den verwendeten Datensätzen, im Prognosehorizont sowie in der räumlichen Größenordnung der Vorhersagen. Während einige Arbeiten lokale Vorhersagen betrachten, analysieren andere regionale Modelle. Diese Unterschiede erfordern eine genaue Prüfung, um sicherzustellen, dass Prognosen unter vergleichbaren Bedingungen gegenübergestellt werden. Insgesamt zeigt sich, dass ein pauschaler Vergleich der Fehlerangaben aus den verschiedenen wissenschaftlichen Arbeiten komplex ist und eine differenzierte Betrachtung der jeweiligen Methodik und Randbedingungen erfordert.

2.3 Zusammenfassung

In Kapitel 2 erfolgt eine Definition des Begriffs ML sowie eine Aufteilung in verschiedene Lernverfahren. ML ist ein Teilbereich der KI, der Algorithmen entwickelt, die aus Daten lernen, um Muster zu erkennen oder Vorhersagen zu treffen. Es wird bei komplexen Prozessen eingesetzt und umfasst überwachtes, unüberwachtes und bestärkendes Lernen. In Abschnitt 2.1 werden die für diese Arbeit relevanten ML-Methoden ANN, LR, KNN, DT, Clustering sowie Ensemble Learning näher erläutert. In Abschnitt 2.2 wird ein Überblick über den aktuellen Stand der Wissenschaft zum Thema Methoden der PV-Prognose gegeben. Für die Prognose der erzeugten PV-Leistung wird eine Vielzahl von Modellen angewendet, die sich in unterschiedliche Gruppen unterteilen lassen. In der wissenschaftlichen Literatur gibt es keine einheitliche Klassifizierung, jedoch werden Modelle oft nach Merkmalen wie Vorhersagemethodik, Prognosehorizont,

räumlicher Größenordnung und Ausgabeart gruppiert. Die Prognosemodelle lassen sich anhand der verwendeten Vorhersagemethodik in physikalische, statistische und hybride Modelle einteilen. Physikalische Modelle simulieren PV-Anlagen, indem sie Anlagenparameter und meteorologische Vorhersagen in mathematische Berechnungen zur Bestimmung der Leistung integrieren. Statistische Modelle umfassen sowohl klassische Zeitreihenvorhersagen als auch ML-basierte Methoden. Sie nutzen historische Daten, um Muster im Datensatz zu erkennen. Hybride Modelle kombinieren verschiedene Modelle, wie physikalische und statistische Modelle, oder mehrere Methoden eines Modells als Ensemble. Der Prognosehorizont von PV-Prognosen wird in sehr kurzfristig, kurzfristig, mittelfristig und langfristig unterteilt. Dies hängt von verfügbaren Daten und dem Anwendungskontext im Energiesystem ab. Prognosemodelle lassen sich nach räumlicher Größenordnung in regionale und lokale Modelle unterteilen. Die Prognoseausgabe kann deterministisch oder probabilistisch erfolgen. Es folgt ein Überblick über verschiedene wissenschaftliche Arbeiten zum Thema PV-Prognose, die anhand dieser Merkmale eingeteilt werden. Die wissenschaftliche Literatur zu Prognosemethoden für PV-Anlagen ist sehr umfangreich. Allerdings mangelt es der Mehrzahl der Arbeiten an Anwendbarkeit und Praktikabilität. Viele Arbeiten gehen von idealen Daten aus, während in der Realität in Deutschland historische Leistungsdaten und kostenlose Wettervorhersagen für die Bestrahlungsstärke begrenzt verfügbar sind. Die Aussagekraft einiger Arbeiten bezüglich verschiedener Wetterbedingungen wird durch die Verwendung kurzer Zeiträume zur Validierung eingeschränkt. Zudem werden reale Wettervorhersagen selten genutzt. Fehlerquellen wie die Verwendung von Wettervorhersagen und die Alterung von Anlagen sowie Verschattung werden oft nicht berücksichtigt, sodass Methoden zur Fehlerminimierung in der wissenschaftlichen Literatur selten behandelt werden. Die Definition eines standardisierten Fehlerwertes für PV-Prognosen ist schwierig, da in wissenschaftlichen Arbeiten unterschiedliche Zeiträume, Testtage und Fehlermetriken verwendet werden. Zudem wird in einigen Arbeiten der Fehler während der Nachtstunden nicht entfernt, was die Daten verzerrt. Unterschiede in Datensätzen, Prognosehorizonten und räumlicher Größenordnung erschweren ebenfalls den Vergleich. Daher erfordert der Vergleich der Fehlerwerte eine differenzierte Betrachtung der jeweiligen Methodik und Randbedingungen.

3 PV-Prognose

Das PV-Leistungsprognosemodell besteht aus einer Prognose der solaren Bestrahlungsstärke und einem physikalischen PV-Modell, das die prognostizierte Bestrahlungsstärke in die erzeugte Leistung umrechnet. Eine Übersicht über die Funktionsweise der PV-Leistungsprognose im Betrieb ist in Abbildung 3.1 dargestellt. Über eine NWP für einen bestimmten Ort werden numerische Wettervorhersagedaten für einen bestimmten Zeitraum und eine bestimmte Auflösung in Abhängigkeit von der NWP bezogen. Diese werden als Eingangsdaten einem trainierten ANN zugeführt, welches als Ausgabe die Bestrahlungsstärke hat. Die Ausgabe ist entsprechend dem Zeitraum und der Auflösung der NWP die Vorhersage der Bestrahlungsstärke für einen bestimmten Ort.

Die für den Standort der PV-Anlage prognostizierte Bestrahlungsstärke wird im physikalischen PV-Modell über mathematische Beziehungen in die von der PV-Anlage erzeugte Leistung umgerechnet. Da die Bestrahlungsstärke für einen ausgewählten Standort prognostiziert wird, kann das PV-Leistungsprognosemodell auf alle PV-Anlagen an diesem Standort angewendet werden und ist nicht auf eine einzelne Anlage beschränkt. Indem die technischen Anlagenparameter in das physikalische PV-Modell integriert werden, lässt sich die Bestrahlungsstärke in die entsprechende erzeugte Leistung der betrachteten PV-Anlage umrechnen.

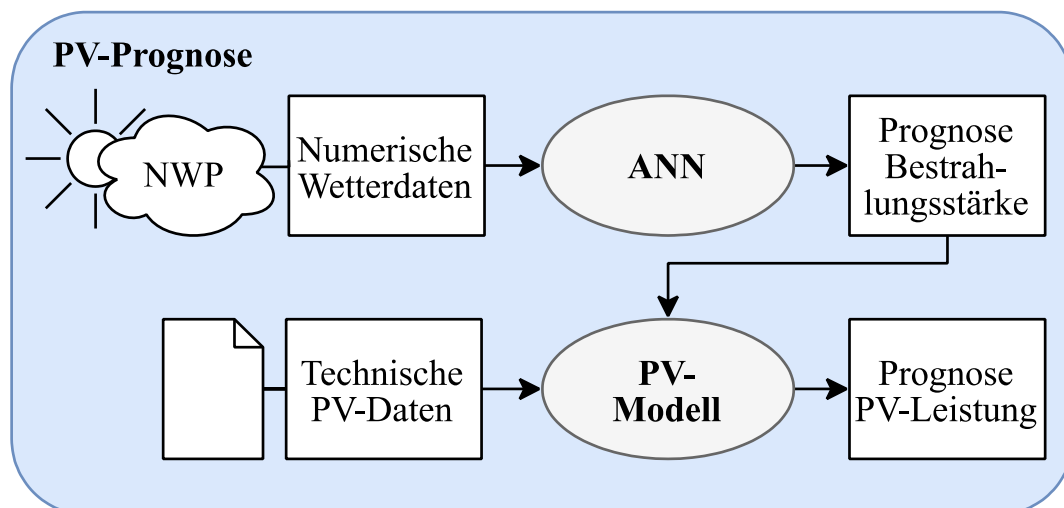


Abbildung 3.1: Übersicht über die Funktionsweise der PV-Prognose im Betrieb

Im Folgenden wird die Methode zur Prognose der Bestrahlungsstärke erläutert. Dazu werden zunächst die dafür genutzten Daten dargestellt und das Training sowie die Validierung des ANN erklärt. Im Anschluss wird die Methode für das physikalische PV-Modell präsentiert. Die Ergebnisse der PV-Leistungsprognose werden anhand einer realen PV-Anlage verifiziert.

3.1 Methode der Prognose der Bestrahlungsstärke

Die Prognose der Bestrahlungsstärke basiert auf einem ANN, welches mit historischen Wetterdaten trainiert wird. Die ANN-Methode erzielt in wissenschaftlichen Arbeiten wie in [10, 36] für die Prognose der Bestrahlungsstärke vielversprechende Ergebnisse, weshalb diese Methode auch in dieser Arbeit zur Prognose der Bestrahlungsstärke eingesetzt wird. Im Gegensatz zu den bisherigen Arbeiten überträgt diese Dissertation die Methode auf eine reale PV-Anlage. Dabei wird eine Wettervorhersage berücksichtigt und die Prognosequalität wird über einen längeren Zeitraum hinweg evaluiert. Nach der Darstellung der Methode der Prognose der Bestrahlungsstärke wird darauf im weiteren Verlauf von diesem Kapitel eingegangen. Die Prognose in dieser Arbeit ist keine Zeitreihenprognose, sondern basiert auf historischen, gemessenen Wetterdaten zum Messzeitpunkt. Für den vorliegenden Anwendungsfall erweist sich ein Multi-Layer-Perzeptron als ausreichend, was einen weiteren Argumentationspunkt für die Nutzung eines ANNs darstellt. Darüber hinaus zeigen auch andere Arbeiten wie in [12, 46, 48, 79] dargestellt, dass sich ANN-Modelle erfolgreich zur Prognose der PV-Leistung auf Basis von Leistungsmessdaten einsetzen lassen.

In Abbildung 3.2 sind die Wetterdaten dargestellt, die als Eingangs- und Zieldaten für das ANN genutzt werden. Diese Wetterdaten werden eingesetzt, da diese in klassischen NWP als Wettervorhersagedaten vorhanden sind. Zudem werden die Wetterdaten in vielen wissenschaftlichen Arbeiten in verschiedenen Kombinationen betrachtet (siehe Kapitel 2.2). Die Vorhersage der Bestrahlungsstärke ist dagegen in der Regel nicht kostenfrei über klassische NWP verfügbar. Aus diesem Grund erfolgt in dieser Arbeit eine Prognose der Bestrahlungsstärke.

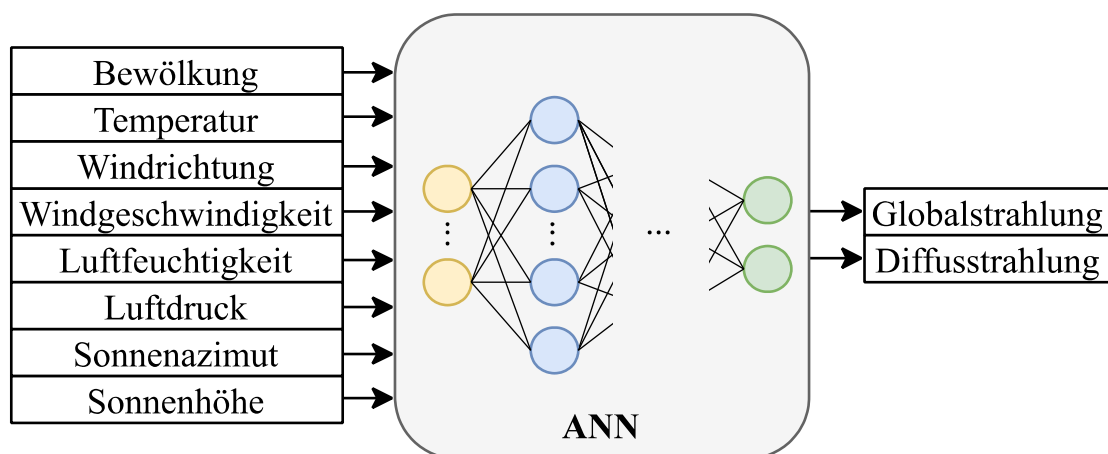


Abbildung 3.2: Darstellung der Methode der Prognose der Bestrahlungsstärke mittels eines ANN

Die Bestrahlungsstärke der Sonne setzt sich aus der direkten und der diffusen Strahlung zusammen. Während die direkte Strahlung direkt aus der Sonnenrichtung kommt, entsteht die diffuse Strahlung durch Streuung in der Atmosphäre und besitzt damit keine definierte Richtung. Die Summe aus direkter Bestrahlungsstärke $E_{\text{Dir,h}}^t$ und diffuser Bestrahlungsstärke $E_{\text{Diff,h}}^t$ ergibt die gesamte bzw. globale Bestrahlungsstärke $E_{\text{G,h}}^t$ auf die horizontale Erdoberfläche zum Zeitpunkt t gemäß Formel (3.1) [80]. In dieser Arbeit wird für das Training eines ANN lediglich die Globalstrahlung und die Diffusstrahlung berücksichtigt, da die Direktstrahlung aus diesen beiden Komponenten nach Formel (3.1) direkt berechnet werden kann.

$$E_{\text{G,h}}^t = E_{\text{Dir,h}}^t + E_{\text{Diff,h}}^t \quad (3.1)$$

Das Ziel des ANN ist die Herstellung einer Beziehung zwischen Eingangs- und Zieldaten, d. h. zwischen verschiedenen Wetterbedingungen und der Bestrahlungsstärke. Das ANN wird dahingehend trainiert, dass es lernt, zu welchen Wetterdaten, die zu einem bestimmten Zeitpunkt vorliegen, welche Bestrahlungsstärke zu diesem Zeitpunkt gehört. Die Prognose ergibt sich durch den Einsatz der NWP mit Wettervorhersagedaten als Input in das trainierte ANN. Für die Entwicklung des ANN für die Prognose der Bestrahlungsstärke wird die Programmiersprache Python genutzt. Dabei wird die Bibliothek Keras für die Entwicklung von Deep-Learning-Modellen mit Tensorflow als Backend verwendet.

Im weiteren Verlauf dieses Kapitels erfolgt eine detaillierte Erläuterung der verwendeten Daten sowie eine Beschreibung der Aufbereitung und Aufteilung dieser Daten für das Training und die Validierung des ANN. Im Anschluss erfolgt eine Darstellung des Trainings und der Validierung des ANN.

3.1.1 Datenaufbereitung und -aufteilung

Die historischen Wetterdaten werden vom Deutschen Wetterdienst (DWD) bezogen. Das Climate Data Center des DWD stellt eine Vielzahl verschiedener gemessener Wetterdaten in unterschiedlichen Auflösungen von verschiedenen Wetterstationen kostenlos zur Verfügung [81]. Da die Anzahl der Wetterstationen begrenzt ist und nicht direkt an jeder PV-Anlage eine Wetterstation vorhanden ist, muss eine geeignete Wetterstation aus der Umgebung ausgewählt werden. Für diese Arbeit wurden verschiedene Wetterstationen auf ihre Eignung für das Training eines ANN geprüft und die Wetterstation in Braunschweig aufgrund der umfangreichen Datenbasis ausgewählt. Die Wetterstation in Braunschweig befindet sich auf einer Höhe von 81,2 m ü. NN mit einem Längengrad von 10,45 und einem Breitengrad von 52,29. Diese Wetterstation befindet sich in der Nähe von Bielefeld, wo sich die PV-Anlage befindet, die in dieser Arbeit zum Testen der Methoden verwendet wird. Beide Städte liegen in einem Gebiet, in dem die mittlere

Jahressumme der Globalstrahlung von 1991 bis 2020 ähnlich ist (siehe Anhang A Abbildung A. 1). Damit kann das PV-Prognose-Modell auf jede beliebige PV-Anlage angewendet werden, die sich in der Nähe der verwendeten Wetterstation und in einem Gebiet mit ähnlichen Bedingungen bezüglich der gemessenen Bestrahlungsstärke befindet.

Bei der Betrachtung von PV-Anlagen in anderen Regionen, in denen die mittlere Jahressumme der Globalstrahlung deutlich höher ist, wie zum Beispiel im Süden von Deutschland, wird die Prognose ebenfalls möglich sein. Voraussichtlich wird der Fehler größer sein, da das ANN für die Prognose der Bestrahlungsstärke nicht alle dort auftretenden Wetterbedingungen, wie zum Beispiel höhere Bestrahlungsstärkewerte, gelernt hat.

Tabelle 3.1 zeigt alle für das Training eines ANN für die Prognose der Bestrahlungsstärke verwendeten Daten. Die Daten sind in Eingangs- und Zieldaten unterteilt. Außerdem sind die Herkunft und die Einheit jedes Merkmals angegeben. Alle historischen Wetterdaten, Bedeckungsgrad aller Wolken, Lufttemperatur, Windrichtung, Windgeschwindigkeit, relative Feuchte, Luftdruck, Globalstrahlung und Diffusstrahlung werden vom DWD bezogen. Neben den Wetterdaten wird auch der Stand der Sonne, ausgedrückt als Sonnenazimut und Sonnenhöhe, für das Training verwendet, da der Sonnenstand einen großen Einfluss auf die Bestrahlungsstärke an einem Ort hat [82]. Der Stand der Sonne wird nach Kapitel 3.2.1 mit Kenntnis des Ortes und der Zeitangabe berechnet. In diesem Fall werden der Längen- und Breitengrad von Braunschweig sowie das Datum und die Uhrzeit des jeweiligen Messzeitpunktes von der Wetterstation zur Berechnung der Sonnenposition genutzt.

Tabelle 3.1: Übersicht über die Daten für das Training eines ANN für die Prognose der Bestrahlungsstärke

Datentyp	Merkmale	Herkunft	Einheit
Eingangsdaten	Bedeckungsgrad aller Wolken	DWD	Achtel
	Lufttemperatur	DWD	°C
	Windrichtung	DWD	°
	Windgeschwindigkeit	DWD	m/s
	Relative Feuchte	DWD	%
	Luftdruck auf Stationshöhe	DWD	hPA
	Sonnenazimut	Berechnet	°
	Sonnenhöhe	Berechnet	°
Zieldaten	Globalstrahlung	DWD	J/cm ²
	Diffusstrahlung	DWD	J/cm ²

Der Bedeckungsgrad aller Wolken wird in Achteln angegeben. Dabei gilt die folgende Definition:

- 0 Achtel entsprechen keine Wolken
- 4 Achtel entsprechen einer Wolkendecke, die maximal zur Hälfte des Himmels bedeckt
- 8 Achtel entsprechen einem vollständig bedeckten Himmel

Die Windrichtung wird als Winkel in Grad angegeben. Dabei gilt für die Richtung des Windes die folgende Definition:

- Windrichtung aus Osten entspricht 90°
- Windrichtung aus Süden entspricht 180°
- Windrichtung aus Westen entspricht 270°
- Windrichtung aus Norden entspricht 360°

Die Angabe von Sonnenazimut und Sonnenhöhe erfolgt als Winkel in Grad. Dabei gilt für den Sonnenazimut die folgende Definition:

- Norden entspricht 0°
- Osten entspricht 90°
- Süden entspricht 180°
- Westen entspricht 270°

Alle Wetterdaten werden vom DWD mit einer Auflösung von einer Stunde bezogen, da in dieser Auflösung alle Daten vorliegen und den gleichen Zeitstempel besitzen. Dabei sind der Bedeckungsgrad, die Temperatur, die relative Feuchte und der Luftdruck als Stundenwerte angegeben. Die Windrichtung und die Windgeschwindigkeit sind als Stundenmittel und die Globalstrahlung und die Diffusstrahlung als Stundensumme angegeben. Die Globalstrahlung und die Diffusstrahlung sind auf eine horizontale Fläche bezogen und werden in die Einheit W/m^2 umgerechnet, um die entsprechende Leistung auf der horizontalen Fläche zu erhalten.

Für die Aufbereitung der Daten werden aus den stündlich aufgelösten Datensätzen die Stunden entfernt, in denen keine Sonneneinstrahlung vorhanden ist, wie zum Beispiel nachts. Fehlerhafte Werte oder fehlende Daten werden entfernt. Ebenso werden die Daten der anderen Eingangsgrößen zum entsprechenden Zeitpunkt entfernt, so dass die Länge des Datensatzes für alle Eingangsgrößen gleichbleibt.

Für das Training und die Validierung werden die Daten aus den fünf Jahren 2012 bis 2016 genutzt. Nach der Aufbereitung ergibt sich ein Datensatz von 20.949 Datenpunkten mit zehn Merkmalen, wobei acht Merkmale die Eingangsdaten und zwei Merkmale die Zieldaten darstellen. Es erfolgt eine Aufteilung der Daten in Trainings-, Validie-

rungs- und Testdaten. Die Trainingsdaten dienen zum Training des ANN. Die Validierungsdaten, die dem ANN unbekannt sind, werden zur Bewertung der Genauigkeit des Modells während des Trainings eingesetzt. Sie helfen bei der Anpassung der Hyperparameter und tragen dazu bei, Overfitting zu vermeiden. Der Testdatensatz wird genutzt, um ein fertig trainiertes ANN mit völlig unbekannten Daten zu bewerten. Aufgrund der Verfügbarkeit einer ausreichend großen Menge an Daten kann eine Aufteilung mit der Holdout-Methode erfolgen [21]. Dazu wird aus dem gesamten Datensatz ein Datensatz für die Validierung zurückgehalten. Der Datensatz wird im Verhältnis 80 % Trainingsdaten und 20 % Validierungsdaten aufgeteilt, wobei die Aufteilung zufällig erfolgt. Diese Aufteilung ist in dieser oder ähnlicher Form üblich, jedoch können bei zum Beispiel kleineren Datensätzen andere Aufteilungen erforderlich oder sinnvoll sein [18, 19]. Der Testdatensatz besteht aus Daten eines ganzen Jahres. Damit wird sichergestellt, dass die Prognose der Bestrahlungsstärke für alle Jahreszeiten und Wetterbedingungen innerhalb eines Jahres getestet werden kann. Zum Testen wird das Jahr 2017 genutzt, was einem Testdatensatz von 4.370 Datenpunkten entspricht. Diese Daten sind dem trainierten ANN unbekannt.

Da die Eingangsdaten sowohl unterschiedliche Dimensionen haben als auch große Werte annehmen, erfolgt vor dem Training eine Normalisierung der Daten. Dies ist wichtig, da das Gradientenverfahren beim Training Schwierigkeiten beim Konvergieren aufweist, sofern die Eingangsdaten unterschiedliche Größenordnungen und große Werte besitzen [20, 21]. Dazu wird die z-Transformation verwendet, wobei von jedem Eingangswert x_n eines Merkmals m der Mittelwert μ aller Eingangsdaten eines Merkmals abgezogen und durch die Standardabweichung σ dividiert wird [21, 22]. Damit haben die Daten eines Eingangsmerkmals den Mittelwert 0 und die Standardabweichung 1. Die normierten Daten $z_{m,n}$ werden nach diesem Verfahren gemäß Formel (3.2) berechnet. Diese Normalisierung wird sowohl für die Trainingsdaten und die Validierungsdaten als auch für die Testdaten durchgeführt.

$$z_{m,n} = \frac{x_{m,n} - \mu_m}{\sigma_m} \quad (3.2)$$

Beim Training werden verschiedene Wetterdaten miteinander kombiniert. Dabei wird untersucht, welche Wetterdaten für die Prognose wichtig sind und mit welcher Kombination von Wetterdaten der geringste Fehler bei der Prognose der Bestrahlungsstärke erreicht wird. In Tabelle 3.2 sind sieben verschiedene Kombinationen von Eingangsdaten dargestellt, die jeweils beim Training untersucht werden.

Tabelle 3.2: Kombinationen von Eingangsdaten

Kombination	Eingangsdaten
Kombination 1	Bewölkung
Kombination 2	Sonnenhöhe, Sonnenazimut
Kombination 3	Bewölkung, Sonnenhöhe, Sonnenazimut
Kombination 4	Bewölkung, Sonnenhöhe, Sonnenazimut, Temperatur
Kombination 5	Bewölkung, Sonnenhöhe, Sonnenazimut, Temperatur, Windrichtung, Windgeschwindigkeit
Kombination 6	Bewölkung, Sonnenhöhe, Sonnenazimut, Temperatur, Windrichtung, Windgeschwindigkeit, Luftdruck
Kombination 7	Bewölkung, Sonnenhöhe, Sonnenazimut, Temperatur, Windrichtung, Windgeschwindigkeit, Luftdruck, Luftfeuchtigkeit

3.1.2 Training und Validierung eines ANN

Das Training eines ANN ist komplex. Ein ANN ist ein nichtlineares System, das durch verschiedene Hyperparameter definiert wird. Dazu gehören unter anderem die Anzahl versteckter Schichten und Neuronen, die Aktivierungsfunktionen und die Lernrate, die bestimmt werden müssen. Beim Training werden die Gewichte im ANN so eingestellt, dass der Fehler der Ausgabe minimiert wird. Die Einstellungen der Hyperparameter beeinflussen sich gegenseitig, sodass eine optimale Einstellung der Hyperparameter schwierig ist. Zudem kann es beim Training zu Overfittung oder Underfitting kommen und das Modell darf nicht zu komplex aber auch nicht zu einfach ausgelegt werden [19]. Die Auswahl der optimalen Hyperparameter kann iterativ in Experimenten erfolgen. Die optimale Einstellung der Hyperparameter kann jedoch viel Zeit und Rechenleistung in Anspruch nehmen, da das ANN für jede neue Kombination von Hyperparametern erneut trainiert und validiert werden muss.

Es ist sinnvoll, die Hyperparameter automatisch bestimmen zu lassen. In dieser Arbeit wird die Optimierung der Hyperparameter mit der Bibliothek Keras Tuner durchgeführt, die verschiedene Tuner zur Verfügung stellt. Der BayesianOptimization-Tuner kommt zum Einsatz, da er im Gegensatz zu anderen Tunern, wie dem RandomSearch-Tuner, nicht einfach zufällige Hyperparameterkombinationen auswählt und meist bessere Ergebnisse zeigt als zufallsbasierte Tuner [19]. Stattdessen erlernt der BayesianOptimization-Tuner schrittweise die Bereiche des Hyperparameterraums, die die optimalen Ergebnisse liefern, und nähert sich den besten Hyperparametern schrittweise an [20]. Dieser Tuner nutzt ein Vorhersagemodell, das auf einem Gaußschen Prozess basiert, um die Leistung der bisher getesteten Hyperparameter darzustellen [19, 20]. Die Auswahl der

nächsten Hyperparameter erfolgt in einem Bereich, in dem besonders hohe Leistungswerte bezüglich der Genauigkeit der Ausgabe eines ANN zu erwarten sind. Auf diese Weise werden die Hyperparameter-Kombinationen nicht einfach zufällig ausgewählt, sondern es wird ein Gleichgewicht zwischen der Suche nach neuen Hyperparametern und der Konzentration auf die besten Hyperparameter-Kombinationen geschaffen.

Um die Rechenzeit in einem vertretbaren Rahmen zu halten, können nicht alle Varianten ausprobiert werden. Zusätzlich müssen die Werte, auf die die Hyperparameter eingestellt werden können, begrenzt werden. Diese Grenzen werden bei der Bayesschen Optimierung im Vorfeld definiert. Dabei wird die Orientierung an vergleichbaren Untersuchungen in der Literatur als vorteilhaft erachtet [19]. Diesem Ansatz folgt auch diese Arbeit. In den wissenschaftlichen Arbeiten, die ebenfalls ANNs für die Prognose der Bestrahlungsstärke oder der PV-Leistung einsetzen, werden die Hyperparameter sehr unterschiedlich angegeben, von über 1000 Neuronen [46] bis zu 100 Neuronen [12, 79] bis hin zu einer niedrigen einstelligen Anzahl von Neuronen [36]. Meist werden mehrere versteckte Schichten von bis zu vier Schichten [10, 12, 36, 46], teilweise bis zu sieben Schichten [79] genutzt. Um die optimalen Hyperparameter des ANN für die Prognose der Bestrahlungsstärke zu bestimmen, werden zwei Optimierungen durchgeführt. Zunächst wird ein großer Wertebereich für die Anzahl der Neuronen untersucht, da in der vergleichbaren Literatur auch sehr hohe Neuronenanzahlen angegeben werden. Parallel dazu wird ein kleinerer Wertebereich berücksichtigt, weil in der Literatur ebenfalls niedrige Werte angegeben sind. Die zweistufige Vorgehensweise ist notwendig, da eine feine Suche im gesamten großen Wertebereich zu viel Rechenzeit beanspruchen würde. Stattdessen erfolgt erst eine grobe Variation im großen Bereich und anschließend eine feinere Suche im kleinen Bereich. Auf diese Weise lässt sich die beste Kombination an Hyperparametern bestimmen.

Die Anzahl der Eingangsneuronen liegt zwischen einem Neuron und acht Neuronen und ergibt sich aus der Kombination der betrachteten Eingangsdaten. Die Anzahl der Ausgangsneuronen beträgt zwei und bleibt konstant. Das eine Ausgangsneuron entspricht der Globalstrahlung und das andere Ausgangsneuron der Diffusstrahlung. Es werden eine bis drei versteckte Schichten untersucht. Die Anzahl der Neuronen in den versteckten Schichten wird in der ersten Optimierung von 10 bis 1000 Neuronen variiert, wobei die Anzahl in Abständen von 20 Neuronen erhöht wird. In der zweiten Optimierung wird die Anzahl der Neuronen in den versteckten Schichten von zwei auf 100 Neuronen in Abständen von zwei Neuronen erhöht. Die Aktivierungsfunktion in der Ausgangsschicht ist als linear definiert und fest vorgegeben, da die Ausgabe kontinuierliche Werte ohne Beschränkung annehmen können muss. Während der Optimierung wird in den versteckten Schichten zwischen den Aktivierungsfunktionen ReLU, Sigmoid und Linear variiert. Die Lernrate wird zwischen einer niedrigen Lernrate von $1 \cdot 10^{-4}$ und einer

hohen Lernrate von $1 \cdot 10^{-2}$ wie in [20] mit einer Schrittweite von $2 \cdot 10^{-4}$ variiert. Als Optimierer wird der Adam-Optimierer eingesetzt. Eine Möglichkeit zur Reduktion von Overfitting stellt die Methode des Dropout dar [21, 24]. Im Rahmen dieser Methode erfolgt eine zufällige Deaktivierung bzw. Löschung von Neuronen einer Schicht. Die Dropout-Quote definiert, in welchem Umfang Neuronen zu Null gesetzt werden. In der Regel liegt dieser Wert zwischen 0,2 und 0,5. Dadurch wird ein Auswendiglernen der Trainingsdaten verhindert, indem in jeder Iteration zufällig Neuronen gelöscht werden. Auf diese Weise werden Störungen bzw. Rauschen im Training berücksichtigt. In dieser Arbeit werden die Grenzen für den Dropout mit 0 und 0,5 in 0,1-Schritten definiert.

Eine weitere Möglichkeit, die Wahrscheinlichkeit eines Overfittings zu reduzieren, stellt der Ansatz des frühen Stoppens (*engl. Early Stopping*) dar [20, 24]. In der ersten Trainingsphase ist eine Reduktion des Fehlers der Trainings- sowie der Validierungsdaten zu beobachten. Zu einem bestimmten Zeitpunkt lernt das ANN die Trainingsdaten auswendig. Ab diesem Zeitpunkt ist eine weitere Reduktion des Fehlers der Trainingsdaten zu verzeichnen, während der Fehler der Validierungsdaten ansteigt. An dieser Stelle wird das Training abgebrochen. Diese Funktion wird in Keras über die EarlyStopping-Callback-Funktion implementiert.

Als Verlustfunktion wird der MSE gewählt, der meist für Regressionen verwendet wird [24]. Dieser wird während des Trainings minimiert und ist ein Maß für den Erfolg des Trainings [21]. Wie in Kapitel 2.1.1 beschrieben, werden große Abweichungen durch das Quadrieren stärker gewichtet. Die während der Optimierung eingestellten Hyperparameter des ANN mit der Eingangsdatenkombination 7 und zwei Schichten, welches das ANN mit dem geringsten Prognosefehler ist, sind in Tabelle 3.3 dargestellt. In Tabelle A. 1, Tabelle A. 2 und Tabelle A. 3 im Anhang A sind alle Hyperparameter aller ANNs mit allen Eingangsdatenkombinationen und allen Schichten aufgeführt.

Tabelle 3.3: Hyperparameter des ANN mit der Eingangsdatenkombination 7 und zwei versteckten Schichten

Hyperparameter		Angabe
Anzahl versteckter Schichten		2
Anzahl Neuronen	Versteckte Schicht 1	990
	Versteckte Schicht 2	990
Aktivierungsfunktion	Versteckte Schicht 1	ReLU
	Versteckte Schicht 2	ReLU
	Ausgabeschicht	Linear
Dropout	Versteckte Schicht 1	0
	Versteckte Schicht 2	0,5
Lernrate		0,0007

Während des Trainings werden mehrere Epochen durchlaufen, in denen die Verlustfunktion durch Anpassen der Gewichte im ANN minimiert wird. Wie in Abbildung 3.3 dargestellt, nimmt der MSE sowohl im Training als auch in der Validierung mit jeder Epoche ab. Wenn der Trainingsverlust weiter abnimmt, der Validierungsverlust aber wieder zunimmt, tritt ein Overfitting auf. In diesem Fall passt sich das ANN zu stark an die Trainingsdaten an, sodass eine Verallgemeinerung auf unbekannte Daten, wie beispielsweise die Validierungsdaten, nicht möglich ist. Um dies zu verhindern, wird das Training überwacht und im Falle eines Anstiegs des Validierungsfehlers gestoppt. Dieser Fall tritt in der Abbildung 3.3 nach 75 Epochen ein.

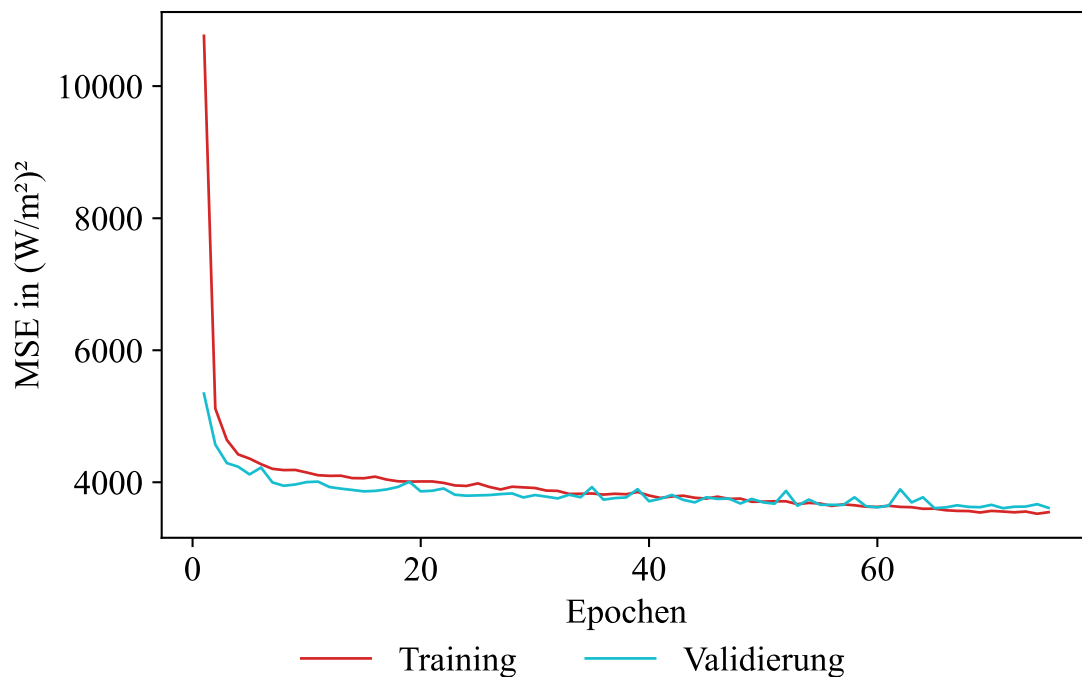


Abbildung 3.3: Verlauf der Verlustfunktion als MSE während des Trainings und der Validierung des ANN mit der Eingangsdatenkombination 7 und zwei versteckten Schichten

Zur Bewertung der Genauigkeit der Ausgaben von ANNs gibt es verschiedene statistische Kennzahlen. In dieser Arbeit werden der MAE, nMAE, RMSE und nRMSE als Fehlermetriken berechnet und zur Bewertung der Prognosegenauigkeit dargestellt.

Der MAE wird nach Formel (3.3) bestimmt, indem die Differenz zwischen dem Vorhersagewert eines ANN und dem tatsächlichen Messwert gebildet wird [83]. Er beschreibt die mittlere Abweichung über einen betrachteten Zeitraum. Dabei ist $\hat{y}_i$ der Vorhersagewert, y_i der tatsächliche Messwert und N die Anzahl der betrachteten Vorhersagen. Der MAE gewichtet alle Fehler gleich [84].

$$\text{MAE} = \frac{1}{N} \sum_{i=1}^N |\hat{y}_i - y_i| \quad (3.3)$$

Der RMSE ist die Wurzel aus dem mittleren quadratischen Fehler. Er gibt an, um welchen Betrag die Prognose im Durchschnitt vom Messwert abweicht. Er ist gemäß folgender Formel (3.4) definiert [27]. Das Quadrieren von Fehlern führt dazu, dass größere Fehler ein größeres Gewicht erhalten. Dadurch ist der RMSE empfindlich gegenüber großen Fehlern und Ausreißern [83, 84].

$$\text{RMSE} = \sqrt{\frac{1}{N} \sum_{i=1}^N (\hat{y}_i - y_i)^2} \quad (3.4)$$

Im Rahmen der Bewertung der Prognose sollte demnach der RMSE nicht isoliert betrachtet werden, da dieser maßgeblich von den Ausreißern beeinflusst wird. Die MAE-Metrik stellt eine robustere Bewertungsmethode dar, die weniger von Ausreißern beeinflusst wird, da sie durch die absoluten Differenzen zwischen vorhergesagten und beobachteten Werten charakterisiert ist.

Aufgrund der Tatsache, dass die Werte für die Bestrahlungsstärke über den Tag eine breite Spanne an Werten aufweisen, werden der MAE und der RMSE darüber hinaus normalisiert dargestellt. Diese Vorgehensweise ermöglicht einen Vergleich zwischen PV-Anlagen unterschiedlicher Größe. Darüber hinaus ist eine relative Angabe des Fehlers einfacher zu interpretieren als absolute Werte des MAE und RMSE, da die wahren Werte mit eingezogen werden. Der nMAE und der nRMSE werden gemäß Formel (3.5) und (3.6) berechnet [33, 85]. Dabei wird der MAE bzw. der RMSE durch die Differenz zwischen dem maximal und minimal gemessenen Wert geteilt. In diesem Fall ist der maximal gemessene Wert die Bestrahlungsstärke bzw. die PV-Leistung.

$$\text{nMAE} = \frac{\text{MAE}}{y_{\max} - y_{\min}} \quad (3.5)$$

$$\text{nRMSE} = \frac{\text{RMSE}}{y_{\max} - y_{\min}} \quad (3.6)$$

Nachfolgend sind die Ergebnisse des Trainings, der Validierung und des Tests für das ANN mit dem geringsten Fehler dargestellt, das im weiteren Verlauf der Arbeit für die Prognose verwendet wird. Das beste ANN hat sich für zwei versteckte Schichten beim Training der ersten Optimierung ergeben, wobei die Anzahl der Neuronen in den Schichten zwischen 10 und 1000 Neuronen variieren. Im Folgenden werden daher die Ergebnisse für ein ANN mit zwei versteckten Schichten aus der ersten Optimierung dargestellt. In Tabelle 3.4 sind der MAE, nMAE, RMSE und nRMSE für das Training, die Validierung und den Test für das ANN mit zwei Schichten für alle sieben getesteten

Kombinationen dargestellt. Die Fehlermetriken für die ANNs mit einer versteckten Schicht und mit drei versteckten Schichten sind in Tabelle A. 4 und Tabelle A. 6 im Anhang A aufgeführt. Zudem sind im Anhang A die Fehlermetriken für die ANNs aus der zweiten Optimierung, bei der die Anzahl der Neuronen in den versteckten Schichten zwischen zwei und 100 Neuronen variiert wird, für eine versteckte Schicht in Tabelle A. 7, zwei versteckte Schichten in Tabelle A. 8 und drei versteckte Schichten in Tabelle A. 9 zu finden.

Tabelle 3.4: Fehlermetriken der Prognose der Bestrahlungsstärke für Trainings-, Validierungs- und Testdaten für ANNs mit zwei versteckten Schichten

Kombination	MAE in W/m ²	nMAE in %	RMSE in W/m ²	nRMSE in %
Training				
1	115,81	14,45	155,15	19,36
2	68,46	8,54	101,88	12,71
3	47,85	5,97	73,05	9,11
4	46,38	5,79	70,82	8,84
5	46,17	5,76	69,50	8,67
6	42,77	5,34	64,59	8,06
7	37,37	4,66	57,34	7,15
Validierung				
1	115,73	14,44	155,04	19,34
2	67,92	8,47	100,83	12,58
3	47,66	5,95	73,25	9,14
4	46,79	5,84	72,32	9,02
5	47,32	5,90	72,14	9,00
6	45,41	5,67	69,72	8,70
7	38,71	4,83	60,05	7,49
Test				
1	122,99	15,35	171,87	21,44
2	70,48	8,79	105,16	13,12
3	62,96	7,85	100,17	12,50
4	60,30	7,52	94,01	11,73
5	59,61	7,44	92,19	11,50
6	59,23	7,39	92,07	11,49
7	47,34	5,91	73,14	9,13

Der Fehler des ANN, welches mit der Kombination 7 der Eingangsdaten trainiert wird, ist am geringsten und der nMAE der Testdaten beträgt 5,91 %. Der nMAE der Trainingsdaten beträgt 4,66 % und der der Validierungsdaten 4,83 %. Da die Testdaten dem ANN unbekannt sind, ist der Fehler der Testdaten folglich höher. Der nRMSE der Testdaten bei Kombination 7 beträgt 9,13 %. Dieser ist größer als der nMAE, was darauf hindeutet, dass die Fehlerwerte Ausreißer enthalten, die durch das Quadrieren der Fehler in der RMSE-Metrik stärker gewichtet werden. Der Fehler, der beim Training des ANN mit der Kombination 1 der Eingangsdaten auftritt, ist am höchsten und der nMAE der Testdaten beträgt 15,35 %. Die Eingangsdaten der Kombination 7 umfassen sämtliche Wetterdaten sowie den Stand der Sonne, während die Kombination 1 lediglich die Bewölkung berücksichtigt. Daher ist es sinnvoll, alle verfügbaren Wetterdaten sowie den Stand der Sonne als Eingangsdaten zu verwenden, um die optimale Prognose der Bestrahlungsstärke zu erhalten. Im weiteren Verlauf dieser Arbeit wird das ANN, welches mit den Daten der Kombination 7 trainiert wird, für die Prognose der Bestrahlungsstärke verwendet.

In Abbildung 3.4 ist die Verteilung des nMAE in einem Boxplot für die verschiedenen Eingangsdatenkombinationen für ANNs mit zwei versteckten Schichten für die Testdaten dargestellt. Die Globalstrahlung und die Diffusstrahlung sind getrennt dargestellt. Ein Boxplot zeigt die Verteilung von Daten [25]. In der Box werden die mittleren 50 % der Daten dargestellt. Die Länge der Box wird als Interquartilsabstand (IQA) bezeichnet. Die horizontale Linie innerhalb der Box ist der Median. Dieser teilt die Daten in genau die Hälfte. Die Antennen, auch Whiskers genannt, außerhalb der Box betragen 1,5-mal dem IQA. Hier werden der kleinste und größte Wert des nMAE innerhalb dieser Grenze angezeigt. Liegen Werte unter- oder oberhalb dieser Grenze, wird dieser Wert mit der Antenne dargestellt. Werte außerhalb werden als Ausreißer bezeichnet und sind in dieser Abbildung für die Übersichtlichkeit nicht dargestellt.

Es zeigt sich erneut, dass die Eingangsdatenkombination 7 den geringsten Fehler sowohl in der Prognose bei der Globalstrahlung als auch bei der Diffusstrahlung hervorruft. Der Fehler für die Globalstrahlung ist höher als für die Diffusstrahlung, weil sich die Globalstrahlung aus Direkt- und Diffusstrahlung zusammensetzt. Es lässt sich feststellen, dass mit einer Zunahme der beim Training verwendeten Eingangsdaten eine Reduktion des nMAE sowie eine Steigerung der Prognosegenauigkeit einhergeht. Die Verwendung lediglich der Bewölkung als Eingangsmerkmal, wie in Kombination 1, resultiert in dem höchsten Fehler. Selbst bei einer Beschränkung auf den Sonnenstand, definiert durch den Sonnenazimut und die Sonnenhöhe, wie in Kombination 2, ist die Vorhersage besser als bei ausschließlicher Berücksichtigung der Bewölkung. Eine weitere Reduktion des Fehlers wird durch die Kombination von Sonnenstand und Bewölkung wie in Kombination 3 erzielt. Mit der Berücksichtigung weiterer Wetterbedingungen wie der Temperatur in Kombination 4, der Windrichtung und Windgeschwindigkeit in Kombination 5

sowie des Luftdrucks in Kombination 6 zeigt sich lediglich eine minimale Verringerung des Fehlers. Diese vier Eingangsdatenkombinationen weisen demnach eine hohe Ähnlichkeit auf. Eine weitere Reduzierung des Fehlers ist zu verzeichnen, wenn zusätzlich wie in Kombination 7 die Luftfeuchtigkeit berücksichtigt wird.

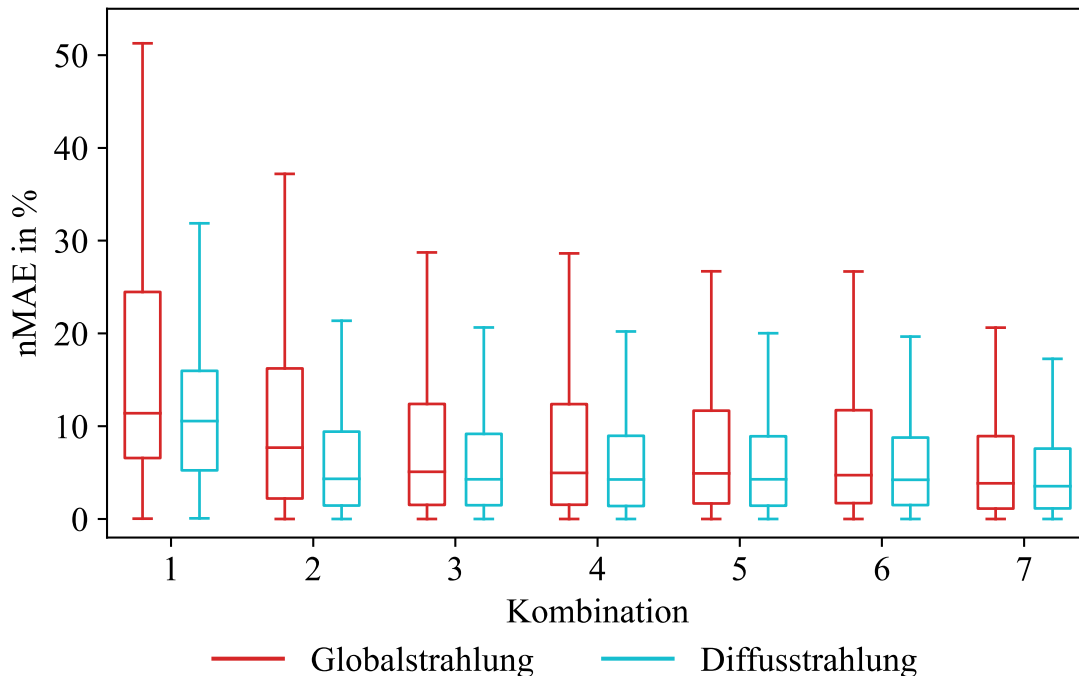


Abbildung 3.4: Darstellung der Verteilung des nMAE in einem Boxplot für verschiedene Eingangsdatenkombinationen für ANNs mit zwei versteckten Schichten für die Testdaten

Die genauen Werte für alle Fehlermetriken aufgeteilt nach Globalstrahlung und Diffusstrahlung sowie als Mittelwerte für Trainings-, Validierungs- und Testdaten für ANNs mit verschiedenen Eingangsdatenkombinationen und verschiedenen Anzahlen an versteckten Schichten sind für die erste Optimierung in Tabelle A. 4, Tabelle A. 5 und Tabelle A. 6 und für die zweite Optimierung in Tabelle A. 7, Tabelle A. 8 und Tabelle A. 9 im Anhang A dargestellt.

In Abbildung 3.5 sind die gemessene und prognostizierte Globalstrahlung, welche vom ANN ausgegeben wird, an zwei Beispieltagen dargestellt. Die zwei Beispieltage entstammen dem Testdatensatz, welcher dem ANN unbekannt ist. Am ersten Beispieltag (06.01.2017) zeigt sich eine hohe Übereinstimmung zwischen Messung und Ausgabe des ANN. Der nMAE liegt bei 0,88 % (siehe Tabelle 3.5) und liegt damit unter dem Durchschnitt von 5,91 % aller Testdaten. Am zweiten Beispieltag (11.07.2017) sind Abweichungen zwischen Messung und Ausgabe aus dem ANN zu erkennen. An diesem Tag liegt der Fehler mit einem nMAE von 15,98 % (siehe Tabelle 3.5) deutlich über dem Durchschnitt. Dies veranschaulicht einen Tag, an dem die Daten durch das ANN

inkorrekt wiedergegeben werden und die zugehörige Globalstrahlung durch die Wetterdaten an diesem Tag unzureichend abgebildet wird. Der 11.07.2017 war ein eher bedeckter Tag, weshalb das ANN vermutlich eine niedrigere Globalstrahlung ausgibt als gemessen wurde.

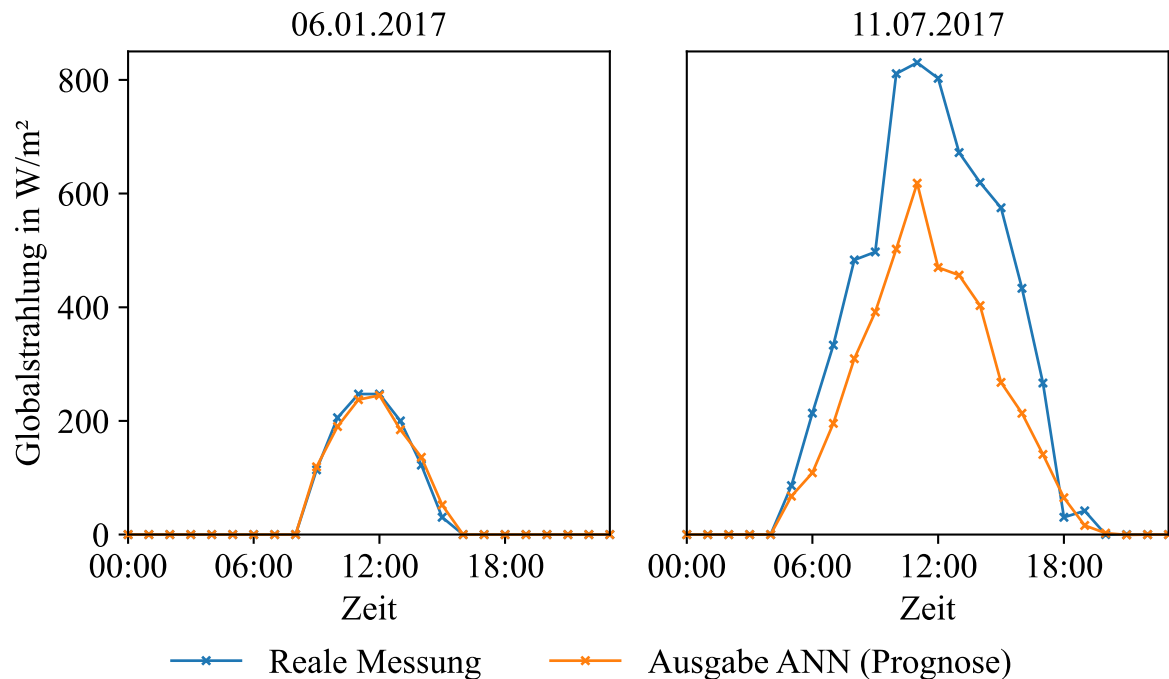


Abbildung 3.5: Gemessene und prognostizierte Globalstrahlung an zwei Beispieltagen der Testdaten

Tabelle 3.5: Fehlermetriken der gemessenen und prognostizierten Globalstrahlung von zwei Beispieltagen aus den Testdaten

Tag	MAE in W/m^2	nMAE in %	RMSE in W/m^2	nRMSE in %
06.01.2017	8,07	0,88	9,95	1,08
11.07.2017	147,30	15,98	179,84	19,51

3.2 Physikalisches PV-Modell

Um die PV-Leistung aus der vom ANN vorhergesagten Bestrahlungsstärke zu berechnen, müssen die technischen Anlagenparameter und Standortgegebenheiten der entsprechenden PV-Anlage, für die die Prognose erstellt werden soll, berücksichtigt werden. Die Leistung, die aus der auf die PV-Anlage treffenden Bestrahlungsstärke resultiert, wird durch diese Anlagenparameter und Standortgegebenheiten beeinflusst. In Abbildung 3.6 sind die im physikalischen PV-Modell betrachteten Parameter dargestellt.

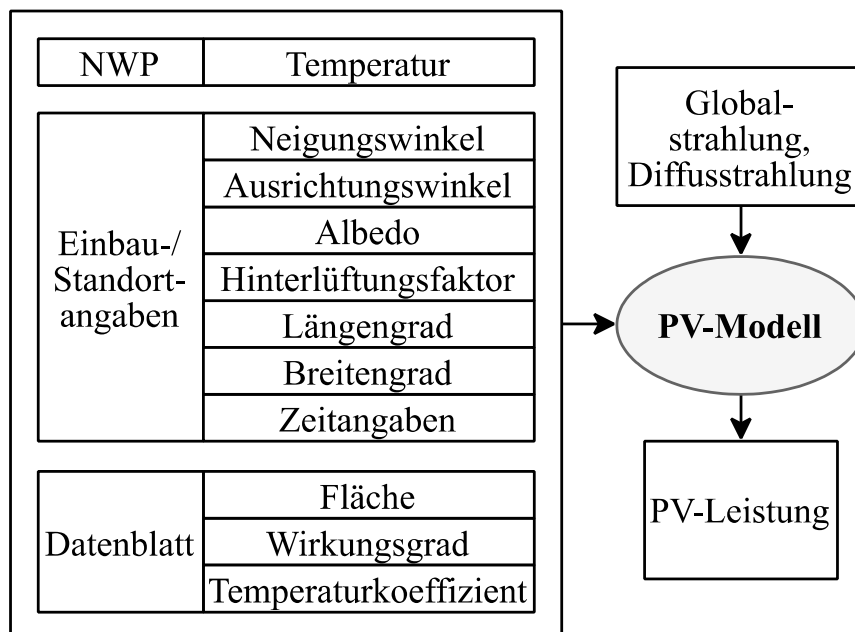


Abbildung 3.6: Übersicht über das physikalische PV-Modell

Die Leistungserzeugung einer PV-Anlage ist von der Umgebungstemperatur abhängig, da die Leistung ab einer bestimmten Temperatur der Module abnimmt. Zusammen mit einem Hinterlüftungsfaktor, der über die Art des Einbaus der PV-Anlage bestimmt ist, ergibt sich eine entsprechende Modultemperatur. Der Temperaturkoeffizient einer PV-Anlage dient der Umrechnung der erzeugten Leistung entsprechend der Modultemperatur. Je größer die Fläche einer PV-Anlage ist, desto größer ist die Leistung. Der Wirkungsgrad einer PV-Anlage gibt an, wie viel der Bestrahlungsstärke in Leistung umgewandelt wird und ist je nach Art der PV-Anlage unterschiedlich. Die Leistung einer PV-Anlage ist von der Ausrichtung und dem Neigungswinkel der PV-Anlage abhängig, da die PV-Anlage dadurch zur Sonne entsprechend ausgerichtet ist. Mittels Zeit- und Standortangaben können der Sonnenazimut und die Sonnenhöhe für den Standort der PV-Anlage bestimmt werden. Diese Angaben sind für die Berechnung der Bestrahlungsstärke auf eine entsprechend geneigte Fläche notwendig. Der Albedo ist definiert durch die Umgebung, in der sich die PV-Anlage befindet und beeinflusst die vom Boden reflektierte Strahlung auf eine geneigte Fläche.

Die Umgebungstemperatur kann über eine NWP bezogen werden. Der Neigungs- und Ausrichtungswinkel sowie der Hinterlüftungsfaktor sind über Einbauangaben bekannt. Die Angabe zum Albedo kann über den Standort abgeleitet werden. Der Längs- und Breitengrad sowie die notwendigen Zeitangaben wie die Uhrzeit, der Tag des Jahres, die Anzahl der Tage im Jahr und die Zeitzone sind ebenfalls über den Standort bekannt. Der Wirkungsgrad, die Fläche und der Temperaturkoeffizient können dem Datenblatt der PV-Anlage entnommen werden.

Im weiteren Verlauf dieses Kapitels erfolgt eine Erläuterung der Berechnung der Sonnenposition, da diese Angabe für die Berechnung der PV-Leistung am Standort der PV-

Anlage erforderlich ist. Darüber hinaus wird die Berechnung der PV-Leistung aus der vorhergesagten Bestrahlungsstärke dargelegt. Des Weiteren wird anhand von Messdaten der Bestrahlungsstärke am Ort einer realen PV-Anlage sowie Leistungsmessdaten dieser PV-Anlage aufgezeigt, wie groß der Fehler des physikalischen PV-Modells ist.

3.2.1 Berechnung der Sonnenposition

Die Position der Sonne ist durch die Sonnenhöhe γ_S und den Sonnenazimut α_S definiert [80]. Die Sonnenhöhe beschreibt den Winkel zwischen dem Sonnenmittelpunkt und dem betrachteten Horizont [80]. Durch die Neigung der Erdachse ist die Sonnenhöhe an einem Standort über das Jahr gesehen unterschiedlich [86]. Der Sonnenazimut ist der Winkel zwischen der Nordrichtung und dem Vertikalkreis durch den Sonnenmittelpunkt [80]. Steht die Sonne im Norden beträgt der Azimutwinkel α_S nach der DIN 5034 0° , im Osten 90° , im Süden 180° und im Westen 270° [80, 86]. Sowohl die Sonnenhöhe also auch der Sonnenazimut sind durch die Position der PV-Anlage sowie das Datum als auch die Uhrzeit definiert [80]. Folglich wird die Position der Sonne gemäß den Formeln (3.7) und (3.8) entsprechend des Tages, der Uhrzeit und des Standortes der PV-Anlage in Abhängigkeit von der Sonnendeklination δ , dem Stundenwinkel ω und dem Breitengrad der PV-Anlage φ_{PV} berechnet. Zur Berechnung wird der Algorithmus der DIN 5034 „Tageslicht in Innenräumen“ verwendet [80, 87].

$$\gamma_S = \arcsin(\cos \omega \cdot \cos \varphi_{PV} \cdot \cos \delta + \sin \varphi_{PV} \cdot \sin \delta) \quad (3.7)$$

$$\alpha_S = \begin{cases} 180^\circ - \arccos\left(\frac{\sin \gamma_S \cdot \sin \varphi_{PV} - \sin \delta}{\cos \gamma_S \cdot \cos \varphi_{PV}}\right) & \text{für WOZ} \leq 12:00 \text{ Uhr} \\ 180^\circ + \arccos\left(\frac{\sin \gamma_S \cdot \sin \varphi_{PV} - \sin \delta}{\cos \gamma_S \cdot \cos \varphi_{PV}}\right) & \text{für WOZ} > 12:00 \text{ Uhr} \end{cases} \quad (3.8)$$

Die Position der PV-Anlage ist durch den Breitengrad φ_{PV} festgelegt, welcher den Winkel zwischen der Verbindung des Erdmittelpunktes zum Standort und dem Äquator angibt [88]. Dieser wird auf der nördlichen Erdhalbkugel positiv und auf der südlichen Erdhalbkugel negativ angegeben. Die Position der Sonne in Abhängigkeit von der Jahreszeit sowie der Uhrzeit wird maßgeblich durch den Deklinationswinkel δ sowie den Stundenwinkel ω bestimmt. Die Erdachse bzw. Rotationsachse der Erde ist um $23,45^\circ$ geneigt, was die Ursache für die Temperatur -und Strahlungsänderungen in einem Jahr ist [88]. Der Sonnendeklinationswinkel gibt den Winkel zwischen dem Sonnenmittelpunkt und Himmeläquator an [80]. Dieser verändert sich im Verlauf eines Jahres zwischen $23,45^\circ$ und $-23,45^\circ$ [80, 82, 89]. Die Nordhalbkugel ist um den 21. Dezember um $-23,45^\circ$ von der Sonne weggeneigt. Dies ist auf der Südhalbkugel die Sommersonnenwende und auf der Nordhalbkugel die Wintersonnenwende. Um den 21. Juni ist die

Nordhalbkugel um $23,45^\circ$ von der Sonne weggeneigt. Auf der Nordhalbkugel ist dann Sommersonnenwende [82, 89]. Die Sonnendeklination δ wird nach Formel (3.9) bestimmt [80]. Die Einheit der Sonnendeklination δ ist Grad.

$$\begin{aligned} \delta = & (0,3948 - 23,2559 \cdot \cos(J' + 9,1^\circ) \\ & - 0,3915 \cdot \cos(2 \cdot J' + 5,4^\circ) \\ & - 0,1764 \cdot \cos(3 \cdot J' + 26^\circ))^\circ \end{aligned} \quad (3.9)$$

Die Berechnung des für die Bestimmung der Sonnendeklination benötigten Parameters J' erfolgt gemäß Formel (3.10) aus der Nummer für den Tag des Jahres (*engl. Day of Year, DOY*) und der Anzahl der Tage im Jahr (*engl. Days in Year, DIY*) [80, 82]. Dieser Parameter gibt die gemittelte Winkelbewegung der Erde um die Sonne an [82].

$$J' = 360^\circ \cdot \frac{\text{DOY}}{\text{DIY}} \quad (3.10)$$

Der Stundenwinkel ω wird nach Formel (3.11) aus der wahren Ortszeit (WOZ) berechnet [80, 90]. Die WOZ richtet sich nach der Bahn der Sonne. Bei dieser Zeit erreicht die Sonne ihren Höchststand um genau 12 Uhr mittags, was einem Stundenwinkel von $\omega = 0^\circ$ entspricht [90]. In einem Zeitraum von 24 Stunden durchläuft die Erde eine vollständige Drehung (360°) mit einer Winkelgeschwindigkeit von 15° pro Stunde [90].

$$\omega = (12 \text{ h} - \text{WOZ}) \cdot \frac{15^\circ}{\text{h}} \quad (3.11)$$

Die WOZ unterscheidet sich von der lokalen Zeit (LZ), die an einem bestimmten Ort in einer Zeitzone gilt [88]. Der Längengrad λ ist der Winkel zwischen der Verbindung vom Erdmittelpunkt zur betrachteten Position und der Äquatorialebene und definiert die Position eines Ortes östlich oder westlich vom Nullmeridian [91]. Dieser verläuft durch Greenwich in England. Östlich vom Nullmeridian ist $\lambda > 0$ und westlich ist $\lambda < 0$. Jede Zeitzone besitzt einen Bezugsmeridian λ_0 . Für die Zeitzone Mitteleuropäische Zeit (MEZ) ist dieser Bezugsmeridian der 15. östliche Längengrad ($\lambda_0 = 15^\circ$) und für die Mitteleuropäische Sommerzeit (MESZ) beträgt $\lambda_0 = 30^\circ$ [88]. In der Regel entspricht der Längengrad λ eines betrachteten Ortes nicht dem Bezugsmeridian [86, 88]. Zudem unterscheidet sich die LZ von der WOZ um vier Minuten pro Längengradabweichung vom Bezugsmeridian [88, 90]. Die LZ an einem Ort wird gemäß Formel (3.12) über die Zeitzone (MEZ = 1 h, MESZ = 2 h) und den Längengrad λ_{PV} des Ortes der PV-Anlage in die mittlere Ortszeit (MOZ) umgerechnet [80]. Aus der MOZ lässt sich gemäß Formel (3.13) die WOZ ableiten, wobei die Zeitgleichung (ZGL) berücksichtigt wird [80].

$$\text{MOZ} = \text{LZ} - \text{Zeitzone} + 4 \cdot \lambda_{\text{PV}} \cdot \frac{\text{min}}{\circ} \quad (3.12)$$

$$\text{WOZ} = \text{MOZ} + \text{ZGL} \quad (3.13)$$

Die ZGL nach Formel (3.14) stellt einen weiteren Korrekturfaktor der LZ von der WOZ dar, welcher aufgrund der ungleichmäßigen Geschwindigkeit der Erde um die Sonne sowie der Neigung der Erdachse erforderlich ist [80, 88]. Die ZGL wird in Minuten angegeben.

$$\begin{aligned} \text{ZGL} = & (0,0066 + 7,3525 \cdot \cos(J' + 85,9^\circ) \\ & + 9,9359 \cdot \cos(2 \cdot J' + 108,9^\circ) \\ & + 0,3387 \cdot \cos(3 \cdot J' + 105,2^\circ)) \text{min} \end{aligned} \quad (3.14)$$

3.2.2 Berechnung der PV-Leistung

Die Berechnung der erzeugten PV-Leistung aus der prognostizierten Bestrahlungsstärke auf eine geneigte und in eine Himmelsrichtung ausgerichtete PV-Anlage erfordert die Berücksichtigung weiterer PV-Anlagen-spezifischer Parameter. Dabei haben die Globalstrahlung auf die geneigte Fläche der PV-Anlage $E_{\text{G,g}}^t$ zum Zeitpunkt t , der Wirkungsgrad und die Fläche der PV-Anlage Einfluss auf die erzeugte Leistung [80, 86]. Die von einer PV-Anlage zum Zeitpunkt t erzeugte PV-Leistung P_{PV}^t wird nach Formel (3.15) berechnet. Der Wirkungsgrad wird dabei als konstant angenommen, da der im Datenblatt angegebene Wirkungsgrad einer PV-Anlage in der Regel unter Standardtestbedingungen (*engl. Standard Test Conditions, STC*) mit einer Bestrahlungsstärke von 1000 W/m^2 , einer Modultemperatur der PV-Anlage von 25°C und einer Luftmasse (*engl. Air Mass, AM*) von 1,5 gemessen angegeben wird [80, 86, 91]. Um den Einfluss der Modultemperatur der PV-Anlage auf die erzeugte Leistung berücksichtigen zu können, wird ein Korrekturfaktor für die Modultemperatur ϑ^t zum Zeitpunkt t eingeführt. Mit steigender Temperatur von PV-Modulen sinken sowohl die Leistung als auch der Wirkungsgrad [80, 91]. Im Folgenden wird die Berechnung der Globalstrahlung auf die geneigte Fläche der PV-Anlage $E_{\text{G,g}}^t$ sowie des Korrekturfaktors für die Modultemperatur ϑ^t näher erläutert. Der Wirkungsgrad und die Fläche können direkt dem Datenblatt der PV-Anlage entnommen werden.

$$P_{\text{PV}}^t = E_{\text{G,g}}^t \cdot \eta_{\text{PV}} \cdot A_{\text{PV}} \cdot \vartheta^t \quad (3.15)$$

Die Messwerte der Global- und Diffusstrahlung vom DWD beziehen sich auf eine horizontale Fläche. Da PV-Anlagen jedoch in der Regel eine Neigung aufweisen und in eine bestimmte Richtung ausgerichtet sind, muss die Bestrahlungsstärke auf einer horizontalen Fläche auf eine beliebig orientierte Fläche umgerechnet werden. Die Globalstrahlung auf eine geneigte Oberfläche $E_{G,g}^t$ zum Zeitpunkt t setzt sich nach Formel (3.16) aus der direkten Bestrahlungsstärke $E_{Dir,g}^t$ und der diffusen Bestrahlungsstärke $E_{Diff,g}^t$ auf eine geneigte Oberfläche sowie einem Teil, der vom Boden reflektiert wird $E_{R,g}^t$, zusammen [80, 86, 91].

$$E_{G,g}^t = E_{Dir,g}^t + E_{Diff,g}^t + E_{R,g}^t \quad (3.16)$$

Die direkte Bestrahlungsstärke auf eine beliebig orientierte PV-Anlage $E_{Dir,g}^t$ ergibt sich gemäß Formel (3.17) aus der direkten horizontalen Bestrahlungsstärke $E_{Dir,h}^t$ multipliziert mit einem geometrischen Faktor aus Sonnenhöhe γ_S^t und Einfallswinkel θ_g^t [80].

$$E_{Dir,g}^t = E_{Dir,h}^t \cdot \frac{\cos \theta_g^t}{\sin \gamma_S^t} \quad (3.17)$$

Die Sonnenhöhe wird entsprechend der Berechnungen in Kapitel 3.2.1 ermittelt. Der Einfallswinkel θ_g^t ist der Winkel zwischen der direkten Sonneneinstrahlung und der Flächennormalen der PV-Anlage [80, 88]. Die Berechnung erfolgt gemäß Formel (3.18) [80, 88]. Der Einfallswinkel ist abhängig vom Sonnenstand sowie vom Neigungswinkel γ_{PV} und der Ausrichtung der PV-Anlage. Die Ausrichtung der PV-Anlage wird durch den Azimutwinkel α_{PV} beschrieben. Dabei entspricht α_{PV} bei einer Südausrichtung 0° , bei einer Westausrichtung 90° und bei einer Ostausrichtung ist -90° . Der Sonnenstand wird durch die Sonnenhöhe γ_S^t und den Sonnenazimut α_S^t definiert. Steht die Sonne im Norden entspricht α_S^t 0° , im Osten 90° , im Süden 180° und im Westen 270° . Der Sonnenazimut wird ebenfalls entsprechend der Berechnungen in Kapitel 3.2.1 ermittelt.

$$\theta_g^t = \arccos(-\cos \gamma_S^t \cdot \sin \gamma_{PV} \cdot \cos(\alpha_S^t - \alpha_{PV}) + \sin \gamma_S^t \cdot \cos \gamma_{PV}) \quad (3.18)$$

Die diffuse Bestrahlungsstärke auf eine geneigte Fläche $E_{Diff,g}^t$ zum Zeitpunkt t wird gemäß der Formel (3.19) berechnet, wobei der Neigungswinkel der PV-Anlage γ_{PV} und die horizontale diffuse Bestrahlungsstärke $E_{Diff,h}^t$ berücksichtigt werden. Dabei wird ein isotroper Ansatz angenommen, was bedeutet, dass die Diffusstrahlung aus allen Himmelsrichtungen gleich ist und somit eine gleichmäßige Verteilung aufweist [80, 86, 88]. Dieser Ansatz stellt eine Näherung dar, da die Verteilung der Strahlung in Abhängigkeit von der Himmelsrichtung vor allem bei klarem Himmel variiert. Mit einem anisotropen Ansatz kann dies berücksichtigt werden [80]. Die Ergebnisse von Untersuchungen, die

im Rahmen dieser Arbeit durchgeführt wurden, zeigen jedoch, dass dies in keiner Verbesserung der Ergebnisse resultiert [S21].

$$E_{\text{Diff,g}}^t = E_{\text{Diff,h}}^t \cdot \frac{1}{2} \cdot (1 + \cos \gamma_{\text{PV}}) \quad (3.19)$$

Die Berechnung der reflektierten Bestrahlungsstärke auf eine geneigte Fläche $E_{\text{R,g}}^t$ aus der globalen horizontalen Bestrahlungsstärke $E_{\text{G,h}}^t$ gemäß Formel (3.20) erfordert die Berücksichtigung der Bodenreflexion. Diese wird über den Albedo-Koeffizienten ρ angegeben. Auch hier wird ein isotroper Ansatz verwendet [80, 86, 88].

$$E_{\text{R,g}}^t = E_{\text{G,h}}^t \cdot \rho \cdot \frac{1}{2} \cdot (1 - \cos \gamma_{\text{PV}}) \quad (3.20)$$

Wenn der Albedo am Standort nicht gemessen werden kann, kann dieser der Tabelle 3.6 für verschiedene Umgebungen entnommen werden.

Tabelle 3.6: Albedo ρ für verschiedene Umgebungen nach [80, 86, 88]

Untergrund	Albedo ρ	Untergrund	Albedo ρ
Gras (Juli, August)	0,25	Asphalt	0,09 - 0,18
Rasen	0,18 - 0,23	Wälder	0,05 - 0,18
Trockenes Gras	0,28 - 0,32	Heide- und Sandflächen	0,10 - 0,25
Nicht bestellte Felder	0,26	Wasserfläche ($\gamma_s > 45^\circ$)	0,05 - 0,08
Nackter Boden	0,17	Wasserfläche ($\gamma_s > 30^\circ$)	0,08
Schotter	0,18	Wasserfläche ($\gamma_s > 20^\circ$)	0,12
Beton, verwittert	0,20	Wasserfläche ($\gamma_s > 10^\circ$)	0,22
Beton, sauber	0,30	FrISChe Schneedecke, sauber	0,80 - 0,90
Zement, sauber	0,55	Alte Schneedecke	0,45 - 0,70
Sandboden, trocken	0,21 - 0,43	Erdboden, trocken	0,12 - 0,15
Sandboden, feucht	0,09	Erdboden, feucht	0,07 - 0,12
Dachziegel, rot	0,33	Holz	0,22
Dachpappe, schwarz	0,12 - 0,13	Kupfer	0,47
Schiefer	0,10 - 0,14	Stahl	0,80

Der Wirkungsgrad einer PV-Anlage ist von der Modultemperatur abhängig. Mit steigender Temperatur sinkt die erzeugte PV-Leistung. Über einen Korrekturfaktor für die Modultemperatur η^t gemäß Formel (3.21) kann die Leistung entsprechend der vorherrschenden Temperatur der PV-Anlage T_{PV}^t angepasst werden. Der Temperaturkoeffizient

der PV-Anlage β_{PV} gibt den Prozentsatz an, um den die Leistung basierend auf der Temperatur unter STC von 25 °C sinkt oder steigt [80, 86]. Bei einer Modultemperatur von 25 °C erfolgt eine Multiplikation der PV-Leistung mit dem Faktor 1, wodurch es keine Anpassung der erzeugten Leistung gibt. Bei einer höheren Modultemperatur als 25 °C wird der Korrekturfaktor kleiner als 1, was zu einer Verringerung der Leistung führt. Im umgekehrten Fall erfolgt eine Erhöhung des Korrekturfaktors, was eine Steigerung der Leistung zur Folge hat.

$$g^t = 1 + (25^\circ\text{C} - T_{PV}^t) \cdot \beta_{PV} \quad (3.21)$$

Die Temperatur des PV-Moduls T_{PV}^t ergibt sich nach Formel (3.22) aus der Umgebungstemperatur T_U zum Zeitpunkt t, der Globalstrahlung auf eine geneigte Oberfläche $E_{G,g}^t$ und einer Proportionalitätskonstanten σ für die Installationsart des PV-Moduls [80]. Die Art des Einbaus der PV-Anlage ist wegen der Hinterlüftung der PV-Anlage von entscheidender Bedeutung. Tabelle 3.7 gibt einen Überblick über verschiedene Einbauvarianten mit den zugehörigen Hinterlüftungsfaktoren.

$$T_{PV}^t = T_U + \frac{\sigma \cdot E_{G,g}^t}{1000 \frac{\text{W}}{\text{m}^2}} \quad (3.22)$$

Tabelle 3.7: Hinterlüftungsfaktor σ für die Berechnung der Modultemperatur für verschiedene Einbauvarianten von PV-Anlagen nach [80]

Art des Einbaus	σ in °C
völlig freie Aufständering	22
auf dem Dach, großer Abstand	28
auf dem Dach bzw. dachintegriert, gute Hinterlüftung	29
auf dem Dach bzw. dachintegriert, schlecht Hinterlüftung	32
an der Fassade bzw. fassadenintegriert, gute Hinterlüftung	35
an der Fassade bzw. fassadenintegriert, schlechte Hinterlüftung	39
Dachintegration, ohne Hinterlüftung	43
Fassadenintegration, ohne Hinterlüftung	55

3.2.3 Darstellung der PV-Anlage und Validierung des PV-Modells

Im Rahmen dieser Arbeit erfolgt eine Betrachtung einer PV-Anlage mit Standort in Bielefeld. Die technischen Daten der PV-Anlage sind in Tabelle 3.8 dargestellt. Die Anlage ist auf einer völlig freien Aufständerung auf dem Dach der Hochschule Bielefeld installiert. Das Institut für Technische Energie-Systeme der Hochschule Bielefeld stellt die Messdaten der erzeugten PV-Leistung sowie weitere Wettermessdaten wie die Temperatur und die Bestrahlungsstärke zur Verfügung, welche für die Validierung der PV-Prognose in dieser Arbeit verwendet werden.

Die PV-Anlage besteht aus acht Modulen mit einer Leistung von jeweils 280 Wp. Die technischen Daten bezüglich der Fläche, des Temperaturkoeffizienten der Leistung sowie der Nennleistung können dem Datenblatt im Anhang A aus der Abbildung A. 2 entnommen werden. Der Wirkungsgrad ist nicht im Datenblatt angegeben, kann jedoch nach Formel (3.23) aus der Leistung P_{STC} und der Bestrahlungsstärke E_{STC} unter STC sowie der Fläche der PV-Anlage berechnet werden [86].

$$\eta_{PV} = \frac{P_{STC}}{E_{STC} \cdot A_{PV}} \quad (3.23)$$

Die PV-Anlage weist eine Neigung von 30° sowie eine Ausrichtung nach Süden auf. Der Albedo-Wert für sauberen Beton kann Tabelle 3.6 und der Faktor für die Hinterlüftung für eine völlig freie Aufständerung kann Tabelle 3.7 entnommen werden.

Tabelle 3.8: Technische Daten der betrachteten PV-Anlage

Quelle	Parameter	Angabe
Einbau-/Standortangaben	Neigungswinkel	30°
	Ausrichtungswinkel	0° (Südausrichtung)
	Albedo	0,3 (Beton, sauber)
	Hinterlüftungsfaktor	22 °C (völlig freie Aufständerung)
	Längengrad	8,49°
	Breitengrad	52,04°
Datenblatt	Fläche	11,68 m ²
	Wirkungsgrad	19,18 %
	Temperaturkoeffizient Leistung	0,41 %/K
	Nennleistung	2,24 kWp

Für die Validierung des physikalischen PV-Modells wird der Fehler zwischen Messung und Berechnung der PV-Leistung der betrachteten PV-Anlage für einen Zeitraum von

mindestens einem Jahr bestimmt. Dadurch werden alle relevanten Wetterbedingungen und Sonnenstände innerhalb eines Jahres berücksichtigt. Dazu stehen im Zeitraum von 2019 bis 2022 eine ausreichende Anzahl an Temperaturmessdaten sowie Messungen der Global- und Diffusstrahlung an der PV-Anlage mit einer Auflösung von zehn Minuten zur Verfügung. Aus der am Standort der PV-Anlage gemessenen Bestrahlungsstärke wird mit den technischen Anlagenparametern und Standortinformationen die resultierende PV-Leistung nach den vorangegangenen Kapiteln 3.2.1 und 3.2.2 berechnet und mit der Messung zum entsprechenden Zeitpunkt verglichen. Die verwendeten Messdaten müssen vorab bereinigt werden, indem die Nachtstunden entfernt werden. Hierbei werden diejenigen Zeitpunkte aus dem Messdatensatz entfernt, bei denen die Sonnenhöhe kleiner Null ist. Auf diese Weise werden lediglich die Zeitpunkte in die Fehlerberechnung miteinbezogen, bei denen eine Erzeugung durch die PV-Anlage stattfindet. Fehlerhafte Werte werden entfernt. Damit ergeben sich für die Validierung 805 Tage im Zeitraum von 2019 bis 2022.

In Tabelle 3.9 sind die Fehlermetriken für die Validierung des PV-Modells dargestellt. Für den gesamten betrachteten Zeitraum beträgt der nMAE zwischen berechneter und gemessener Leistung 2,62 %. Dieser Fehler muss bei der PV-Leistungsprognose berücksichtigt werden, da sich der Fehler der Prognose aus dem Fehler des ANN für die Prognose der Bestrahlungsstärke und dem Fehler des physikalischen PV-Modells zusammensetzt.

Tabelle 3.9: Fehlermetriken der Validierung des physikalischen PV-Modells für verschiedene Zeiträume

Zeitraum	MAE in W	nMAE in %	RMSE in W	nRMSE in %
Gesamter	54,98	2,62	117,23	5,58
02.07.2019	31,11	1,48	45,31	2,16
27.12.2019	389,57	18,55	427,89	20,38

In Abbildung 3.7 sind die berechnete und die zugehörige gemessene PV-Leistung an zwei Beispieltagen dargestellt. Am ersten Beispieltag (02.07.2019) zeigt sich eine hohe Übereinstimmung zwischen berechneter und gemessener Leistung, was auf einen geringen Fehler hindeutet. Der Tag ist ein Sommertag im Juli, an dem die Bewölkung sehr wechselhaft war, was sich in den sprunghaften Leistungsschwankungen widerspiegelt. Insgesamt ist die erzeugte Leistung an diesem Tag hoch und erreicht einen Spitzenwert von über 2 kW. Der nMAE beträgt an diesem Tag 1,48 % (vgl. Tabelle 3.9) und liegt unterhalb der durchschnittlichen Abweichung im gesamten Zeitraum. Dieser Tag stellt einen Tag dar, an dem das PV-Modell sehr genau ist. Der zweite Tag (27.12.2019) ist

ein Wintertag im Dezember. Es zeigt sich eine hohe Abweichung zwischen der berechneten und der gemessenen PV-Leistung. Der nMAE beträgt an diesem Tag 18,55 % (vgl. Tabelle 3.9), was deutlich über dem nMAE für den gesamten Zeitraum liegt. Es sind folglich auch Tage zu verzeichnen, an denen das PV-Modell signifikante Fehler aufweist, was sich nachteilig auf die PV-Leistungsprognose auswirkt, da sich der Fehler der PV-Leistungsprognose aus dem Fehler des ANN und dem PV-Modell ergibt.

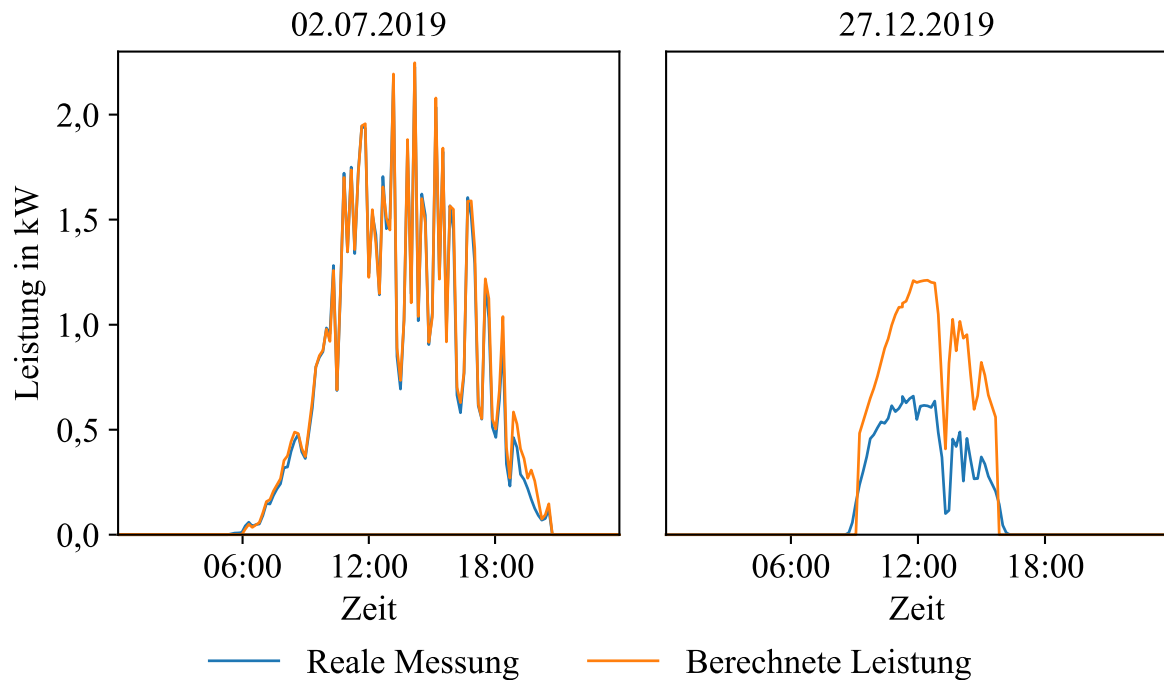


Abbildung 3.7: Berechnete und gemessene PV-Leistung an zwei Beispieltagen

Die berechnete Leistung liegt an diesem Tag deutlich über der gemessenen Leistung. Das PV-Modell bestimmt an diesem Tag demnach eine deutliche höhere Erzeugung für diese PV-Anlage. Aufgrund des niedrigeren Standes der Sonne im Winter deutet dies auf eine Verschattung der PV-Anlage hin. Dies wird auch in der Abbildung 3.8 deutlich, welche die gemessene und berechnete Leistung über eine gesamte Woche darstellt. Die Woche ist die Kalenderwoche 52 des Jahres 2019, in der ebenfalls der zweite Beispieltag aus Abbildung 3.7 dargestellt ist. Es wird ersichtlich, dass insbesondere am 26., 27., 28. und 29. Dezember eine signifikant höhere Leistung berechnet wird als tatsächlich gemessen wurde. Eine Berücksichtigung von Verschattung, wie sie bei dieser PV-Anlage im Winter auftritt, ist im PV-Modell nicht vorgesehen. Das Modell betrachtet keine komplexen Aspekte wie Alterung oder Verschattung.

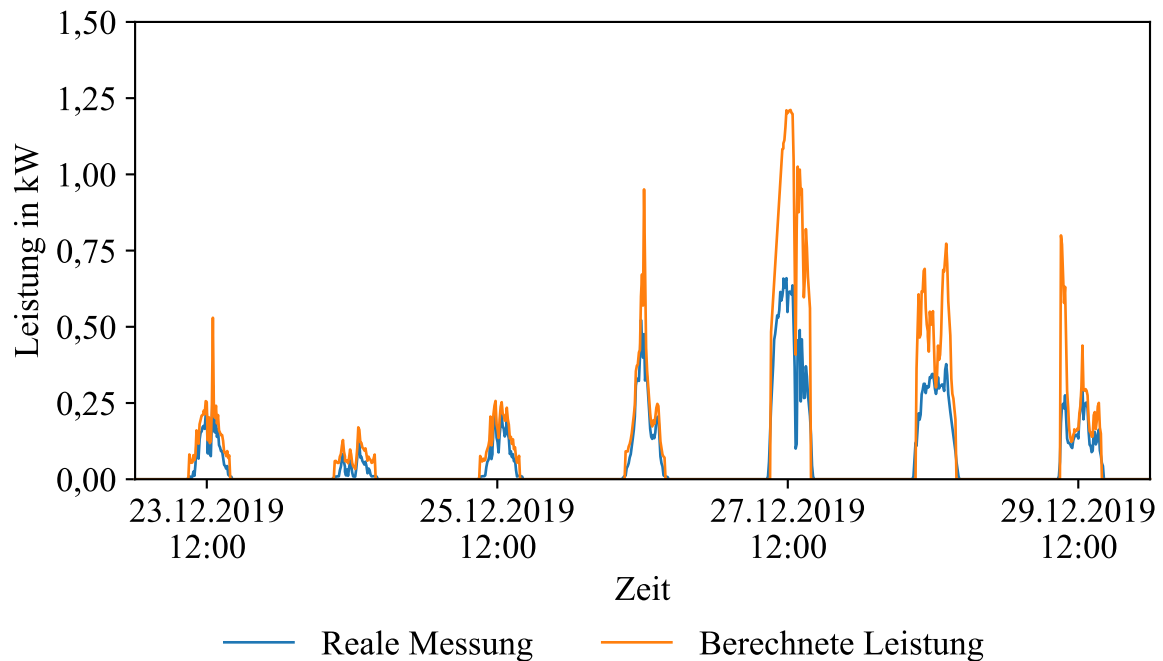


Abbildung 3.8: Berechnete und gemessene PV-Leistung in der Kalenderwoche 52 des Jahres 2019

In Abbildung 3.9 ist der nMAE des physikalischen PV-Modells für jeden Monat in einem Balkendiagramm dargestellt. Grundsätzlich ist zu erkennen, dass der Fehler im Winter höher ist als im Sommer, was auf die bereits erwähnte Verschattung zurückzuführen ist. Trotzdem variieren die Fehler zwischen den Monaten aufgrund unterschiedlicher Wetterbedingungen. Verschattungen treten besonders häufig im sonnigen Winter auf, was zu größeren Fehlern führt. Diese Fehler sind deutlich ausgeprägter als in einem Monat mit durchgehend bedecktem Himmel. Dadurch fallen die Fehler etwas unterschiedlich aus. Im Winter kann der Fehler wie zum Beispiel im Dezember 2020 ebenfalls geringer ausfallen. In diesem Monat war das Wetter überwiegend bedeckt, wodurch der diffuse Anteil der Bestrahlungsstärke höher war und die Verschattung weniger ins Gewicht fiel. Im Gegensatz dazu fällt bei sonnigem Wetter der direkte Anteil höher aus und hat einen größeren Effekt auf die Verschattung. Des Weiteren ist zu berücksichtigen, dass sich die Anzahl der Tage in den Monaten unterscheidet, was sich auf die Aussagekraft des Fehlers im jeweiligen Monat auswirkt. Die unterschiedliche Anzahl an Tagen in den Monaten resultiert daraus, dass einige Tage aufgrund fehlender oder fehlerhafter Daten vorab in der Datenaufbereitung entfernt wurden. Trotz dieser Einflussfaktoren lässt sich eine Tendenz der Fehler im Winter und im Sommer erkennen.

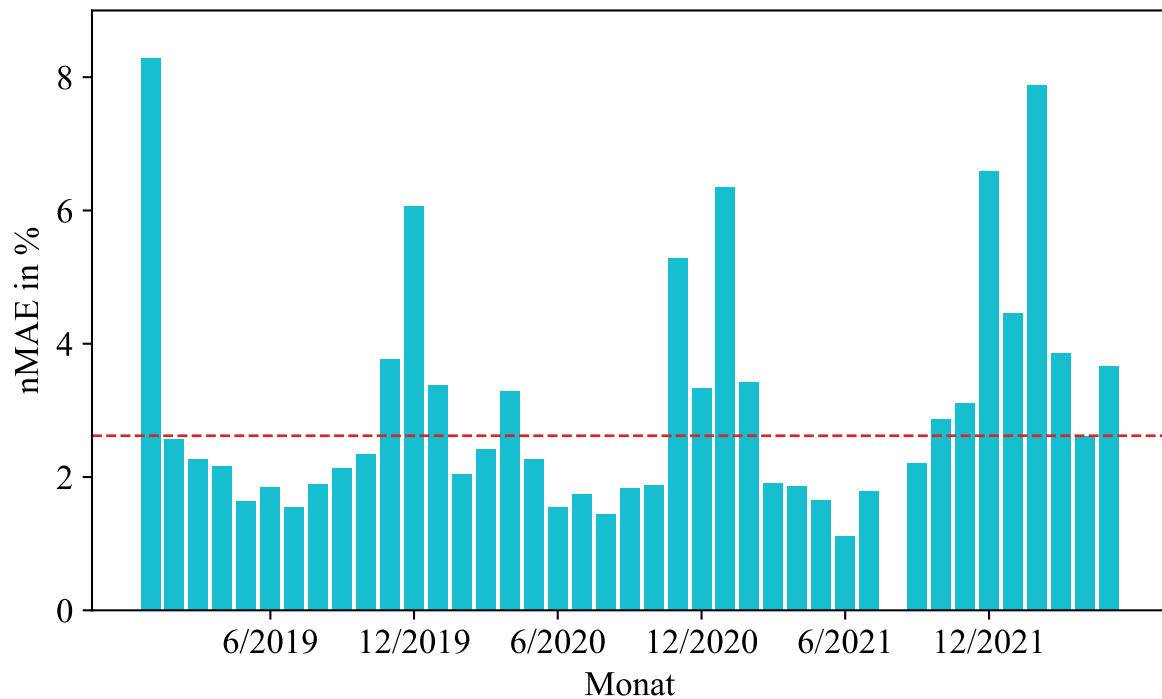


Abbildung 3.9: Darstellung des nMAE des physikalischen PV-Modells als Balkendiagramm pro Monat von 2019 bis 2022

3.3 Ergebnisse der Day-Ahead-Prognose einer PV-Anlage

In diesem Kapitel wird eine reale Vorhersage der erzeugten PV-Leistung unter Verwendung von Wettervorhersagedaten präsentiert. Dies unterscheidet sich von der Validierung eines ANN im Kapitel 3.1.2, bei der untersucht wird, wie gut das ANN anhand von Wetterdaten die Bestrahlungsstärke zum aktuellen Zeitpunkt abbilden kann.

Das PV-Leistungsprognosemodell wird anhand der in Kapitel 3.2.3 dargestellten PV-Anlage validiert. Die Bestrahlungsstärke am Standort der PV-Anlage der Hochschule Bielefeld wird mithilfe einer Wettervorhersage für die nächsten 24 Stunden prognostiziert. Das für die Prognose der Bestrahlungsstärke verwendete ANN ist dasjenige, welches beim Training mit der Eingangsdatenkombination 7 und zwei versteckten Schichten die besten Ergebnisse erzielt hat. Als Grundlage für die Wettervorhersage dient die kostenfreie Vorhersage des Anbieters *wetter.com GmbH*. Der Grund für die Wahl von dieser Wettervorhersage liegt in der guten Datenbasis, da alle benötigten Wettervorhersagedaten zur Eingabe für das ANN vorhanden sind, und der kostenfreien Verfügbarkeit. Der Datensatz für die 24-Stunden-Prognose umfasst eine stündliche Auflösung der Wettervorhersagedaten für den Zeitraum von 2019 bis 2023. Eine Aufbereitung des Datensatzes der Wettervorhersage ist erforderlich. Im Rahmen der Datenbereinigung werden Zeitpunkte, bei denen die Sonnenhöhe kleiner als null ist, somit ohne Erzeugung durch die PV-Anlage, aus dem Datensatz entfernt. Tage, an denen einzelne Datenpunkte

fehlen, werden komplett entfernt. Die Wettervorhersagedaten werden entsprechend der Trainings- und Validierungsdaten vor dem Einlesen in das ANN normiert. Die PV-Anlage wird mithilfe des physikalischen PV-Modells aus Kapitel 3.2 abgebildet, um aus der vorhergesagten Bestrahlungsstärke die erzeugte PV-Leistung zu berechnen.

Für den Vergleich von Prognose und gemessener Leistung wird der gleiche Datensatz der PV-Messdaten verwendet, der auch bei der Validierung des PV-Modells genutzt wird. Die 10-Minuten-Werte der gemessenen PV-Leistung werden auf eine stündliche Auflösung gemittelt, da auch die Wettervorhersage eine stündliche Auflösung aufweist. Insgesamt ergeben sich somit 629 Tage im Zeitraum von 2019 bis 2022, die für den Vergleich zwischen vorhergesagter und gemessener PV-Leistung der betrachteten PV-Anlage herangezogen werden.

Die Fehlermetriken der Validierung der PV-Prognose sind anhand der PV-Anlage für den gesamten betrachteten Zeitraum in Tabelle 3.10 dargestellt. Bei einer Prognose der PV-Leistung für die nächsten 24 Stunden beträgt der nMAE 9,33 %. Eine stündliche Aktualisierung der PV-Prognose, bei der das ANN eine stündlich aktualisierte Wettervorhersage erhält, führt zu einer Reduzierung des Fehlers auf einen nMAE von 9,01 %. Dadurch kann der Fehler in der verwendeten NWP reduziert werden. Der nRMSE beträgt für eine Day-Ahead-Prognose 12,25 % und für eine stündlich aktualisierte Prognose 11,76 %. Es zeigt sich, dass weiterhin Ausreißer unter den Fehlern vorhanden sind. Durch eine stündliche Aktualisierung ist die Differenz zwischen dem nMAE und dem nRMSE leicht gesunken, was auf eine Reduktion der Fehlerausreißer durch die Aktualisierung hinweist.

Tabelle 3.10: Fehlermetriken der Validierung der PV-Prognose anhand einer realen PV-Anlage

Prognosehorizont Leistung		MAE in W	nMAE in %	RMSE in W	nRMSE in %
Leistung	24 h Prognose	195,85	9,33	257,27	12,25
	1 h Prognose	189,20	9,01	246,95	11,76
Prognosehorizont Bestrahlungs- stärke		MAE in W/m²	nMAE in %	RMSE in W/m²	nRMSE in %
Globalstrahlung	24 h Prognose	84,70	9,19	106,37	11,54
	1 h Prognose	82,50	8,95	102,88	11,16
Diffusstrahlung	24 h Prognose	54,80	8,05	70,09	10,29
	1 h Prognose	54,51	8,01	69,60	10,22

Der nMAE der Globalstrahlung beträgt 9,19 % und der Diffusstrahlung 8,05 % für eine 24-Stunden-Prognose bevor die Bestrahlungsstärke im PV-Modell in die entsprechende erzeugte Leistung umgerechnet wird. Auch diese Fehler können durch eine stündliche Aktualisierung leicht reduziert werden. Hier wird nochmal deutlich, dass sich der Fehler des physikalischen PV-Modells auf den Prognosefehler auswirkt, da der Fehler der Prognose der PV-Leistung etwas höher ist als der der Bestrahlungsstärke.

In Abbildung 3.10 sind die gemessene und prognostizierte PV-Leistung für einen Prognosehorizont von 24 Stunden sowie einer Stunde an zwei Beispieltagen abgebildet. Am 27.07.2019 zeigt sich eine hohe Übereinstimmung zwischen Prognose und Messung der PV-Leistung. Eine noch exaktere Prognose kann durch eine stündliche Aktualisierung der Vorhersage erzielt werden. Der nMAE für eine 24-Stunden-Prognose beträgt 3,11 % während er für eine Prognose von einer Stunde bei 3,04 % liegt (siehe Tabelle 3.11). Der Fehler liegt somit deutlich unter dem Durchschnitt. Am 20.09.2019 zeigt sich eine Abweichung zwischen Prognose und Messung der PV-Leistung, die durch eine stündliche Aktualisierung der Prognose jedoch signifikant reduziert werden kann. Die 24-Stunden-Prognose weicht von den gemessenen Werten ab, da sie deutlich niedrigere Werte vorhersagt. Durch eine stündliche Aktualisierung wird die NWP angepasst. Folglich ergibt die Prognose analog zu den Messungen nun ebenfalls eine höhere erzeugte Leistung. Der nMAE beträgt 15,45 % für eine 24-Stunden-Prognose und bei einer stündlichen Aktualisierung beträgt der nMAE 10,55 % (siehe Tabelle 3.11), welcher damit noch etwas über dem Durchschnitt liegt. Durch eine stündliche Aktualisierung können demnach Fehler aus der verwendeten NWP reduziert werden.

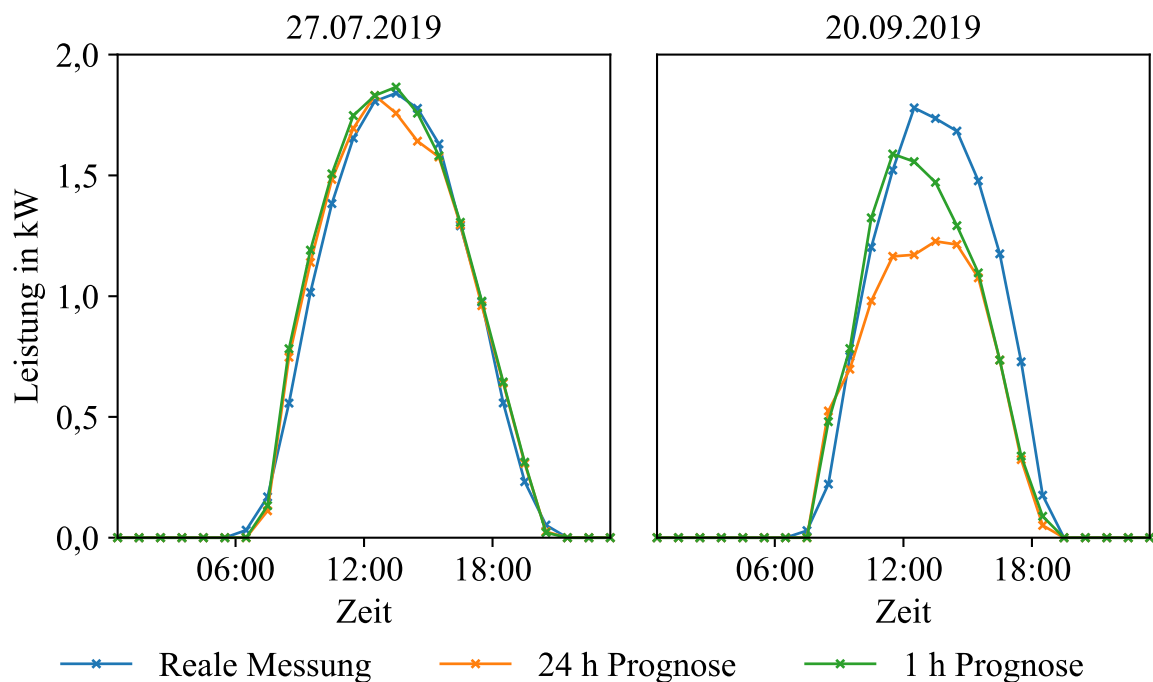


Abbildung 3.10: Gemessene und prognostizierte PV-Leistung an zwei Beispieltagen

Tabelle 3.11: Fehlermetriken der Validierung der PV-Prognose anhand einer realen PV-Anlage für verschiedene Beispieltage

Datum	Prognosehorizont	MAE in W	nMAE in %	RMSE in W	nRMSE in %
27.07.2019	24 h Prognose	65,29	3,11	82,81	3,94
	1 h Prognose	63,74	3,04	89,07	4,24
20.09.2019	24 h Prognose	324,37	15,45	371,36	17,68
	1 h Prognose	221,54	10,55	267,60	12,74
27.12.2019	24 h Prognose	362,75	17,27	430,47	20,50
	1 h Prognose	308,83	14,71	374,51	17,83

In der folgenden Abbildung 3.11 sind die gemessenen sowie prognostizierten Werte der PV-Leistung für einen Prognosehorizont von 24 Stunden sowie einer Stunde in der Kalenderwoche 52 des Jahres 2019 dargestellt. Dies entspricht der gleichen Woche, die auch bei der Validierung des PV-Modells in Kapitel 3.2.3 betrachtet wird. In dieser Woche wird demonstriert, dass die PV-Anlage im Winter verschattet ist und dadurch eine geringere Leistung erzeugt. Dieses Verhalten zeigt sich auch hier in der zweiten Wochenhälfte, in der die Prognose die gemessene PV-Leistung deutlich übersteigt. Das PV-Prognosemodell geht von einer nicht verschatteten PV-Anlage aus und prognostiziert anhand der vorhergesagten Bestrahlungsstärke die zu erwartende PV-Leistung. Diese müsste jedoch aufgrund der Verschattung im Winter geringer ausfallen. Am 27.12.2019 beträgt der nMAE für eine 24-Stunden-Prognose 17,27 % (siehe Tabelle 3.11).

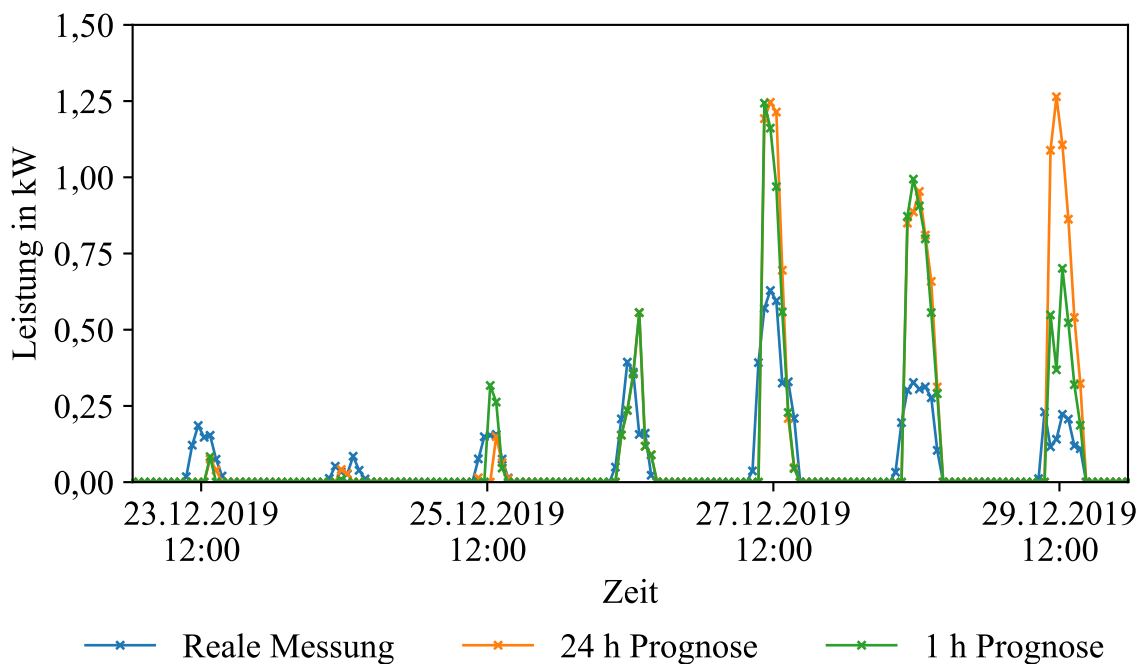


Abbildung 3.11: Gemessene und prognostizierte PV-Leistung in der Kalenderwoche 52 des Jahres 2019

Zur Veranschaulichung des Prognosefehlers bei verschiedenen Wetterbedingungen sind in Tabelle 3.12 die Fehlermetriken für verschiedene Wertebereiche des Clear-Sky Index (CSI) für eine 24-Stunden-Prognose dargestellt. Der CSI gibt an, wie groß die aktuelle Bestrahlungsstärke im Vergleich zur theoretisch möglichen Bestrahlungsstärke bei klarem Himmel ist [33, 68, 92]. Wenn der $CSI = 1$ ist, ist die aktuelle Bestrahlungsstärke gleich der theoretisch möglichen Bestrahlungsstärke bei klarem Himmel. Ist der $CSI < 1$, dann ist die Bestrahlungsstärke geringer als die theoretisch mögliche bei klarem Himmel, was auf eine Bewölkung des Himmels hinweist. Der CSI wird nach [92] aus der Globalstrahlung und Orts- und Zeitangaben wie dem Längengrad und Breitengrad sowie der Uhrzeit berechnet. In vielen wissenschaftlichen Arbeiten werden Prognosefehler auch anhand von CSI-Werte klassifiziert. In den meisten Fällen ist die Vorhersage an Tagen mit sehr hohen CSI-Werten, d. h. nahe bei 1, am besten, wie z. B. in [33], da an diesen Tagen der Himmel klar und nicht mit Wolken bedeckt ist und es keine Änderung der Wetterbedingungen aufgrund von Bewölkung gibt.

In dieser Arbeit wird der nMAE anhand des CSI in drei Bereiche unterteilt. Der CSI wird dabei in Intervalle von 0 bis 1 aufgeteilt. Der erste Bereich, mit einem CSI von 0,66 bis 1, wird als sonnig bezeichnet, da hier die höchsten CSI-Werte vorliegen. Der zweite Bereich, mit einem CSI von 0,33 bis 0,66, wird als wolkig definiert, da er in der Mitte liegt und wechselnde Bewölkung aufweist. Der dritte Bereich, mit einem CSI von 0 bis 0,33, wird als bedeckt klassifiziert, da hier die niedrigsten CSI-Werte vorhanden sind. Die Tage werden basierend auf ihrem mittleren täglichen CSI in diese drei Kategorien eingeteilt und der nMAE für jede Kategorie berechnet. Der niedrigste nMAE von 7,34 % (siehe Tabelle 3.12) tritt an bedeckten Tagen auf, da es hier weniger wechselnde Wetterbedingungen und fast keine Verschattung der PV-Anlage gibt, da der Himmel bedeckt ist und der diffuse Anteil der Sonneneinstrahlung am höchsten ist. Der höchste nMAE wird mit 10,89 % (siehe Tabelle 3.12) an wolkigen Tagen verzeichnet, was auf die ständig wechselnden Wetterbedingungen durch die Wolkenbewegungen zurückzuführen ist, die die PV-Leistungsvorhersage erschweren. An sonnigen Tagen beträgt der nMAE 10,52 % (siehe Tabelle 3.12), was der niedrigste nMAE-Wert aller Bereiche ist. Er unterscheidet sich jedoch nur geringfügig von den wolkigen Tagen. In vielen wissenschaftlichen Arbeiten wird gezeigt, dass an sonnigen Tagen der Prognosefehler am geringsten und die Prognose am einfachsten ist. In dieser Arbeit ist dagegen nur ein geringfügiger Unterschied zwischen sonnigen und wolkigen Tagen zu erkennen. Der Grund für diesen Unterschied liegt darin, dass bei einem hohen Anteil direkter Sonneneinstrahlung, wie an einem sonnigen Tag, eine Verschattung der PV-Anlage stattfindet, die in der Prognose nicht berücksichtigt wird. Dies führt zu einem höheren Prognosefehler an sonnigen Tagen, die in dieser Arbeit im Winter auftreten. Des Weiteren impliziert ein CSI von 0,66 bis 1 nicht zwangsläufig einen vollständig wolkenfreien Himmel, sondern kann auch eine gewisse Bewölkung umfassen.

Die PV-Prognose, welche ausschließlich Wetterdaten sowie physikalische Daten der PV-Anlage nutzt, ist nicht in der Lage, Prognosefehler, die durch Verschattung entstehen, zu berücksichtigen. Daher ist die Implementierung eines zusätzlichen Korrekturmodells zur Reduzierung von Prognosefehlern erforderlich. Da die Prognosefehler in hohem Maße individuell und standortspezifisch sind, können sie nicht aus Wetterdaten gelernt werden. Leistungsmessdaten einer PV-Anlage ermöglichen jedoch eine Korrektur solcher Prognosefehler. Eine detaillierte Untersuchung dieser Thematik erfolgt in Kapitel 4.

Tabelle 3.12: Fehlermetriken der Validierung der PV-Prognose anhand einer realen PV-Anlage für verschiedene CSI-Wertebereiche für eine 24 Stunden Prognose

CSI-Bereiche	MAE in W	nMAE in %	RMSE in W	nRMSE in %
1 - 0,66 (sonnig)	220,89	10,52	280,42	13,35
0,66 - 0,33 (wolkig)	228,78	10,89	299,07	14,24
0,33 – 0 (bedeckt)	154,14	7,34	205,94	9,81

In einer Vielzahl wissenschaftlicher Arbeiten erfolgt ein Vergleich der dargestellten Methoden mit der Persistenzmethode, um das Potenzial der in der Arbeit entwickelten Methoden aufzuzeigen. Das zugrundeliegende Modell ist relativ einfach und besagt, dass sich der zu prognostizierende Tag wie der Tag davor verhält. Das bedeutet, dass die PV-Leistung des zu prognostizierenden Tages den Messdaten der Leistung am Tag zuvor entspricht. Darauf wird in dieser Arbeit verzichtet, da bereits in einer Reihe von anderen wissenschaftlichen Arbeiten, wie in [10, 36, 72, 76, 79], aufgezeigt wird, dass ML-Methoden für die PV-Leistungsprognose besser geeignet sind.

3.4 Zusammenfassung

In Kapitel 3 wird zunächst die Funktionsweise des PV-Prognosemodells im Betrieb vorgestellt und anschließend die Methodik der einzelnen Elemente des Modells erläutert. Zudem werden die Ergebnisse der PV-Prognose für eine reale PV-Anlage dargelegt und diskutiert. Die PV-Prognose besteht aus einer Prognose der Bestrahlungsstärke und einem physikalischen PV-Modell, das die prognostizierte Bestrahlungsstärke in die erzeugte Leistung umrechnet. Aus den Wettervorhersagedaten einer NWP als Eingangsdaten wird mittels eines trainierten ANN die Bestrahlungsstärke vorhergesagt. Die für den Standort der PV-Anlage prognostizierte Bestrahlungsstärke wird im physikalischen PV-Modell in die von der PV-Anlage erzeugte Leistung umgerechnet. Da die Bestrahlungsstärke für einen ausgewählten Standort prognostiziert wird, kann das PV-Leistungsprognosemodell auf alle PV-Anlagen an diesem Standort angewendet werden und

ist nicht auf eine einzelne Anlage beschränkt. In Abschnitt 3.1 wird die Methode für die Erstellung der Prognose der Bestrahlungsstärke erläutert. Die Prognose der Bestrahlungsstärke basiert auf einem ANN, welches mit historischen Wetterdaten trainiert wird. Das ANN stellt beim Training eine Beziehung zwischen den Wetterdaten, die zu einem bestimmten Zeitpunkt vorliegen, und der Bestrahlungsstärke zu diesem Zeitpunkt her. Im Rahmen des Trainings werden verschiedene Kombinationen von Wetterdaten für die Eingangsdaten untersucht, um die Kombination mit dem geringsten Fehler bei der Prognose der Bestrahlungsstärke zu ermitteln. Die Optimierung der Hyperparameter des ANN erfolgt mittels der Bayesschen Optimierung, welche eine schrittweise Annäherung an die optimalen Hyperparameter ermöglicht. Der Fehler des ANN, das mit allen verfügbaren Wetterdaten sowie dem Stand der Sonne als Eingangsdaten trainiert wird, ist am geringsten und der nMAE der Testdaten beträgt 5,91 %. In Abschnitt 3.2 wird das physikalische PV-Modell präsentiert, welches die PV-Leistung aus der vom ANN vorhergesagten Bestrahlungsstärke auf Basis von technischen Anlagenparameter und Standortgegebenheiten der entsprechenden PV-Anlage über mathematische Beziehungen berechnet. Die Validierung der PV-Prognose erfolgt anhand einer PV-Anlage mit einer Nennleistung von 2,24 kWp, die sich auf dem Dach der Hochschule Bielefeld befindet. Für die reine Berechnung der PV-Leistung aus der Bestrahlungsstärke im Rahmen des PV-Modells beträgt der nMAE 2,62 %. Die Ergebnisse für die PV-Prognose der PV-Anlage werden in Abschnitt 3.3 präsentiert. Bei einer Day-Ahead-Prognose der PV-Leistung beträgt der nMAE 9,33 % für einen Testzeitraum von 629 Tagen im Zeitraum von 2019 bis 2022. Eine stündliche Aktualisierung der Prognose, bei der die Wettervorhersage stündlich neu in das ANN eingegeben wird, führt zu einer leichten Reduzierung des Fehlers. Die Ergebnisse zeigen, dass die PV-Anlage im Winter verschattet ist. In einer Beispielwoche im Winter ist die Prognose deutlich höher als die gemessene PV-Leistung. Das Prognosemodell berücksichtigt nur technische Parameter und geht von einer nicht verschatteten Anlage aus. Verschattungseffekte bleiben daher unberücksichtigt. Die Einbindung von Leistungsmessdaten der PV-Anlage in Kapitel 4 ermöglicht es, solche Faktoren zu erkennen und die Prognose zu verbessern.

4 Korrekturmodell der PV-Prognose

Das Korrekturmodell basiert auf einem Stacking-Ensemble von verschiedenen ML-Modellen. Die prognostizierte Leistung eines Tages einer PV-Anlage sowie weitere zeitliche Angaben werden als Eingabe für das Korrekturmodell verwendet. Durch die zeitlichen Angaben werden sowohl saisonale als auch tägliche Trends des entsprechenden Tages berücksichtigt. Als Ausgabe erstellt das Korrekturmodell eine korrigierte prognostizierte PV-Leistung für den jeweiligen Tag. Die Funktionsweise des Korrekturmodells im Betrieb ist in Abbildung 4.1 dargestellt.

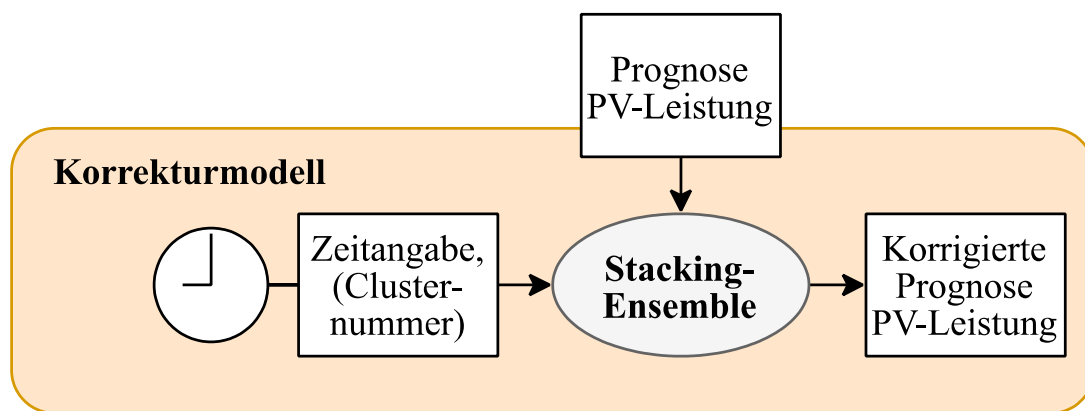


Abbildung 4.1: Übersicht über die Funktionsweise des Korrekturmodells im Betrieb

Das Ziel besteht in der Reduktion von wiederkehrenden Prognosefehlern, welche auf die spezifischen Gegebenheiten der jeweiligen PV-Anlage zurückzuführen sind. Solche Fehler können beispielsweise durch eine Verschattung oder durch eine fortgeschrittene Alterung der PV-Anlage bedingt sein. Diese Gegebenheiten, die je nach PV-Anlage sehr individuell sind, können nicht über eine PV-Prognose abgebildet werden, die nur Wetterdaten für die Vorhersage nutzt. Daher ist es erforderlich, auch die Leistungsmessdaten der jeweiligen PV-Anlage zu berücksichtigen. Obwohl eine Betrachtung derartiger Gegebenheiten in einem mathematischen PV-Modell zur Umrechnung der Bestrahlungsstärke in die entsprechende erzeugte Leistung prinzipiell möglich ist, erweist sich beispielsweise die Ermittlung einer Verschattung einer PV-Anlage als komplex.

Im Folgenden werden die Methodik sowie die Ergebnisse des Korrekturmodells für PV-Prognosen erläutert. Im Rahmen der Methodik erfolgt eine Darlegung der Aufbereitung und Aufteilung der Daten sowie der Auswahl relevanter Merkmale für den Trainingsprozess. Die Trainingsphase des Korrekturmodells basiert auf Prognosedaten und den dazugehörigen Messwerten der PV-Leistung, um die Zuordnung der jeweiligen Vorhersage zu den entsprechenden Messdaten zu lernen. Die Methodik umfasst zudem die Auswahl der Basismodelle für das Stacking-Ensemble, die Festlegung der optimalen

Hyperparameter sowie die Erstellung eines Metamodells für die Kombination der Vorhersagen aller Basismodelle zu einer finalen Vorhersage. Im Anschluss erfolgt die Ergebnispräsentation, in der die Vorhersagen mit und ohne Korrektur miteinander verglichen werden, um zu analysieren, ob die Korrektur eine Reduktion der Prognosefehler bewirkt hat. Des Weiteren erfolgt eine Ergebnisdarstellung in Bezug auf die Trainingsdatenmenge, um aufzuzeigen, wie viele Daten erforderlich sind, bis eine Verbesserung der Prognosegenauigkeit erreicht wird.

4.1 Methode des Stacking-Ensembles

Die Methode des Korrekturmodells basiert auf einem Stacking-Ensemble von verschiedenen ML-Modellen. Dies sind die Basismodelle des Stacking-Ensembles. Die Basismodelle zeigen jeweils unterschiedliche Stärken und Schwächen bezüglich eines bestimmten Datensatzes. Die zugrundeliegende Idee ist, dass die verschiedenen Modelle in Abhängigkeit von den Wetterbedingungen oder den Daten im Datensatz unterschiedliche Verhaltensweisen aufweisen. Im Rahmen der Modellkombination wird ein Metamodell erstellt, welches eine sinnvolle Kombination der Modelle darstellt und dadurch eine optimierte Prognose ermöglicht.

In Abbildung 4.2 ist eine Übersicht über die Methodik für das Korrekturmodell dargestellt. Als Basismodelle des Stacking-Ensembles finden die Modelle auf Basis der Methoden KNN, RF, Gradient Boosting, XGBoost und Catboost Regressor Anwendung. Die Auswahl dieser Modelle wird im weiteren Verlauf dieses Kapitels näher erläutert. Die genannten Modelle werden jeweils einzeln mit der Prognose der PV-Leistung sowie zeitlichen Angaben als Eingangsdaten und der zugehörigen Messung der PV-Leistung trainiert. Im Anschluss geben sie am Ausgang parallel Prognosen für die korrigierte PV-Prognose aus. Die finale Prognose wird durch das Metamodell unter Zuhilfenahme der LR-Methode generiert. Zu diesem Zweck wird das Metamodell mit den Prognosen der Basismodelle als Eingangsdaten sowie den realen Messdaten der PV-Leistung als Ziel-
daten trainiert. Auf diese Weise lernt das Metamodell, wie die Vorhersagen der Basismodelle kombiniert werden müssen, um die Prognosefehler zur realen Messung der PV-Leistung zu reduzieren. Die Stärken der einzelnen Basismodelle werden optimal miteinander kombiniert.

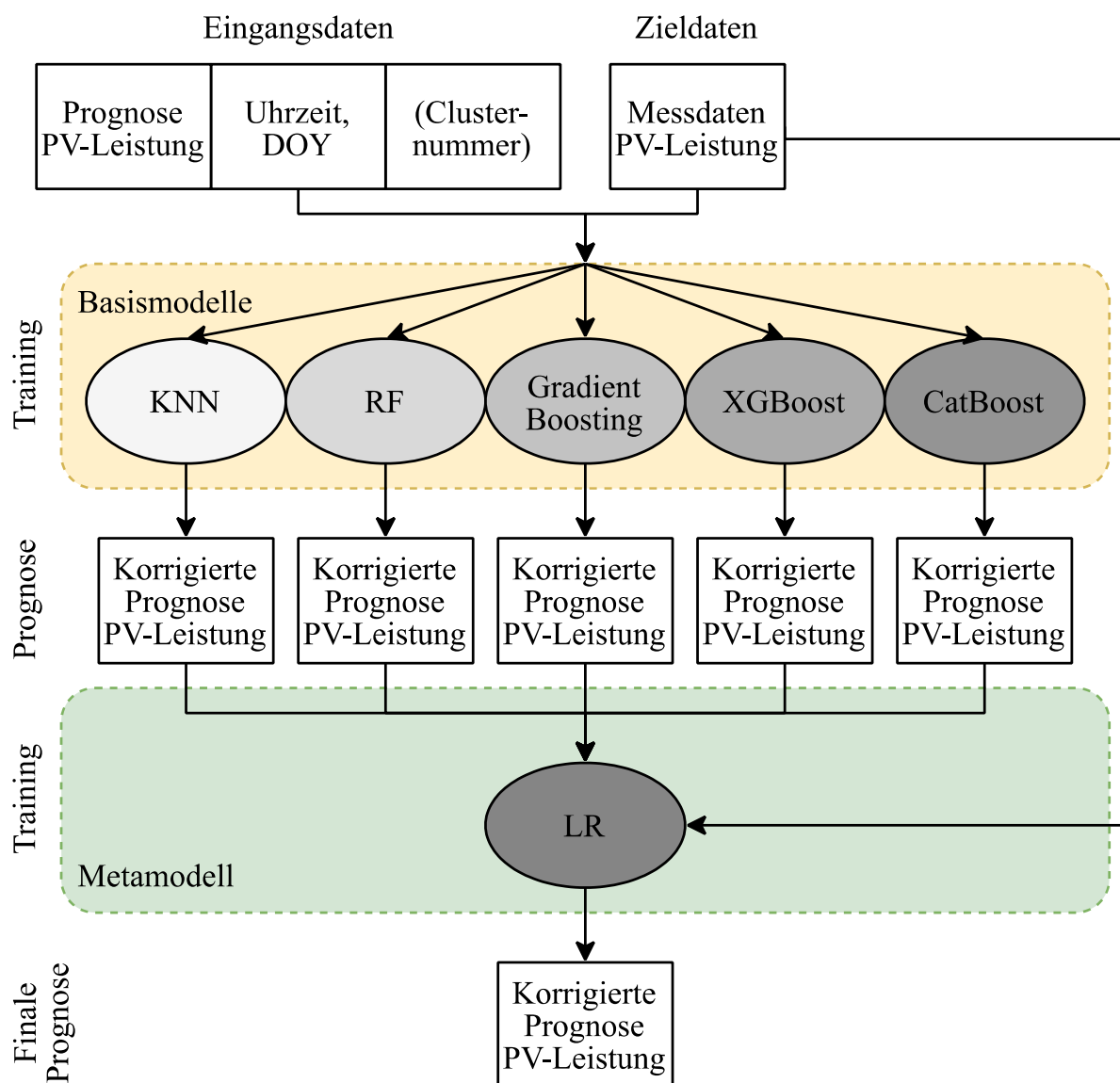


Abbildung 4.2: Darstellung der Methode des Korrekturmodells mittels eines Stacking-Ensembles

Im Folgenden erfolgt eine detaillierte Erläuterung der zum Training und der Validierung der Modelle im Stacking-Ensemble verwendeten Daten. Zudem wird die Auswahl der Modelle sowie die Auswahl der Hyperparameter im Rahmen des Trainings dargelegt. Für die Entwicklung des Stacking-Ensemble für die Korrektur der PV-Prognose findet die Programmiersprache Python Anwendung. Dabei werden die ML-Bibliotheken Sci-kit-learn, XGBoost und CatBoost genutzt.

4.1.1 Datenaufbereitung und -aufteilung

Die für die Korrektur der PV-Prognose herangezogenen Daten entsprechen denen, die in Kapitel 3 beschrieben sind. Die vorliegenden Daten stellen einerseits die Leistungsmessung der PV-Anlage auf dem Dach der Hochschule Bielefeld dar. Die Daten sind im gleichen Format wie in Kapitel 3 beschrieben und auf die gleiche Weise aufbereitet.

Dies beinhaltet die Entfernung der Nachtstunden, eine Datennormalisierung für die ML-Methoden, für die es sinnvoll ist wie die KNN-Methode und die Ridge Regression sowie eine Anpassung der zeitlichen Auflösung auf eine stündliche Auflösung. Des Weiteren sind PV-Leistungsprognosedaten für den gleichen Zeitraum verfügbar, welche ebenfalls eine stündliche Auflösung aufweisen. Die Messdaten der PV-Leistung dienen als Ausgangs- bzw. Zieldaten, während die entsprechenden Prognosedaten die Eingabedaten für das Modell sind. Das Ziel besteht in der Herstellung einer Beziehung zwischen den Eingangs- und Zieldaten, um die zugehörige reale Messung für jede Prognose zu ermitteln. Auf diese Weise soll eine Verringerung des Fehlers des ersten Prognosemodells erreicht werden. Die Nutzung von zusätzlichen Merkmalen als Eingabedaten, welche das zeitliche Verhalten der Daten widerspiegeln, können vorteilhaft sein. Die Erzeugung einer PV-Anlage unterliegt im Jahresverlauf Schwankungen, die auf saisonale Unterschiede zurückzuführen sind. In Deutschland zeigt die Sonne im Sommer eine frühere Auf- und spätere Untergangszeit, während die Tage im Winter kürzer sind. Des Weiteren lässt sich ein Unterschied im zeitlichen Auftreten von Erzeugungsspitzen zwischen den Jahreszeiten beobachten. Das Lernen solcher temporalen Muster kann dazu beitragen, die Prognosequalität des Modells zu verbessern. Aus diesem Grund erfolgt eine Auswahl zusätzlicher Merkmale, darunter der Tag des Jahres, der DOY, und die Uhrzeit. Der DOY repräsentiert saisonale Trends, wobei beispielsweise der 1. Januar als $\text{DOY} = 1$ definiert ist. Die Uhrzeit zeigt tägliche Trends an, wobei beispielsweise 14 Uhr der Zahl 14 entspricht. Die Berücksichtigung beider Merkmale ist von entscheidender Bedeutung, da sie in ihrer Gesamtheit die zeitliche Struktur der Daten beschreiben. Des Weiteren kann es hilfreich sein, weitere Merkmale wie beispielsweise meteorologische Parameter wie Windrichtung und Wolkenbedeckung als zusätzliche Merkmale zu berücksichtigen. Allerdings stehen diese Daten in der Regel lediglich als Wettervorhersage zur Verfügung und sind bereits in der ursprünglichen PV-Prognose berücksichtigt.

Der Datensatz umfasst den Zeitraum von Januar 2019 bis April 2023 und beinhaltet insgesamt 18.504 Datenpunkte, die 771 Tagen entsprechen. Hierbei ist zu beachten, dass die Nachtstunden in dieser Angabe miteinbezogen sind. Die Nachtstunden werden aus dem Datensatz für das Training und die Validierung entfernt. Für das Training und die anschließende Bewertung des trainierten Modells erfolgt eine Aufteilung im Verhältnis von 70 % zu 30 % in Trainings- und Testdaten, um eine hinreichend große Menge an Testdaten zu gewährleisten. Die Trainingsdaten umfassen zufällig ausgewählte 540 Tage, was eine ausreichend große Datenmenge darstellt, um Muster zu erkennen. Die Anzahl von 231 Testtagen ist ausreichend groß und zufällig verteilt, sodass eine Prüfung des trainierten Modells auf unterschiedliche Tage gewährleistet ist. Die zufällige Verteilung der Trainings- und Testtage ist von entscheidender Bedeutung, um zu gewährleisten, dass das Modell nicht lediglich auf einen spezifischen Zeitraum, beispielsweise

den Sommer, trainiert wird. Eine solche Einschränkung könnte dazu führen, dass das Modell bei neuen Daten schlechtere Leistungen zeigt.

Die Daten für das Training des Stacking-Ensembles für die Korrektur der PV-Prognose liegen entsprechend der Abbildung 4.3 in einer Matrix vor. Die Eingangsdaten bestehen aus einer Matrix mit drei Spalten, welche die prognostizierte Leistung $P_{\text{Prog},i}$ am Tag des Jahres DOY_i und die Uhrzeit T_i sowie die Werte für den DOY_i und die Uhrzeit T_i selbst enthalten. Die Zieldaten werden in einem Array mit einer Spalte bereitgestellt. Dieses Array enthält die zugehörige gemessenen Leistung $P_{\text{Mess},i}$ in Abhängigkeit vom DOY_i und der Uhrzeit T_i . Der DOY_i kann Werte zwischen 1 und 365 bzw. 366 in einem Schaltjahr annehmen. Die Uhrzeit T_i , angegeben in ganzen Stunden, ist eine Ganzzahl zwischen 0 und 23. Bestimmte Werte für T_i werden nicht vorkommen, da die Nachtstunden in dem Datensatz entfernt sind. Alle Datenpunkte über mehrere Tage des Zeitraums werden vertikal in der Matrix (für die Eingangsdaten) bzw. im Array (für die Zieldaten) angeordnet. Die Zuordnung der Datenpunkte zu einem bestimmten Tag erfolgt über den Wert des DOY. Die maximale Anzahl an Datenpunkten ist durch n definiert und der Index i liegt im Bereich $[1, 2, \dots, n]$.

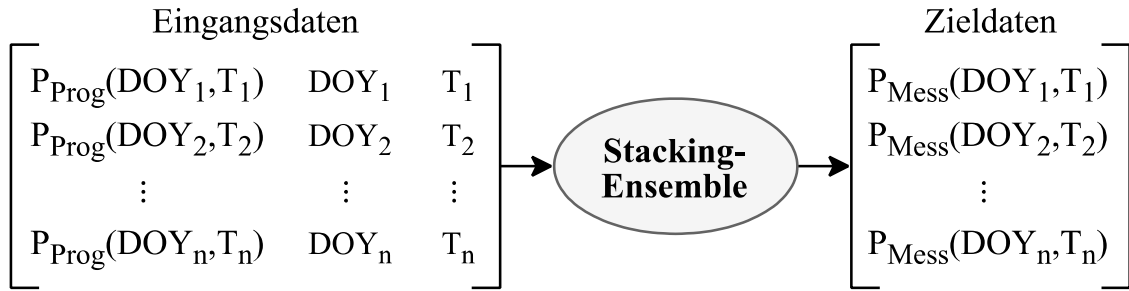


Abbildung 4.3: Eingangs- und Zieldaten für das Training eines Stacking-Ensembles

Neben den Merkmalen, die direkt aus dem Datensatz abgeleitet werden können, besteht die Möglichkeit, neue Merkmale zu erstellen, um weitere in den Daten verborgene Trends zu erfassen. In anderen wissenschaftlichen Arbeiten wie in [64, 79] sowie in einer im Rahmen dieser Dissertation wissenschaftlichen Veröffentlichung in [B10] wird gezeigt, dass die PV-Prognose durch den Einbezug von geclusterten Daten optimiert werden kann. Die Gewinnung von Merkmalen, die ähnliche Tage beinhalten, kann durch Methoden wie Clustering erfolgen. Die Idee besteht darin, ähnliche Tagesverläufe zu clustern, um Tage zu identifizieren, die aufgrund ähnlicher Wetterbedingungen und saisonaler Muster vergleichbare Tagesverläufe und dadurch wiederkehrende Prognosefehler aufweisen. Zur Generierung neuer Eingangsdaten für das Korrekturmodell mittels Clustering wird der Datensatz der PV-Prognosedaten mithilfe eines Clustering-Algorithmus in Cluster eingeteilt. Jedem Cluster wird eine Nummer zugewiesen, sodass jede PV-Prognose eine Clusternummer für das Training des Korrekturmodells erhält. Diese Clusternummer K wird als zusätzliches Merkmal für das Training des Korrekturmodells, wie in Abbildung 4.4 dargestellt, integriert. Im Rahmen der Prognose erfolgt

die Zuordnung der PV-Prognose zu einem Cluster, dem die Prognose am ähnlichsten ist. Die Clusternummer dieses Clusters wird als zusätzlicher Eingangswert in das trainierte Korrekturmodell integriert.

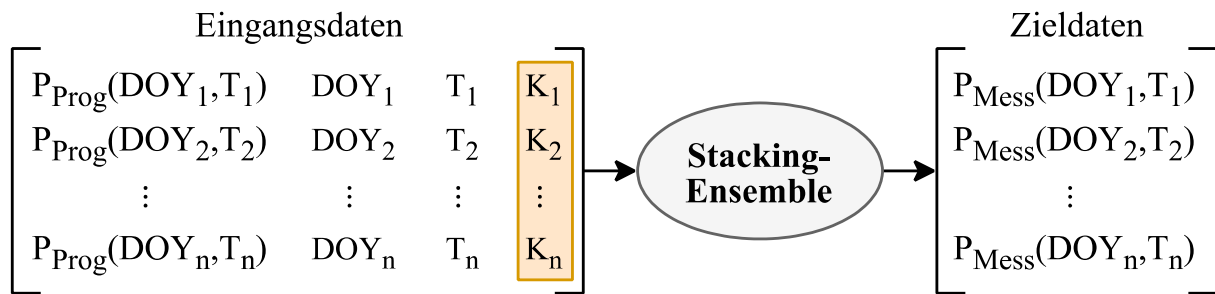


Abbildung 4.4: Eingangs- und Zieldaten für das Training eines Stacking-Ensembles mit dem hinzugefügten Merkmal Clusternummer

Für das Clustern der PV-Prognosedaten findet das weit verbreitete K-Means-Clustering-Verfahren Anwendung, welches auch in [64, 79] zum Einsatz kommt. Die Bestimmung der Clusteranzahl erfolgt mittels der sogenannten Ellenbogen-Methode. Dabei wird die Trägheit in Abhängigkeit von der Clusteranzahl dargestellt [20]. Die Trägheit wird durch die Summe der quadrierten Abweichungen (*engl. Sum of Squared Errors, SSE*) beschrieben und dient als Maß für den Abstand der Datenpunkte zum Centroid eines Clusters. Eine Erhöhung der Clusteranzahl resultiert in einer Reduktion des Abstandes und damit in einer Abnahme der Trägheit. Zunächst ist ein schnelles Sinken der Trägheit zu beobachten, gefolgt von einem Knickpunkt, der als Ellenbogen bezeichnet wird. Ab diesem Punkt verringert sich die Trägheit nur noch langsam. Am Ellenbogen kann die optimale Clusteranzahl abgelesen werden, da ab diesem Punkt die Trägheit sich mit einer höheren Anzahl von Clustern nicht mehr stark verringert. Das Ellenbogendiagramm ist in Abbildung 4.5 dargestellt. Da die Zeitreihen der PV-Prognosen starke Schwankungen aufweisen können, wird ein großer Bereich von Clusteranzahlen zwischen 2 und 30 analysiert. Es ist zu erkennen, dass die Kurve sehr glatt ohne ausgeprägte Ellenbogen verläuft. Dies kann auf die Komplexität des Datensatzes zurückzuführen sein. Die optimale Clusteranzahl kann beispielsweise zwischen 7 und 20 Clustern liegen. Die Ellenbogen-Methode stellt eine eher grobe Methode dar und ist zudem subjektiv [20].

Da die Ellenbogen-Methode keine eindeutige Bestimmung der Clusteranzahl ermöglicht, wird die optimale Clusteranzahl mittels eines empirischen Ansatzes bestimmt. Dieser Ansatz umfasst die Berechnung des Prognosefehlers der PV-Prognose nach dem Korrekturmodell für verschiedene Clusteranzahlen. Das Korrekturmodell wird mehrfach für einen Bereich von 2 bis 30 Clustern durchgeführt. Der nMAE wird für jede Prognose berechnet und gegen die Anzahl der Cluster aufgetragen, wie in Abbildung 4.6 dargestellt. Es wird eine leichte Tendenz ersichtlich, dass mit einer erhöhten Zahl an Clustern der Prognosefehler sinkt. Gleichzeitig weist die Kurve signifikante Schwankungen auf. Dabei ist anzumerken, dass der dargestellte Bereich des nMAE zwischen

7,67 % und 7,88 % lediglich eine starke Vergrößerung aufweist und die Änderung des Prognosefehlers lediglich in einem geringen Bereich bei einer Änderung der Clusteranzahl erfolgt. Der niedrigste nMAE-Wert von 7,67 % wird bei einer Clusteranzahl von $K = 19$ erreicht. Aus diesem Grund wird eine Clusteranzahl von 19 für die Gruppierung der PV-Prognosedaten in Cluster festgelegt.

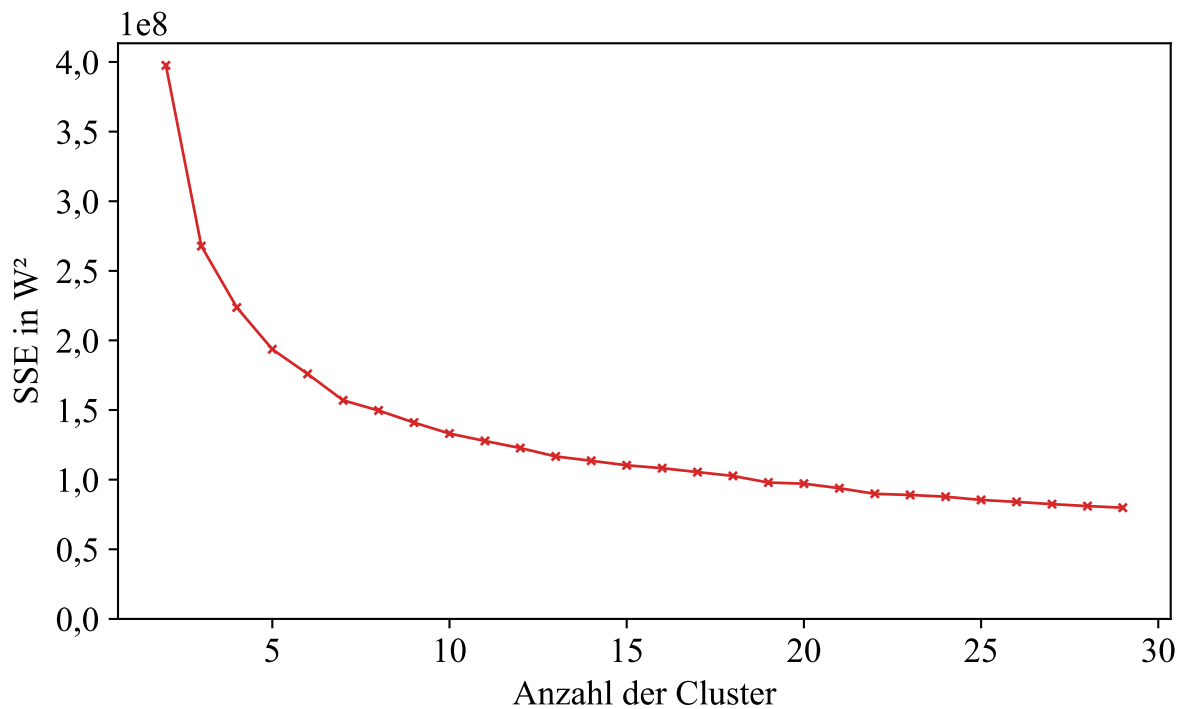


Abbildung 4.5: Trägheit als SSE in Abhängigkeit der Clusteranzahl

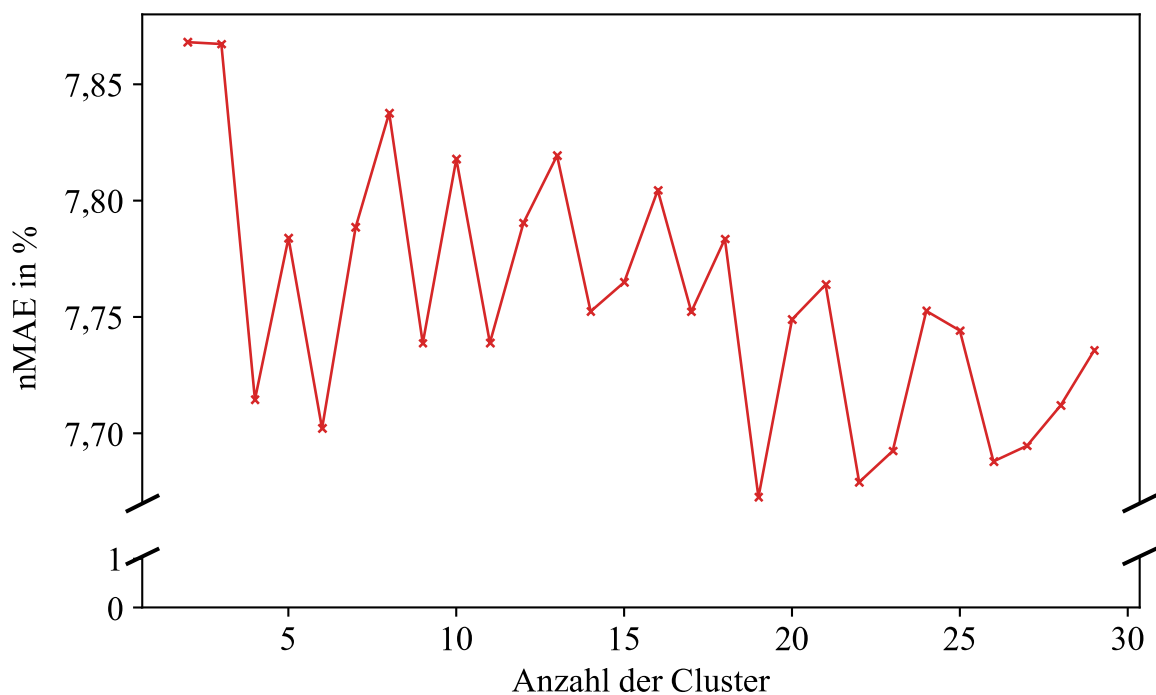


Abbildung 4.6: Prognosefehler als nMAE in Abhängigkeit der Clusteranzahl

Die Einstellung der Hyperparameter für den Clustering-Algorithmus erfolgt mit dem Ziel, den niedrigsten nMAE des Korrekturmodells durch die Nutzung der Clusternummern im Korrekturmodell zu erreichen. Der häufig verwendete Lloyd-Algorithmus wird eingesetzt, der auf Basis des euklidischen Abstands die Zuordnung der Datenpunkte zu den Centroids berechnet und die Position der Centroids iterativ optimiert [18]. Die Toleranz wird sehr niedrig gewählt, damit der Algorithmus stoppt, wenn die Trägheit zwischen den Centroids in zwei aufeinander folgenden Iterationen kleiner als $1 \cdot 10^{-12}$ ist. Die Kombination einer hohen Iterationsgrenze mit einer niedrigen Toleranz hilft dem Algorithmus, genauere Centroids zu ermitteln.

In Abbildung 4.7 werden die mittlere Zeitreihe der PV-Leistung in Rot sowie alle Zeitreihen in Grau von Cluster 1 und Cluster 6 dargestellt. Eine Übersicht über alle 19 Cluster sind im Anhang B in Abbildung B. 1 zu finden. Cluster 1 repräsentiert eher einen sonnigen Tag im Sommer, da die Tage im Gegensatz zu Cluster 6 länger sind und die Kurve den typischen Verlauf der PV-Erzeugung an einem sonnigen Tag ohne Leistungseinbrüche durch Wolken zeigt. Dennoch zeigen die grauen Kurven, dass es innerhalb von Cluster 1 auch Tage mit geringen Leistungseinbrüchen gibt. Cluster 6 zeigt ein Cluster, welches eher den Wintertagen zuzuordnen ist, da die Tage kürzer sind und die Erzeugung geringer ist. Zudem zeigt dieses Cluster keinen typischen Glockenverlauf, wie er an einem sonnigen Tag üblich ist, sondern eine höhere Erzeugung vormittags und eine niedrigere Erzeugung nachmittags. Diese Beobachtung könnte auf Tage zurückzuführen sein, an denen die PV-Anlage nachmittags im Winter verschattet ist und daher nachmittags eine geringere Leistung erzeugt. Diese Hypothese lässt sich jedoch nicht eindeutig aus den Clustern ableiten. Eine weitere mögliche Erklärung könnte in einer veränderten Bewölkung zur Nachmittagszeit liegen.

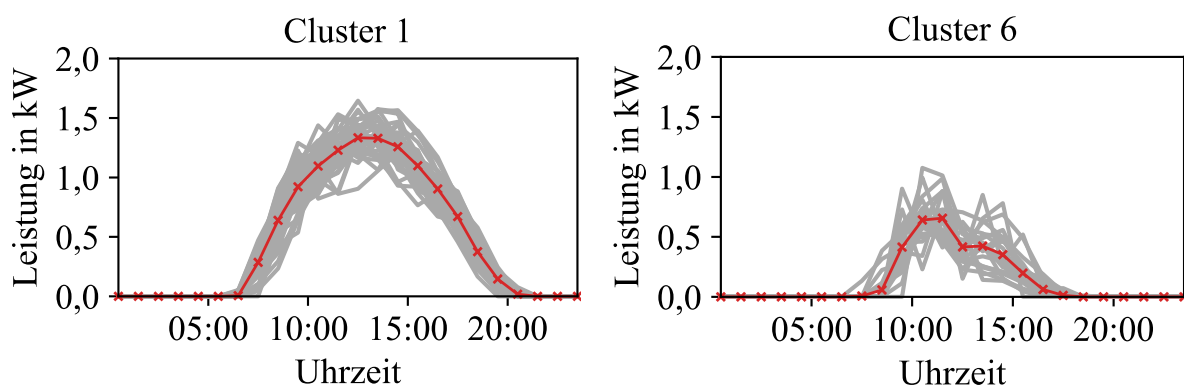


Abbildung 4.7: Mittlere Zeitreihe der PV-Leistung in Rot sowie alle Zeitreihen in Grau von Cluster 1 und Cluster 6

Trotzdem wird ersichtlich, dass Cluster mit saisonalen Mustern, wie sie etwa die Jahreszeiten aufweisen, oder wiederkehrende Muster, wie sie durch eine bestimmte Erzeugungsstruktur über den Tag hinweg entstehen, mit der Clustering-Methode dargestellt

werden können. Die Cluster können als Eingangsmerkmal für das Korrekturmodell dienen, um diesem zu ermöglichen, entsprechende Tage zu identifizieren. Tägliche Schwankungen in der Erzeugung aufgrund der Bewölkung sind dagegen schwieriger abzubilden, da dieses Verhalten sehr zufällig ist und nicht an bestimmte Muster gekoppelt ist. Dies wird an den beiden dargestellten Clustern deutlich, da die grauen Kurven auch Ausreißer zeigen, die sich deutlich von der roten mittleren Leistungskurve unterscheiden.

4.1.2 Auswahl und Training der ML-Modelle

Für das Training der Basismodelle wird der im vorherigen Kapitel beschriebene Trainingsdatensatz verwendet. Der Testdatensatz wird im Rahmen von Kapitel 0 genutzt, um die Leistung des Korrekturmodells bei unbekannten Daten zu bewerten. Für die Basismodelle wird zudem zum Trainingsdatensatz ein Validierungsdatensatz benötigt. Um die Prognosegenauigkeit der Basismodelle zu evaluieren und Overfitting zu erkennen, wird die Methode der k -fachen Kreuzvalidierung angewendet. [17, 28]. Zudem gewährleistet diese Methode die Bereitstellung einer ausreichenden Anzahl an Daten für das Training des Metamodells. Diesbezüglich sei angemerkt, dass bei einem Datensatz, der zeitabhängige Daten beinhaltet, die herkömmliche Aufteilung von Daten in einen Trainings- und einen Validierungssatz nicht zwangsläufig ausreichend ist, um sicherzustellen, dass jedes Zeitintervall, beispielsweise jede Jahreszeit, als Validierungsdatensatz verwendet wird. Die k -fache Kreuzvalidierung gewährleistet die Generalisierbarkeit der Modelle. Dabei erfolgt eine Unterteilung der Trainingsdatenmenge in k Teilmengen. Im Rahmen der k -fachen Kreuzvalidierung wird jedes Basismodell k -mal trainiert. Dabei wird in jeder Iteration ein unterschiedlicher Teil des Datensatzes entfernt. Diese Teilmenge wird für die Validierung verwendet. Folglich erfolgt in jeder Iteration ein Training und eine Validierung mit einem unterschiedlichen Trainings- und Validierungsdatensatz. Jede Teilmenge wird dabei einmal als Validierungsdatensatz genutzt. Dieses Verfahren ist in Abbildung 4.8 dargestellt.

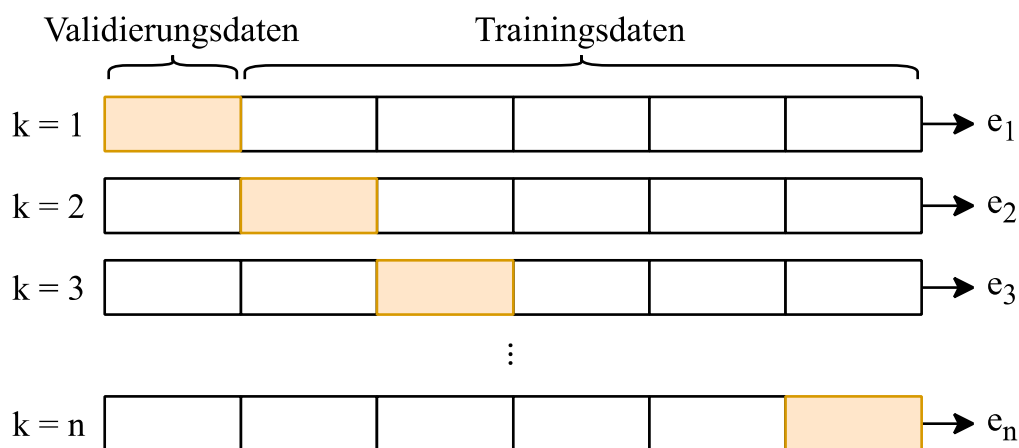


Abbildung 4.8: Methode der k -fachen Kreuzvalidierung nach [17, 18]

Im Anschluss an die Iterationen erfolgt eine Mittelung der Fehler der Prognosen aus der Validierung aus sämtlichen Iterationen gemäß Formel (4.1) [18]. Hierbei bezeichnet e_i den Fehler für jede Iteration i . Ein Nachteil der k -fachen Kreuzvalidierung besteht in dem mit ihr einhergehenden erhöhten Rechenaufwand, da jedes Modell k -mal trainiert und validiert werden muss [17].

$$E_m = \frac{1}{n} \sum_{i=1}^n e_i \quad (4.1)$$

Im Rahmen dieser Arbeit erfolgt eine k -fache Kreuzvalidierung für fünf Iterationen. In der Regel wird entweder ein Wert von $k = 5$ oder von $k = 10$ gewählt [93]. Allerdings existiert keine allgemein anwendbare Methode zur exakten Ermittlung des optimalen k -Wertes für eine spezifische Datenmenge. In der Praxis wird häufig ein größerer k -Wert, beispielsweise $k = 10$, präferiert, da dies eine Aufteilung der Trainingsdaten in zehn gleich große Teile erlaubt. Dies resultiert in einem Trainingsdatensatz, der nahezu die gleiche Größe und Vollständigkeit wie der ursprüngliche Datensatz aufweist, was zu einer präziseren Modellanpassung führen kann. Bei geringeren k -Werten, wie beispielsweise $k = 5$, werden lediglich vier Fünftel der Daten für das Training verwendet, während ein Fünftel als Validierungsdaten dient. Dies ist bei einem Stacking-Ensemble jedoch nicht von entscheidender Bedeutung, da das Metamodell die Schwächen der Basismodelle durch eine gewichtete Kombination der Vorhersagen kompensiert. Des Weiteren lässt sich die Rechenzeit im Vergleich zu höheren k -Werten signifikant reduzieren. Aus den genannten Gründen wird in der vorliegenden Arbeit der Wert $k = 5$ für die k -fache Kreuzvalidierung als Kompromiss zwischen Rechenaufwand, Modellkomplexität und Generalisierbarkeit verwendet.

Die Prognosen, welche auf Basis des Validierungsdatensatzes aus jeder Iteration erstellt werden, werden zusammengefasst, sodass am Ende die gleiche Datenmenge wie im Trainingsdatensatz für das Metamodell zur Verfügung steht. Da fünf Basismodelle genutzt werden, entsteht ein Trainingsdatensatz mit fünf Merkmalen für das Metamodell. Für das Training des Metamodells wird der gesamte erstellte Trainingsdatensatz genutzt. Im Anschluss an die Trainingsphase erfolgt eine Evaluierung des Stacking-Ensembles auf Basis des Testdatensatzes, welcher für alle Modelle bislang unbekannt ist.

Für die Festlegung der ML-Modelle für das Stacking-Ensemble ist es notwendig, Modelle auszuwählen, die verschiedene Algorithmen verwenden. Die Modelle sollen möglichst unterschiedliche Fehler machen, die sich in einem Ensemble gegenseitig ausgleichen. Auf diese Weise können die Stärken von jedem Modell kombiniert werden, was zu einer Optimierung der Gesamtvorhersage führt. Es besteht die Möglichkeit, in einem Stacking-Ensemble andere ensemblebasierte Modelle wie RF (Bagging) und Gradient Boosting zu integrieren. Dies erlaubt die Erkennung komplexer Muster in den Daten.

Als mögliche Basismodelle des Stacking-Ensembles finden die Modelle auf Basis der Methoden KNN, LR, DT, RF, Gradient Boosting, XGBoost und Catboost Regressor Anwendung. Die genannten Modelle lassen sich sowohl in einfache, wie die KNN-Methode, als auch in komplexe Methoden, wie die Boosting-Algorithmen, unterteilen. Auf die Nutzung von ANNs wird verzichtet, da die Datenbasis nicht ausreichend ist, um ein ANN zu trainieren. Die Methodik des Korrekturmodells basiert auf der Idee, möglichst wenig Daten der gemessenen PV-Leistung zu sammeln, um die Prognose zu korrigieren. Daher ist der Einsatz von ANNs innerhalb eines Stacking-Ensembles nicht zielführend.

Es folgt eine Erläuterung der Gründe für die Nutzung der Modelle im Stacking-Ensemble sowie die Auswahl der Hyperparameter der einzelnen Modelle beim Training. Für die Auswahl der Modelle ist es erforderlich, dass diese eine Diversität aufweisen, um die Stärken unterschiedlicher Modelle zu kombinieren und um verschiedene Aspekte des Datensatzes abzudecken. Es gibt kein spezifisches Verfahren, um die Modelle für das Stacking-Ensemble auszuwählen. Die Auswahl der Modellkombinationen erfolgt durch eine systematische Untersuchung verschiedener Modelle, indem verschiedene Modelle getestet und miteinander verglichen werden.

Im Rahmen der Auswahl werden mehrere ML-Modelle mit der zuvor beschriebenen Methode der k-fachen Kreuzvalidierung in einem Stacking-Ensemble zusammengefasst. Die Bestimmung der Hyperparameter ist wie beim Training von ANNs zeitaufwendig und ungenau. Die Bestimmung der Hyperparameter für die Modelle erfolgt daher mittels der GridSearchCV-Methode. GridSearchCV ist eine Methode der ML-Bibliothek Scikit-learn zur Optimierung von Hyperparametern von ML-Modellen. Dabei werden in einer Gitter-Suche in Kombination mit einer k-fachen Kreuzvalidierung alle möglichen Kombinationen von Hyperparametern je Modell evaluiert, wobei diejenige Kombination ausgewählt wird, bei der das Modell die besten Ergebnisse erzielt [20]. Das ist der Fall, wenn der Fehler bei den Validierungsdaten am geringsten ist. Zu diesem Zweck wird vorab ein Hyperparameter-Gitter definiert, innerhalb dessen die Hyperparameter variieren. Die Methode des GridSearchCV zielt darauf ab, eine Metrik zu maximieren. Daher findet der negative MAE als Fehlermetrik Verwendung, um den Fehler der Vorhersage zu minimieren. Der MAE als Fehlermetrik weist eine höhere Robustheit gegenüber Ausreißern auf [94]. Das Quadrieren der Fehlerwerte beim MSE beispielsweise führt zu einer stärkeren Gewichtung von größeren Werten. Die Nutzung des MAE gewährleistet, dass Modelle nicht zu stark an Ausreißern ausgerichtet werden. Als Beispiel für derartige Ausreißer können Tage mit stark wechselnder Bewölkung angeführt werden, welche mit einer Day-Ahead-Prognose nur schwer zu erfassen sind. Das Ziel der vorliegenden Vorhersage besteht in der Reduktion von Fehlern, die unter anderem durch Verschattung oder Alterung bedingt sein können.

Die Auswahl der möglichen Hyperparameterwerte je Modell erfolgt grundsätzlich mit dem Ziel, dass die einfachen Modelle wie KNN und RF zur Verallgemeinerung beitragen. Die Modelle weisen demnach eine geringere Komplexität auf. Die Prognose von Tagen mit wechselnder Bewölkung, welche zu Leistungsspitzen und -einbrüchen über den Tag führt, erweist sich als besonders herausfordernd. Dies ist darauf zurückzuführen, dass sich die Mehrzahl der Modelle auf die allgemeine Generalisation konzentriert, um eine Überanpassung der Daten zu vermeiden. Die Hyperparameter der Boosting-Methoden werden daher so gewählt, dass sie komplexe Muster erfassen können. Diese Vorgehensweise kann beispielsweise durch eine große Tiefe und eine hohe Anzahl von Iterationen erreicht werden.

Da alle möglichen Hyperparameterkombinationen in einem Gitter getestet werden, wird der Hyperparameter-Raum begrenzt und möglichst einfach gehalten, um den zeitlichen Aufwand gering zu halten. Für die Auswahl der möglichen Werte eines Hyperparameters wird ein großer Bereich rund um den voreingestellten Standardwert von den ML-Bibliotheken in Python wie Scikit-learn angekommen. Im ersten Trainingsdurchlauf erfolgt die Bestimmung der optimalen Hyperparameter für alle Modelle unter Verwendung der Methode GridSearchCV. Sofern die dort bestimmten Hyperparameter sich am Rand des definierten Wertebereichs befinden, erfolgt eine Anpassung des Bereichs für den jeweiligen Hyperparameter. Die Anpassung erfolgt derart, dass der neue Wertebereich den Rand des vorherigen Bereichs übersteigt, wodurch ein größerer Wertebereich als zuvor betrachtet wird. Wenn ein Wert innerhalb des definierten Bereichs liegt, wird ein höherer Detaillierungsgrad des Wertebereichs betrachtet, um sicherzustellen, dass das Modell fein genug dimensioniert ist. In einem nächsten Schritt erfolgt ein erneutes Training mit den modifizierten Wertebereichen. Dieses Vorgehen wird wiederholt, bis die optimalen Hyperparameter identifiziert sind. Darüber hinaus werden die Hyperparameter der Boosting-Algorithmen nach den Durchläufen bewusst so angepasst, dass die Boosting-Algorithmen unterschiedlich arbeiten. Da XGBoost und Gradient Boosting vom Prinzip her ähnliche Arbeitsweisen aufweisen, wird beispielsweise die Lernrate beim XGBoost-Modell signifikant kleiner gesetzt im Gegensatz zum Gradient Boosting. In der Konsequenz wird das XGBoost-Modell dazu angehalten, langsamer und in kleineren Schritten zu lernen. Das soll dazu führen, dass das XGBoost-Modell Overfittung vermeidet, da die Boosting-Modelle schnell zu Overfitting neigen. Damit wird sichergestellt, dass XGBoost anders arbeitet. Das CatBoost-Modell wird dahingehend angepasst, die Komplexitäten des Datensatzes zu erfassen. Ein Beispiel für diese Anpassung ist die deutliche Erhöhung der Anzahl der DTs im CatBoost-Modell. Diese Maßnahme zielt darauf ab, eine Variabilität in der Arbeitsweise der Boosting-Algorithmen zu erzeugen, die durch die Einstellung der Hyperparameter bedingt ist. Dabei ist zu berücksichtigen, dass der MAE-Wert im Vergleich zu anderen Hyperparametereinstellungen möglicherweise eine Verschlechterung aufweist. Das Ziel dieses Vorgehens besteht da-

rin, Modelle zu erstellen, die jeweils spezifische Vorteile aufweisen und in einem Stacking-Ensemble kombiniert werden können. Das Stacking-Verfahren zielt darauf ab, die Stärken der einzelnen Modelle zu kombinieren, um die Schwächen abzuschwächen und die Vorhersagegenauigkeit zu verbessern, indem die spezifischen Vorteile der einzelnen Modelle genutzt werden. Die für jedes Modell definierten Hyperparameterräume sind im Anhang B in Tabelle B. 1, Tabelle B. 2, Tabelle B. 3, Tabelle B. 4, Tabelle B. 5, Tabelle B. 6 und Tabelle B. 7 dargestellt.

Für die Auswertung des Trainings werden für jedes einzelne Modell die ebenfalls in Kapitel 3.1.2 betrachteten statistischen Kennzahlen berechnet. Diese Metriken werden verwendet, um die Modelle für das Stacking-Ensemble auszuwählen, wobei ein Vergleich mit den Metriken der ursprünglichen PV-Prognose erfolgt. Im Anschluss erfolgt eine Auswahl der Modelle, welche sich durch eine signifikante Reduktion des nMAE auszeichnen. Die Tabelle 4.1 zeigt die Fehlermetriken der Testdaten aller betrachteten Basismodelle mit der Clustering-Methode. Im Anhang B in Tabelle B. 9 sind die Fehlermetriken aller Basismodelle ohne die Clustering-Methode dargestellt.

Die tabellarische Übersicht der Fehlermetriken zeigt, dass das XGBoost-Modell mit einem nMAE von 17,93 % den größten Prognosefehler aufweist. Im Vergleich zur ursprünglichen Prognose mit einem nMAE von 8,76 % ist eine deutliche Verschlechterung der Prognose zu verzeichnen. Es ist jedoch anzumerken, dass das XGBoost-Modell bewusst angepasst wurde, um einen Unterschied zu Gradient Boosting zu erreichen. Dabei wird ein höherer Prognosefehler in Kauf genommen. Das Modell wird zunächst nicht aussortiert und im Stacking-Ensemble berücksichtigt. Auch das CatBoost-Modell wird im Stacking-Ensemble berücksichtigt. Es wurde gezielt angepasst, um die Komplexitäten im Datensatz zu erfassen. Zudem führt das CatBoost-Modell mit einem nMAE von 8,06 % zu einer Verbesserung des Prognosefehlers. Eine Verbesserung ist hingegen beim DT nicht erkennbar, dessen nMAE bei 9,59 % liegt. Demgegenüber weist der RF einen nMAE von 7,70 % auf, was eine Verbesserung der Prognose bedeutet. Da das Funktionsprinzip von DTs sehr ähnlich zu dem von einem RF ist, welches ein Bagging von mehreren DTs ist, findet die Methode des DT im weiteren Verlauf keine Berücksichtigung im Rahmen des Stacking-Ensembles. Für ein optimales Stacking ist die Auswahl von Basismodellen mit unterschiedlichen Algorithmen erforderlich. Das LR-Modell weist einen nMAE von 8,67 % auf. Die Verbesserung im Vergleich zur ursprünglichen Prognose ist demnach nur sehr gering. Nichtsdestotrotz wird das LR-Modell in der weiteren Betrachtung berücksichtigt. Die Prognosefehler des Gradient Boosting Modells und des KNN-Modells betragen 7,80 % bzw. 7,83 %. Aus diesem Grund werden diese Modelle auch in einem Stacking-Ensemble berücksichtigt.

Tabelle 4.1: Fehlermetriken für die Korrektur der PV-Prognose von verschiedenen Basismodellen für die Testdaten mit der Clustering-Methode

ML-Modell	MAE in W	nMAE in %	RMSE in W	nRMSE in %
LR	183,01	8,67	270,94	12,84
KNN	165,16	7,83	249,19	11,82
DT	202,36	9,59	295,51	14,01
RF	162,44	7,70	245,73	11,65
CatBoost	169,88	8,06	260,91	12,37
XGBoost	378,25	17,93	534,10	25,33
Gradient Boosting	164,44	7,80	245,70	11,65
Ursprüngliche Prognose	184,66	8,76	267,49	12,68

Die ausgewählten Modelle werden einer weiteren Untersuchung unterzogen. Wie bei den Methoden DT und RF ist es weiterhin wichtig, die Modellvielfalt zu gewährleisten. Allerdings erweist sich eine Bewertung der Diversität allein auf der Grundlage von Fehlermetriken als schwierig. Daher werden verschiedene Kombinationen von Basismodellen in einem Stacking-Ensemble untersucht. Im Folgenden wird ermittelt, ob sechs Modelle als Basismodelle in einem Stacking-Ensemble erforderlich sind. Dazu werden die Vorhersagen der Basismodelle in verschiedenen Kombinationen in einem Stacking-Ensemble kombiniert und in ein Metamodell überführt. Das Metamodell basiert auf der LR-Methode und kombiniert die Vorhersagen der Basismodelle zu einer finalen Vorhersage. Für das Stacking-Ensemble wird die LR-Methode als Metamodell eingesetzt. Die LR-Methode ist eine einfache und nachvollziehbare Methode. Es lässt sich erkennen, wie die einzelnen Vorhersagen der Basismodelle gewichtet werden. Diese Regressionskoeffizienten geben an, in welchem Maß die einzelnen Basismodelle zur endgültigen Vorhersage beitragen. In Tabelle 4.2 sind die Fehlermetriken des Korrekturmodells von verschiedenen Kombinationen von Basismodellen in einem Stacking-Ensemble mit der Clustering-Methode dargestellt. Die Fehlermetriken des Korrekturmodells von verschiedenen Kombinationen von Basismodellen in einem Stacking-Ensemble ohne die Clustering-Methode sind in Anhang B in Tabelle B. 10 dargestellt. Alle Kombinationen von Basismodellen führen zu einer Verbesserung der Prognose, obwohl nicht alle Basismodelle alleine den Prognosefehler reduzieren. Dies zeigt, dass in einem Stacking-Ensemble die Stärken der einzelnen Basismodelle kombiniert werden. Der niedrigste Prognosefehler von 7,67 % wird mit einem Stacking-Ensemble aus den fünf Basismodellen KNN, RF, CatBoost, XGBoost und Gradient Boosting erreicht. Dabei wird nur das LR-Modell nicht berücksichtigt. Wird das LR-Modell hingegen berücksichtigt, so beträgt der Prognosefehler auf 7,74 %. Der Unterschied zu einem Stacking-Ensemble aus den Basismodellen KNN, RF, CatBoost und Gradient Boosting ist sehr gering und

der nMAE beträgt 7,68 %. In diesem Ensemble findet neben dem LR-Modell auch das XGBoost-Modell keine Berücksichtigung. Auch bei einem Ensemble, das ausschließlich die Modelle KNN und RF umfasst, wird ein niedriger nMAE von 7,69 % erreicht. Die Boosting-Modelle werden in der Regel für die komplexen Elemente im Datensatz eingesetzt. Dies lässt darauf schließen, dass komplexe Muster im Datensatz möglicherweise nicht zuverlässig erkannt und verbessert werden können. Als Beispiel für derartige komplexe Muster können Fehler genannt werden, die nicht wiederkehrend sind und aufgrund des lokal sich schnell veränderten Wetters sowie durch Fehler in der NWP entstehen. Dennoch wird im weiteren Verlauf das Stacking-Ensemble mit den fünf Basismodellen KNN, RF, CatBoost, XGBoost und Gradient Boosting, welches den geringsten Prognosefehler aufweist, ausgewählt. Insgesamt zeigt sich zwar nur ein sehr geringer Unterschied zu den Fehlern von den anderen Stacking-Ensembles, an einzelnen Tagen können jedoch erhebliche Unterschiede in den Prognosen auftreten, die auf die spezifischen Eigenschaften des gewählten Basismodells zurückzuführen sind.

Der Hyperparameterraum vom Metamodell, welches durch eine Ridge Regression umgesetzt wird, ist im Anhang B in Tabelle B. 8 dargestellt. Die Hyperparameter der fünf ausgewählten Basismodelle sowie die Hyperparameter des Metamodells sind in Anhang B in Tabelle B. 11 aufgeführt.

Tabelle 4.2: Fehlermetriken für die Korrektur der PV-Prognose von verschiedenen Basismodellkombinationen eines Stacking-Ensembles für die Testdaten mit der Clustering-Methode

Basismodellkombinationen	MAE in W	nMAE in %	RMSE in W	nRMSE in %
KNN, RF, LR, CatBoost, XGBoost, Gradient Boosting	163,29	7,74	246,15	11,68
KNN, RF, CatBoost, XGBoost, Gradient Boosting	161,82	7,67	244,44	11,59
KNN, RF, Gradient Boosting	162,60	7,71	244,46	11,59
KNN, RF, CatBoost, Gradient Boosting	161,92	7,68	244,44	11,59
KNN, RF	162,25	7,69	244,58	11,60
CatBoost, XGBoost, Gradient Boosting	166,57	7,90	245,78	11,66

4.2 Ergebnisse der korrigierten Day-Ahead-Prognose einer PV-Anlage

In diesem Kapitel werden die Ergebnisse des Korrekturmodells der PV-Leistungsprognose vorgestellt und diskutiert. Die Entwicklung und Validierung des Korrekturmodells erfolgt anhand der in Kapitel 3 betrachteten PV-Anlage der Hochschule Bielefeld. Im Folgenden werden die Ergebnisse des Stacking-Ensembles mit den Ergebnissen der einzelnen ML-Modelle in Bezug auf die Testdaten verglichen. Es wird erörtert, ob sich die Anwendung einer Stacking-Ensemble-Methode als vorteilhaft gegenüber der Verwendung einzelner ML-Modelle erweist. Des Weiteren wird diskutiert, ob die Anwendung der Clustering-Methode zur Generierung zusätzlicher Merkmale zu einer Optimierung der Prognosegenauigkeit führt. Darüber hinaus wird eine Visualisierung der Daten anhand von Stichproben vorgenommen. Zudem wird untersucht, wie viele PV-Leistungsmessdaten vorab zu sammeln sind, um das Korrekturmodell im laufenden Betrieb zu optimieren. Die Methode ist nur praktikabel, wenn wenige Messdaten ausreichen, um das Modell zu trainieren. Andernfalls wäre ein direktes Training mit den PV-Leistungsmessdaten möglich.

In Tabelle 4.3 sind die Fehlermetriken des Korrekturmodells der PV-Leistungsprognose für die Testdaten mit und ohne der Clustering-Methode dargestellt. Die Fehlermetriken sind sowohl für die einzelnen ML-Modelle des Stacking-Ensembles als auch für ein Stacking-Ensemble der ML-Modelle aufgeführt. Im Vergleich dazu ist ebenfalls der Fehler für die ursprüngliche Prognose der Testdaten aufgeführt. In diesem Zusammenhang ist anzumerken, dass sich der Prognosefehler der ursprünglichen Prognose aus Kapitel 3 zu diesem leicht unterscheidet, da in diesem Kapitel andere Testdaten verwendet werden. In Kapitel 3 wird der vollständige Datensatz für die Berechnung des Prognosefehlers herangezogen, während in diesem Kapitel die Testdaten einen Anteil von 30 % aus diesem Datensatz ausmachen. Die mittels eines Stacking-Ensembles ohne Clustering korrigierte PV-Prognose weist einen nMAE von 7,89 % auf. Im Vergleich dazu beträgt der nMAE der PV-Prognose ohne Korrektur 8,76 %. Dies bedeutet eine Verbesserung um 9,93 % bzw. um 0,87 Prozentpunkte, die durch die Korrektur erzielt wird. Mit der Clustering-Methode werden die besten Ergebnisse durch ein Stacking-Ensemble erzielt. Der nMAE beträgt 7,67 %. Dies bedeutet eine Verbesserung des nMAE im Gegensatz zur ursprünglichen Prognose von 12,44 % bzw. 1,09 Prozentpunkten. Die Ergebnisse zeigen, dass ein vorangehendes Clustering der Daten sowie eine Berücksichtigung der Clusternummer als Merkmal eine sinnvolle Vorgehensweise darstellt, um dem Stacking-Ensemble ein zusätzliches Merkmal zur Erkennung von Tagen mit ähnlichen Mustern zur Verfügung zu stellen.

Tabelle 4.3: Fehlermetriken des Korrekturmodells der PV-Prognose für verschiedene ML-Modelle und ein Stacking-Ensemble der ML-Modelle für die Testdaten

Korrekturmethode	MAE in W/m ²	nMAE in %	RMSE in W/m ²	nRMSE in %
ohne Clustering-Methode				
KNN	167,63	7,95	254,84	12,08
RF	166,72	7,91	252,68	11,98
CatBoost	177,82	8,43	273,80	12,98
XGBoost	378,27	17,94	534,12	25,33
Gradient Boosting	168,14	7,97	251,46	11,92
Stacking-Ensemble	166,28	7,89	251,27	11,92
Ursprüngliche Prognose	184,66	8,76	267,49	12,68
mit Clustering-Methode				
KNN	165,16	7,83	249,19	11,82
RF	162,44	7,70	245,73	11,65
CatBoost	169,88	8,06	260,91	12,37
XGBoost	378,25	17,93	534,10	25,33
Gradient Boosting	164,44	7,80	245,70	11,65
Stacking-Ensemble	161,82	7,67	244,44	11,59
Ursprüngliche Prognose	184,66	8,76	267,49	12,68

Die Anwendung der Clustering-Methode führt auch beim RF zu vielversprechenden Ergebnissen mit einem nMAE von 7,70 %. Diese Ergebnisse sind die zweitbesten, die mittels Clustering erzielt werden. In diesem Zusammenhang stellt sich die Frage, ob ein solches Stacking-Ensemble überhaupt erforderlich ist, wenn ein ML-Modell alleine gute Ergebnisse erzielt. Bei dem RF handelt es sich ebenfalls um ein Ensemble-Modell, genauer gesagt um ein Bagging-Modell. Folglich ist ersichtlich, dass der RF ebenfalls in der Lage ist alleine gute Ergebnisse zu liefern. Im Vergleich zur Entwicklung eines Stacking-Modells kann die Verwendung eines RF sogar zu einer Einsparung von Aufwand führen. Darüber hinaus zeichnet sich die RF-Methode durch ihre Einfachheit und die Vermeidung von Overfitting durch die Kombination mehrerer DTs aus. Demnach ist die alleinige Verwendung von RF zur Korrektur der PV-Prognose als durchaus geeignet zu betrachten, um lange Rechenzeiten zu vermeiden im Gegensatz zu einem Stacking-Ensemble mit verschiedenen ML-Modellen. Obgleich die Fehlermetriken der korrigierten Prognose nur geringfügig unterschiedlich ausfallen, erweist sich das Stacking-Ensemble mit Clustering als überlegen. Dies ist darauf zurückzuführen, dass die Vorzüge sämtlicher ML-Modelle in einem einzigen Modell vereint werden können. Das Merkmal der Clusternummer unterstützt das Stacking-Ensemble dabei, Muster innerhalb der Tage zu

erkennen und die Stärken jedes Modells im Ensemble auf die spezifischen Muster anzuwenden.

Der Effekt der Korrektur lässt sich besonders in Abbildung 4.9 erkennen. In dieser Abbildung sind die gemessene, die prognostizierte sowie die korrigierte prognostizierte PV-Leistung in der Kalenderwoche 52 des Jahres 2019 dargestellt. Die vorliegende Woche ist bereits in Kapitel 3 dargestellt, wo aufgezeigt wird, dass insbesondere an den letzten drei Tagen der Woche, bedingt durch eine Verschattung der PV-Anlage im Winter, ein signifikanter Prognosefehler zu verzeichnen ist, der den Durchschnitt übersteigt.

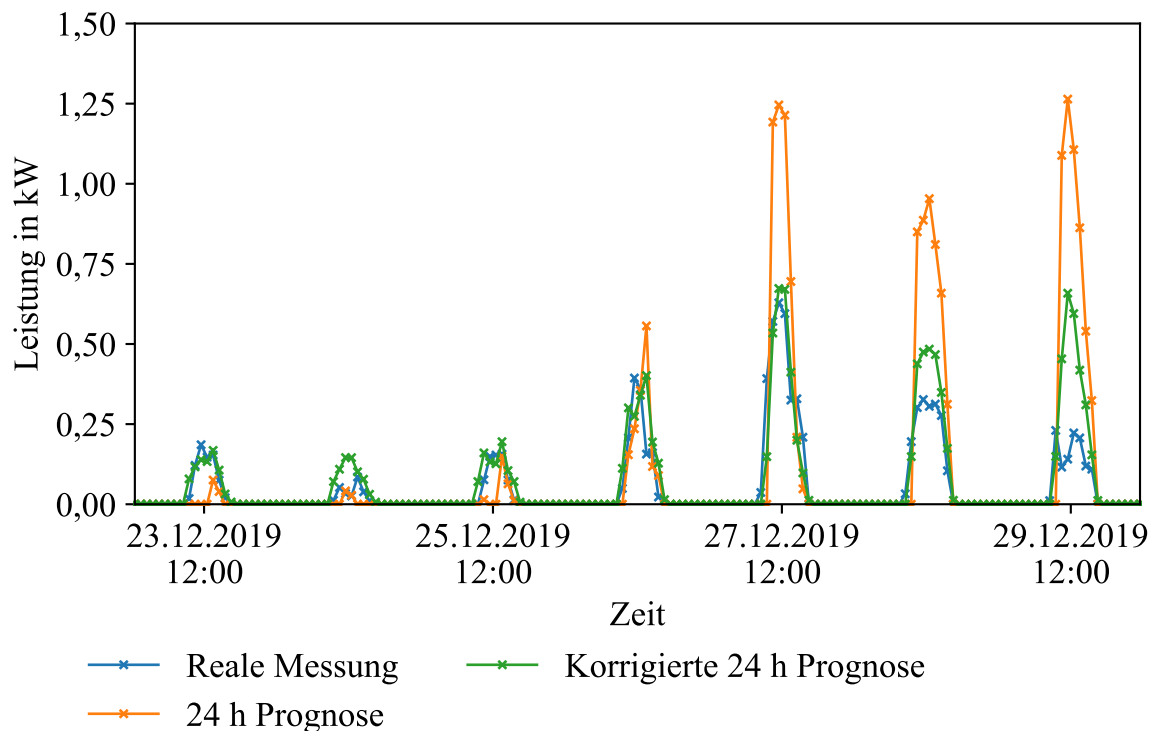


Abbildung 4.9: Gemessene, prognostizierte PV-Leistung und korrigierte prognostizierte PV-Leistung in der Kalenderwoche 52 des Jahres 2019

Die Day-Ahead-Prognose weist für den 27.12.2019 einen nMAE von 17,27 % auf. Durch eine stündliche Aktualisierung der Prognose kann der Fehler auf 14,71 % verringert werden. Die Abbildung verdeutlicht, dass durch das Korrekturmodell auf Basis eines Stacking-Ensembles eine optimierte Prognose möglich ist. Das Korrekturmodell prognostiziert eine geringere PV-Leistung über den Tag, da das durch die Verschattung bedingte Muster erkannt wird. Am 27.12.2019 kann der Prognosefehler somit auf einen nMAE von 4,41 % verringert werden (siehe Tabelle 4.4).

Tabelle 4.4: Fehlermetriken der PV-Prognose und des Korrekturmodells der PV-Prognose für den 27.12.2019

Prognose	MAE in W	nMAE in %	RMSE in W	nRMSE in %
24 h Prognose	362,75	17,27	430,47	20,50
1 h Prognose	308,83	14,71	374,51	17,83
Korrigierte 24 h Prognose	92,56	4,41	114,76	5,46

In Abbildung 4.10 sind die gemessene, die prognostizierte PV-Leistung und korrigierte prognostizierte PV-Leistung an zwei Beispieltagen aus den Testdaten dargestellt. Die Abbildung zeigt, dass die Korrektur an diesen beiden Tagen nicht zu einer Verringerung, sondern zu einer Verschlechterung des Prognosefehlers geführt hat. Die Fehlermetriken beider Tage sind in Tabelle 4.5 aufgeführt. Eine Verschlechterung des nMAE ist am 28.04.2019 von 9,77 % auf 10,10 % und am 29.04.2019 von 8,20 % auf 9,66 % zu verzeichnen. Die Verschlechterung ist insbesondere an der Leistungsspitze am 28.04.2019 zu beobachten. Die Prognose der PV-Leistung sowie das Korrekturmodell prognostizieren eine Leistungsspitze in der Tagesmitte, obwohl dort ein Leistungseinbruch zu verzeichnen ist. Die Korrektur führt sogar zu einer Erhöhung der prognostizierten Spitze.

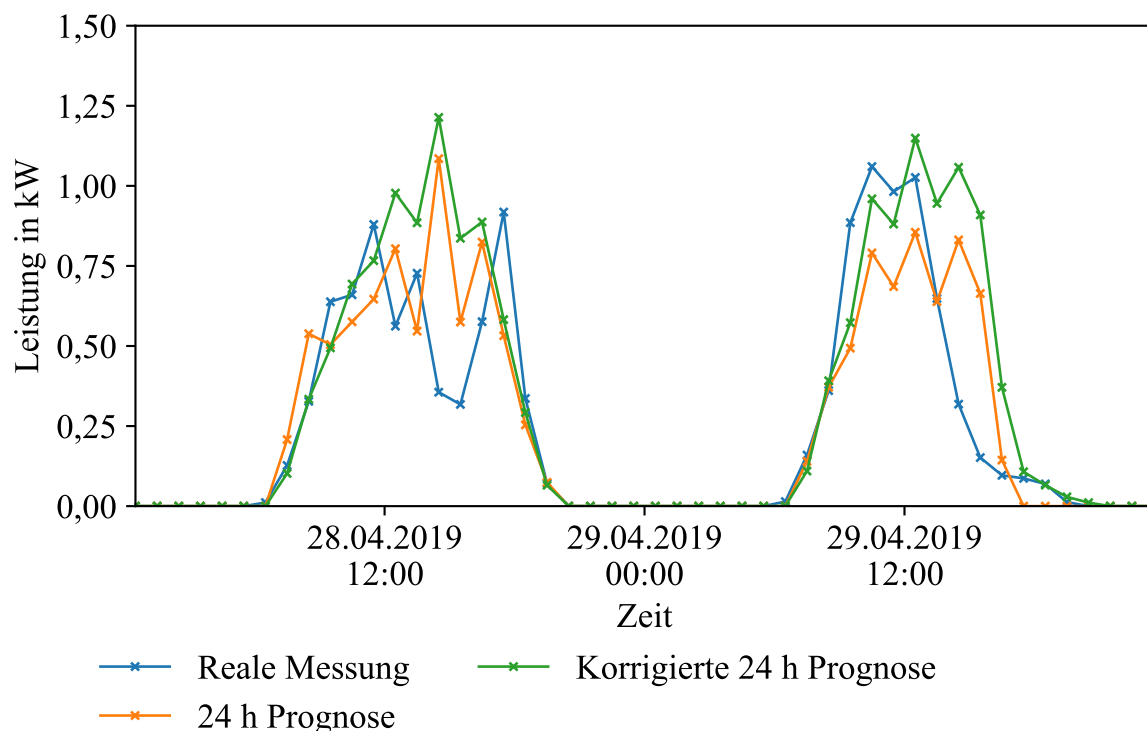


Abbildung 4.10: Gemessene, prognostizierte PV-Leistung und korrigierte prognostizierte PV-Leistung an zwei Beispieltagen

Die Messung der erzeugten PV-Leistung weist zu Beginn des Tages am 29.04.2019 einen höheren Wert auf als prognostiziert. Dies wird durch das Korrekturmodell korrigiert, wobei eine höhere Leistung vorhergesagt wird. Allerdings wird auch der Leistungseinbruch am Ende dieses Tags nicht erkannt, sodass eine noch höhere Leistung vorhergesagt wird als vorher. Die Leistungseinbrüche können durch die im Rahmen der PV-Prognose genutzte Wettervorhersage offenbar nicht identifiziert werden, sodass die PV-Prognose diesbezüglich keine Vorhersage treffen kann. Dies ist auf Fehler in der Wettervorhersage zurückzuführen. Prognosefehler, die durch eine Wettervorhersage bedingt sind, können durch eine Korrektur nicht identifiziert werden, da es sich hierbei um zufällige Fehler handelt, die nicht in einem bestimmten Muster wiederkehren. Gerade an Tagen, an denen eine wechselnde Bewölkung zu beobachten ist, zeigt das Korrekturmodell Schwächen. In diesem Kontext sind Prognosemodelle erforderlich, die einen kürzeren Prognosezeitraum als eine 24-Stunden-Vorhersage abdecken. Dabei kann beispielsweise die Bewegung von Wolkenzügen am Himmel berücksichtigt werden. In der vorliegenden Arbeit werden diese Aspekte nicht behandelt, da durch die sehr kurzfristigen Prognosen, welche lediglich eine Vorhersage für wenige Minuten treffen, ganz andere Anwendungsfälle betrachtet werden. In einer im Rahmen dieser Arbeit getätigten wissenschaftlichen Veröffentlichung wird ein Prognosehorizont von zehn Minuten auf Basis von Satellitenbildern untersucht [B13].

Tabelle 4.5: Fehlermetriken der PV-Prognose und des Korrekturmodells der PV-Prognose für zwei Beispieltage

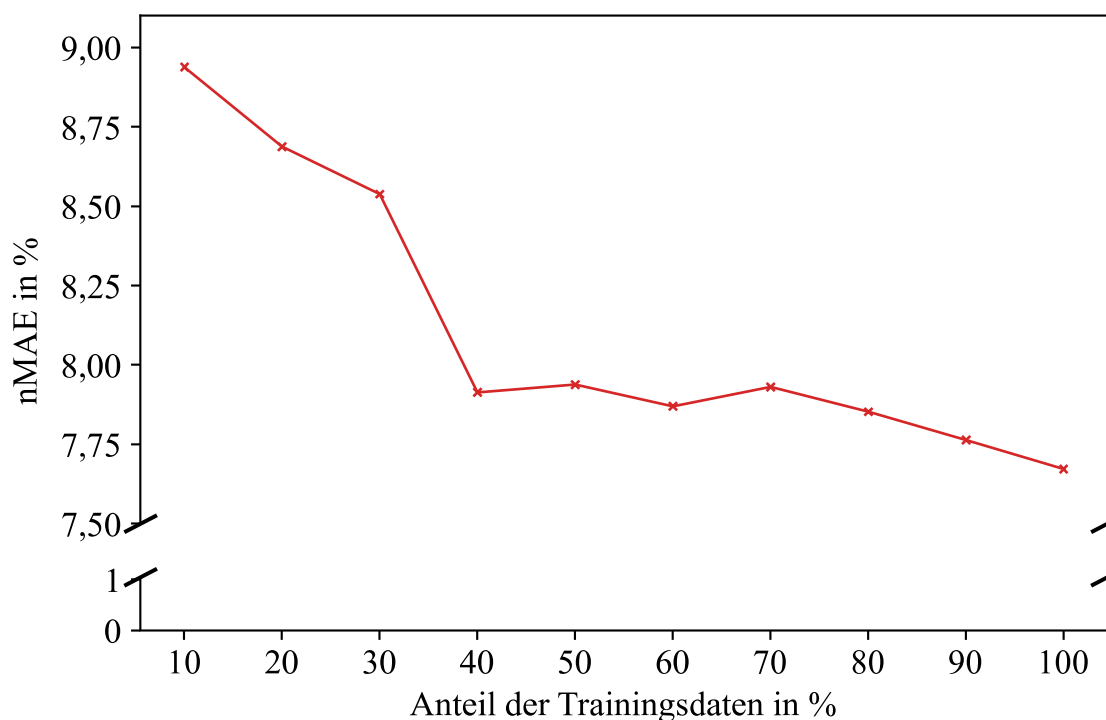
Tag	Prognose	MAE in W	nMAE in %	RMSE in W	nRMSE in %
28.04.2019	24 h Prognose	205,22	9,77	272,22	12,96
	Korrigierte 24 h Prognose	212,09	10,10	321,81	15,32
29.04.2019	24 h Prognose	172,18	8,20	251,34	11,97
	Korrigierte 24 h Prognose	202,77	9,66	318,71	15,18

Die Leistungsfähigkeit des Korrekturmodells zu verschiedenen Wetterbedingungen, insbesondere an wolkigen Tagen, wird auch anhand der in Tabelle 4.6 dargestellten Ergebnisse ersichtlich. In der vorliegenden Tabelle werden die Fehlermetriken des Korrekturmodells für verschiedene CSI-Wertebereiche dargestellt. Der nMAE erreicht an wolkigen Tagen mit 8,74 % seinen Höchstwert, während er an bedeckten Tagen mit 5,93 % am niedrigsten ausfällt. An sonnigen Tagen, wobei in diesem Bereich auch Tage mit geringfügiger Bewölkung enthalten sind, beträgt der nMAE nach der Korrektur 7,90 %.

Tabelle 4.6: Fehlermetriken des Korrekturmodells mit der Clustering-Methode für verschiedene CSI-Wertebereiche für die Testdaten

CSI-Bereiche	MAE in W	nMAE in %	RMSE in W	nRMSE in %
1 - 0,66 (sonnig)	165,90	7,90	216,19	10,29
0,66 - 0,33 (wolkig)	183,62	8,74	242,55	11,55
0,33 – 0 (bedeckt)	124,49	5,93	161,15	7,67

Die Anwendung der Methodik des Korrekturmodells ist nur dann sinnvoll, wenn sie mit einer geringen Anzahl gesammelter Messdaten der PV-Leistung arbeitet. Ziel ist es, die PV-Prognose bereits mit wenigen Messdaten zu verbessern. Dazu muss das Stacking-Ensemble mit einer geringen Trainingsmenge trainiert werden. Daher zeigt Abbildung 4.11 den Prognosefehler der Testdaten als nMAE in Bezug auf den Anteil der Trainingsdaten für das Stacking-Ensemble mit der Clustering-Methode. Somit kann aufgezeigt werden, ab welcher Datenmenge zum Training eine signifikante Optimierung der Prognose erreicht werden kann. Zu diesem Zweck erfolgt eine sukzessive Erhöhung der Trainingsdaten in chronologischer Reihenfolge für das Training in 10-Prozent-Schritten.

**Abbildung 4.11:** Prognosefehler der Testdaten in Bezug auf den Anteil der Trainingsdaten für das Stacking-Ensemble mit der Clustering-Methode

Bei einer Trainingsdatenmenge von 10 %, was einem Zeitraum von etwas weniger als zwei Monaten bzw. 54 Tagen entspricht, zeigt sich, dass der nMAE mit einem Wert von 8,94 % etwas höher ist als der nMAE ohne Korrektur. Dies ist darauf zurückzuführen, dass die Trainingsdaten lediglich 54 chronologisch aufeinanderfolgende Tage umfassen,

während die Testdaten Daten von Tagen aus allen Jahreszeiten enthalten. Das Korrekturmodell ist zu diesem Zeitpunkt noch nicht in der Lage, eine optimierte Prognose zu erzeugen, sondern liefert sogar eine geringfügig schlechtere Prognose. Bei 20 % der Trainingsdaten liegt der nMAE mit einem Wert von 8,69 % niedriger als der nMAE-Wert der ursprünglichen PV-Prognose. Bis zu einem Anteil von 100 % der Trainingsdaten sinkt der nMAE auf 7,67 %. Bei einem Anteil von 40 % der Trainingsdaten ist eine Veränderung des nMAE nur noch in geringfügigen Maßen zu verzeichnen. Eine optimierte Prognose kann folglich bereits mit 40 % der Trainingsdatenmenge erstellt werden, was einem Datensatz von 216 Tagen bzw. in etwa sieben Monaten entspricht. In der Praxis erfolgt in der Regel keine Messung der erzeugten Leistung der dezentralen PV-Anlagen, sodass eine ausreichende Menge an historischen Messdaten für ein sofortiges Training nicht zur Verfügung steht. Sobald jedoch Messdaten in der Anwendung gesammelt werden können, können sie zur Optimierung der Prognose verwendet werden. Nach einer Erfassungsdauer von etwas mehr als einem halben Jahr ist eine signifikante Verbesserung der Prognose möglich. Die Testdaten werden nach dem Zufallsprinzip aus dem Datensatz ausgewählt, sodass eine Abdeckung unterschiedlicher Wetterbedingungen und Jahreszeiten gewährleistet ist. Die Ergebnisse belegen, dass eine geringe Menge an chronologisch gesammelten PV-Leistungsmessdaten ausreicht, um die Vorhersage für verschiedene Wetterbedingungen zu optimieren. Die in dieser Arbeit präsentierte Methode zur Verbesserung von Prognosen der PV-Leistung ist praktikabel, da eine geringe Menge an Messdaten ausreicht, die nach einer kurzen Zeitspanne in der Anwendung gesammelt werden können. Für das unmittelbare Training eines PV-Vorhersagemodells selbst sind keine großen Mengen an PV-Leistungsmessdaten erforderlich.

4.3 Zusammenfassung

Kapitel 4 beschreibt zunächst die Funktionsweise des Korrekturmodells im Betrieb. Es basiert auf einem Stacking-Ensemble verschiedener ML-Modelle. Die prognostizierte Leistung eines Tages einer PV-Anlage sowie weitere zeitliche Angaben werden als Eingabe für das Korrekturmodell verwendet, um saisonale und tägliche Trends zu berücksichtigen. Die Ausgabe des Korrekturmodells ist eine korrigierte prognostizierte PV-Leistung. Das Ziel ist die Reduzierung wiederkehrender Prognosefehler, z. B. durch Verschattung oder Alterung der PV-Anlage, die allein durch die Prognose auf Basis von Wetterdaten nicht erfassbar sind. Dafür werden zusätzlich Leistungsmessdaten der PV-Anlage herangezogen. Abschnitt 4.1 beschreibt die Methodik des Korrekturmodells, einschließlich der Datenaufbereitung, Auswahl der Basismodelle sowie der Hyperparameter. Als Basismodelle des Stacking-Ensembles finden die Methoden KNN, RF, Gradient Boosting, XGBoost und Catboost Regressor Anwendung. Die Basismodelle werden jeweils einzeln mit der Prognose der PV-Leistung sowie zeitlichen Angaben als

Eingangsdaten und der zugehörigen Messung der PV-Leistung als Zieldaten trainiert. Darüber hinaus wird untersucht, ob die Anwendung der Clustering-Methode zur Generierung zusätzlicher Eingangsmerkmale durch Clustern der PV-Prognosedaten zu einer Optimierung der Prognosegenauigkeit führt. Das Metamodell, basierend auf der LR-Methode, kombiniert die Prognosen der Basismodelle zu einer finalen korrigierten PV-Prognose. Zu diesem Zweck wird das Metamodell mit den Prognosen der Basismodelle als Eingangsdaten sowie den realen Messdaten der PV-Leistung als Zieldaten trainiert. Die Stärken der einzelnen Basismodelle werden optimal miteinander kombiniert. In Abschnitt 0 werden die Ergebnisse des Korrekturmodells für die Korrektur der PV-Prognose für die betrachtete PV-Anlage diskutiert. Das Clustering der PV-Prognosedaten und die Verwendung der Clusternummer als Eingangsmerkmal ermöglichen dem Stacking-Ensemble, ähnliche Tagesmuster zu erkennen und Prognosefehler zu reduzieren. Das Stacking-Ensemble mit Clustering erzielt einen nMAE von 7,67 % gegenüber 8,76 % bei der PV-Prognose ohne Korrektur für einen Testzeitraum von 231 Tagen im Zeitraum von 2019 bis 2023. Der Effekt der Korrektur lässt sich vor allem an einer Beispielwoche im Winter erkennen, in der die PV-Anlage verschattet ist und folglich eine geringere Leistung erzeugt. Das Korrekturmodell prognostiziert in dieser Woche eine geringere PV-Leistung über den Tag als die ursprüngliche PV-Prognose ohne Korrektur. Dies lässt darauf schließen, dass das Korrekturmodell den wiederkehrenden Fehler durch die Verschattung identifiziert und diesen reduziert. Prognosefehler, die durch die Wettervorhersage bedingt sind, sind aufgrund ihres zufälligen Charakters schwerer durch das Korrekturmodell zu identifizieren. Besonders bei wechselnder Bewölkung zeigt das Korrekturmodell Schwächen. Hierfür sind Modelle erforderlich, die kürzere Prognosezeiträume als 24 Stunden abdecken und beispielsweise Wolkenzüge am Himmel betrachten. Bereits etwa sechs Monate chronologisch gesammelter PV-Leistungsmessdaten für das Training des Korrekturmodells reichen aus, um die Prognose für verschiedene Wetterbedingungen zu optimieren. Die Methode ist aufgrund der kurzen Erfassungszeit von Messdaten praktikabel.

5 Anwendung der PV-Prognose für das Lademanagement von Elektrofahrzeugen

Kurzfristige Prognosen, welche die erzeugte Leistung einer PV-Anlage in einem Zeitraum von einer Stunde bis zu einer Woche vorhersagen, finden unter anderem Anwendung in Energiemanagementsystemen [11, 12]. In [10] findet die Prognose Anwendung im Energiemanagement von Batteriespeichern im Haushalt, wobei die Netzstabilität gewährleistet und der Eigenverbrauch erhöht werden soll. Dies erfordert eine Anpassung und Verschiebung der Haushaltslasten sowie eine Steuerung der Batterieladung auf Basis der PV-Prognose über den Tag. Ein weiteres Anwendungsgebiet ist die Steuerung der Ladung von Elektrofahrzeugen auf Basis einer Prognose der PV-Erzeugung über den Tag. In [34] wird die Prognose für ein intelligentes Ladesystem auf Haushaltsebene auf Basis einer Optimierung genutzt, um den durch die PV-Anlage erzeugten Strom optimal für das Laden von Elektrofahrzeugen zu nutzen, den Eigenverbrauch zu erhöhen und Lastspitzen zu reduzieren. Dabei wird ein Eigenverbrauch von bis zu 89 % mit fehlerfreier Prognose erreicht. Die Prognose in dieser Arbeit weist einen Prognosehorizont von 24 Stunden auf und ist somit den kurzfristigen Prognosen zuzuordnen. Um den Nutzen der PV-Prognose zu zeigen, wird in diesem Kapitel die Anwendung der PV-Leistungsprognose für das Lademanagement von Elektrofahrzeugen dargestellt. Diese wurde im Rahmen dieser Dissertation in [B15] wissenschaftlich veröffentlicht.

Elektrofahrzeuge, die über längere Zeiträume geparkt sind, wie beispielsweise private Fahrzeuge von Mitarbeitern, die morgens am Arbeitsplatz ankommen und den gesamten Arbeitstag dort verbleiben, sollten nicht gleichzeitig mit hoher Leistung geladen werden. Hohe Ladeleistungen können zu Lastspitzen im elektrischen Netz führen. Für Unternehmen kann dies zu einem höheren Leistungspreis führen, der durch die Lastspitzen entsteht. Zudem würden die Ladestationen im weiteren Verlauf des Tages nicht optimal genutzt, da die Fahrzeuge aufgrund der hohen Ladeleistung bereits früh vollständig geladen wären. Es ist daher sinnvoll, die Leistung einer Ladestation auf mehrere Elektrofahrzeuge zu verteilen und den Ladevorgang zu steuern. Dies kann mithilfe der in dieser Arbeit vorgestellten PV-Prognose erfolgen. In diesem Kapitel werden die Methode der linearen Optimierung sowie die Anwendung der PV-Prognose für den Einsatz als Ladesteuerungssystem für mehrere Elektrofahrzeuge vorgestellt. Die Optimierung zielt darauf ab, mehrere Elektrofahrzeuge innerhalb eines vorgegebenen Zeitraums auf einen gewünschten Ladezustand (*engl. State of Charge, SoC*) zu laden. Dabei wird die von einer PV-Anlage erzeugte Leistung in der Ladeleistung maximiert und Ladespitzen werden minimiert. Die Elektrofahrzeuge werden dabei mit der Leistung einer einzelnen La-

destation geladen. Im Forschungsprojekt Power2Load wurde eine Ladestation mit mehreren Ladepunkten entwickelt, die in der Lage ist, die Leistung auf mehrere Elektrofahrzeuge aufzuteilen [B14]. Die Ladeleistung der an die Ladestation angeschlossenen Fahrzeuge wird entsprechend der durch die Optimierung erstellten Ladepläne angepasst. Indem mehrere Elektrofahrzeuge entweder nicht gleichzeitig oder mit reduzierter Ladeleistung geladen werden, können Ladespitzen reduziert und die Ladeleistung entsprechend der PV-Prognose gesteuert werden. In dieser Arbeit wird ein Lademanagementsystem vorgestellt, welches die Steuerung durch die Anwendung der PV-Prognose ermöglicht.

In Abbildung 5.1 ist eine Übersicht über das Lademanagementsystem für Elektrofahrzeuge dargestellt. Das Modell zur Vorhersage der PV-Leistung basiert auf einer Wettervorhersage sowie technischen Daten der PV-Anlage und prognostiziert die zu erwartende PV-Leistung für den kommenden Tag. Die Ladesteuerung benötigt Nutzereingaben über die geplante Abfahrtszeit, den SoC des Elektrofahrzeugs bei Ankunft, den gewünschten SoC bei Abfahrt sowie das Fahrzeugmodell. Daraus können weitere notwendige Informationen aus einer Datenbank abgeleitet werden wie die maximale und minimale Ladeleistung, der Ladewirkungsgrad und die Batteriekapazität des Elektrofahrzeuges. Die Übermittlung der genannten Daten kann beispielsweise mittels einer App an das Lademanagementsystem erfolgen.

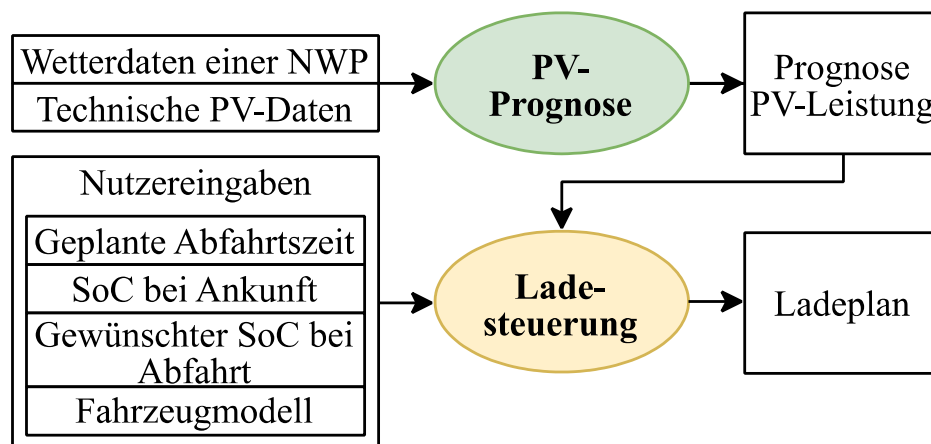


Abbildung 5.1: Übersicht über das Lademanagementsystem für Elektrofahrzeuge

5.1 Definition einer Optimierungsaufgabe mittels linearer Optimierung

Im Folgenden wird zunächst ein Überblick über die Herkunft und Definition der Optimierung gegeben, bevor das konkrete Optimierungsmodell vorgestellt wird. Zu den zentralen Methoden des Operations Research (OR) gehören die Optimierung und Simulation [95]. Das OR befasst sich mit der Entwicklung und Anwendung mathematischer

Modelle sowie deren numerischer Lösung zur Unterstützung von Entscheidungsprozessen [95, 96]. Beispiele für Methoden aus dem Bereich des OR sind die lineare und die nichtlineare Optimierung [96]. Diese Ansätze finden breite Anwendung in betriebs- und volkswirtschaftlichen Fragestellungen, wie bei der Kostenminimierung oder Gewinnmaximierung, aber auch im natur- und ingenieurwissenschaftlichen Bereich [96].

Im Allgemeinen lässt sich eine Optimierungsaufgabe durch eine Zielfunktion beschreiben, die maximiert oder minimiert werden soll. Die Freiheitsgrade des Systems werden durch Entscheidungsvariablen definiert, die während der Optimierung entsprechend der Zielfunktion angepasst werden [95, 97]. Das System unterliegt jedoch gewissen Einschränkungen, die aus den physikalischen Gesetzen resultieren, denen es unterliegt. Dies wird über Nebenbedingungen in Form von Gleichungen und Ungleichungen formuliert. Die zugrundeliegenden Gleichungen setzen sich aus Variablen und festen Parametern zusammen. Die Variablen werden so angepasst, dass ein optimaler Wert gemäß der Zielfunktion und den Nebenbedingungen erreicht wird. Die Lösung solcher Optimierungsprobleme erfolgt in der Regel durch iterative numerische Verfahren. Dabei wird ausgehend von einer anfänglichen Schätzung der Lösung schrittweise eine Verbesserung der Lösung generiert [97]. Ein Optimierungsmodell wird als lineare Optimierung (*engl. Linear Programming, LP*) bezeichnet, wenn sowohl die Zielfunktion als auch alle Nebenbedingungen lineare Kombinationen der Entscheidungsvariablen sind [95, 97]. Enthalten die Zielfunktion und/oder die Nebenbedingungen Nichtlinearitäten, wird die Optimierung als nichtlineare Optimierung (*engl. Nonlinear Programming, NLP*) bezeichnet. Da in dieser Arbeit sowohl die Zielfunktion als auch die Nebenbedingungen linear sind, wird ein LP-System definiert.

Ein definiertes Optimierungsproblem besitzt in der Regel mehrere gültige Lösungen. Eine mögliche Lösung besteht in der Auswahl einer Wertekombination für die Entscheidungsvariablen, wobei alle Nebenbedingungen eingehalten werden. Das Ziel besteht in der Findung einer Lösung, welche im Sinne der Zielfunktion optimal ist [95]. Die Simplex-Methode stellt die älteste und bekannteste Lösungsmethode für LP-Modelle dar [95]. Bei dieser Methode erfolgt die Suche nach einer Lösung des Optimierungsproblems innerhalb des durch die Nebenbedingungen definierten Zulässigkeitsbereichs. Die Suche startet an einem Startpunkt, von dem aus eine Folge von Eckpunkten mit wachsenden Zielfunktionswerten gebildet wird [96]. Das Resultat ist ein zulässiges Polyeder, wobei sich die optimale Lösung in einer Ecke befindet [95]. Nach einer definierten Anzahl von Schritten soll eine optimale Lösung des Optimierungsproblems gefunden worden sein [96]. Allerdings zeigt sich beim Simplex-Verfahren eine exponentielle Zunahme der Laufzeit mit der Größe des Problems [95, 96]. Meistens können Innere-Punkte-Verfahren bei LP-Problemen schneller eine Lösung finden als das Simplex-Verfahren [95]. Im Gegensatz zum Simplex-Verfahren erfolgt die Lösungsfindung beim Innere-Punkte-Verfahren iterativ durch das Innere des Polyeders [96]. Allerdings kann

keine Methode als die beste bezeichnet werden und beide Methoden haben ihre Berechtigung entsprechend verschieden formulierter Optimierungsprobleme [95].

Das in dieser Arbeit behandelte Optimierungsproblem wird mithilfe des Python-basierenden Software-Pakets Pyomo in Kombination mit dem Solver GLPK definiert. Pyomo ermöglicht die strukturierte Definition und Analyse einer Vielzahl von Optimierungsmethoden. GLPK (kurz für GNU Linear Programming Kit) ist ein frei verfügbares sowie weit verbreitetes Software-Paket, welches die Lösung von linearen Optimierungsproblemen ermöglicht [98]. Das Lösungsverfahren wird entsprechend dem definierten Optimierungsproblem automatisch ausgewählt.

Die in der Formulierung des Optimierungsproblems verwendeten Parameter sind in Tabelle 5.1 dargestellt. Die Optimierung erfolgt für jeden Zeitschritt i über die Länge der PV-Prognose T . Des Weiteren erfolgt die Optimierung für jedes Elektrofahrzeug ω innerhalb der Menge aller Elektrofahrzeuge Ω . Die gewünschte Dauer der Ladung jedes Elektrofahrzeugs ω wird mit jedem Zeitschritt i in der Menge $Y(\omega)$ definiert. Diese Zeitschritte befinden sich innerhalb des Rahmens von der Menge T , da die Elektrofahrzeuge innerhalb der Länge der PV-Prognose, welche für einen Zeitraum von 24 Stunden erstellt wird, geladen werden sollen.

Tabelle 5.1: Parameter der Optimierung

Parameter	Beschreibung	Einheit
$P_{\text{Prog,PV}}$	Prognostizierte PV-Leistung	kW
$P_{\text{Ladesäule}}$	Leistung einer Ladesäule	kW
$P_{\text{Lade,min}}$	Minimale Ladeleistung	kW
$P_{\text{Lade,max}}$	Maximale Ladeleistung	kW
K_{ap}	Batteriekapazität	kWh
SoC_0	SoC zu Beginn der Ladung	%
SoC_{Ziel}	Gewünschter SoC nach Ende der Ladung	%
η_{Lade}	Ladewirkungsgrad	1
Ω	Menge der ω angeschlossenen Elektrofahrzeuge	-
T	Menge der Zeitschritte i in der Länge der PV-Prognose	-
$Y(\omega)$	Menge der Zeitschritte i für die Länge der Ladung des jeweiligen Elektrofahrzeuges ω innerhalb der Länge der PV-Prognose	-
$\Delta\tau$	Zeitauflösung	h
Φ_1	Gewicht 1	1
Φ_2	Gewicht 2	1

Die maximale und minimale Ladeleistung, der Ladewirkungsgrad sowie die Batteriekapazität sind über das entsprechende Elektrofahrzeug bekannt. Die Kenntnis des SoC der Batterie zu Beginn der Ladung sowie des gewünschten SoC zum Ende der Ladung und des Elektrofahrzeugs sowie der Aufenthaltsdauer ist erforderlich, damit diese Daten der Ladesteuerung übermittelt werden können. Diese Daten müssen über den Nutzer beispielsweise über eine Eingabe in eine App dem Ladesystem mitgeteilt werden.

In Tabelle 5.2 sind die Variablen dargestellt, die während der Optimierung der Zielfunktion entsprechend berechnet werden. Im Rahmen der vorliegenden Optimierungsaufgabe werden die Variablen $P_{PVinBatt}$ und $P_{NetzinBatt}$ entsprechend berechnet, um die Zielfunktion zu optimieren. Die Ladeleistung der Elektrofahrzeuge setzt sich aus den beiden Komponenten $P_{PVinBatt}$ und $P_{NetzinBatt}$ zusammen. Sie geben den Anteil der PV-Leistung sowie der Leistung aus dem angeschlossenen elektrischen Netz an der Ladeleistung der Elektrofahrzeuge an. Die beiden Variablen sind sowohl zeit- als auch fahrzeugabhängig und bekommen in jedem Zeitschritt i sowie für jedes Elektrofahrzeug ω einen Wert zugewiesen.

Tabelle 5.2: Variablen der Optimierung

Variable	Beschreibung	Einheit
$P_{PVinBatt}$	Anteil PV-Leistung an der Ladeleistung	kW
$P_{NetzinBatt}$	Anteil Leistung aus dem Netz an der Ladeleistung	kW

Das Ziel der Optimierung besteht in der Maximierung des Anteils der von PV-Anlagen erzeugten Leistung an der Ladeleistung. Ein Bezug von Leistung aus dem elektrischen Netz ist nur dann erforderlich, wenn die Leistung der PV-Anlage nicht ausreichend ist. Die nachfolgende Formel (5.1) zeigt die Zielfunktion, welche die beiden Variablen $P_{PVinBatt}$ und $P_{NetzinBatt}$ beinhaltet. Die Zielfunktion sieht eine Maximierung der Gesamtenergie vor, die erforderlich ist, um die angeschlossenen Elektrofahrzeuge innerhalb der vorgegebenen Zeitspanne auf den gewünschten SoC zu laden. Dies gewährleistet, dass die Elektrofahrzeuge unter den gegebenen Bedingungen möglichst bis zum erzielten SoC geladen werden. In manchen Fällen kann es vorkommen, dass die Elektrofahrzeuge nicht in der vorgegebenen Zeit auf den gewünschten SoC geladen werden können. Dennoch wird durch diese Zielfunktion eine Maximierung der Gesamtenergie zum Laden der Elektrofahrzeuge angestrebt. Die Gewichte Φ_1 und Φ_2 sorgen dafür, dass der Anteil der Energie, welche durch die PV-Anlage erzeugt wird, gleichzeitig maximiert wird. Das Gewicht Φ_1 muss sehr groß und das Gewicht Φ_2 sehr klein gewählt werden.

$$\max \left(\sum_{i \in T} \sum_{\omega \in \Omega} (\Phi_1 \cdot P_{PVinBatt}^{i,\omega} + \Phi_2 \cdot P_{NetzinBatt}^{i,\omega}) \cdot \Delta\tau \right) \quad (5.1)$$

Formel (5.2) stellt die erste Nebenbedingung dar. Die Ladeleistung für jedes Elektrofahrzeug, welche sich aus dem PV-Anteil sowie dem Netzanteil zusammensetzt, wird über die Ladezeit aufsummiert und mit der Zeitauflösung $\Delta\tau$ multipliziert. Im Rahmen dieser Arbeit erfolgt die Verwendung von Zeitreihen mit einer Auflösung von zehn Minuten. Dies resultiert aus der Tatsache, dass die Messdaten der PV-Anlage mit einer zeitlichen Auflösung von zehn Minuten vorliegen. Daher wird auch die PV-Prognose auf eine Auflösung von zehn Minuten linear interpoliert. Die Zeitauflösung $\Delta\tau$ ist demnach 1/6 h. Die Energie muss der Kapazität der Batterie Kap der Elektrofahrzeuge entsprechen oder diese unterschreiten, damit die Elektrofahrzeuge nicht über die Kapazität der Fahrzeugbatterie hinaus geladen werden. Der SoC zu Beginn des Ladevorgangs SoC_0 sowie der Ziel-SoC SoC_{Ziel} werden gleichermaßen berücksichtigt genauso wie der Ladewirkungsgrad η_{Lade} .

$$\sum_{i \in T} \eta_{\text{Lade}}^{\omega} \cdot (P_{\text{PVinBatt}}^{i,\omega} + P_{\text{NetzinBatt}}^{i,\omega}) \cdot \Delta\tau \leq \text{Kap}^{\omega} - \text{Kap}^{\omega} \cdot \text{SoC}_0^{\omega} - \text{Kap}^{\omega} \cdot (1 - \text{SoC}_{\text{Ziel}}^{\omega}) \quad \forall \omega \in \Omega \quad (5.2)$$

Die zweite Nebenbedingung gewährleistet, dass der Anteil der PV-Leistung an der Ladeleistung aller Elektrofahrzeuge in jedem Zeitschritt i kleiner oder gleich der vorhergesagten PV-Leistung $P_{\text{Prog,PV}}$ ist. Dies ist in Formel (5.3) dargestellt.

$$\sum_{\omega \in \Omega} P_{\text{PVinBatt}}^{i,\omega} \leq P_{\text{Prog,PV}}^i \quad \forall i \in T \quad (5.3)$$

Die Ladeleistung jedes angeschlossenen Elektrofahrzeugs darf die Leistung der Ladesäule insgesamt nicht überschreiten. Dies wird durch die Nebenbedingung in (5.4) sichergestellt. Die Begrenzung der Ladeleistung aller Elektrofahrzeuge auf die Leistung der Ladesäule bedingt, dass die resultierende Ladespitze maximal der Leistung der Ladesäule entspricht. Die angeschlossenen Elektrofahrzeuge teilen sich die Leistung einer Ladesäule.

$$P_{\text{PVinBatt}}^{i,\omega} + P_{\text{NetzinBatt}}^{i,\omega} \leq P_{\text{Ladesäule}} \quad \forall \omega \in \Omega, \forall i \in T \quad (5.4)$$

Die maximale Ladeleistung $P_{\text{Lade,max}}$ eines jeden Elektrofahrzeuges darf nicht überschritten werden, wenn sich die Ladung innerhalb der gewünschten Ladedauer befindet. Dies ist in der Nebenbedingung in (5.5) dargestellt. Die Ladeleistung ist gleich Null, wenn der Zeitschritt i nicht in der Menge $Y(\omega)$ für die gewünschte Ladedauer zu finden ist.

$$\begin{aligned}
 \sum_{i \in T} P_{PVinBatt}^{i,\omega} + P_{NetzinBatt}^{i,\omega} &\leq P_{Lade,max}^{\omega}, \forall i \in Y(\omega) \\
 \sum_{i \in T} P_{PVinBatt}^{i,\omega} + P_{NetzinBatt}^{i,\omega} &= 0, \text{ sonst} \\
 &\forall \omega \in \Omega
 \end{aligned} \tag{5.5}$$

Des Weiteren ist zu berücksichtigen, dass die Elektrofahrzeuge über eine minimale Ladeleistung $P_{Lade,min}$ verfügen, unterhalb derer eine Ladung der Fahrzeuge nicht mehr möglich ist. Im Falle einer Ladeleistung von Null in der Mitte des Ladevorgangs ist eine selbstständige Wiederaufnahme des Ladevorgangs durch die Elektrofahrzeuge möglich. Dies setzt voraus, dass die Ladeleistung entweder größer oder gleich der minimalen Ladeleistung ist oder gleich Null. Die Erfüllung dieser Nebenbedingung ist durch Formel (5.6) gewährleistet.

$$\begin{aligned}
 \sum_{i \in T} P_{PVinBatt}^{i,\omega} + P_{NetzinBatt}^{i,\omega} &\geq P_{Lade,min}^{\omega} \quad \forall \\
 \sum_{i \in T} P_{PVinBatt}^{i,\omega} + P_{NetzinBatt}^{i,\omega} &= 0 \\
 &\forall \omega \in \Omega
 \end{aligned} \tag{5.6}$$

Die in (5.7) und (5.8) formulierten Nicht-Negativitätsbedingungen gewährleisten, dass der Anteil der PV-Leistung sowie der Anteil der Leistung aus dem elektrischen Netz zu keinem Zeitpunkt negative Werte annehmen können.

$$\begin{aligned}
 P_{PVinBatt}^{i,\omega} &\geq 0 \\
 &\forall \omega \in \Omega, \forall i \in T
 \end{aligned} \tag{5.7}$$

$$\begin{aligned}
 P_{NetzinBatt}^{i,\omega} &\geq 0 \\
 &\forall \omega \in \Omega, \forall i \in T
 \end{aligned} \tag{5.8}$$

5.2 Ergebnisse der Elektrofahrzeug-Ladesteuerung auf Basis der PV-Prognose

Zur Validierung des Lademanagementsystems auf Basis eines LP-Systems und der PV-Prognose wird die Ladesteuerung mit sechs Elektrofahrzeugen getestet. Das untersuchte Szenario umfasst die Ladung privater Elektrofahrzeuge von Mitarbeitern eines Unternehmens auf dem Unternehmensgelände während der Arbeitszeit. Die Auswahl von sechs Elektrofahrzeugen erfolgt aufgrund der Ladesäule aus dem Forschungsprojekt

Power2Load, welche eine Verteilung der Leistung der Ladesäule von 22 kW auf sechs Fahrzeuge ermöglicht. Für die Definition eines Arbeitstages wird die Ankunft der Mitarbeiter zwischen 07:00 und 8:30 Uhr am Morgen gewählt, wobei eine Aufenthaltsdauer zwischen 8 und 9,5 Stunden angenommen wird. In diesem Zeitraum können die Elektrofahrzeuge über den Tag verteilt geladen werden. Bei Ankunft weisen die Elektrofahrzeuge unterschiedliche SoC-Werte auf, wobei angenommen wird, dass die Werte zwischen 50 % und 80 % liegen. Die angestrebte Zielvorgabe für den SoC-Wert am Ende des Ladevorgangs ist bei allen Elektrofahrzeugen 100 %. Die Angaben zu den einzelnen Elektrofahrzeugen sind in Tabelle 5.3 aufgeführt. Zur Vereinfachung wird angenommen, dass die Elektrofahrzeuge alle dem gleichen Fahrzeugtyp angehören. Hier werden die Fahrzeugwerte vom Renault Zoe genutzt. Die nutzbare Kapazität der Batterie beträgt 52 kWh und es wird ein Ladewirkungsgrad von 90 % angenommen [34, 99]. Die Ladeleistung liegt zwischen 2,3 kW und 22 kW [99].

Tabelle 5.3: Übersicht über die Angaben der sechs betrachteten Elektrofahrzeuge

Elektrofahrzeug	SoC ₀ in %	Uhrzeit Ankunft	Aufenthaltsdauer in min
EV 1	60	07:30	540
EV 2	60	07:50	480
EV 3	80	08:20	540
EV 4	50	07:10	570
EV 5	80	08:10	540
EV 6	70	08:20	480

Im Folgenden wird das gesteuerte Laden mit dem ungesteuerten Laden für verschiedene Szenarien verglichen. Beim gesteuerten Laden gibt es lediglich eine Ladesäule, die eine Leistung von 22 kW an die sechs angeschlossenen Elektrofahrzeuge verteilt. Beim ungesteuerten Laden gibt es jeweils sechs Ladesäulen für sechs Elektrofahrzeuge. Die Leistung jeder Ladesäule $P_{\text{Ladesäule}}$ unterscheidet sich in den Szenarien hinsichtlich der Höhe der Leistung, die die Ladesäule zur Verfügung stellen kann. In Bezug auf das einphasige und dreiphasige Laden sowie einen maximalen Ladestrom von 16 A bzw. 32 A resultieren unterschiedliche Leistungen [100, 101]. Die einphasige Ladeleistung von 7,2 kW sowie die dreiphasigen Ladeleistungen von 11 kW und 22 kW werden in dieser Arbeit betrachtet. Die Anzahl der Ladesäulen im jeweiligen Szenario wird durch n definiert. Es werden folgende Szenarien betrachtet und miteinander verglichen:

- Szenario 1: Gesteuertes Laden, $n = 1$, $P_{\text{Ladesäule}} = 22 \text{ kW}$
- Szenario 2: Ungesteuertes Laden, $n = 6$, $P_{\text{Ladesäule}} = 22 \text{ kW}$
- Szenario 3: Ungesteuertes Laden, $n = 6$, $P_{\text{Ladesäule}} = 11 \text{ kW}$
- Szenario 4: Ungesteuertes Laden, $n = 6$, $P_{\text{Ladesäule}} = 7,2 \text{ kW}$

Im Rahmen der Validierung des Lademanagementsystems mit sechs Elektrofahrzeugen wird die in dieser Arbeit präsentierte PV-Prognose herangezogen. Ein Beispieltag wird dargestellt. Als Beispieltag wird der 27.12.2019 genutzt, da dieser Tag bereits in den vorherigen Untersuchungen berücksichtigt wird und an diesem Tag der Prognosefehler nach der Korrektur mit 4,41 % unter dem Durchschnitt liegt. Die PV-Anlage, welche sich auf dem Dach der Hochschule Bielefeld befindet, weist eine Leistung von 2,24 kWp auf. In den betrachteten Szenarien werden die Elektrofahrzeuge mit dem von der PV-Anlage erzeugten Strom geladen. Zu diesem Zweck werden die PV-Messdaten sowie die prognostizierten Daten an diesem Tag normiert. Zudem ist der 27.12.2029 ein Wintertag und die Erzeugung entsprechend zu einem Tag, wo die Sonne höher steht, geringer. In der Folge wird die normierte PV-Leistung über den Tag mit einer anderen Nennleistung für eine größere PV-Anlage multipliziert. Es wird eine Nennleistung von 120 kWp angenommen. Diese Annahmen gewährleisten, dass der an diesem Tag durch die PV-Anlage generierte Strom zur Ladung von sechs Elektrofahrzeugen ausreicht. Zur verbesserten Visualisierung und detaillierteren Darstellung wird die Auflösung der PV-Prognosedaten und Messdaten von einer Stunde auf zehn Minuten interpoliert. Folglich werden auch die durch die Ladesteuerung generierten Ladepläne der sechs Elektrofahrzeuge eine Auflösung von zehn Minuten aufweisen.

Die Abbildung 5.2 stellt den SoC-Wert beim gesteuerten Laden von sechs Elektrofahrzeugen an einer Ladesäule mit 22 kW und beim ungesteuerten Laden der Elektrofahrzeuge an sechs Ladesäulen mit 22 kW dar. Dabei werden die Elektrofahrzeuge gemäß der englischen Bezeichnung für Elektrofahrzeuge (*engl. Electric Vehicle, EV*) abgekürzt. In beiden Fällen werden die Fahrzeuge auf einen SoC-Wert von 100 % geladen. Es ist zu sehen, dass der SoC-Wert beim gesteuerten Laden erst nach einer längeren Zeitspanne erreicht wird. Eine Reduktion der Ladeleistung beim gesteuerten Laden bedingt, dass die Elektrofahrzeuge über einen längeren Zeitraum geladen werden. Beim gesteuerten Laden ist um 16:20 Uhr der SoC-Wert von 100 % aller Fahrzeuge erreicht. Im Gegensatz dazu werden die Elektrofahrzeuge beim ungesteuerten Laden mit einer Leistung von 22 kW bis zur vollständigen Ladung geladen, was zu einer schnelleren Erreichung des SoC-Wertes von 100 % führt. Bereits um 9 Uhr sind alle Elektrofahrzeuge beim ungesteuerten Laden auf einen SoC von 100 % geladen, während beim gesteuerten Laden zu diesem Zeitpunkt nicht einmal alle Fahrzeuge mit der Ladung begonnen haben. Es zeigt sich, dass beim gesteuerten Laden eine niedrigere Ladeleistung ausreichend ist, um alle Fahrzeuge während der Arbeitszeit vollständig zu laden. Bei einer langen Parkdauer, wie sie während der Arbeitszeit in Unternehmen auftritt, ist eine Ladestation ausreichend, um sechs Fahrzeuge vollständig aufzuladen.

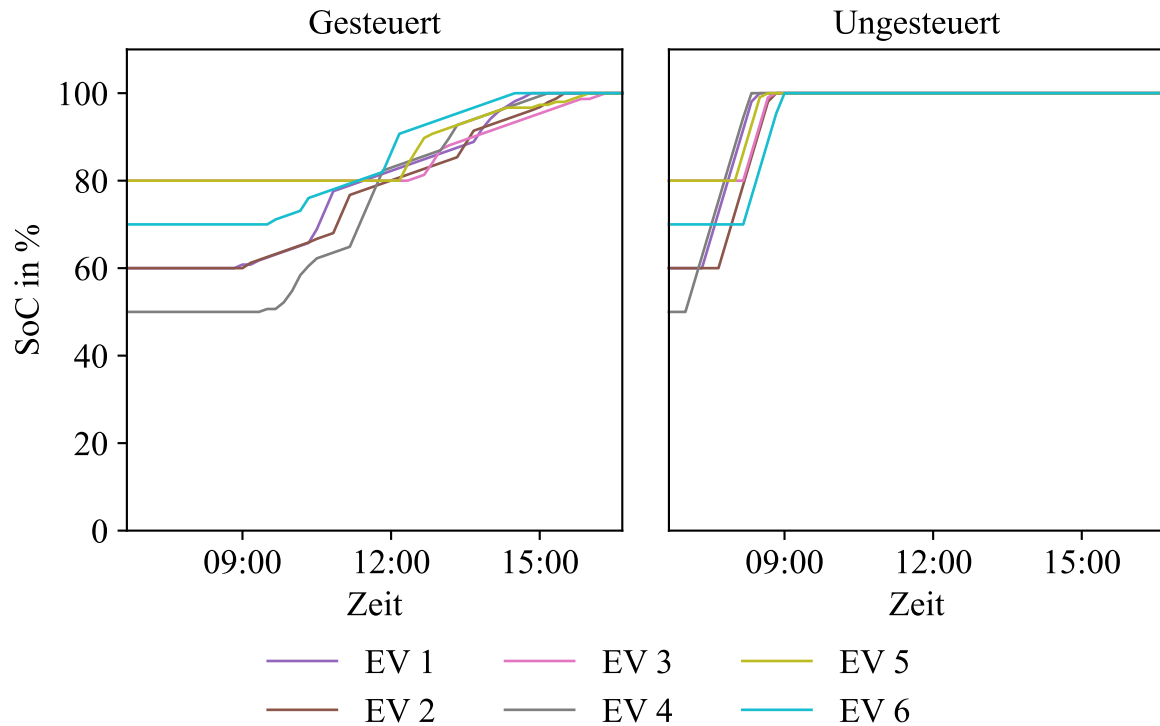


Abbildung 5.2: Ladezustand von sechs Elektrofahrzeugen beim Szenario 1) Gesteuert und 2) Ungesteuert, $P_{\text{Ladesäule}} = 22 \text{ kW}$

In Abbildung 5.3 wird der Gesamtladeplan für alle sechs Elektrofahrzeuge in Bezug auf die vier untersuchten Szenarien dargestellt. Zusätzlich zum Ladeplan werden die korrigierte prognostizierte sowie die gemessene PV-Leistung dargestellt. Die Anteile der Ladeenergie, die während des Ladevorgangs von der PV-Anlage und durch das angeschlossene Stromnetz abgedeckt werden, sind farblich markiert. Beim gesteuerten Laden wird die Leistung der Ladestation über den Tag auf alle Fahrzeuge aufgeteilt, so dass die PV-Leistung im Ladestrom maximiert und die Leistung der Ladesäule von 22 kW nicht überschritten wird (Szenario 1). Dies resultiert in einer Reduktion der Ladelastspitzen. Im Gegensatz dazu resultiert das ungesteuerte Laden jedes Elektrofahrzeugs an einer Ladesäule mit einer Leistung von 22 kW in einer hohen Ladelastspitze von 132 kW (Szenario 2). Zudem ist der Anteil der aus dem Stromnetz bezogenen Energie während des Ladevorgangs im Vergleich zum gesteuerten Laden deutlich höher. Der PV-Anteil ist während des Ladevorgangs gering, da die Elektrofahrzeuge direkt bei der Ankunft geladen werden und die von der PV-Anlage erzeugte Leistung in der Mittagszeit nicht genutzt wird. Im Vergleich dazu erzeugen sechs einzelne Ladesäulen hohe Ladelastspitzen. Auch bei geringeren Leistungen der Ladesäule, wie 11 kW und 7,2 kW, führt das ungesteuerte Laden zu hohen Ladelastspitzen und die erzeugte PV-Leistung wird nicht optimal genutzt (Szenario 3 und 4).

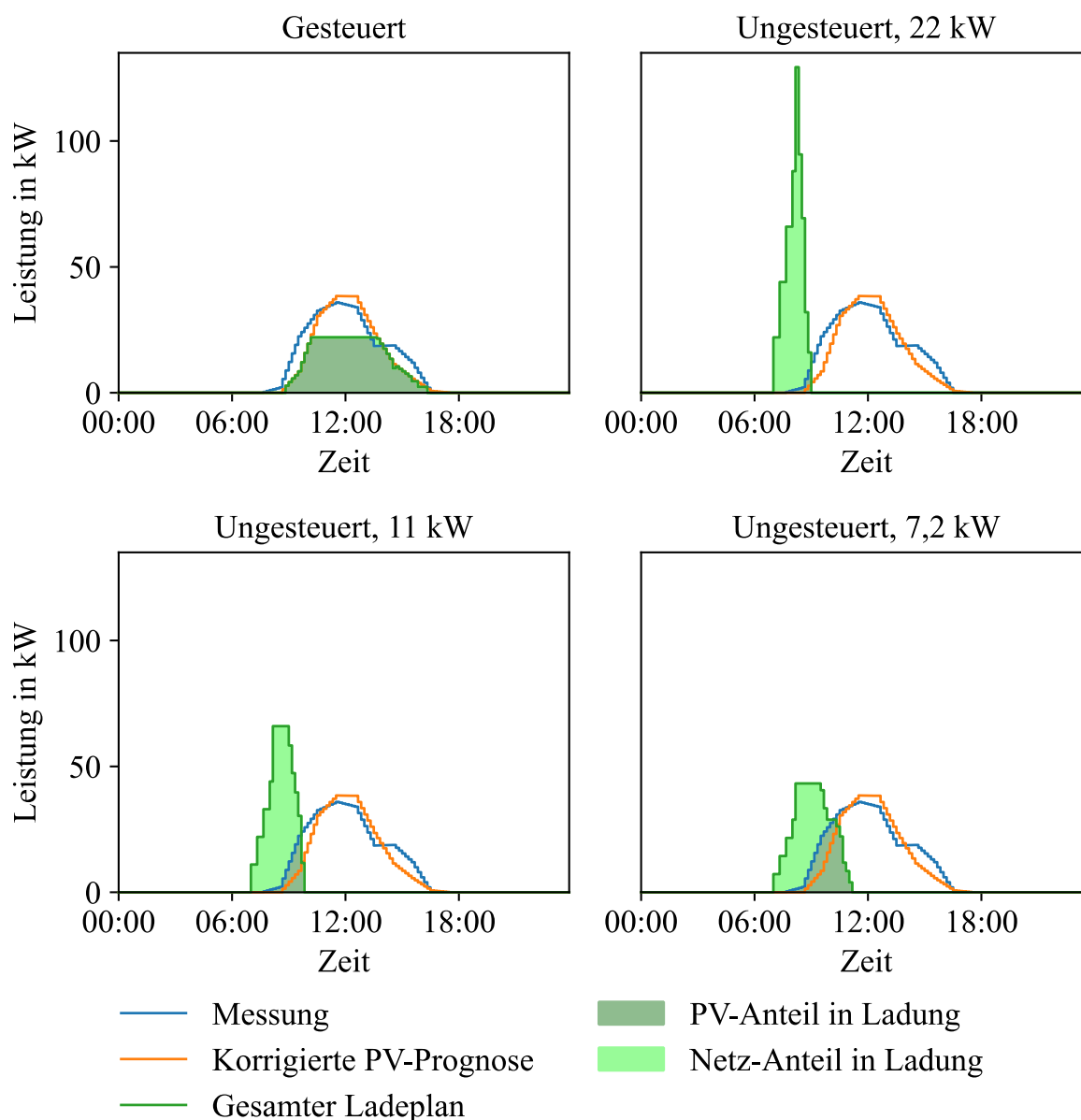


Abbildung 5.3: PV-Messung, korrigierte PV-Prognose und gesamter Ladeplan für sechs Elektrofahrzeuge für die vier verschiedenen Szenarien 1) Gesteuert, 2) Ungesteuert, $P_{\text{Ladesäule}} = 22 \text{ kW}$, 3) Ungesteuert, $P_{\text{Ladesäule}} = 11 \text{ kW}$, 4) Ungesteuert, $P_{\text{Ladesäule}} = 7,2 \text{ kW}$

Die in der vorangegangenen Abbildung dargestellten Erkenntnisse werden in Tabelle 5.4 in numerischer Form nochmals verdeutlicht. Wenn die Elektrofahrzeuge gemäß den optimierten Ladeplänen geladen werden, beträgt der PV-Anteil an der Ladung 98,84 %. Der Anteil der Ladung, der aus dem Stromnetz bezogen wird, beträgt dementsprechend nur 1,16 %. Der Eigenverbrauch des erzeugten Stroms aus der PV-Anlage beim Laden der Elektrofahrzeuge liegt bei 64,75 %. Der verbleibende PV-Strom wird demnach nicht zur vollständigen Aufladung der Elektrofahrzeuge genutzt. Ein minimaler Teil muss aus dem Netz bezogen werden, da die Prognose der PV-Leistung einen gewissen Prognosefehler aufweist. Die Ladepläne werden anhand der prognostizierten PV-Leistung er-

stellt. Wie in der Abbildung 5.3 zu sehen ist, wird die PV-Leistung zu hoch prognostiziert. Dies führt dazu, dass die Ladepläne so erstellt werden, dass diese Leistung auch zur Verfügung steht. Die tatsächliche Messung zeigt jedoch eine geringere Leistung als prognostiziert. Diese Differenz muss durch das Netz zur Verfügung gestellt werden.

Tabelle 5.4: PV-Anteil und Netz-Anteil in der Ladung sowie PV-Eigenverbrauch und Ladespitze in den verschiedenen Szenarien

Szenarien		PV-Anteil in Ladung in %	Netz-An- teil in La- dung in %	PV-Eigen- verbrauch in %	Lade- spitze in kW
Unge- steuert	22 kW	3,13	96,87	2,05	132
	11 kW	14,81	85,19	9,70	66
	7,2 kW	38,87	61,13	25,47	43,20
Ge- steuert	Korrigierte PV-Prognose	98,84	1,16	64,75	22
	PV-Prognose	96,25	3,75	63,07	22

Beim ungesteuerten Laden sind die PV-Anteile an der Ladung deutlich geringer, während der Anteil des Stromnetzes deutlich höher ist. Bei der Nutzung von Ladesäulen mit einer Leistung von 22 kW zum Laden der sechs Elektrofahrzeuge beträgt der PV-Anteil lediglich 3,13 %, während der Anteil aus dem Stromnetz 96,87 % beträgt. Der PV-Anteil steigt bei abnehmender Leistung der Ladesäulen. Dieser Zusammenhang lässt sich durch die längere Dauer des Ladens mit geringer Leistung erklären, wodurch ein größerer Anteil der Ladung des durch die PV-Anlage erzeugten Stroms abgedeckt werden kann. Bei einer Ladung von sechs Elektrofahrzeugen an Ladesäulen mit je 7,2 kW beträgt der PV-Anteil 38,87 %, wobei dieser Wert ebenfalls signifikant unter dem beim gesteuerten Laden liegt. Zudem fällt die Ladespitze beim ungesteuerten Laden mit 132 kW deutlich höher aus als beim gesteuerten Laden mit 22 kW. Selbst bei einer reduzierten Leistung der Ladesäule von 7,2 kW erreicht der ungesteuerte Ladevorgang eine Ladespitze von 43,20 kW, was etwa einer Verdopplung der Ladespitze im Vergleich zum gesteuerten Laden entspricht.

In Abbildung 5.4 ist der gesamte Ladeplan für alle sechs Elektrofahrzeuge dargestellt, der einmal auf Basis der ursprünglichen PV-Prognose sowie einmal auf Basis der korrigierten PV-Prognose erzeugt wird. Der Prognosefehler ist bei der ursprünglichen PV-Prognose größer, was wiederum dazu führt, dass der Anteil, der beim Laden aus dem Netz gezogen wird, größer ist. In diesem Fall beträgt der Netz-Anteil bei der Ladung 3,75 %. Obwohl der Effekt in diesem Beispiel nur geringfügig ist, wird der Nutzen der Korrektur der PV-Prognose an diesem Beispieltag demonstriert. In beiden Fällen des gesteuerten Szenarios beträgt die Ladespitze 22 kW, da die maximale Leistung durch die Ladesäule begrenzt wird.

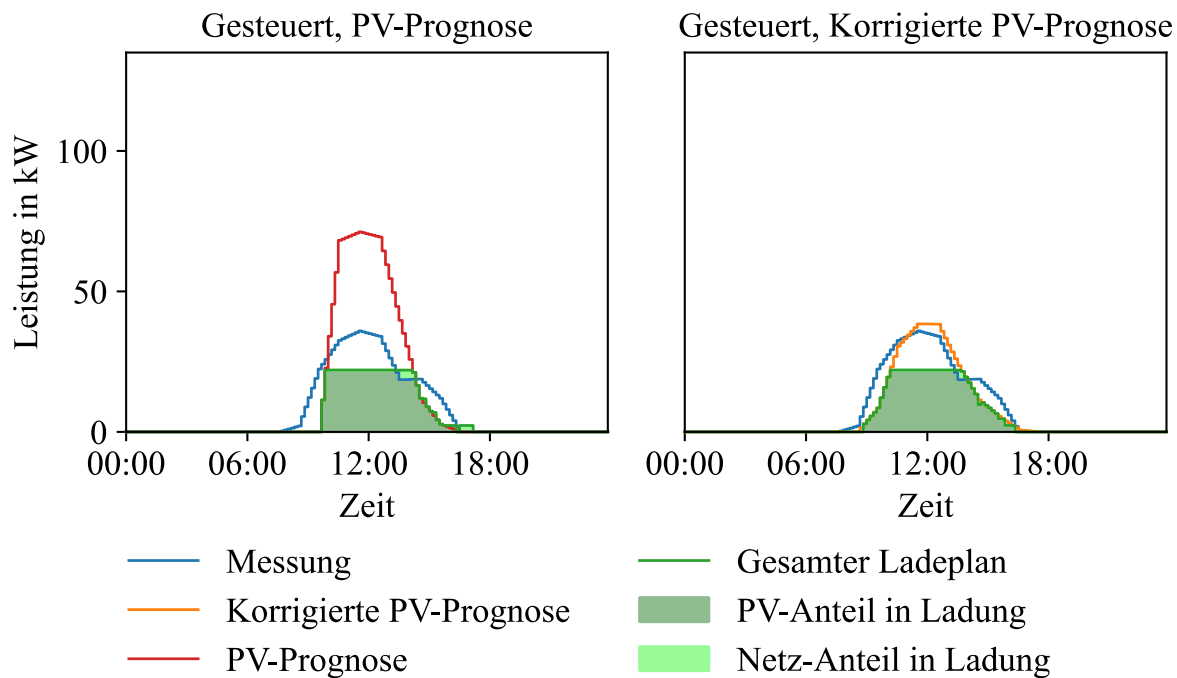


Abbildung 5.4: PV-Messung, PV-Prognose und gesamter Ladeplan für sechs Elektrofahrzeuge im Vergleich mit korrigierter PV-Prognose und ursprünglicher PV-Prognose für das Szenario 1) Gesteuert

5.3 Zusammenfassung

Kurzfristige Prognosen, die die PV-Leistung für Zeiträume von einer Stunde bis zu einer Woche vorhersagen, finden in Energiemanagementsystemen Anwendung. In Kapitel 5 wird die Anwendung der Day-Ahead-PV-Prognose für das Lademanagement von Elektrofahrzeugen aufgezeigt. In Abschnitt 5.1 wird die Methode der linearen Optimierung zur Definition einer Optimierungsaufgabe für die Erstellung von Ladeplänen für Elektrofahrzeuge im Rahmen des Lademanagements auf Basis der PV-Prognose präsentiert. Die Zielfunktion sowie die Nebenbedingungen der Optimierung werden erläutert. Die Optimierung zielt darauf ab, mehrere Elektrofahrzeuge innerhalb eines vorgegebenen Zeitraums auf einen gewünschten SoC zu laden. Dabei wird die von einer PV-Anlage erzeugte Leistung in der Ladeleistung maximiert und Ladespitzen werden minimiert. In Abschnitt 5.2 werden die Ergebnisse der Validierung der Ladesteuerung auf Basis der PV-Prognose präsentiert. Das untersuchte Szenario umfasst die Ladung privater Elektrofahrzeuge von Mitarbeitern eines Unternehmens auf dem Unternehmensgelände während der Arbeitszeit. Eine PV-Anlage auf dem Unternehmensgelände erzeugt den PV-Strom zum Laden der Elektrofahrzeuge. Reicht dieser nicht aus, wird Strom aus dem elektrischen Netz bezogen. Die Untersuchung dieses Szenarios ist von besonderer Relevanz, da Elektrofahrzeuge, die über einen längeren Zeitraum nicht in Gebrauch sind, wie beispielsweise Mitarbeiterfahrzeuge an ihrem Arbeitsplatz, nicht gleichzeitig mit hoher Leistung geladen werden sollten. Andernfalls können Lastspitzen im elektrischen

Netz entstehen. Zudem würden die Ladestationen im weiteren Verlauf des Tages nicht voll ausgelastet sein, da die Fahrzeuge aufgrund der hohen Ladeleistung schnell vollständig geladen wären. Stattdessen wird eine Verteilung der Ladeleistung auf mehrere Fahrzeuge sowie eine Steuerung des Ladevorgangs untersucht. Beim gesteuerten Laden auf Basis einer PV-Prognose für einen Beispieltag mit einem nMAE von 4,41 % wird die Leistung einer Ladestation auf sechs Elektrofahrzeuge so verteilt, dass die PV-Leistung optimal genutzt wird und ein PV-Anteil an der Ladung von 98,84 % erreicht wird. Die maximale Leistung von 22 kW wird dabei nicht überschritten, wodurch Ladespitzen reduziert werden. Demgegenüber resultiert ungesteuertes Laden mit einer Leistung von 22 kW pro Fahrzeug an sechs Ladestationen in einer hohen Ladespitze von 132 kW und einem deutlich geringen PV-Anteil an der Ladung von 3,13 %. Dieser Umstand ist darauf zurückzuführen, dass die PV-Leistung zur Mittagszeit ungenutzt bleibt und die Elektrofahrzeuge direkt bei Ankunft morgens mit der maximalen Leistung geladen werden.

6 Zusammenfassung und Ausblick

Im Folgenden wird eine Zusammenfassung der gesamten Arbeit präsentiert. Zu diesem Zweck wird ein Überblick über das Ziel der Arbeit sowie die Neuheiten in Bezug auf die aktuelle wissenschaftliche Literatur gegeben. Darüber hinaus werden die Methodik sowie die wichtigsten erzielten Ergebnisse dargelegt. Schließlich wird ein Bezug zu den anfangs definierten Forschungsthemen hergestellt, um diese zusammenfassend zu belegen. Das Kapitel endet mit einem auf den Ergebnissen aufbauenden Ausblick auf weitere Forschungsarbeiten.

6.1 Zusammenfassung

Die vorliegende Arbeit adressiert zentrale Herausforderungen, die sich im Zuge der Energiewende aus dem starken Zubau dezentraler PV-Anlagen in Niederspannungsnetzen ergeben. In Regionen mit hoher PV-Dichte kann es insbesondere zur Mittagszeit zunehmend zu Netzengpässen kommen, wenn die lokale Stromerzeugung den Verbrauch übersteigt und überschüssiger Strom zurück ins Mittelspannungsnetz eingespeist wird. Gleichzeitig verläuft der dringend benötigte Ausbau und die Digitalisierung der Verteilnetze nur langsam, wodurch kurzfristig flexible, datenbasierte Lösungsansätze erforderlich sind. Eine vorausschauende Steuerung auf Basis lokaler Einspeiseprognosen von PV-Anlagen bietet hier eine vielversprechende Möglichkeit, Netzengpässe zu vermeiden und Flexibilitätspotenziale durch Lastverschiebung zu nutzen. Vor diesem Hintergrund wird in dieser Arbeit ein neuartiges zweistufiges PV-Prognosemodell entwickelt und untersucht. Trotz eingeschränkter Datenverfügbarkeit ermöglicht es eine zuverlässige lokale Day-Ahead-Prognose der PV-Leistung und ist somit praktikabel einsetzbar. Das PV-Prognosemodell besteht aus einer PV-Prognose und einem Korrekturmodell. Die PV-Prognose kombiniert eine Vorhersage der Bestrahlungsstärke und ein physikalisches PV-Modell. Die Bestrahlungsstärke wird mithilfe eines ANN prognostiziert, das mit historischen Wetterdaten trainiert ist und Wettervorhersagedaten als Eingabe verwendet. Der Prognosehorizont beträgt 24 Stunden. Das PV-Modell berechnet aus der prognostizierten Bestrahlungsstärke auf Basis von anlagenspezifischen Parametern die PV-Leistung. Die PV-Prognose ist dadurch für verschiedene PV-Anlagen anwendbar. Ein Stacking-Ensemble aus ML-Modellen dient als Korrekturmodell, das Leistungsmessdaten berücksichtigt, um wiederkehrende Prognosefehler, die beispielsweise durch Verschattung oder Alterung der PV-Anlage entstehen, zu reduzieren.

Die wissenschaftliche Literatur zu Prognosemethoden für PV-Anlagen ist umfangreich, jedoch oft wenig praxisnah. Ein großer Anteil der wissenschaftlichen Arbeiten basiert

auf idealen Daten, obwohl in Deutschland historische Leistungsmessdaten begrenzt verfügbar sind. Darüber hinaus findet die Nutzung realer Wettervorhersagen nur selten statt. Häufig werden kurze Validierungszeiträume genutzt, was die Aussagekraft bei unterschiedlichen Wetterbedingungen einschränkt. Fehlerquellen von Prognosen wie die Alterung von Anlagen sowie Verschattung werden oft nicht berücksichtigt, sodass Methoden zur Fehlerminimierung in der wissenschaftlichen Literatur selten behandelt werden. In dieser Arbeit wird daher erstmals ein zweistufiges PV-Prognosemodell entwickelt und untersucht, das ausschließlich frei verfügbare Daten wie historische Wetterdaten und kostenlose Wettervorhersagen nutzt. Dies gewährleistet eine praktikable Methodik, die sofort anwendbar ist, ohne dass zuvor Leistungsmessdaten der PV-Anlage gesammelt werden müssen. Durch schrittweises Sammeln von wenigen Leistungsmessdaten können Prognosefehler reduziert werden. Die Validierung des Modells erfolgt anhand eines umfangreichen Datensatzes, welcher einen signifikanten Zeitraum und eine Vielzahl unterschiedlicher Wetterbedingungen abdeckt.

Die Validierung des zweistufigen PV-Prognosemodells erfolgt anhand einer PV-Anlage mit einer Nennleistung von 2,24 kWp, die sich auf dem Dach der Hochschule Bielefeld befindet. Bei einer Day-Ahead-Prognose der PV-Leistung beträgt der nMAE 9,33 % für einen Testzeitraum von 629 Tagen im Zeitraum von 2019 bis 2022. Eine stündliche Aktualisierung der Prognose, bei der die Wettervorhersage stündlich neu in das ANN eingegeben wird, führt zu einer leichten Reduzierung des Fehlers. Das Korrekturmodell reduziert den nMAE auf 7,67 % im Vergleich zu 8,76 % bei einer PV-Prognose ohne Korrektur für einen Testzeitraum von 231 Tagen im Zeitraum von 2019 bis 2023. Der Effekt der Korrektur zeigt sich insbesondere in einer Beispielwoche im Winter, in der die PV-Anlage verschattet ist und folglich eine geringere Leistung erzeugt. Das Korrekturmodell prognostiziert in dieser Woche eine geringere PV-Leistung über den Tag als die ursprüngliche PV-Prognose ohne Korrektur. Dies lässt darauf schließen, dass das Korrekturmodell den wiederkehrenden Fehler durch die Verschattung identifiziert und diesen reduziert. Prognosefehler, die durch die Wettervorhersage bedingt sind, sind aufgrund ihres zufälligen Charakters schwerer durch das Korrekturmodell zu identifizieren. Bereits etwa sechs Monate chronologisch gesammelter PV-Leistungsmessdaten für das Training des Korrekturmodells reichen aus, um die Prognose für verschiedene Wetterbedingungen zu optimieren.

Der Nutzen der PV-Prognose wird über die Anwendung der Day-Ahead-PV-Prognose für das Lademanagement von Elektrofahrzeugen demonstriert. Die auf Optimierung basierende Ladesteuerung unter der Nutzung der PV-Prognose wird für ein Szenario validiert, das die Ladung privater Elektrofahrzeuge von Mitarbeitern eines Unternehmens auf dem Unternehmensgelände während der Arbeitszeit umfasst. Aufgrund der langen Parkzeiten der Elektrofahrzeuge in diesem Szenario bietet sich eine Verteilung der La-

deistung auf mehrere Elektrofahrzeuge über den Tag an. Die Untersuchung dieses Szenarios ist von besonderer Relevanz, da ein gleichzeitiges Laden von Elektrofahrzeugen mit hoher Leistung zu Lastspitzen im elektrischen Netz führen kann. Beim gesteuerten Laden auf Basis einer PV-Prognose für einen Beispieltag mit einem nMAE von 4,41 % wird die Leistung einer Ladestation auf sechs Elektrofahrzeuge so verteilt, dass die PV-Leistung optimal genutzt wird und ein PV-Anteil an der Ladung von 98,84 % erreicht wird. Die maximale Leistung von 22 kW wird dabei nicht überschritten, wodurch Ladespitzen reduziert werden. Demgegenüber resultiert ungesteuertes Laden mit einer Leistung von 22 kW pro Fahrzeug an sechs Ladestationen in einer hohen Ladespitze von 132 kW und einem deutlich geringen PV-Anteil an der Ladung von 3,13 %. Die PV-Leistung zur Mittagszeit bleibt in diesem Fall ungenutzt und die Elektrofahrzeuge werden morgens direkt mit der maximalen Leistung geladen.

Das in dieser Arbeit entwickelte Prognosemodell zielt auf die eingangs dargestellten Herausforderungen der Energiewende ab. Sie bietet eine praktikable Lösung zur Integration volatiler PV-Erzeugung in das Energiesystem, ohne dass flächendeckend vorhandene PV-Leistungsmessdaten aus Smart Meter erforderlich sind. Damit leistet das zweistufige PV-Prognosemodell einen Beitrag zur Digitalisierung der Verteilnetze und kann die Umsetzung von Anwendungen zur Lastverschiebung wie das gesteuerte Laden von Elektrofahrzeugen ermöglichen. Dies ist ein zentraler Baustein für ein intelligentes, flexibles Energiesystem der Zukunft.

6.2 Bezugnahme zu den Forschungsthese

In der vorliegenden Arbeit werden die zu Beginn definierten Forschungsthese untersucht. Im Folgenden werden die Ergebnisse der Arbeit mit den Thesen in Beziehung gesetzt, um eine Prüfung dieser vorzunehmen.

Die lokal erzeugte PV-Leistung einer einzelnen PV-Anlage kann mit Hilfe von ML-Methoden unter der Verwendung von frei verfügbaren Wetterdaten prognostiziert werden.

Die erste Forschungsthese kann durch einen Vergleich der PV-Prognose dieser Arbeit mit verwandten wissenschaftlichen Arbeiten belegt werden. Ziel dieses Vergleichs ist die Einordnung der Prognosequalität der PV-Prognose im aktuellen Stand der Wissenschaft. Die PV-Prognose in dieser Arbeit weist für eine PV-Anlage auf dem Dach der Hochschule Bielefeld einen Prognosefehler als nMAE von 8,76 % ohne Korrektur und von 7,67 % mit Korrektur auf. Um diese Fehlerwerte einzuordnen, ist ein Vergleich mit anderen wissenschaftlichen Arbeiten sinnvoll. Da kein allgemeiner Prognosefehler für den aktuellen Stand der Wissenschaft aus den in Kapitel 2.2 genannten Gründen defi-

niert werden kann, werden einige der untersuchten wissenschaftlichen Arbeiten ausgewählt, um einen Vergleich mit diesen zu ziehen. Zu diesem Zweck werden jene wissenschaftlichen Arbeiten ausgewählt, die ebenfalls den nMAE als Fehlermetrik angeben. In [38] wird ein nMAE von 7,7 % als Prognosefehler für eine Day-Ahead-PV-Prognose angegeben. Die Datenbasis umfasst 157 Tage mit PV-Leistungsmessdaten sowie Wettervorhersagedaten, welche die Bestrahlungsstärke beinhalten. Es ist unklar, ob der Fehler nur für einen Sommertag angegeben ist. Der nMAE in [70] beträgt 7,74 %, während in [76] ein nMAE von 7,72 % angegeben wird. In beiden wissenschaftlichen Arbeiten wird die PV-Prognose ebenfalls auf Basis von PV-Leistungsmessdaten sowie Wettervorhersagedaten wie die Bestrahlungsstärke erstellt, wobei die Daten für ein Jahr vorliegen. In [70] umfasst der Testzeitraum die drei Monate April, Mai und Juni. In [76] werden verschiedene Tage aus unterschiedlichen Jahreszeiten mit unterschiedlichen Wetterbedingungen getestet und der Fehler von jedem Tag angegeben. Der angegebene nMAE von 7,72 % stellt den minimalen nMAE dar, der an einem Tag erreicht wird. In [33] wird für eine Day-Ahead-PV-Prognose ein nMAE von 3,07 % angegeben. Der Datensatz umfasst PV-Leistungsmessdaten sowie Wetterdaten, darunter die Bestrahlungsstärke, über einen Zeitraum von über zwei Jahren. Der Testdatensatz deckt ein ganzes Jahr ab, wobei der Fehler für 32 große PV-Parks zusammengefasst wird. Es ist anzumerken, dass die Nachtstunden in dieser Arbeit nicht entfernt wurden. Die Prognose der Globalstrahlung in [61] erreicht einen nMAE von 6,7 % für eine Prognose für eine Stunde im Voraus unter der Nutzung von Himmelsbildern. Der Testdatensatz umfasst 113 Tage. Für eine PV-Prognose für einen Prognosehorizont von zehn Tagen wird in [67] ein nMAE von unter 7,5 % angegeben. Der Datensatz umfasst einen Zeitraum von einem Monat und beinhaltet PV-Leistungsmessdaten sowie Wetterdaten wie die Bestrahlungsstärke. Der Fehler wird für einen Testdatensatz von zehn Tagen angegeben. Der in der vorliegenden Arbeit ermittelte Prognosefehler ist mit den in anderen Arbeiten ermittelten Fehlern vergleichbar und befindet sich in derselben Größenordnung. Dies stützt die Forschungsthese, die besagt, dass die lokal erzeugte PV-Leistung einer einzelnen PV-Anlage mit ML-Methoden unter der Verwendung von Wetterdaten prognostiziert werden kann. Eine direkte Vergleichbarkeit ist aufgrund der Verwendung unterschiedlicher Testdatensätze nur eingeschränkt möglich. Darüber hinaus gehen die Arbeiten von der Annahme aus, dass alle erforderlichen Daten vorhanden sind, wie beispielsweise die Prognose der Bestrahlungsstärke. Im Gegensatz zu den vorangegangenen Arbeiten wird in dieser Arbeit die Bestrahlungsstärke nicht als gegeben vorausgesetzt. Stattdessen erfolgt eine Prognose mittels eines ANNs über eine klassische Wettervorhersage. Die Anwendbarkeit der PV-Prognose ist dadurch direkt zu Beginn gegeben, ohne dass PV-Messdaten vorab gesammelt werden müssen. Mit einer geringen Anzahl an Messdaten der PV-Leistung wird eine Korrektur der PV-Prognose durchgeführt. Der Vergleich mit anderen wissenschaftlichen Arbeiten zeigt auch, dass bei einer

kurzfristigen Prognose mit einem Prognosehorizont von 24 Stunden stets ein Prognosefehler auftritt, wie auch in der vorliegenden Arbeit. Die Methodik der vorliegenden Arbeit erweist sich im Gegensatz zu anderen wissenschaftlichen Arbeiten als praktikabel.

Der Fehler der PV-Prognose kann durch die Verwendung weniger Messdaten der erzeugten Leistung der PV-Anlage reduziert werden.

Die zweite Forschungsthese kann durch einen Vergleich der PV-Prognose der vorliegenden Arbeit bezüglich der Menge der verwendeten PV-Leistungsmessdaten mit verwandten wissenschaftlichen Arbeiten belegt werden. Ziel dieses Vergleichs ist die Einordnung der Datenmenge für das Korrekturmodell im aktuellen Stand der Wissenschaft. Dadurch kann festgestellt werden, ob die Datenmenge im Vergleich zur Datenmenge anderer wissenschaftlicher Arbeiten gering ist und dadurch einen Vorteil bietet, anstatt direkt mit der geringen Datenmenge ein Prognosemodell zu trainieren. In einer Vielzahl von wissenschaftlichen Arbeiten, wie beispielsweise in [31, 35, 51, 64, 69, 70, 72, 76, 79], wird für die Entwicklung eines Prognosemodells für PV-Anlagen ein Jahr an PV-Leistungsmessdaten als Datenmenge eingesetzt. In [42, 63] liegt die Datenmenge zwischen einem und zwei Jahren und in [33, 65, 73, 74] beträgt die Datenmenge zwischen zwei und drei Jahren. Darüber hinaus finden sich in der Literatur Arbeiten, die eine größere Datenmenge aufweisen, wie in [12] mit mehr als drei Jahren, in [75] mit vier Jahren, in [62] mit fünf Jahren und in [50] mit sieben Jahren. Im Vergleich dazu genügen in der vorliegenden Arbeit etwa sechs Monate chronologisch gesammelter PV-Leistungsmessdaten für das Training des Korrekturmodells aus, um die Prognose für verschiedene Wetterbedingungen zu optimieren. Es gibt vereinzelt weitere Publikationen, die ähnliche geringe Datenmengen nutzen. In [78] werden knapp acht Monate an PV-Leistungsmessdaten verwendet. Es ist jedoch anzumerken, dass sich die Arbeit nicht mit der lokalen PV-Prognose, sondern mit der regionalen Prognose beschäftigt. In [67] wird ein Monat an PV-Leistungsmessdaten genutzt und der Testzeitraum beträgt zehn Tage. In [38] werden 157 Tage an Messdaten der PV-Leistung eingesetzt. In beiden Arbeiten wird davon ausgegangen, dass die Bestrahlungsstärke über eine Wettervorhersage immer vorhanden ist und als Eingangsdaten in das Prognosemodell einfließt. In [36] werden Messdaten der Bestrahlungsstärke von sechs Monaten genutzt. In dieser Arbeit wird ebenfalls das Problem der geringen Datenmenge betitelt, welches auftritt, wenn ein System gerade erst seinen Betrieb aufnimmt. In allen drei Arbeiten wird jeweils nur ein kurzer Testzeitraum betrachtet. Dieser Zeitraum deckt demnach nicht alle Wetterbedingungen und Jahreszeiten ab. Ob die Datenmenge ausreicht, um ein Prognosemodell zu entwickeln, welches verschiedene Wetterbedingungen abdeckt, bleibt fraglich. Zudem müssen auch die wenigen Daten vorab gesammelt werden und die Modelle in den Arbeiten sind nicht direkt zu Beginn einsetzbar. Dieser Vergleich zeigt, dass die These durch die Ergebnisse der vorliegenden Arbeit gestützt werden kann und ein geringer

Datensatz an PV-Leistungsmessdaten für eine Verbesserung der PV-Prognose ausreichend ist. Bereits etwas mehr als sechs Monate chronologisch gesammelter PV-Leistungsmessdaten für das Training des Korrekturmodells reichen aus, um die Prognose für verschiedene Wetterbedingungen zu verbessern. Die Datenmenge ist im Vergleich zur wissenschaftlichen Literatur gering und im Vergleich zu Arbeiten, die ebenfalls geringe Daten verwenden, praktikabel. Das zweistufige Prognosemodell kann direkt angewendet werden und die Prognose der Bestrahlungsstärke wird nicht als gegeben angenommen. Kostenfreie Wettervorhersagen bieten in der Regel keine Prognosen der Bestrahlungsstärke an. Zudem werden die PV-Prognose und das Korrekturmodell an ausreichend Testdaten für verschiedene Wetterbedingungen und Jahreszeiten validiert.

Die Anwendung der PV-Prognose für das Lademanagement von Elektrofahrzeugen ermöglicht eine Maximierung des Anteils von PV-Leistung in der Ladeleistung.

Die dritte Forschungsthese kann durch die Ergebnisse des untersuchten Szenarios der Ladung von privaten Elektrofahrzeugen von Mitarbeitern eines Unternehmens auf dem Unternehmensgelände während der Arbeitszeit belegt werden. Die Verteilung der Ladeleistung auf mehrere Elektrofahrzeuge sowie eine Steuerung des Ladevorgangs erfolgt dabei mittels eines Optimierungsalgorithmus. Beim gesteuerten Laden auf Basis einer PV-Prognose für einen Beispieltag mit einem geringen Prognosefehler wird die Leistung einer Ladestation auf sechs Elektrofahrzeuge so verteilt, dass die PV-Leistung optimal genutzt wird und ein PV-Anteil an der Ladung von 98,84 % erreicht wird. Die maximale Leistung von 22 kW wird dabei nicht überschritten, wodurch Ladespitzen reduziert werden. Demgegenüber resultiert ungesteuertes Laden mit einer Leistung von 22 kW pro Elektrofahrzeug an sechs Ladestationen in einer hohen Ladespitze von 132 kW und einem deutlich geringen PV-Anteil an der Ladung von 3,13 %. Dies ist darauf zurückzuführen, dass die PV-Leistung zur Mittagszeit ungenutzt bleibt und die Elektrofahrzeuge direkt bei Ankunft morgens mit der maximalen Leistung geladen werden. Die Anwendung der PV-Prognose für das Lademanagement von Elektrofahrzeugen zeigt, dass es möglich ist, den Anteil von PV-Leistung in der Ladeleistung zu maximieren. Unter dieser Annahme eines geringen Prognosefehlers und eines sonnigen Tages zeigt sich der entscheidende Vorteil der PV-Prognose für das Lademanagement von Elektrofahrzeugen. Darüber hinaus kann die durch Elektrofahrzeuge bereitgestellte Flexibilität gezielt genutzt werden, um Lasten zu verschieben. Dadurch lässt sich die Netzauslastung besser steuern, insbesondere zur Mittagszeit, wenn viele PV-Anlagen gleichzeitig hohe Leistungen einspeisen. Dies kann dazu beitragen, potenzielle Überlastungen im Netz zu reduzieren oder sogar zu vermeiden.

6.3 Ausblick

Die Optimierung und Weiterentwicklung von PV-Leistungsprognosen eröffnet ein breites Feld für zukünftige Forschungsarbeiten. Basierend auf den Ergebnissen bietet sich die Untersuchung der folgenden Aspekte an.

Um die universelle Anwendbarkeit der PV-Prognose zu überprüfen, sollten Tests mit weiteren PV-Anlagen durchgeführt werden, die sich an verschiedenen Standorten befinden. Dies würde die Robustheit des Modells gegenüber unterschiedlichen klimatischen Bedingungen und geografischen Gegebenheiten belegen. In der vorliegenden Arbeit wird die Annahme getroffen, dass ähnliche Wetterbedingungen an der Wetterstation, an der die historischen Wetterdaten für das Training des ANN bezogen werden, und dem Standort der betrachteten PV-Anlage herrschen, da beide Standorte in einem Gebiet liegen, in dem die mittlere Jahressumme der Globalstrahlung von 1991 bis 2020 ähnlich ist. Es gilt zu prüfen, inwiefern die Nähe zur Wetterstation die Prognose bzw. den Prognosefehler beeinflusst. Darüber hinaus kann auch untersucht werden, ob eine Anwendbarkeit auch in Gebieten, wo die Jahressumme der Globalstrahlung höher oder niedriger ist, möglich ist. Es ist davon auszugehen, dass eine Prognose ebenfalls möglich ist, aber voraussichtlich mit größeren Fehlern, da das ANN für die Prognose der Bestrahlungsstärke nicht alle dort auftretenden Wetterbedingungen, wie zum Beispiel höhere Bestrahlungsstärkewerte, gelernt hat. Um die Wirksamkeit des Korrekturmodells weiter zu belegen, sollten auch PV-Anlagen untersucht werden, die älter oder stärker von Verschattung betroffen sind, als die PV-Anlage in der vorliegenden Arbeit. Insbesondere bei der Korrektur von wiederkehrenden Fehlern wie Alterung und Verschattung zeigt sich der Effekt des Korrekturmodells.

In der vorliegenden Arbeit wird die Anwendbarkeit der PV-Leistungsprognose im Zusammenhang mit dem Lademanagement von Elektrofahrzeugen untersucht. In zukünftigen Arbeiten sollte das Prognosemodell unter realen Bedingungen zur Ladung von Elektrofahrzeugen getestet werden, da in der Simulation von idealen Ladeverläufen ausgegangen wird. Dies ermöglicht die Ermittlung des Anteils des von PV-Anlagen erzeugten Stroms in der Ladeleistung von Elektrofahrzeugen unter realen Bedingungen, insbesondere an Tagen mit hohen Prognosefehlern. Die Ergebnisse liefern Aufschluss darüber, wie hoch der Prognosefehler in der Anwendung ausfallen darf und ob eine präzisere Gestaltung des Prognosemodells erforderlich ist. Darüber hinaus könnte der Beitrag zu Vermeidung von Netzüberlastungen unter realen Bedingungen untersucht werden. Neben dem Lademanagement für Elektrofahrzeuge könnten weitere Anwendungsbereiche der PV-Prognose im Bereich des Energiemanagements untersucht werden. Ein konkretes Beispiel ist die Steuerung von Batteriespeichern in Haushalten. Ins-

besondere im Hinblick auf netzdienliche Zwecke erweist es sich als sinnvoll, die Batteriespeicher gezielt vor der Mittagsspitze zu entladen, die durch die gleichzeitige Einspeisung mehrerer PV-Anlagen entsteht. Dadurch können die Batteriespeicher die Spitzen abfangen und zu Zeiten mit höchsten Einspeisungen geladen werden. Auf diese Weise wird das Niederspannungsnetz vor Erzeugungsspitzen geschützt. Dafür muss die erzeugte PV-Leistung über den Tag prognostiziert werden. Dieser Ansatz wurde auch im Forschungsprojekt AI4DG betrachtet, in dem KI-basierte Prognoseverfahren zur netzdienlichen Integration dezentraler Batteriespeicher in Haushalten entwickelt und erprobt wurden [B8, B9, B11, B12]. In diesem Zusammenhang bietet sich auch die Untersuchung des Einsatzes der PV-Prognose angesichts der aktuellen regulatorischen Entwicklungen, wie den Vorgaben aus §14a des EnWG. Dabei könnten PV-Prognosen zur netzdienlichen Steuerung von Verbrauchseinrichtungen beitragen beispielsweise im Zusammenhang mit der Steuerung von Elektrofahrzeugen, Wärmepumpen und Batteriespeichern, um Netzüberlastungen zu vermeiden. Die Anwendungen demonstrieren ein breites Spektrum an Möglichkeiten für eine optimierte Integration von PV-Anlagen in das Energiesystem sowie für neue Lasten, wie beispielsweise Elektrofahrzeuge.

Literaturverzeichnis

- [1] Fraunhofer ISE Harry Wirth, *Aktuelle Fakten zur Photovoltaik in Deutschland*, 04.01.2025. [Online]. Verfügbar unter: www.pv-fakten.de (Zuletzt geprüft am: 20.01.2025).
- [2] Bundesnetzagentur, *Ausbau Erneuerbarer Energien 2024*, 08.01.2025. [Online]. Verfügbar unter: https://www.bundesnetzagentur.de/SharedDocs/Pressemitteilungen/DE/2025/20250108_EE.html (Zuletzt geprüft am: 20.01.2025).
- [3] Bundesministerium für Wirtschaft und Klimaschutz (BMWK), *Aktualisierung des integrierten nationalen Energie- und Klimaplanes*, August 2024. [Online]. Verfügbar unter: https://www.bmwk.de/Redaktion/DE/Publikationen/Energie/20240820-aktualisierung-necp.pdf?__blob=publicationFile&v=8 (Zuletzt geprüft am: 20.01.2025).
- [4] Umweltbundesamt, *Photovoltaik-Dachanlagen*, 26.03.2024. [Online]. Verfügbar unter: <https://www.umweltbundesamt.de/themen/klima-energie/erneuerbare-energien/photovoltaik/photovoltaik-dachanlagen> (Zuletzt geprüft am: 20.01.2025).
- [5] Agora Energiewende, *Die Energiewende in Deutschland: Stand der Dinge 2024: Rückblick auf die wesentlichen Entwicklungen sowie Ausblick auf 2025*. [Online]. Verfügbar unter: https://www.agora-energiewende.de/fileadmin/Projekte/2025/2024-18_DE_JAW24/A-EW_351_JAW24_WEB.pdf (Zuletzt geprüft am: 21.01.2025).
- [6] M. Tretschock, M. Greve, F. Probst, S. Kippelt, C. Wagner, C. Rehtanz, A. Moser, N. Wehbring, T. Offergeld, F. Schmidtke, M. Wahl, M. Zdrallek, M. Popp, K. Kotthaus und R. Schmidt, *Gutachten zur Weiterentwicklung der Strom-Verteilnetze in Nordrhein-Westfalen auf Grund einer fortschreitenden Sektorenkopplung und neuer Verbraucher: Für das Ministerium für Wirtschaft Innovation, Digitalisierung und Energie des Landes Nordrhein-Westfalen*, 2021. [Online]. Verfügbar unter: https://www.wirtschaft.nrw/sites/default/files/documents/210609_nrw_verteilnetzstudie_final.pdf (Zuletzt geprüft am: 21.01.2025).
- [7] Bundesnetzagentur, *Bericht zum Zustand und Ausbau der Verteilernetze 2022*, Juli 2023. [Online]. Verfügbar unter: https://www.bundesnetzagentur.de/SharedDocs/Downloads/DE/Sachgebiete/Energie/Unternehmen_Institutionen/NetzentwicklungUndSmartGrid/ZustandAusbauVerteilernetze2022.pdf?__blob=publicationFile&v=1 (Zuletzt geprüft am: 21.01.2025).

- [8] Fraunhofer-Institut für Energiewirtschaft und Energiesystemtechnik IEE, *Künstliche Intelligenz im Energiesektor*. [Online]. Verfügbar unter: https://kognitive-energie-systeme.de/wp-content/uploads/2023/04/FraunhoferIEE_K-ES_KI_im_Energiesektor_web.pdf (Zuletzt geprüft am: 11.07.2025).
- [9] Bernd Michael Buchholz und Zbigniew Antoni Styczynski, *Smart Grids: Grundlagen und Technologien der elektrischen Netze der Zukunft*, 2. Aufl. Berlin: VDE VERLAG, 2018.
- [10] S. P. Durrani, S. Balluff, L. Wurzer and S. Krauter, "Photovoltaic yield prediction using an irradiance forecast model based on multiple neural networks," *Journal of Modern Power Systems and Clean Energy*, vol. 6, no. 2, pp. 255–267, 2018, doi: 10.1007/s40565-018-0393-5.
- [11] B. D. Dimd, S. Voller, U. Cali and O.-M. Midtgard, "A Review of Machine Learning-Based Photovoltaic Output Power Forecasting: Nordic Context," *IEEE Access*, vol. 10, pp. 26404–26425, 2022, doi: 10.1109/ACCESS.2022.3156942.
- [12] D. Rangelov, M. Boerger, N. Tcholtchev, P. Lämmel and M. Hauswirth, "Design and Development of a Short-Term Photovoltaic Power Output Forecasting Method Based on Random Forest, Deep Neural Network and LSTM Using Readily Available Weather Features," *IEEE Access*, vol. 11, pp. 41578–41595, 2023, doi: 10.1109/ACCESS.2023.3270714.
- [13] Bundesnetzagentur, *Roll-out intelligente Messsysteme: Quartalsweise Erhebungen*. [Online]. Verfügbar unter: <https://www.bundesnetzagentur.de/DE/Fachthemen/ElektrizitaetundGas/NetzzugangMesswesen/Mess-undZaehlwesen/iMSys/start.html> (Zuletzt geprüft am: 21.01.2025).
- [14] Agora Energiewende, *Die Energiewende in Deutschland: Stand der Dinge 2023: Rückblick auf die wesentlichen Entwicklungen sowie Ausblick auf 2024*. [Online]. Verfügbar unter: <https://www.agora-energiewende.de/publikationen/die-energiewende-in-deutschland-stand-der-dinge-2023> (Zuletzt geprüft am: 21.01.2025).
- [15] Bundesnetzagentur, *Messeinrichtungen / Intelligente Messsysteme*. [Online]. Verfügbar unter: <https://www.bundesnetzagentur.de/DE/Vportal/Energie/Metering/start.html> (Zuletzt geprüft am: 21.02.2025).
- [16] S. Stock, D. Babazadeh and C. Becker, "Applications of Artificial Intelligence in Distribution Power System Operation," *IEEE Access*, vol. 9, pp. 150098–150119, 2021, doi: 10.1109/ACCESS.2021.3125102.
- [17] B. Botsch, *Maschinelles Lernen - Grundlagen und Anwendungen: Mit Beispielen in Python*. Berlin, Germany: Springer Spektrum, 2023.
- [18] J. Frochte, *Maschinelles Lernen: Grundlagen und Algorithmen in Python*, 3. Aufl. München: Hanser, 2021.

-
- [19] G. Paaß und D. Hecker, *Künstliche Intelligenz: Was steckt hinter der Technologie der Zukunft?* Wiesbaden: Springer Vieweg, 2020.
- [20] A. Géron, *Praxiseinstieg Machine Learning mit Scikit-Learn, Keras und TensorFlow: Konzepte, Tools und Techniken für intelligente Systeme*, 3. Aufl. Heidelberg: O'Reilly, 2023.
- [21] F. Chollet, *Deep learning mit Python und Keras: Das Praxis-Handbuch vom Entwickler der Keras-Bibliothek*, 1. Aufl. Frechen: mitp, 2018.
- [22] R. Kruse, S. Mostaghim, C. Borgelt, C. Braune und M. Steinbrecher, *Computational Intelligence: A Methodological Introduction*. Switzerland: Springer, 2022.
- [23] Z. A. Styczynski, *Einführung in Expertensysteme*. Berlin, Heidelberg: Springer Berlin Heidelberg, 2017.
- [24] C. Albon, *Machine Learning Kochbuch: Praktische Lösungen mit Python: von der Vorverarbeitung der Daten bis zum Deep Learning*. Heidelberg: O'Reilly, 2019.
- [25] L. Fahrmeir, C. G. Heumann, R. Künstler, I. Pigeot und G. Tutz, *Statistik: Der Weg zur Datenanalyse*, 9. Aufl. Berlin, Heidelberg: Springer, 2024.
- [26] J. Cleve und U. Lämmel, *Data Mining: Datenanalyse für Künstliche Intelligenz*, 4. Aufl. Berlin, Boston: De Gruyter, 2024.
- [27] M. von der Hude, *Predictive Analytics und Data Mining: Eine Einführung mit R*, 1. Aufl. Wiesbaden: Springer Vieweg, 2020.
- [28] W. Ertel, *Grundkurs Künstliche Intelligenz: Eine praxisorientierte Einführung*, 5. Aufl. Wiesbaden: Springer Vieweg, 2021.
- [29] A. V. Dorogush, V. Ershov and A. Gulin, "CatBoost: gradient boosting with categorical features support," Oct. 2018. [Online]. Verfügbar unter: <http://arxiv.org/pdf/1810.11363v1>
- [30] A. Kumar und M. Jain, *Ensemble Learning for AI Developers*. Berkeley, CA: Apress, 2020.
- [31] Q. Li, Y. Xu, B. S. H. Chew, H. Ding and G. Zhao, "An Integrated Missing-Data Tolerant Model for Probabilistic PV Power Generation Forecasting," *IEEE Transactions on Power Systems*, vol. 37, no. 6, pp. 4447–4459, 2022, doi: 10.1109/TPWRS.2022.3146982.
- [32] E. Vrettos and C. Gehbauer, "A Hybrid Approach for Short-Term PV Power Forecasting in Predictive Control Applications," *2019 IEEE Milan PowerTech*, Milan, Italy, 2019, pp. 1–6, doi: 10.1109/PTC.2019.8810672.
- [33] L. Gigoni *et al.*, "Day-Ahead Hourly Forecasting of Power Generation From Photovoltaic Plants," *IEEE Transactions on Sustainable Energy*, vol. 9, no. 2, pp. 831–842, 2018, doi: 10.1109/TSTE.2017.2762435.
- [34] R. Fachrizal, D. van der Meer and J. Munkhammar, "Direct forecast of solar irradiance for EV smart charging scheme to improve PV self-consumption at

- home," *2021 IEEE PES Innovative Smart Grid Technologies Europe (ISGT Europe)*, Espoo, Finland, 2021, pp. 1–5, doi: 10.1109/ISGTEurope52324.2021.9640149.
- [35] M. Massaoudi *et al.*, "An Effective Hybrid NARX-LSTM Model for Point and Interval PV Power Forecasting," *IEEE Access*, vol. 9, pp. 36571–36588, 2021, doi: 10.1109/ACCESS.2021.3062776.
- [36] G. Moreno, P. Martin, C. Santos, F. J. Rodriguez and E. Santiso, "A Day-Ahead Irradiance Forecasting Strategy for the Integration of Photovoltaic Systems in Virtual Power Plants," *IEEE Access*, vol. 8, pp. 204226–204240, 2020, doi: 10.1109/ACCESS.2020.3036140.
- [37] C.-J. Huang and P.-H. Kuo, "Multiple-Input Deep Convolutional Neural Network Model for Short-Term Photovoltaic Power Forecasting," *IEEE Access*, vol. 7, pp. 74822–74834, 2019, doi: 10.1109/ACCESS.2019.2921238.
- [38] J. Dumas, C. Cointe, X. Fettweis and B. Cornelusse, "Deep learning-based multi-output quantile forecasting of PV generation," *2021 IEEE Madrid PowerTech*, Madrid, Spain, 2021, pp. 1–6, doi: 10.1109/PowerTech46648.2021.9494976.
- [39] F. E. Atencio Espejo, S. Grillo and L. Luini, "Photovoltaic Power Production Estimation Based on Numerical Weather Predictions," *2019 IEEE Milan PowerTech*, Milan, Italy, 2019, pp. 1–6, doi: 10.1109/PTC.2019.8810897.
- [40] A. Tascikaraoglu *et al.*, "Compressive Spatio-Temporal Forecasting of Meteorological Quantities and Photovoltaic Power," *IEEE Transactions on Sustainable Energy*, vol. 7, no. 3, pp. 1295–1305, 2016, doi: 10.1109/TSTE.2016.2544929.
- [41] M. Kelker *et al.*, "Development of a forecast model for the prediction of photovoltaic power using neural networks and validating the model based on real measurement data of a local photovoltaic system," *2019 IEEE Milan PowerTech*, Milan, Italy, 2019, pp. 1–6, doi: 10.1109/PTC.2019.8810719.
- [42] X. G. Agoua, R. Girard and G. Kariniotakis, "Short-Term Spatio-Temporal Forecasting of Photovoltaic Power Production," *IEEE Transactions on Sustainable Energy*, vol. 9, no. 2, pp. 538–546, 2018, doi: 10.1109/TSTE.2017.2747765.
- [43] L. H. Dissawa *et al.*, "On-Site Solar Power Forecasting Using Sky-Images," *2020 Australasian Universities Power Engineering Conference (AUPEC)*, Hobart, Australia, 2020, pp. 1–8.
- [44] M. Lee, W. Lee and J. Jung, "24-Hour photovoltaic generation forecasting using combined very-short-term and short-term multivariate time series model," *2017 IEEE Power & Energy Society General Meeting*, Chicago, IL, USA, 2017, pp. 1–5, doi: 10.1109/PESGM.2017.8274605.

-
- [45] M. Ueshima, T. Babasaki, K. Yuasa and I. Omura, "Examination of Correction Method of Long-term Solar Radiation Forecasts of Numerical Weather Prediction," *2019 8th International Conference on Renewable Energy Research and Applications (ICRERA)*, Brasov, Romania, 2019, pp. 113–117, doi: 10.1109/ICRERA47325.2019.8997070.
- [46] L. P. Franco, P. Hirmer and L. H. Thom, "Regional PV Energy Forecasting using Distributed Data and Deep Neural Networks," *2022 IEEE Power & Energy Society Innovative Smart Grid Technologies Conference (ISGT)*, New Orleans, LA, USA, 2022, pp. 1–5, doi: 10.1109/ISGT50606.2022.9817479.
- [47] J. Liu, W. Fang, X. Zhang and C. Yang, "An Improved Photovoltaic Power Forecasting Model With the Assistance of Aerosol Index Data," *IEEE Transactions on Sustainable Energy*, vol. 6, no. 2, pp. 434–442, 2015, doi: 10.1109/TSTE.2014.2381224.
- [48] S. Leva, M. Mussetta, A. Nespoli and E. Ogliari, "PV power forecasting improvement by means of a selective ensemble approach," *2019 IEEE Milan PowerTech*, Milan, Italy, 2019, pp. 1–5, doi: 10.1109/PTC.2019.8810921.
- [49] K. P. Moustris, K. A. Kavvadias, A. I. Kokkosis and A. G. Paliatsos, "One day-ahead forecasting of mean hourly global solar irradiation for energy management systems purposes using artificial neural network modeling," *Mediterranean Conference on Power Generation, Transmission, Distribution and Energy Conversion (MedPower 2016)*, Belgrade, 2016, 1-6, doi: 10.1049/cp.2016.1093.
- [50] M. S. Hossain and H. Mahmood, "Short-Term Photovoltaic Power Forecasting Using an LSTM Neural Network and Synthetic Weather Forecast," *IEEE Access*, vol. 8, pp. 172524–172533, 2020, doi: 10.1109/ACCESS.2020.3024901.
- [51] D. Kothona, I. P. Panapakidis and G. C. Christoforidis, "An Hour-Ahead Photovoltaic Power Forecasting Based on LSTM Model," *2021 IEEE Madrid PowerTech*, Madrid, Spain, 2021, pp. 1–5, doi: 10.1109/PowerTech46648.2021.9494841.
- [52] V. Suresh, P. Janik, J. Rezmer and Z. Leonowicz, "Forecasting Solar PV Output Using Convolutional Neural Networks with a Sliding Window Algorithm," *Energies*, vol. 13, no. 3, p. 723, 2020, doi: 10.3390/en13030723.
- [53] H. Zhou, J. Wang, F. Ouyang, C. Cui and X. Li, "A Two-Stage Method for Ultra-Short-Term PV Power Forecasting Based on Data-Driven," *IEEE Access*, vol. 11, pp. 41175–41189, 2023, doi: 10.1109/ACCESS.2023.3267515.
- [54] X. Zhang *et al.*, "Prediction Interval Estimation and Deterministic Forecasting Model Using Ground-Based Sky Image," *IEEE Transactions on Industry Applications*, vol. 59, no. 2, pp. 2210–2224, 2023, doi: 10.1109/TIA.2022.3218758.

- [55] Z. Si, Y. Yu, M. Yang and P. Li, "Hybrid Solar Forecasting Method Using Satellite Visible Images and Modified Convolutional Neural Networks," *IEEE Transactions on Industry Applications*, vol. 57, no. 1, pp. 5–16, 2021, doi: 10.1109/TIA.2020.3028558.
- [56] S. Tiwari, R. Sabzehgar and M. Rasouli, "Short Term Solar Irradiance Forecast based on Image Processing and Cloud Motion Detection," *2019 IEEE Texas Power and Energy Conference (TPEC)*, College Station, TX, USA, 2019, pp. 1–6, doi: 10.1109/TPEC.2019.8662134.
- [57] T. Leelaruji and N. Teerakawanich, "Short Term Prediction of Solar Irradiance Fluctuation Using Image Processing with ResNet," *2020 8th International Electrical Engineering Congress (iEECON)*, Chiang Mai, Thailand, 2020, pp. 1–4, doi: 10.1109/iEECON48109.2020.229573.
- [58] Y. Fu *et al.*, "Sky Image Prediction Model Based on Convolutional Auto-Encoder for Minutely Solar PV Power Forecasting," *IEEE Transactions on Industry Applications*, vol. 57, no. 4, pp. 3272–3281, 2021, doi: 10.1109/TIA.2021.3072025.
- [59] M. Penteliuc and M. Frincu, "Prediction of Cloud Movement from Satellite Images Using Neural Networks," *2019 21st International Symposium on Symbolic and Numeric Algorithms for Scientific Computing (SYNASC)*, Timisoara, Romania, 2019, pp. 222–229, doi: 10.1109/SYNASC49474.2019.00038.
- [60] M. E. Penteliuc and M. Frincu, "Short Term Cloud Motion Forecast based on Boid's Algorithm for use in PV Output Prediction," *2021 IEEE PES Innovative Smart Grid Technologies Europe (ISGT Europe)*, Espoo, Finland, 2021, pp. 1–5, doi: 10.1109/ISGTEurope52324.2021.9640027.
- [61] C. Feng *et al.*, "Short-term global horizontal irradiance forecasting based on sky imaging and pattern recognition," *2017 IEEE Power & Energy Society General Meeting*, Chicago, IL, USA, 2017, pp. 1–5, doi: 10.1109/PESGM.2017.8274480.
- [62] J. Liu *et al.*, "Research of Photovoltaic Power Forecasting Based on Big Data and mRMR Feature Reduction," *2018 IEEE Power & Energy Society General Meeting (PESGM)*, Portland, OR, USA, 2018, pp. 1–5, doi: 10.1109/PESGM.2018.8585967.
- [63] U. K. Das *et al.*, "Optimized Support Vector Regression-Based Model for Solar Power Generation Forecasting on the Basis of Online Weather Reports," *IEEE Access*, vol. 10, pp. 15594–15604, 2022, doi: 10.1109/ACCESS.2022.3148821.
- [64] F. Najibi, D. Apostolopoulou and E. Alonso, "Clustering Sensitivity Analysis for Gaussian Process Regression Based Solar Output Forecast," *2021 IEEE Madrid PowerTech*, Madrid, Spain, 2021, pp. 1–6, doi: 10.1109/PowerTech46648.2021.9495007.

-
- [65] J. Simeunovic, B. Schubnel, P.-J. Alet and R. E. Carrillo, "Spatio-Temporal Graph Neural Networks for Multi-Site PV Power Forecasting," *IEEE Transactions on Sustainable Energy*, vol. 13, no. 2, pp. 1210–1220, 2022, doi: 10.1109/TSTE.2021.3125200.
- [66] N. Zhou, X. Xu, Z. Yan and M. Shahidehpour, "Spatio-Temporal Probabilistic Forecasting of Photovoltaic Power Based on Monotone Broad Learning System and Copula Theory," *IEEE Transactions on Sustainable Energy*, vol. 13, no. 4, pp. 1874–1885, 2022, doi: 10.1109/TSTE.2022.3174012.
- [67] M. S. Potyka, H. Seefluth and P. Schegner, "Prognosemodell einer PV-Anlage basierend auf einem Kurzzeitmesssystem, Wetterdaten und Machine-Learning Verfahren," *17. Symposium Energieinnovation*, Graz, Austria, 2022.
- [68] A. Aliberti, D. Fucini, L. Bottaccioli, E. Macii, A. Acquaviva and E. Patti, "Comparative Analysis of Neural Networks Techniques to Forecast Global Horizontal Irradiance," *IEEE Access*, vol. 9, pp. 122829–122846, 2021, doi: 10.1109/ACCESS.2021.3110167.
- [69] S. M. Miraftabzadeh, C. G. Colombo, M. Longo and F. Foiadelli, "A Day-Ahead Photovoltaic Power Prediction via Transfer Learning and Deep Neural Networks," *Forecasting*, vol. 5, no. 1, pp. 213–228, 2023, doi: 10.3390/forecast5010012.
- [70] Y. Zhang, M. Beaudin, R. Taheri, H. Zareipour and D. Wood, "Day-Ahead Power Output Forecasting for Small-Scale Solar Photovoltaic Electricity Generators," *IEEE Transactions on Smart Grid*, vol. 6, no. 5, pp. 2253–2262, 2015, doi: 10.1109/TSG.2015.2397003.
- [71] B. P. Mukhoty, V. Maurya and S. K. Shukla, "Sequence to sequence deep learning models for solar irradiation forecasting," *2019 IEEE Milan PowerTech*, Milan, Italy, 2019, pp. 1–6, doi: 10.1109/PTC.2019.8810645.
- [72] G. Li, S. Xie, B. Wang, J. Xin, Y. Li and S. Du, "Photovoltaic Power Forecasting With a Hybrid Deep Learning Approach," *IEEE Access*, vol. 8, pp. 175871–175880, 2020, doi: 10.1109/ACCESS.2020.3025860.
- [73] M. Abuella and B. Chowdhury, "Random forest ensemble of support vector regression models for solar power forecasting," *2017 IEEE Power & Energy Society Innovative Smart Grid Technologies Conference (ISGT)*, Washington, DC, USA, 2017, pp. 1–5, doi: 10.1109/ISGT.2017.8086027.
- [74] S. Pretto, E. Ogliari, A. Niccolai and A. Nespoli, "A New Probabilistic Ensemble Method for an Enhanced Day-Ahead PV Power Forecast," *IEEE Journal of Photovoltaics*, vol. 12, no. 2, pp. 581–588, 2022, doi: 10.1109/JPHOTOV.2021.3138223.

- [75] M. Alaraj, A. Kumar, I. Alsaidan, M. Rizwan and M. Jamil, "Energy Production Forecasting From Solar Photovoltaic Plants Based on Meteorological Parameters for Qassim Region, Saudi Arabia," *IEEE Access*, vol. 9, pp. 83241–83251, 2021, doi: 10.1109/ACCESS.2021.3087345.
- [76] M. Q. Raza, M. Nadarajah and C. Ekanayake, "A multivariate ensemble framework for short term solar photovoltaic output power forecast," *2017 IEEE Power & Energy Society General Meeting*, Chicago, IL, USA, 2017, pp. 1–5, doi: 10.1109/PESGM.2017.8274676.
- [77] A. Bracale, G. Carpinelli and P. de Falco, "A Probabilistic Competitive Ensemble Method for Short-Term Photovoltaic Power Forecasting," *IEEE Transactions on Sustainable Energy*, vol. 8, no. 2, pp. 551–560, 2017, doi: 10.1109/TSTE.2016.2610523.
- [78] Y. Li *et al.*, "A machine-learning approach for regional photovoltaic power forecasting," *2016 IEEE Power and Energy Society General Meeting (PESGM)*, Boston, MA, USA, 2016, pp. 1–5, doi: 10.1109/PESGM.2016.7741991.
- [79] S. Theocharides, G. Makrides, A. Livera, M. Theristis, P. Kaimakis and G. E. Georghiou, "Day-ahead photovoltaic power production forecasting methodology based on machine learning and statistical post-processing," *Applied Energy*, vol. 268, p. 115023, 2020, doi: 10.1016/j.apenergy.2020.115023.
- [80] V. Quaschnig, *Regenerative Energiesysteme: Technologie, Berechnung, Klimaschutz*, 12. Aufl. München: Hanser, 2024.
- [81] DWD Climate Data Center (CDC). [Online]. Verfügbar unter: https://open-data.dwd.de/climate_environment/CDC/ (Zuletzt geprüft am: 18.06.2024).
- [82] H. WATTER, *Regenerative Energiesysteme*. Wiesbaden: Springer Fachmedien Wiesbaden, 2022.
- [83] M. Feindt und U. Kerzel, *Prognosen bewerten*. Berlin, Heidelberg: Springer Berlin Heidelberg, 2015.
- [84] A. Ng und K. Soo, *Data Science – was ist das eigentlich?!* Berlin, Heidelberg: Springer Berlin Heidelberg, 2018.
- [85] L. P. Dinesh, N. A. Khafaf and B. McGrath, "A short term multistep forecasting model for photovoltaic generation using deep learning model," *Sustainable Operations and Computers*, vol. 6, pp. 34–46, 2025, doi: 10.1016/j.susoc.2024.11.003.
- [86] K. Mertens, *Photovoltaik: Lehrbuch zu Grundlagen, Technologie und Praxis*, 6. Aufl. München: Hanser, Carl, 2022.
- [87] *DIN 5034-2, Tageslicht in Innenräumen – Teil 2: Grundlagen*, DIN Deutsches Institut für Normung e. V., 2021-08-00.
- [88] V. Wesselak, J. Fischer, T. Link und T. Schabbach, *Handbuch Regenerative Energietechnik*, 3. Aufl. Berlin: Springer Vieweg, 2017.

-
- [89] S. Mousavi Maleki, H. Hizam and C. Gomes, "Estimation of Hourly, Daily and Monthly Global Solar Radiation on Inclined Surfaces: Models Re-Visited," *Energies*, vol. 10, no. 1, p. 134, 2017, doi: 10.3390/en10010134.
- [90] A. Wagner, *Photovoltaik Engineering*, 5. Aufl.: Springer Berlin Heidelberg, 2019.
- [91] P. Konstantin und M. Konstantin, *Praxisbuch Energiewirtschaft: Energieumwandlung, -transport und -beschaffung, Übertragungsnetzausbau und Kernenergieausstieg*, 5. Aufl. Berlin, Heidelberg: Springer Vieweg, 2023.
- [92] N. A. Engerer and F. P. Mills, "KPV: A clear-sky index for photovoltaics," *Solar Energy*, vol. 105, pp. 679–693, 2014, doi: 10.1016/j.solener.2014.04.019.
- [93] M. Kuhn und K. Johnson, *Applied Predictive Modeling*. New York, NY: Springer New York, 2013.
- [94] T. O. Hodson, "Root-mean-square error (RMSE) or mean absolute error (MAE): when to use them or not," *Geosci. Model Dev.*, vol. 15, no. 14, pp. 5481–5487, 2022, doi: 10.5194/gmd-15-5481-2022.
- [95] L. Suhl und T. Mellouli, *Optimierungssysteme*, 3. Aufl.: Springer Berlin Heidelberg, 2013.
- [96] A. Koop und H. Moock, *Lineare Optimierung - eine anwendungsorientierte Einführung in Operations Research: Mit Python-Programmen*, 3. Aufl. Berlin: Springer Berlin; Springer Spektrum, 2023.
- [97] W.-S. Lu und A. Antoniou, *Practical Optimization: Algorithms and Engineering Applications*: Springer, 2021.
- [98] D. L. Woodruff, J.-P. Watson, J. D. Siirola, B. L. Nicholson, C. D. Laird, W. E. Hart, G. A. Hackebeil und M. L. Bynum, *Pyomo -- Optimization Modeling in Python*: Springer, 2021.
- [99] Allgemeiner Deutscher Automobil-Club e. V. (ADAC), *Autokatalog*. [Online]. Verfügbar unter: <https://www.adac.de/rund-ums-fahrzeug/autokatalog/marken-modelle/renault/zoe/1generation-facelift/301604/#technische-daten> (Zuletzt geprüft am: 03.01.2025).
- [100] A. Kampker und H. H. Heimes, *Elektromobilität*. Berlin, Heidelberg: Springer Berlin Heidelberg, 2024.
- [101] P. Komarnicki, J. Haubrock und Z. A. Styczynski, *Elektromobilität und Sektorenkopplung*. Berlin, Heidelberg: Springer Berlin Heidelberg, 2020.
- [102] Deutscher Wetterdienst (DWD), *Globalstrahlung (mittlere 30-jährige Monats- und Jahressummen)*. [Online]. Verfügbar unter: https://www.dwd.de/DE/leistungen/solarenergie/strahlungskarten_mvs.html?nn=16102 (Zuletzt geprüft am: 19.06.2024).

Abkürzungs- und Formelzeichenverzeichnis

Abkürzungen

AM	Luftmasse (<i>engl. Air Mass</i>)
ANN	Künstliches Neuronales Netz (<i>engl. Artificial Neural Network</i>)
CNN	Convolutional Neural Network
CSI	Clear-Sky Index
DIY	Anzahl der Tage im Jahr (<i>engl. Days in Year</i>)
DOY	Nummer für den Tag des Jahres (<i>engl. Day of Year</i>)
DT	Entscheidungsbaum (<i>engl. Decision Tree</i>)
DWD	Deutscher Wetterdienst
ECMWF	Europäisches Zentrum für mittelfristige Wettervorhersage (<i>engl. European Centre for Medium-Range Weather Forecast</i>)
EEG	Erneuerbare-Energien-Gesetz
EnWG	Energiewirtschaftsgesetz
EV	Elektrofahrzeug (<i>engl. Electric Vehicle</i>)
FNN	Feedforward Neural Network
IQA	Interquartilsabstand
KI	Künstliche Intelligenz
KNN	K-Nächsten-Nachbarn (<i>engl. K-Nearest Neighbors</i>)
LP	Lineare Optimierung (<i>engl. Linear Programming</i>)
LR	Lineare Regression
LSTM	Long Short-Term Memory
LZ	Lokale Zeit
MAE	Mittlerer Absoluter Fehler (<i>engl. Mean Absolute Error</i>)
MAPE	Mittlerer Absoluter Prozentualer Fehler (<i>engl. Mean Absolute Percentage Error</i>)

MESZ	Mitteleuropäische Sommerzeit
MEZ	Mitteleuropäische Zeit
ML	Maschinelles Lernen (<i>engl. Machine Learning</i>)
MOZ	Mittlere Ortszeit
MsbG	Messstellenbetriebsgesetz
MSE	Mittlerer Quadratischer Fehler (<i>engl. Mean Squared Error</i>)
NLP	Nichtlineare Optimierung (<i>engl. Nonlinear Programming</i>)
nMAE	Normalisierte Mittlere Absolute Fehler (<i>engl. normalized Mean Absolute Error</i>)
nRMSE	normalisierte Quadratwurzel des Mittleren Quadratischen Fehlers (<i>engl. normalized Root Mean Square Error</i>)
NWP	Numerische Wettervorhersage (<i>engl. Numerical Weather Prediction</i>)
OR	Operations Research
PV	Photovoltaik
RF	Zufallswald (<i>engl. Random Forest</i>)
RMSE	Quadratwurzel des Mittleren Quadratischen Fehlers (<i>engl. Root Mean Square Error</i>)
SoC	Ladezustand (<i>engl. State of Charge</i>)
SSE	Summe der Quadrierten Abweichungen (<i>engl. Sum of Squared Errors</i>)
STC	Standardtestbedingungen (<i>engl. Standard Test Conditions</i>)
SVM	Support Vector Machine
SVR	Support Vector Regression
WOZ	Wahre Ortszeit
ZGL	Zeitgleichung

Formelzeichen

A_{PV}	Fläche der PV-Anlage
a_v	Aktivierung des Neurons v
C_k	Cluster mit der Nummer k

$d_{\text{euklidisch}}$	Euklidischer Abstand
$d_{\text{Manhattan}}$	Manhattan-Abstand
E	Fehlervektor
$E(W)$	Gesamtfehler in Abhängigkeit der gesamten Gewichte
$E_{\text{Diff,g}}^t$	Diffuse Bestrahlungsstärke auf die geneigte Oberfläche zum Zeitpunkt t
$E_{\text{Diff,h}}^t$	Diffuse Bestrahlungsstärke auf die horizontale Erdoberfläche zum Zeitpunkt t
$E_{\text{Dir,g}}^t$	Direkte Bestrahlungsstärke auf die geneigte Oberfläche zum Zeitpunkt t
$E_{\text{Dir,h}}^t$	Direkte Bestrahlungsstärke auf die horizontale Erdoberfläche zum Zeitpunkt t
$E_{\text{G,g}}^t$	Globale Bestrahlungsstärke auf die geneigte Oberfläche zum Zeitpunkt t
$E_{\text{G,h}}^t$	Globale Bestrahlungsstärke auf die horizontale Erdoberfläche zum Zeitpunkt t
E_m	Mittlerer Fehler aus allen Iterationen
$E_{\text{R,g}}^t$	Reflektierte Bestrahlungsstärke auf die geneigte Oberfläche zum Zeitpunkt t
E_{STC}	Bestrahlungsstärke unter STC
e_i	Fehler für jede Iteration i
$e(\tau)$	Fehler zum τ -ten Schritt des Trainingsalgorithmus
F	Aktivierungsfunktion bzw. Übertragungsfunktion
J	Kostenfunktion
J'	Parameter für die gemittelte Winkelbewegung der Erde um die Sonne
K	Anzahl der Cluster
K_{ap}	Batteriekapazität
$P_{\text{Ladesäule}}$	Leistung einer Ladesäule
$P_{\text{Lade,max}}$	Maximale Ladeleistung
$P_{\text{Lade,min}}$	Minimale Ladeleistung
$P_{\text{NetzinBatt}}$	Anteil Leistung aus dem Netz an der Ladeleistung

$P_{\text{Prog,PV}}$	Prognostizierte PV-Leistung
P_{PV}^t	Erzeugte PV-Leistung zum Zeitpunkt t
P_{PVinBatt}	Anteil PV-Leistung an der Ladeleistung
P_{STC}	Leistung unter STC
SoC_{Ziel}	Gewünschter SoC nach Ende der Ladung
SoC_0	SoC zu Beginn der Ladung
T	Menge der Zeitschritte in der Länge der PV-Prognose
T_{PV}^t	Temperatur der PV-Anlage zum Zeitpunkt t
T_{U}^t	Umgebungstemperatur zum Zeitpunkt t
u	Gewichtete Summe
u_v	Gewichtete Summe von den Neuronen der Vorgängerschicht in ein nachfolgendes Neuron v
w_j	Gewicht des Eingangssignals j
$w_{\lambda v}$	Gewicht zwischen einem Neuron λ und einem Neuron v
$w_{\lambda v}(\tau)$	Gewicht der Verbindung zwischen dem Neuron λ und dem Ausgabeneuron v zum τ -ten Schritt des Trainingsalgorithmus
x_i	Unabhängige Variable
x_j	Eingangswert des Eingangssignals j
$x_{m,n}$	Eingangswert eines Merkmals m
$x_{\lambda}(\tau)$	Eingangswert des Neurons λ zum τ -ten Trainingsschritt
$Y(\omega)$	Menge der Zeitschritte für die Länge der Ladung des jeweiligen Elektrofahrzeuges ω innerhalb der Länge der PV-Prognose
y	Ausgangswert/ Ausgabevektor
y_i	Abhängige Variable/ Zielwert
y_{max}	Maximal gemessener Wert
y_{min}	Minimal gemessener Wert
$\hat{y}$	Erwarteter Ausgangswert/ Ausgangsvektor
$\hat{y}_i$	Vorhersagewert

$\hat{y}_v(\tau)$	Erwarteter Ausgangswert des Neurons v zum τ -ten Schritt des Trainingsalgorithmus
y_λ	Ausgangswert des Neurons λ
$y_v(\tau)$	Ausgangswert des Neurons v zum τ -ten Schritt des Trainingsalgorithmus
$z_{m,n}$	Normierte Daten auf Basis der z-Transformation
α	Achsenabschnitt
α_{PV}	Azimutwinkel der PV-Anlage
α_S	Sonnenazimut
β	Steigung der Geraden
β_{PV}	Temperaturkoeffizient der PV-Anlage
γ_{PV}	Neigungswinkel der PV-Anlage
γ_S	Sonnenhöhe
ΔW	Änderung der gesamten Gewichte
$\Delta w_{\lambda v}$	Änderung des Gewichtes zwischen den Neuronen λ und v
$\Delta w_{\lambda v}(\tau)$	Änderung des Gewichtes zwischen den Neuronen λ und v zum τ -ten Schritt des Trainingsalgorithmus
$\Delta \tau$	Zeitauflösung
δ	Sonnendeklination
δ_v	δ -Wert
ε_i	Zufälliger Fehlerterm
η	Lernrate
η_{Lade}	Ladewirkungsgrad
η_{PV}	Wirkungsgrad der PV-Anlage
θ_g	Einfallswinkel zwischen der direkten Sonneneinstrahlung und der Flächennormalen der PV-Anlage
ϑ^t	Korrekturfaktor für die Modultemperatur zum Zeitpunkt t
λ	Stärke der Regularisierung
λ_{PV}	Längengrad der PV-Anlage

μ_k	Centroid des Clusters C_k
μ_m	Mittelwert aller Eingangswerte eines Merkmals m
ξ	Bias
ρ	Albedo-Koeffizient
σ	Hinterlüftungsfaktor
σ_m	Standardabweichung aller Eingangswerte eines Merkmals m
Φ_1	Gewicht 1
Φ_2	Gewicht 2
φ_{PV}	Breitengrad der PV-Anlage
Ω	Menge der angeschlossenen Elektrofahrzeuge
ω	Stundenwinkel

Abbildungs- und Tabellenverzeichnis

Abbildungen

Abbildung 1.1:	Übersicht über die Funktionsweise des zweistufigen PV-Prognosemodells im Betrieb	5
Abbildung 2.1:	Übersicht über Lernverfahren und Methoden des ML nach [17, 20]	10
Abbildung 2.2:	Aufbau eines Neurons in einem ANN nach [18, 23]	11
Abbildung 2.3:	Übersicht über häufige Aktivierungsfunktionen nach [20, 22, 23] ..	12
Abbildung 2.4:	Aufbau eines Feed-Forward-Netzes mit zwei Hidden-Layers nach [17, 23]	14
Abbildung 2.5:	Beispiel für die LR-Methode nach [18, 26]	19
Abbildung 2.6:	Beispiel der KNN-Methode für $k = 3$ Nachbarn nach [26, 27]	21
Abbildung 2.7:	Beispiel für die Struktur eines DTs nach [18]	23
Abbildung 2.8:	Beispiel für den K-Means-Algorithmus nach [20]	25
Abbildung 2.9:	Ablauf des Trainings beim Bagging nach [20, 26]	26
Abbildung 2.10:	Ablauf des Trainings bei AdaBoost nach [20]	28
Abbildung 2.11:	Ablauf des Trainings beim Stacking nach [20, 30]	29
Abbildung 2.12:	Übersicht über die Gruppierung von Prognosemodellen der PV-Leistung	30
Abbildung 3.1:	Übersicht über die Funktionsweise der PV-Prognose im Betrieb	43
Abbildung 3.2:	Darstellung der Methode der Prognose der Bestrahlungsstärke mittels eines ANN	44
Abbildung 3.3:	Verlauf der Verlustfunktion als MSE während des Trainings und der Validierung des ANN mit der Eingangsdatenkombination 7 und zwei versteckten Schichten	52
Abbildung 3.4:	Darstellung der Verteilung des nMAE in einem Boxplot für verschiedene Eingangsdatenkombinationen für ANNs mit zwei versteckten Schichten für die Testdaten	56
Abbildung 3.5:	Gemessene und prognostizierte Globalstrahlung an zwei Beispieltagen der Testdaten	57
Abbildung 3.6:	Übersicht über das physikalische PV-Modell	58
Abbildung 3.7:	Berechnete und gemessene PV-Leistung an zwei Beispieltagen	67
Abbildung 3.8:	Berechnete und gemessene PV-Leistung in der Kalenderwoche 52 des Jahres 2019	68
Abbildung 3.9:	Darstellung des nMAE des physikalischen PV-Modells als Balkendiagramm pro Monat von 2019 bis 2022	69

Abbildung 3.10: Gemessene und prognostizierte PV-Leistung an zwei Beispieltagen.....	71
Abbildung 3.11: Gemessene und prognostizierte PV-Leistung in der Kalenderwoche 52 des Jahres 2019	72
Abbildung 4.1: Übersicht über die Funktionsweise des Korrekturmodells im Betrieb.....	77
Abbildung 4.2: Darstellung der Methode des Korrekturmodells mittels eines Stacking-Ensembles.....	79
Abbildung 4.3: Eingangs- und Zieldaten für das Training eines Stacking-Ensembles	81
Abbildung 4.4: Eingangs- und Zieldaten für das Training eines Stacking-Ensembles mit dem hinzugefügten Merkmal Clusternummer	82
Abbildung 4.5: Trägheit als SSE in Abhängigkeit der Clusteranzahl	83
Abbildung 4.6: Prognosefehler als nMAE in Abhängigkeit der Clusteranzahl.....	83
Abbildung 4.7: Mittlere Zeitreihe der PV-Leistung in Rot sowie alle Zeitreihen in Grau von Cluster 1 und Cluster 6	84
Abbildung 4.8: Methode der k-fachen Kreuzvalidierung nach [17, 18].....	85
Abbildung 4.9: Gemessene, prognostizierte PV-Leistung und korrigierte prognostizierte PV-Leistung in der Kalenderwoche 52 des Jahres 2019.....	94
Abbildung 4.10: Gemessene, prognostizierte PV-Leistung und korrigierte prognostizierte PV-Leistung an zwei Beispieltagen.....	95
Abbildung 4.11: Prognosefehler der Testdaten in Bezug auf den Anteil der Trainingsdaten für das Stacking-Ensemble mit der Clustering-Methode	97
Abbildung 5.1: Übersicht über das Lademanagementsystem für Elektrofahrzeuge	102
Abbildung 5.2: Ladezustand von sechs Elektrofahrzeugen beim Szenario 1) Gesteuert und 2) Ungesteuert, $P_{\text{Ladesäule}} = 22 \text{ kW}$	110
Abbildung 5.3: PV-Messung, korrigierte PV-Prognose und gesamter Ladeplan für sechs Elektrofahrzeuge für die vier verschiedenen Szenarien 1) Gesteuert, 2) Ungesteuert, $P_{\text{Ladesäule}} = 22 \text{ kW}$, 3) Ungesteuert, $P_{\text{Ladesäule}} = 11 \text{ kW}$, 4) Ungesteuert, $P_{\text{Ladesäule}} = 7,2 \text{ kW}$	111
Abbildung 5.4: PV-Messung, PV-Prognose und gesamter Ladeplan für sechs Elektrofahrzeuge im Vergleich mit korrigierter PV-Prognose und ursprünglicher PV-Prognose für das Szenario 1) Gesteuert.....	113

Abbildungen Anhang

Abbildung A. 1: Verteilung der mittleren Jahressumme der Globalstrahlung von 1991 bis 2020 mit Markierung der Städte Bielefeld und Braunschweig [102]	145
Abbildung A. 2: Datenblatt der betrachteten PV-Anlage der Hochschule Bielefeld	156

Tabellen

Tabelle 2.1: Einteilung des zeitlichen Prognosehorizonts nach Zeithorizont und Anwendung nach [11, 12]	31
Tabelle 3.1: Übersicht über die Daten für das Training eines ANN für die Prognose der Bestrahlungsstärke	46
Tabelle 3.2: Kombinationen von Eingangsdaten	49
Tabelle 3.3: Hyperparameter des ANN mit der Eingangsdatenkombination 7 und zwei versteckten Schichten	51
Tabelle 3.4: Fehlermetriken der Prognose der Bestrahlungsstärke für Trainings-, Validierungs- und Testdaten für ANNs mit zwei versteckten Schichten	54
Tabelle 3.5: Fehlermetriken der gemessenen und prognostizierten Globalstrahlung von zwei Beispieltagen aus den Testdaten	57
Tabelle 3.6: Albedo ρ für verschiedene Umgebungen nach [80, 86, 88]	63
Tabelle 3.7: Hinterlüftungsfaktor σ für die Berechnung der Modultemperatur für verschiedene Einbauvarianten von PV-Anlagen nach [80]	64
Tabelle 3.8: Technische Daten der betrachteten PV-Anlage	65
Tabelle 3.9: Fehlermetriken der Validierung des physikalischen PV-Modells für verschiedene Zeiträume	66
Tabelle 3.10: Fehlermetriken der Validierung der PV-Prognose anhand einer realen PV-Anlage	70
Tabelle 3.11: Fehlermetriken der Validierung der PV-Prognose anhand einer realen PV-Anlage für verschiedene Beispieltage	72
Tabelle 3.12: Fehlermetriken der Validierung der PV-Prognose anhand einer realen PV-Anlage für verschiedene CSI-Wertebereiche für eine 24 Stunden Prognose	74
Tabelle 4.1: Fehlermetriken für die Korrektur der PV-Prognose von verschiedenen Basismodellen für die Testdaten mit der Clustering-Methode	90
Tabelle 4.2: Fehlermetriken für die Korrektur der PV-Prognose von verschiedenen Basismodellkombinationen eines Stacking-Ensembles für die Testdaten mit der Clustering-Methode	91

Tabelle 4.3:	Fehlermetriken des Korrekturmodells der PV-Prognose für verschiedene ML-Modelle und ein Stacking-Ensemble der ML-Modelle für die Testdaten	93
Tabelle 4.4:	Fehlermetriken der PV-Prognose und des Korrekturmodells der PV-Prognose für den 27.12.2019	95
Tabelle 4.5:	Fehlermetriken der PV-Prognose und des Korrekturmodells der PV-Prognose für zwei Beispieltage	96
Tabelle 4.6:	Fehlermetriken des Korrekturmodells mit der Clustering-Methode für verschiedene CSI-Wertebereiche für die Testdaten.....	97
Tabelle 5.1:	Parameter der Optimierung	104
Tabelle 5.2:	Variablen der Optimierung	105
Tabelle 5.3:	Übersicht über die Angaben der sechs betrachteten Elektrofahrzeuge	108
Tabelle 5.4:	PV-Anteil und Netz-Anteil in der Ladung sowie PV-Eigenverbrauch und Ladespitze in den verschiedenen Szenarien ..	112

Tabellen Anhang

Tabelle A. 1:	Hyperparameter der ANNs mit einer versteckten Schicht mit verschiedenen Eingangsdatenkombinationen	146
Tabelle A. 2:	Hyperparameter der ANNs mit zwei versteckten Schichten mit verschiedenen Eingangsdatenkombinationen	146
Tabelle A. 3:	Hyperparameter der ANNs mit drei versteckten Schichten mit verschiedenen Eingangsdatenkombinationen	147
Tabelle A. 4:	Fehlermetriken der Prognose der Bestrahlungsstärke aufgeteilt in Mittelwert/ Globalstrahlung/ Diffusstrahlung/ für Trainings-, Validierungs- und Testdaten für ANNs mit verschiedenen Eingangsdatenkombinationen und einer versteckten Schicht.....	147
Tabelle A. 5:	Fehlermetriken der Prognose der Bestrahlungsstärke aufgeteilt in Mittelwert/ Globalstrahlung/ Diffusstrahlung/ für Trainings-, Validierungs- und Testdaten für ANNs mit verschiedenen Eingangsdatenkombinationen und zwei versteckten Schichten	149
Tabelle A. 6:	Fehlermetriken der Prognose der Bestrahlungsstärke aufgeteilt in Mittelwert/ Globalstrahlung/ Diffusstrahlung/ für Trainings-, Validierungs- und Testdaten für ANNs mit verschiedenen Eingangsdatenkombinationen und drei versteckten Schichten.....	150
Tabelle A. 7:	Fehlermetriken der Prognose der Bestrahlungsstärke aufgeteilt in Mittelwert/ Globalstrahlung/ Diffusstrahlung/ für Trainings-, Validierungs- und Testdaten für ANNs mit verschiedenen Eingangsdatenkombinationen und einer versteckten Schicht der zweiten Optimierung	151

Tabelle A. 8:	Fehlermetriken der Prognose der Bestrahlungsstärke aufgeteilt in Mittelwert/ Globalstrahlung/ Diffusstrahlung/ für Trainings-, Validierungs- und Testdaten für ANNs mit verschiedenen Eingangsdatenkombinationen und zwei versteckten Schichten der zweiten Optimierung.....	153
Tabelle A. 9:	Fehlermetriken der Prognose der Bestrahlungsstärke aufgeteilt in Mittelwert/ Globalstrahlung/ Diffusstrahlung/ für Trainings-, Validierungs- und Testdaten für ANNs mit verschiedenen Eingangsdatenkombinationen und drei versteckten Schichten der zweiten Optimierung.....	154
Tabelle B. 1:	Hyperparameterraum der LR-Methode	159
Tabelle B. 2:	Hyperparameterraum des DT-Modells	159
Tabelle B. 3:	Hyperparameterraum des KNN-Modells.....	160
Tabelle B. 4:	Hyperparameterraum des RF-Modells.....	160
Tabelle B. 5:	Hyperparameterraum des XGBoost-Modells	161
Tabelle B. 6:	Hyperparameterraum des Gradient Boosting Modells	161
Tabelle B. 7:	Hyperparameterraum des CatBoost-Modells.....	162
Tabelle B. 8:	Hyperparameterraum der Ridge Regression.....	162
Tabelle B. 9:	Fehlermetriken für die Korrektur der PV-Prognose von verschiedenen Basismodellen für die Testdaten ohne der Clustering-Methode	163
Tabelle B. 10:	Fehlermetriken für die Korrektur der PV-Prognose von verschiedenen Basismodellkombinationen eines Stacking-Ensembles für die Testdaten ohne der Clustering-Methode.....	163
Tabelle B. 11:	Hyperparameter der Basismodelle und des Metamodells des Stacking-Ensembles.....	163

A. Anhang PV-Prognose

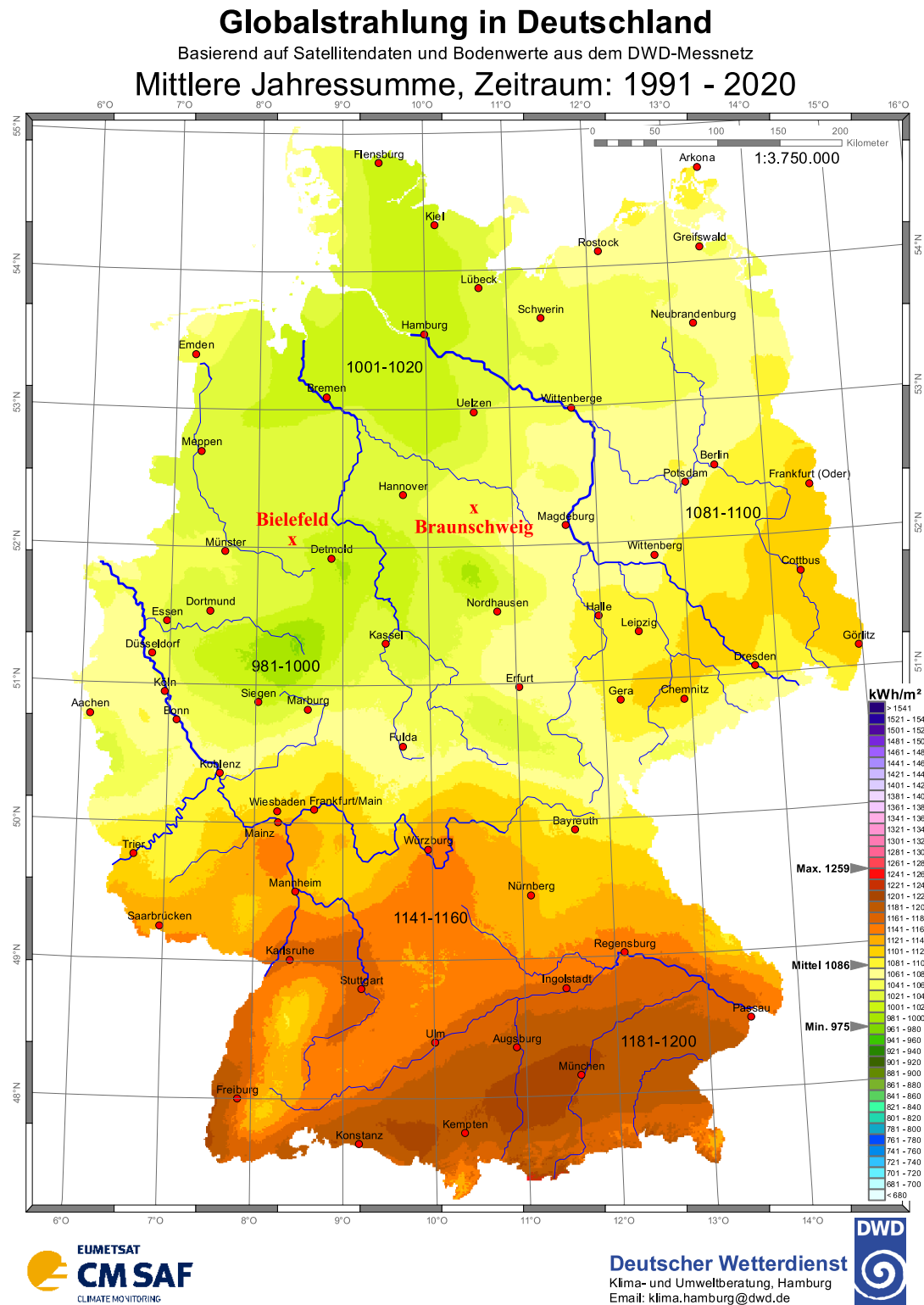


Abbildung A. 1: Verteilung der mittleren Jahressumme der Globalstrahlung von 1991 bis 2020 mit Markierung der Städte Bielefeld und Braunschweig [102]

Tabelle A. 1: Hyperparameter der ANNs mit einer versteckten Schicht mit verschiedenen Eingangsdatenkombinationen

Hyperparameter		Angabe für jede Kombination						
		1	2	3	4	5	6	7
Anzahl Neuronen		930	990	990	990	770	990	990
Aktivierungs-funktion	Versteckte Schicht	ReLU	ReLU	ReLU	ReLU	ReLU	ReLU	ReLU
	Ausgabe-schicht	Linear	Linear	Linear	Linear	Linear	Linear	Linear
Dropout		0,2	0,4	0	0	0	0	0
Lernrate		0,0099	0,0083	0,0025	0,0041	0,0099	0,0079	0,0099

Tabelle A. 2: Hyperparameter der ANNs mit zwei versteckten Schichten mit verschiedenen Eingangsdatenkombinationen

Hyperparameter		Angabe für jede Kombination						
		1	2	3	4	5	6	7
Anzahl Neuro-nen	Versteckte Schicht 1	990	10	990	390	990	470	990
	Versteckte Schicht 2	990	990	990	990	10	490	990
Aktivierungs-funktion	Versteckte Schicht 1	ReLU	ReLU	Linear	ReLU	ReLU	ReLU	ReLU
	Versteckte Schicht 2	ReLU	ReLU	ReLU	ReLU	Linear	ReLU	ReLU
	Ausgabe-schicht	Linear	Linear	Linear	Linear	Linear	Linear	Linear
Dropout	Versteckte Schicht 1	0,2	0	0	0	0	0	0
	Versteckte Schicht 2	0	0,5	0	0	0	0	0,5
Lernrate		0,0001	0,0099	0,0001	0,0059	0,0099	0,0049	0,0007

Tabelle A. 3: Hyperparameter der ANNs mit drei versteckten Schichten mit verschiedenen Eingangsdatenkombinationen

Hyperparameter		Angabe für jede Kombination						
		1	2	3	4	5	6	7
Anzahl Neuro-nen	Versteckte Schicht 1	990	10	990	990	990	990	990
	Versteckte Schicht 2	990	310	670	330	990	230	990
	Versteckte Schicht 3	110	990	390	990	990	490	990
Aktivierungs-funktion	Versteckte Schicht 1	ReLU	ReLU	ReLU	ReLU	ReLU	ReLU	ReLU
	Versteckte Schicht 2	ReLU	ReLU	Sig-moid	ReLU	ReLU	ReLU	ReLU
	Versteckte Schicht 3	ReLU	Linear	Linear	ReLU	Linear	Linear	Sig-moid
	Ausgabe-schicht	Linear	Linear	Linear	Linear	Linear	Linear	Linear
Dropout	Versteckte Schicht 1	0,5	0	0	0	0	0	0,4
	Versteckte Schicht 2	0	0	0	0	0	0	0
	Versteckte Schicht 3	0	0,5	0	0	0	0	0
Lernrate		0,0001	0,0011	0,0001	0,0099	0,0001	0,0039	0,0001

Tabelle A. 4: Fehlermetriken der Prognose der Bestrahlungsstärke aufgeteilt in Mittelwert/ Globalstrahlung/ Diffusstrahlung/ für Trainings-, Validierungs- und Testdaten für ANNs mit verschiedenen Eingangsdatenkombinationen und einer versteckten Schicht

Kombina-tion	MAE in W/m ²	nMAE in %	RMSE in W/m ²	nRMSE in %
Training				
1	116,51/ 155,49/ 77,54	14,54/ 16,86/ 11,39	155,38/ 197,02/ 97,30	19,39/ 21,37/ 14,29
2	68,66/ 95,40/ 41,92	8,57/ 10,35/ 6,16	101,83/ 131,33/ 59,09	12,70/ 14,24/ 8,68

Anhang PV-Prognose

3	47,92/ 58,88/ 36,95	5,98/ 6,39/ 5,43	72,96/ 88,37/ 53,27	9,10/ 9,59/ 7,82
4	47,00/ 57,27/ 36,72	5,86/ 6,21/ 5,39	71,37/ 86,15/ 52,59	8,90/ 9,34/ 7,72
5	46,17/ 56,13/ 36,21	5,76/ 6,09/ 5,32	69,68/ 83,70/ 52,00	8,69/ 9,08/ 7,64
6	44,57/ 54,13/ 35,02	5,56/ 5,87/ 5,14	67,03/ 80,54/ 49,99	8,36/ 8,74/ 7,34
7	38,78/ 46,43/ 31,12	4,84/ 5,04/ 4,57	58,55/ 69,27/ 45,36	7,30/ 7,51/ 6,66
Validierung				
1	116,42/ 154,52/ 78,32	14,53/ 16,76/ 11,50	155,09/ 196,11/ 98,20	19,35/ 21,27/ 14,42
2	68,17/ 94,95/ 41,39	8,51/ 10,30/ 6,08	100,80/ 129,96/ 58,59	12,58/ 14,10/ 8,60
3	47,75/ 58,75/ 36,74	5,96/ 6,37/ 5,40	73,23/ 88,90/ 53,12	9,14/ 9,64/ 7,80
4	47,41/ 58,33/ 36,48	5,91/ 6,33/ 5,36	72,61/ 88,23/ 52,54	9,06/ 9,57/ 7,72
5	47,40/ 57,90/ 36,90	5,91/ 6,28/ 5,42	72,42/ 87,38/ 53,43	9,04/ 9,48/ 7,85
6	46,06/ 56,13/ 35,99	5,75/ 6,09/ 5,29	70,22/ 84,41/ 52,32	8,76/ 9,15/ 7,68
7	39,87/ 47,82/ 31,93	4,97/ 5,19/ 4,69	60,61/ 72,07/ 46,40	7,56/ 7,82/ 6,81
Test				
1	123,27/ 162,34/ 84,21	15,38/ 17,61/ 12,36	170,89/ 217,04/ 106,31	21,32/ 23,54/ 15,61
2	70,95/ 98,23/ 43,68	8,85/ 10,65/ 6,41	105,32/ 135,73/ 61,34	13,14/ 14,72/ 9,01
3	62,96/ 83,14/ 42,78	7,86/ 9,02/ 6,28	100,04/ 127,85/ 60,58	12,48/ 13,87/ 8,90
4	61,04/ 79,74/ 42,33	7,62/ 8,65/ 6,22	95,58/ 121,38/ 59,50	11,93/ 13,16/ 8,74
5	60,37/ 78,62/ 42,11	7,53/ 8,53/ 6,18	93,74/ 118,67/ 59,11	11,70/ 12,87/ 8,68
6	59,41/ 77,03/ 41,79	7,41/ 8,35/ 6,14	91,89/ 116,00/ 58,58	11,47/ 12,58/ 8,60

7	48,04/ 60,01/ 36,08	5,99/ 6,51/ 5,30	73,44/ 90,56/ 50,86	9,16/ 9,82/ 7,47
---	------------------------	---------------------	------------------------	---------------------

Tabelle A. 5: Fehlermetriken der Prognose der Bestrahlungsstärke aufgeteilt in Mittelwert/ Globalstrahlung/ Diffusstrahlung/ für Trainings-, Validierungs- und Testdaten für ANNs mit verschiedenen Eingangsdatenkombinationen und zwei versteckten Schichten

Kombination	MAE in W/m ²	nMAE in %	RMSE in W/m ²	nRMSE in %
Training				
1	115,81/ 154,69/ 76,92	14,45/ 16,78/ 11,30	155,15/ 196,67/ 97,27	19,36/ 21,33/ 14,28
2	68,46/ 95,15/ 41,77	8,54/ 10,32/ 6,13	101,88/ 131,41/ 59,08	12,71/ 14,25/ 8,68
3	47,85/ 58,69/ 37,02	5,97/ 6,37/ 5,44	73,05/ 88,49/ 53,32	9,11/ 9,60/ 7,83
4	46,38/ 56,65/ 36,10	5,79/ 6,14/ 5,30	70,82/ 85,37/ 52,39	8,84/ 9,26/ 7,69
5	46,17/ 56,44/ 35,90	5,76/ 6,12/ 5,27	69,50/ 83,66/ 51,59	8,67/ 9,07/ 7,58
6	42,77/ 51,60/ 33,94	5,34/ 5,60/ 4,98	64,59/ 77,37/ 48,54	8,06/ 8,39/ 7,13
7	37,37/ 44,65/ 30,09	4,66/ 4,84/ 4,42	57,34/ 67,80/ 44,49	7,15/ 7,35/ 6,53
Validierung				
1	115,73/ 153,82/ 77,65	14,44/ 16,68/ 11,40	155,04/ 196,03/ 98,20	19,34/ 21,26/ 14,42
2	67,92/ 94,57/ 41,27	8,47/ 10,26/ 6,06	100,83/ 129,95/ 58,71	12,58/ 14,09/ 8,62
3	47,66/ 58,54/ 36,79	5,95/ 6,35/ 5,40	73,25/ 88,92/ 53,14	9,14/ 9,64/ 7,80
4	46,79/ 57,54/ 36,03	5,84/ 6,24/ 5,29	72,32/ 87,66/ 52,69	9,02/ 9,51/ 7,74
5	47,32/ 58,09/ 36,55	5,90/ 6,30/ 5,37	72,14/ 87,11/ 53,11	9,00/ 9,45/ 7,80
6	45,41/ 54,91/ 35,91	5,67/ 5,96/ 5,27	69,72/ 83,33/ 52,70	8,70/ 9,04/ 7,74
7	38,71/ 46,32/ 31,09	4,83/ 5,02/ 4,57	60,05/ 71,41/ 45,97	7,49/ 7,75/ 6,75

Test				
1	122,99/ 162,04/ 83,94	15,35/ 17,57/ 12,33	171,87/ 218,20/ 107,07	21,44/ 23,67/ 15,72
2	70,48/ 97,65/ 43,32	8,79/ 10,59/ 6,36	105,16/ 135,54/ 61,19	13,12/ 14,70/ 8,99
3	62,96/ 83,00/ 42,91	7,85/ 9,00/ 6,30	100,17/ 127,98/ 60,75	12,50/ 13,88/ 8,92
4	60,30/ 78,90/ 41,71	7,52/ 8,56/ 6,13	94,01/ 119,05/ 59,20	11,73/ 12,91/ 8,69
5	59,61/ 77,48/ 41,74	7,44/ 8,40/ 6,13	92,19/ 116,27/ 58,97	11,50/ 12,61/ 8,66
6	59,23/ 77,06/ 41,41	7,39/ 8,36/ 6,08	92,07/ 116,38/ 58,40	11,49/ 12,62/ 8,58
7	47,34/ 59,00/ 35,69	5,91/ 6,40/ 5,24	73,14/ 89,55/ 51,77	9,13/ 9,71/ 7,60

Tabelle A. 6: Fehlermetriken der Prognose der Bestrahlungsstärke aufgeteilt in Mittelwert/ Globalstrahlung/ Diffusstrahlung/ für Trainings-, Validierungs- und Testdaten für ANNs mit verschiedenen Eingangskombinationen und drei versteckten Schichten

Kombination	MAE in W/m²	nMAE in %	RMSE in W/m²	nRMSE in %
Training				
1	116,26/ 155,25/ 77,26	14,50/ 16,84/ 11,35	155,13/ 196,61/ 97,36	19,36/ 21,32/ 14,30
2	68,56/ 95,43/ 41,69	8,5/ 10,3/ 6,12	101,79/ 131,32/ 58,95	12,70/ 14,24/ 8,66
3	47,96/ 58,60/ 37,31	5,98/ 6,36/ 5,48	72,91/ 88,25/ 53,33	9,10/ 9,57/ 7,83
4	46,07/ 56,18/ 35,95	5,75/ 6,09/ 5,28	70,54/ 84,97/ 52,27	8,80/ 9,22/ 7,68
5	44,93/ 54,62/ 35,24	5,61/ 5,92/ 5,17	68,19/ 81,84/ 51,01	8,51/ 8,88/ 7,49
6	43,50/ 52,76/ 34,24	5,43/ 5,72/ 5,03	66,31/ 79,49/ 49,75	8,27/ 8,62/ 7,31
7	38,35/ 45,62/ 31,09	4,79/ 4,95/ 4,57	58,31/ 68,75/ 45,53	7,27/ 7,46/ 6,69
Validierung				

1	116,20/ 154,39/ 78,01	14,50/ 16,75/ 11,46	155,06/ 196,03/ 98,28	19,35/ 21,26/ 14,43
2	68,15/ 95,10/ 41,19	8,50/ 10,31/ 6,05	100,95/ 130,19/ 58,58	12,59/ 14,12/ 8,60
3	47,87/ 58,56/ 37,19	5,97/ 6,35/ 5,46	73,28/ 88,92/ 53,24	9,14/ 9,64/ 7,82
4	46,74/ 57,43/ 36,04	5,83/ 6,23/ 5,29	72,31/ 87,58/ 52,78	9,02/ 9,50/ 7,75
5	47,14/ 57,82/ 36,47	5,88/ 6,27/ 5,36	72,26/ 87,28/ 53,16	9,02/ 9,47/ 7,81
6	45,40/ 55,05/ 35,75	5,66/ 5,97/ 5,25	70,11/ 83,67/ 53,22	8,75/ 9,07/ 7,81
7	39,78/ 47,66/ 31,89	4,96/ 5,17/ 4,68	61,14/ 72,80/ 46,64	7,63/ 7,90/ 6,85
Test				
1	123,13/ 162,18/ 84,09	15,36/ 17,59/ 12,35	171,20/ 217,36/ 106,65	21,36/ 23,58/ 15,66
2	70,75/ 98,12/ 43,37	8,83/ 10,64/ 6,37	105,24/ 135,70/ 61,14	13,13/ 14,72/ 8,98
3	62,84/ 82,39/ 43,29	7,84/ 8,94/ 6,36	99,67/ 127,05/ 61,05	12,44/ 13,78/ 8,96
4	59,61/ 77,35/ 41,87	7,44/ 8,39/ 6,15	92,44/ 116,30/ 59,71	11,53/ 12,61/ 8,77
5	60,27/ 78,64/ 41,90	7,52/ 8,53/ 6,15	94,62/ 119,88/ 59,45	11,81/ 13,00/ 8,73
6	58,24/ 76,31/ 40,18	7,27/ 8,28/ 5,90	91,34/ 115,83/ 57,16	11,40/ 12,56/ 8,39
7	48,64/ 60,82/ 36,45	6,07/ 6,60/ 5,35	75,23/ 92,69/ 52,21	9,39/ 10,05/ 7,67

Tabelle A. 7: Fehlermetriken der Prognose der Bestrahlungsstärke aufgeteilt in Mittelwert/ Globalstrahlung/ Diffusstrahlung/ für Trainings-, Validierungs- und Testdaten für ANNs mit verschiedenen Eingangskombinationen und einer versteckten Schicht der zweiten Optimierung

Kombination	MAE in W/m ²	nMAE in %	RMSE in W/m ²	nRMSE in %
Training				

Anhang PV-Prognose

1	122,19/ 160,86/ 83,51	15,24/ 17,45/ 12,26	159,53/ 200,79/ 102,87	19,90/ 21,78/ 15,11
2	68,54/ 95,29/ 41,79	8,55/ 10,34/ 6,14	101,46/ 130,83/ 58,92	12,66/ 14,19/ 8,65
3	49,08/ 60,27/ 37,88	6,12/ 6,54/ 5,56	74,15/ 89,81/ 54,13	9,25/ 9,74/ 7,95
4	47,88/ 58,60/ 37,15	5,97/ 6,36/ 5,46	72,25/ 87,33/ 53,04	9,01/ 9,47/ 7,79
5	46,25/ 56,35/ 36,16	5,77/ 6,11/ 5,31	70,09/ 84,40/ 51,98	8,74/ 9,15/ 7,63
6	47,26/ 57,72/ 36,80	5,90/ 6,26/ 5,40	71,07/ 85,76/ 52,41	8,87/ 9,30/ 7,70
7	42,89/ 51,65/ 34,13	5,35/ 5,60/ 5,01	63,09/ 75,11/ 48,15	7,87/ 8,15/ 7,07
Validierung				
1	123,53/ 163,53/ 83,53	15,41/ 17,74/ 12,27	161,77/ 204,45/ 102,65	20,18/ 22,17/ 15,07
2	70,01/ 97,45/ 42,56	8,73/ 10,57/ 6,25	103,31/ 133,59/ 59,16	12,89/ 14,49/ 8,69
3	48,17/ 59,10/ 37,23	6,01/ 6,41/ 5,47	72,34/ 87,83/ 52,46	9,03/ 9,53/ 7,70
4	48,27/ 59,47/ 37,07	6,02/ 6,45/ 5,44	73,28/ 88,99/ 53,11	9,14/ 9,65/ 7,80
5	48,07/ 58,99/ 37,14	6,00/ 6,40/ 5,45	72,08/ 87,39/ 52,48	8,99/ 9,48/ 7,71
6	47,24/ 57,84/ 36,65	5,89/ 6,27/ 5,38	71,77/ 86,85/ 52,53	8,95/ 9,42/ 7,71
7	43,42/ 51,95/ 34,89	5,42/ 5,63/ 5,12	62,84/ 74,50/ 48,45	7,84/ 8,08/ 7,11
Test				
1	126,57/ 167,03/ 86,12	15,79/ 18,12/ 12,65	171,59/ 218,56/ 105,42	21,41/ 23,71/ 15,48
2	71,13/ 98,34/ 43,92	8,87/ 10,67/ 6,45	105,23/ 135,59/ 61,33	13,13/ 14,71/ 9,01
3	63,29/ 83,34/ 43,24	7,90/ 9,04/ 6,35	98,99/ 126,27/ 60,45	12,35/ 13,70/ 8,88
4	61,40/ 80,10/ 42,69	7,66/ 8,69/ 6,27	95,27/ 120,81/ 59,63	11,89/ 13,10/ 8,76

5	60,42/ 78,65/ 42,20	7,54/ 8,53/ 6,20	94,10/ 119,08/ 59,40	11,74/ 12,92/ 8,72
6	59,90/ 77,44/ 42,35	7,47/ 8,40/ 6,22	92,23/ 116,20/ 59,27	11,51/ 12,60/ 8,70
7	51,39/ 63,38/ 39,40	6,41/ 6,87/ 5,79	76,23/ 92,80/ 54,88	9,51/ 10,07/ 8,06

Tabelle A. 8: Fehlermetriken der Prognose der Bestrahlungsstärke aufgeteilt in Mittelwert/ Globalstrahlung/ Diffusstrahlung/ für Trainings-, Validierungs- und Testdaten für ANNs mit verschiedenen Eingangsdatenkombinationen und zwei versteckten Schichten der zweiten Optimierung

Kombination	MAE in W/m ²	nMAE in %	RMSE in W/m ²	nRMSE in %
Training				
1	116,31/ 154,90/ 77,73	14,51/ 16,80/ 11,41	155,04/ 196,17/ 97,94	19,34/ 21,28/ 14,38
2	68,97/ 95,72/ 42,22	8,61/ 10,38/ 6,20	101,26/ 130,43/ 59,11	12,63/ 14,15/ 8,68
3	48,06/ 59,20/ 36,92	6,00/ 6,42/ 5,42	73,27/ 88,75/ 53,49	9,14/ 9,63/ 7,85
4	47,17/ 57,18/ 37,15	5,88/ 6,20/ 5,45	71,53/ 86,10/ 53,10	8,92/ 9,34/ 7,80
5	44,32/ 53,57/ 35,07	5,53/ 5,81/ 5,15	68,07/ 81,66/ 50,96	8,49/ 8,86/ 7,48
6	44,09/ 53,06/ 35,11	5,50/ 5,76/ 5,16	67,04/ 80,30/ 50,41	8,36/ 8,71/ 7,40
7	39,11/ 46,48/ 31,74	4,88/ 5,04/ 4,66	59,04/ 70,04/ 45,46	7,37/ 7,60/ 6,68
Validierung				
1	116,93/ 156,72/ 77,15	14,59/ 17,00/ 11,33	156,10/ 198,27/ 97,07	19,48/ 21,50/ 14,25
2	70,48/ 97,94/ 43,03	8,79/ 10,62/ 6,32	103,38/ 133,58/ 59,43	12,90/ 14,49/ 8,73
3	47,46/ 58,56/ 36,35	5,92/ 6,35/ 5,34	71,93/ 87,37/ 52,10	8,97/ 9,48/ 7,65
4	47,54/ 58,12/ 36,96	5,93/ 6,30/ 5,43	72,48/ 87,78/ 52,92	9,04/ 9,52/ 7,77

5	46,53/ 56,71/ 36,36	5,81/ 6,15/ 5,34	70,72/ 85,43/ 52,00	8,82/ 9,27/ 7,64
6	45,26/ 54,84/ 35,68	5,65/ 5,95/ 5,24	69,82/ 84,04/ 51,83	8,71/ 9,11/ 7,61
7	40,55/ 48,10/ 33,00	5,06/ 5,22/ 4,85	60,62/ 71,83/ 46,79	7,56/ 7,79/ 6,87
Test				
1	123,27/ 162,24/ 84,30	15,38/ 17,60/ 12,38	171,21/ 217,30/ 106,80	21,36/ 23,57/ 15,68
2	71,69/ 99,36/ 44,03	8,95/ 10,78/ 6,47	105,41/ 135,83/ 61,44	13,15/ 14,73/ 9,02
3	62,97/ 83,42/ 42,51	7,86/ 9,05/ 6,24	100,24/ 128,30/ 60,31	12,51/ 13,91/ 8,86
4	60,86/ 79,51/ 42,21	7,59/ 8,62/ 6,20	95,00/ 120,60/ 59,20	11,85/ 13,08/ 8,69
5	60,13/ 78,60/ 41,66	7,50/ 8,53/ 6,12	95,34/ 121,12/ 59,22	11,89/ 13,14/ 8,70
6	58,16/ 74,78/ 41,54	7,26/ 8,11/ 6,10	91,07/ 114,69/ 58,59	11,36/ 12,44/ 8,60
7	50,15/ 62,20/ 38,11	6,26/ 6,75/ 5,60	76,79/ 94,16/ 54,10	9,58/ 10,21/ 7,94

Tabelle A. 9: Fehlermetriken der Prognose der Bestrahlungsstärke aufgeteilt in Mittelwert/ Globalstrahlung/ Diffusstrahlung/ für Trainings-, Validierungs- und Testdaten für ANNs mit verschiedenen Eingangskombinationen und drei versteckten Schichten der zweiten Optimierung

Kombination	MAE in W/m²	nMAE in %	RMSE in W/m²	nRMSE in %
Training				
1	116,47/ 155,61/ 77,32	14,53/ 16,88/ 11,35	155,08/ 196,42/ 97,56	19,35/ 21,30/ 14,33
2	68,70/ 95,24/ 42,16	8,57/ 10,33/ 6,19	101,68/ 130,85/ 59,62	12,69/ 14,19/ 8,76
3	49,16/ 60,28/ 38,03	6,13/ 6,54/ 5,58	73,78/ 89,21/ 54,11	9,21/ 9,68/ 7,95
4	47,25/ 57,89/ 36,60	5,89/ 6,28/ 5,37	71,37/ 86,19/ 52,52	8,90/ 9,35/ 7,71

5	47,60/ 57,88/ 37,31	5,94/ 6,28/ 5,48	72,27/ 87,23/ 53,25	9,02/ 9,46/ 7,82
6	43,78/ 52,83/ 34,74	5,46/ 5,73/ 5,10	67,25/ 80,66/ 50,37	8,39/ 8,75/ 7,40
7	39,96/ 47,26/ 32,66	4,99/ 5,13/ 4,80	60,12/ 71,23/ 46,43	7,50/ 7,73/ 6,82
Validierung				
1	116,98/ 157,01/ 76,95	14,59/ 17,03/ 11,30	155,91/ 198,14/ 96,70	19,45/ 21,49/ 14,20
2	70,16/ 97,35/ 42,96	8,75/ 10,56/ 6,31	103,57/ 133,68/ 59,85	12,92/ 14,50/ 8,79
3	48,54/ 59,54/ 37,54	6,06/ 6,46/ 5,51	72,41/ 87,82/ 52,66	9,03/ 9,53/ 7,73
4	47,70/ 58,95/ 36,44	5,95/ 6,39/ 5,35	72,39/ 87,92/ 52,44	9,03/ 9,54/ 7,70
5	48,67/ 59,51/ 37,84	6,07/ 6,45/ 5,56	73,35/ 89,16/ 53,02	9,15/ 9,67/ 7,79
6	45,33/ 55,01/ 35,65	5,66/ 5,97/ 5,23	70,56/ 84,91/ 52,41	8,80/ 9,21/ 7,70
7	41,37/ 48,84/ 33,89	5,16/ 5,30/ 4,98	61,60/ 72,71/ 47,99	7,69/ 7,89/ 7,05
Test				
1	123,22/ 162,25/ 84,18	15,37/ 17,60/ 12,36	170,52/ 215,96/ 107,31	21,27/ 23,42/ 15,76
2	70,87/ 98,07/ 43,66	8,84/ 10,64/ 6,41	105,18/ 135,39/ 61,60	13,12/ 14,68/ 9,05
3	62,55/ 81,82/ 43,29	7,80/ 8,87/ 6,36	97,23/ 123,43/ 60,60	12,13/ 13,39/ 8,90
4	61,57/ 80,77/ 42,38	7,68/ 8,76/ 6,22	96,31/ 122,55/ 59,44	12,02/ 13,29/ 8,73
5	62,02/ 81,11/ 42,93	7,74/ 8,80/ 6,30	97,41/ 124,13/ 59,74	12,15/ 13,46/ 8,77
6	58,11/ 74,82/ 41,40	7,25/ 8,12/ 6,08	91,21/ 114,63/ 59,14	11,38/ 12,43/ 8,68
7	50,45/ 63,14/ 37,75	6,29/ 6,85/ 5,54	78,22/ 97,03/ 53,11	9,76/ 10,52/ 7,80

Anhang PV-Prognose

Irrtümer und Änderungen vorbehalten | 2016 SOLARWATT GmbH | AZ-TDB-PMS-0262
Dieses Datenblatt entspricht den Vorgaben der DIN EN 50380:2003 | REV 007 | 08/2016



Glas-Glas-Modul: SOLARWATT 60P

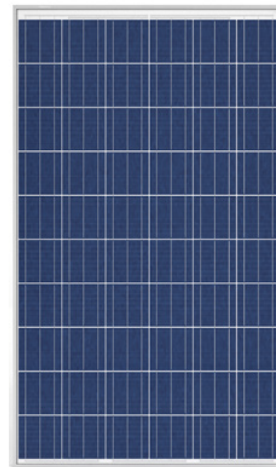
SOLARWATT Solarmodule

DIE INNOVATIVE GLAS-GLAS-GENERATION SOLARWATT 60P

- Super-Leichtgewicht durch 2 mm dünnes Glas
- Höchste Ertragszuverlässigkeit
- 100% Schutz gegen PID
- Höhere mechanische Belastbarkeit
- Höhere Brandsicherheit
- Polykristalline Hochleistungszellens
- 260Wp–285Wp (100% Plussortierung)

Produkteigenschaften

- langlebig
- belastbar
- ertragreich
- innovativ
- sicher
- blendarm
- ammoniakbeständig
- hagelbeständig
- salznebelbeständig



SOLARWATT Service

SOLARWATT Komplettschutz
inklusive (bis 1000 kWp*)

30 Jahre Produkt-Garantie
gemäß „Besondere Garantiebedingungen für SOLARWATT-Solarmodule“

Einfache Finanzierung
ohne zusätzliche Sicherheits-nachweise

30 Jahre Leistungs-Garantie
gemäß „Besondere Garantiebedingungen für SOLARWATT-Solarmodule“

Unkomplizierte Rücknahme
gemäß den Lieferbedingungen für SOLARWATT-Solarmodule

Made in Dresden
Herkunfts-Garantie
Qualität aus Deutschland

* in Italien bis 50 kWp

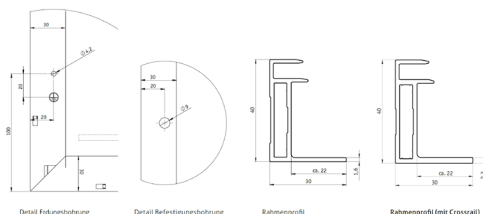
SOLARWATT GmbH | Maria-Reiche-Str. 2a | 01109 Dresden | Germany
Tel. +49 351 8895-333 | Fax +49 351 8895-111 | www.solarwatt.de
Zertifiziert nach DIN EN ISO 9001 und 14001 | BS OHSAS 18001:2007

Irrtümer und Änderungen vorbehalten | 2016 SOLARWATT GmbH | AZ-TDB-PMS-0262
Dieses Datenblatt entspricht den Vorgaben der DIN EN 50380:2003 | REV 007 | 08/2016

Technische Daten | SOLARWATT 60P



ABMESSUNGEN



* Für die Variante mit Zusatzausstattung Crossrail gilt als freigegebene Belastung nach SOLARWATT Montageanleitung: Auflast bei Quermontage¹⁾: bis 5500 Pa
Testbedingungen: Schrägbelastung mit 9000 Pa (Die Bedingungen berücksichtigen Sicherheitsfaktoren für Schneeüberhang und Eislast gemäß Eurocode 1.)

ALLGEMEINE DATEN

Modultechnologie	Glas-Glas-Laminat; Aluminiumrahmen
Deckmaterial	Gehärtetes Solarglas mit Antireflex-Veredelung, 2 mm
Verkapselung	EVA-Solarzellen-EVA, weiß
Rückseitenmaterial	Gehärtetes Solarglas, 2 mm
Solarzellen	60 polykristalline Hochleistungszellens
Maße der Zellen	156 x 156 mm
L x B x D	1680 ⁺¹ x 990 ⁺¹ x 40 ^{+0,3} mm
Gewicht	ca. 22,8 kg / ca. 24 kg mit Crossrail
Anschluss-technik	Kabel 2 x 1,0 m/4 mm ² , HC4-Steckverbinder
Bypass-Dioden	3
Anwendungs-kategorie	A (nach IEC 61730)
Max. Systemspannung	1000 V
Prüfungen zur mechanischen Belastbarkeit	Soglast bis 2400 Pa Auflast bis 5400 Pa
Freigegebene Belastungen nach SOLARWATT Montageanleitung	Auflast bei Quermontage ¹⁾ : 3500 Pa Testbedingungen: Schrägbelastung mit 5400 Pa (Die Bedingungen berücksichtigen Sicherheitsfaktoren für Schneeüberhang und Eislast gemäß Eurocode 1.) 1) Beachten Sie hierzu bitte die Angaben in der Montageanleitung
Qualifikationen	IEC 61215 Ed.2 IEC 61730 (inkl. Schutzklasse II)

ELEKTRISCHE EIGENSCHAFTEN BEI STC

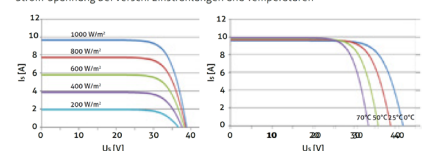
STC: Standard Test Conditions: Bestrahlungsstärke 1000W/m², Spektrale Verteilung AM 1,5 | Temperatur 25±2°C, entsprechend EN 60904-3

	260 Wp	265 Wp	270 Wp	275 Wp	280 Wp	285 Wp
Nennleistung P _n	260 Wp	265 Wp	270 Wp	275 Wp	280 Wp	285 Wp
Nennspannung U _{mp}	30,9 V	31,0 V	31,1 V	31,2 V	31,3 V	31,4 V
Nennstrom I _{mp}	8,50 A	8,63 A	8,76 A	8,89 A	9,02 A	9,15 A
Leerlaufspannung U _{oc}	38,1 V	38,3 V	38,5 V	38,7 V	38,9 V	39,1 V
Kurzschlussstrom I _{sc}	9,20 A	9,32 A	9,44 A	9,56 A	9,68 A	9,80 A

Mess-toleranzen bezogen auf P_{max} ± 5 %;
Reduktion des Modulwirkungs-grades bei Rückgang der Bestrahlungsstärke von 1000 W/m² auf 200 W/m² (bei 25 °C): 4 ± 2 % (relativ) / -0,6 ± 0,3 % (absolut).
Rückstrombelastbarkeit I_r: 20 A, Betrieb der Module mit eingespeistem Fremdstrom ist nur bei Verwendung einer Strangsicherung mit Auslösestrom ≤ 20 A zulässig.

KENNLINIEN (Leistungskategorie 280 Wp)

Strom-Spannung bei versch. Einstrahlungen und Temperaturen



ELEKTRISCHE EIGENSCHAFTEN BEI NOCT

NOCT: Normal Operation Cell Temperature: Bestrahlungsstärke 800 W/m², AM 1,5 | Temperatur 20 °C, Windgeschwindigkeit 1m/s, elektrischer Leerlauf

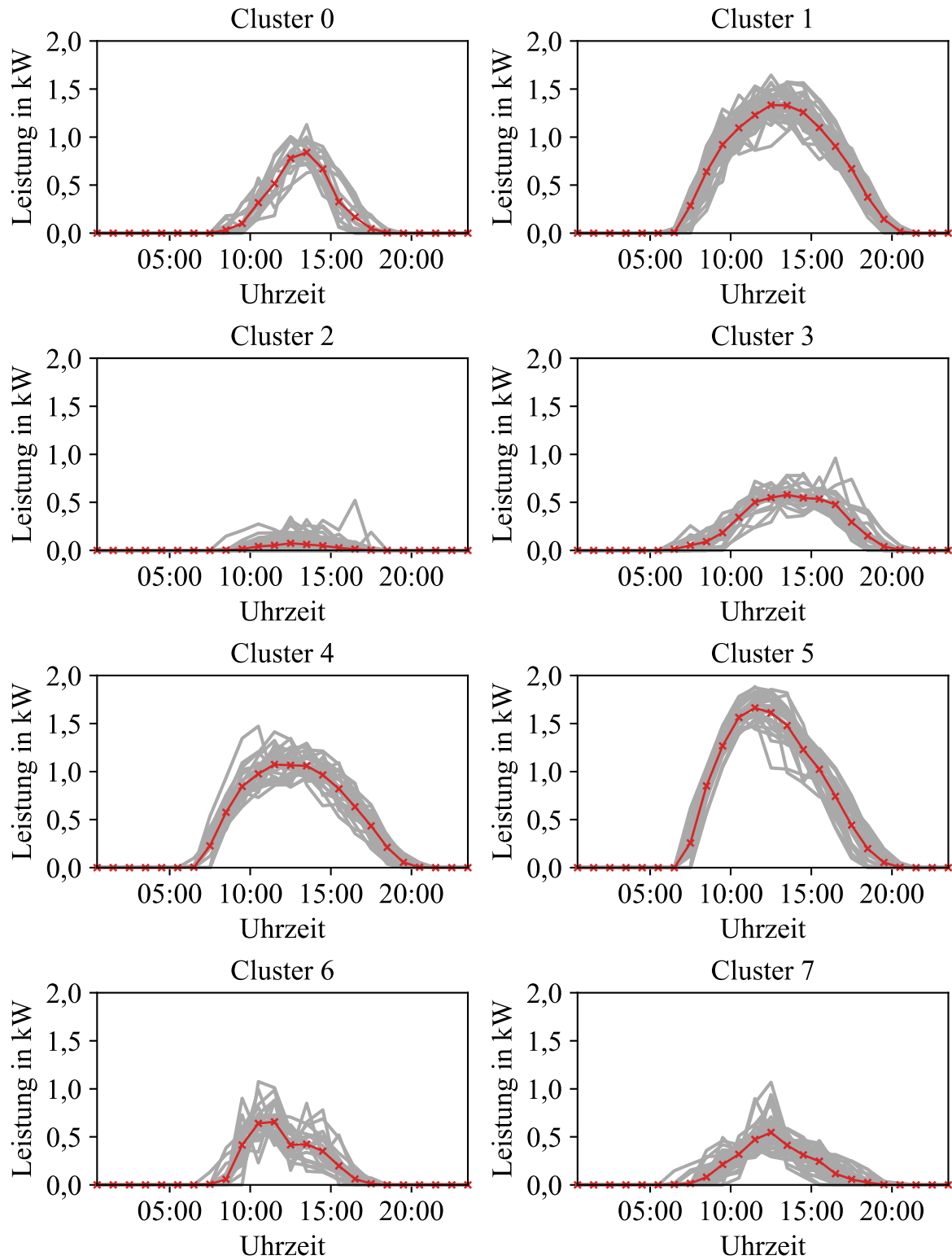
	191 W	195 W	198 W	202 W	206 W	209 W
Nennleistung P _n	191 W	195 W	198 W	202 W	206 W	209 W
Nennspannung U _{mp}	28,5 V	28,6 V	28,7 V	28,8 V	28,9 V	29,0 V
Leerlaufspannung U _{oc}	35,7 V	35,9 V	36,1 V	36,3 V	36,5 V	36,7 V
Kurzschlussstrom I _{sc}	7,43 A	7,53 A	7,63 A	7,72 A	7,82 A	7,92 A

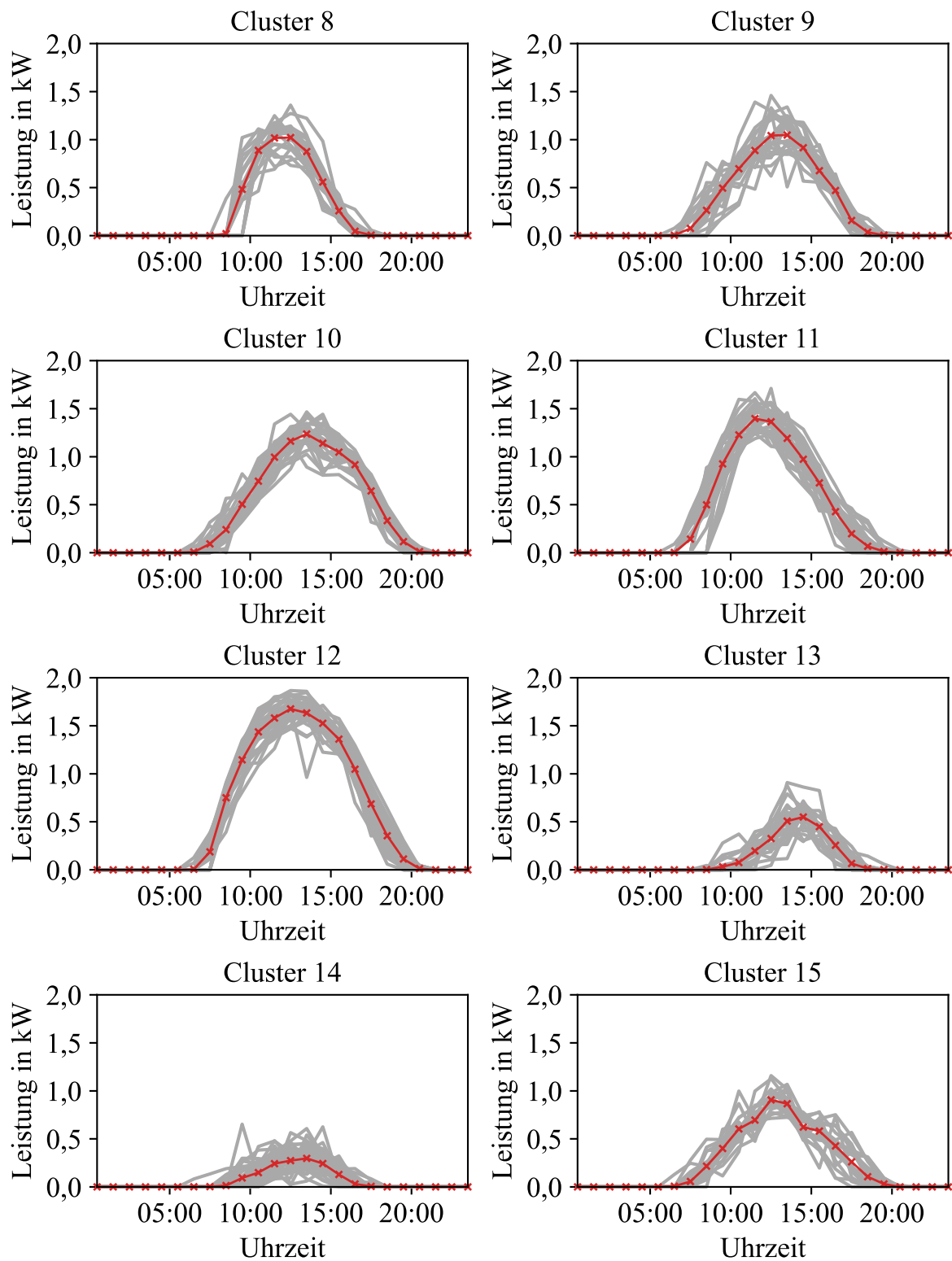
THERMISCHE EIGENSCHAFTEN

Betriebstemperaturbereich	-40 ... +85 °C
Umgebungstemperaturbereich	-40 ... +45 °C
Temperaturkoeffizient P _n	-0,41 %/K
Temperaturkoeffizient U _{oc}	-0,31 %/K
Temperaturkoeffizient I _{sc}	0,05 %/K
NOCT	45 °C

Abbildung A. 2: Datenblatt der betrachteten PV-Anlage der Hochschule Bielefeld

B. Anhang Korrekturmodell





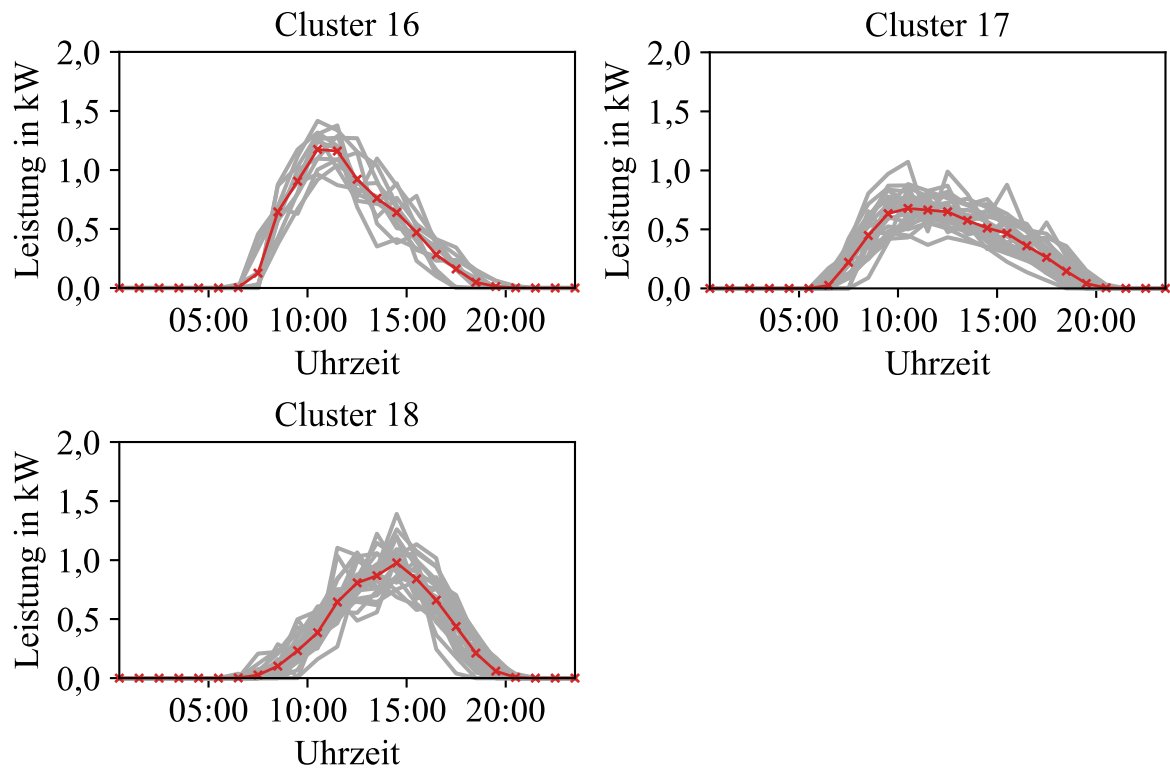


Abbildung B. 1: Mittlere Zeitreihen der PV-Leistung in Rot sowie alle Zeitreihen in Grau der 19 Cluster

Tabelle B. 1: Hyperparameterraum der LR-Methode

Hyperparameter	Durchlauf 1
Achsenabschnitt	[True, False]
Positive Koeffizienten	[True, False]
Beste Hyperparameter	True, True
MAE	120,67

Tabelle B. 2: Hyperparameterraum des DT-Modells

Hyperparameter	Durchlauf 1	Durchlauf 2	Durchlauf 3
Bewertungsmetrik der Splits	['squared_error', 'absolute_error', 'friedman_mse', 'poisson']	['absolute_error']	['absolute_error']
Maximale Tiefe des DTs	[3, 5, 7, 10, None]	[4, 5, 6, 15, 20]	[5, 25, 30, 40]
Minimale Anzahl an Datensätzen pro Split	[2, 5, 10, 20]	[2]	[2]

Minimale Anzahl an Datensätzen pro Blatt	[1, 2, 5, 10]	[10, 15, 20]	[15, 30, 40, 50]
Maximale Anzahl an Merkmalen pro Split	['sqrt', 'log2', None]	['sqrt']	['sqrt']
Beste Hyperparameter	'absolute_error', 5, 2, 10, 'sqrt'	'absolute_error', 5, 2, 15, 'sqrt'	'absolute_error', 5, 2, 15, 'sqrt'
MAE	104,46	103,97	103,97

Tabelle B. 3: Hyperparameterraum des KNN-Modells

Hyperparameter	Durchlauf 1	Durchlauf 2	Durchlauf 3
Anzahl der nächsten Nachbarn	[1, 5, 10, 15, 20, 25]	[23, 24, 25, 30, 35, 40]	[31, 32, 33, 34, 35, 36, 37, 38, 39]
Größe der Blätter für tree-basierte Algorithmen	[15, 20, 25, 30, 35, 40]	[5, 10, 15, 16, 17],	[1, 2, 3, 4, 5, 6, 7, 8, 9, 10]
Distanzmetrik	['manhattan', 'euclidean', 'minkowski']	'manhattan'	'manhattan'
Beste Hyperparameter	25, 15, 'manhattan'	35, 5, 'manhattan'	37, 3, 'manhattan'
MAE	100,47	100,12	99,97

Tabelle B. 4: Hyperparameterraum des RF-Modells

Hyperparameter	Durchlauf 1	Durchlauf 2	Durchlauf 3
Anzahl der DTs im RF	[50, 100, 150]	[125, 150, 200, 300]	[250, 275, 300, 350, 400]
Bewertungsmetrik der Splits	['squared_error', 'absolute_error', 'friedman_mse', 'poisson']	['squared_error']	['squared_error']
Maximale Tiefe des DTs	[None, 10, 20, 30]	[8, 9, 10, 11, 12]	[11, 12, 13, 14, 15, 16, 17, 18, 19]
Minimale Anzahl an Datensätzen pro Split	[2, 5, 10]	[2, 3, 4]	[4]

Minimale Anzahl an Datensätzen pro Blatt	[1, 5, 10]	[1]	[1]
Maximale Anzahl an Merkmalen pro Split	['sqrt', 'log2']	['sqrt']	['sqrt']
Beste Hyperparameter	150, 'squared_error', 10, 2, 1, 'sqrt'	300, 'squared_error', 12, 4, 1, 'sqrt'	275, 'squared_error', 12, 4, 1, 'sqrt'
MAE	103,11	102,85	102,84

Tabelle B. 5: Hyperparameterraum des XGBoost-Modells

Hyperparameter	Durchlauf 1	Durchlauf 2
Anzahl der DTs	[50, 100, 200]	[50, 500]
Lernrate	[0.01, 0.1, 0.3]	[0.001]
Maximale Tiefe des DTs	[3, 5, 7]	[5, 10, 15]
Minimale Summe der Gewichte	[1, 3, 5]	[3]
Mindestverlust für den Split	[0, 0.1, 0.3]	[0]
Anteil Trainingsdaten für das Training jedes DTs	[0.6, 0.8, 1.0]	[0.8]
Anteil Merkmale für Erstellung von jedem DT	[0.6, 0.8, 1.0]	[1.0]
L1-Regularisierungsterm	[0, 0.1, 0.5]	[0]
L2-Regularisierungsterm	[0.1, 1, 10]	[7, 10]
Beste Hyperparameter	50, 0.001, 5, 3, 0, 0.8, 1.0, 0, 10	50, 0.001, 5, 3, 0, 0.8, 1.0, 0, 7
MAE	105,23	348,74

Tabelle B. 6: Hyperparameterraum des Gradient Boosting Modells

Hyperparameter	Durchlauf 1	Durchlauf 2	Durchlauf 3
Anzahl der DTs	[50, 100, 200]	[25, 30, 50, 75, 400]	[50, 1000, 5000]
Lernrate	[0.01, 0.1, 0.3]	[0.1]	[0.1]
Maximale Tiefe des DTs	[3, 4, 5]	[5, 10, 15, 20, 25]	[5, 6, 7, 8, 9]
Anteil Trainingsdaten für das Training jedes DTs	[0.6, 0.8, 1.0]	[0.8]	[0.8]

Maximale Anzahl an Merkmalen pro Split	['sqrt', 'log2']	['sqrt']	['sqrt']
Minimale Anzahl an Datensätzen pro Blatt	[1, 2, 4]	[2]	[2]
Minimale Anzahl an Datensätzen pro Split	[2, 5, 10]	[4, 5, 6]	[5]
Beste Hyperparameter	50, 0.1, 5, 0.8, 'sqrt', 2, 5	50, 0.1, 5, 0.8, 'sqrt', 2, 5	50, 0.1, 6, 0.8, 'sqrt', 2, 5
MAE	105,16	105,16	104,59

Tabelle B. 7: Hyperparameterraum des CatBoost-Modells

Hyperparameter	Durchlauf 1	Durchlauf 2	Durchlauf 3
Anzahl der DTs	[200, 400]	[50, 100, 200]	[700]
Lernrate	[0.01, 0.1]	[0.1]	[0.01]
Maximale Tiefe des DTs	[4, 6, 10]	[5, 6, 7]	[12]
L2-Regularisierung für die Blätter	[1, 3, 10]	[3, 4, 5]	[2, 5]
Zufallsraummethode	[0.7, 0.8]	[0.8, 1.0]	[1.0]
Anzahl der Bins (Grenzwerte) für numerische Merkmale	[32, 64]	[32]	[32, 182]
Zufallsstärke	[0.5, 1.0]	[0.5]	[1.0]
Beste Hyperparameter	200, 0.1, 6, 3, 0.8, 32, 0.5	100, 0.1, 7, 5, 0.8, 32, 0.5	700, 0.01, 12, 2, 1.0, 32, 1.0
MAE	103,15	102,83	110,31

Tabelle B. 8: Hyperparameterraum der Ridge Regression

Hyperparameter	Durchlauf 1	Durchlauf 2
Optimierungsalgorithmus	['auto', 'svd', 'cholesky', 'lsqr', 'sparse_cg', 'saga']	['sparse_cg']
Regulierungsstärke	[1e-4, 1e-3, 1e-2, 1e-1, 1, 10, 100]	[1, 10, 100, 1000, 10000, 100000]
Achsenabschnitt	[True, False]	[True]
Beste Hyperparameter	'sparse_cg', 10, True	'sparse_cg', 10, True
MAE	93,85	93,85

Tabelle B. 9: Fehlermetriken für die Korrektur der PV-Prognose von verschiedenen Basismodellen für die Testdaten ohne der Clustering-Methode

ML-Modell	MAE in W	nMAE in %	RMSE in W	nRMSE in %
LR	182,90	8,67	270,95	12,85
KNN	167,63	7,95	254,84	12,08
DT	237,76	11,27	339,00	16,07
RF	166,72	7,91	252,68	11,98
CatBoost	177,82	8,43	273,80	12,98
XGBoost	378,27	17,94	534,12	25,33
Gradient Boosting	168,14	7,97	251,46	11,92
Ursprüngliche Prognose	184,66	8,76	267,49	12,68

Tabelle B. 10: Fehlermetriken für die Korrektur der PV-Prognose von verschiedenen Basismodellkombinationen eines Stacking-Ensembles für die Testdaten ohne der Clustering-Methode

Basismodellkombinationen	MAE in W	nMAE in %	RMSE in W	nRMSE in %
KNN, RF, LR, CatBoost, XGBoost, Gradient Boosting	167,34	7,94	252,80	11,99
KNN, RF, CatBoost, XGBoost, Gradient Boosting	166,29	7,89	251,27	11,92
KNN, RF, Gradient Boosting	166,60	7,90	251,33	11,92
KNN, RF, CatBoost, Gradient Boosting	166,59	7,90	251,29	11,92
KNN, RF	166,31	7,89	251,71	11,94
CatBoost, XGBoost, Gradient Boosting	169,91	8,06	251,53	11,93

Tabelle B. 11: Hyperparameter der Basismodelle und des Metamodells des Stacking-Ensembles

ML-Modell	Hyperparameter	Angabe
KNN	Anzahl der nächsten Nachbarn	37
	Distanzmetrik	'manhattan'
	Größe der Blätter für tree-basierte Algorithmen	3
RF	Anzahl der DTs im RF	275
	Maximale Tiefe des DTs	12
	Maximale Anzahl an Merkmalen pro Split	'sqrt'

	Minimale Anzahl an Datensätzen pro Blatt	1
	Minimale Anzahl an Datensätzen pro Split	4
	Bewertungsmetrik der Splits	'squared_error'
Gradient Boosting	Anzahl der DTs	50
	Lernrate	0,1
	Maximale Tiefe des DTs	6
	Anteil Trainingsdaten für das Training jedes DTs	0,8
	Maximale Anzahl an Merkmalen pro Split	'sqrt'
	Minimale Anzahl an Datensätzen pro Blatt	2
	Minimale Anzahl an Datensätzen pro Split	5
Catboost	Anzahl der DTs	700
	Lernrate	0,01
	Maximale Tiefe des DTs	12
	L2-Regularisierung für die Blätter	2
	Zufallsraummethode	1,0
	Anzahl der Bins (Grenzwerte) für numerische Merkmale	32
	Zufallsstärke	1,0
XGBoost	Anzahl der DTs	50
	Lernrate	0,001
	Maximale Tiefe des DTs	5
	Minimale Summe der Gewichte	3
	Mindestverlust für den Split	0
	Anteil Trainingsdaten für das Training jedes DTs	0,8
	Anteil Merkmale für Erstellung von jedem DT	1
	L1-Regularisierungsterm	0
	L2-Regularisierungsterm	7
Ridge Regression (Meta-modell)	Optimierungsalgorithmus	'sparse_cg'
	Regulierungsstärke	10
	Achsenabschnitt	True

C. Anhang Wissenschaftlicher Tätigkeitsnachweis

Betreute Masterarbeiten

- [S1] Katrin Handel, *Auswirkungen der Optimierungsfunktionen von Solar-Home-Batteriespeichern an ausgewählten Beispielen*. Masterarbeit, 2025.
- [S2] Majkel Denis Straus, *Optimale Auslegung eines Großbatteriespeichers am Gütersloher Mittelspannungsnetz*. Masterarbeit, 2025.
- [S3] Lars Engel, *Optimierung und Verbesserung einer PV-Prognose mit ML-Algorithmen*. Masterarbeit, 2023.
- [S4] Teresa Stratmann, *Prognose der Energiemenge am Hausanschluss mittels künstlicher neuronaler Netze unter Berücksichtigung der Vorhersageunsicherheit*. Masterarbeit, 2023.
- [S5] Melina Gurcke, *Entwicklung einer autonomen Steuerung für Niederspannungsnetze unter Beachtung der Netzstabilität basierend auf Fuzzy-Logik*. Masterarbeit, 2023.
- [S6] Oliver Runde, *Prognosemodell zur Vorhersage der lokalen Wolkenzüge zur Optimierung der Leistungsprognose einer Photovoltaikanlage*. Masterarbeit, 2021.
- [S7] Michael Pühs, *Lademanagementverfahren zur optimierten Ladung mehrerer Elektrofahrzeuge*. Masterarbeit, 2020.

Betreute Bachelorarbeiten

- [S8] Julius Dresselhaus, *Vergleichende Analyse zwischen Netzausbau und netzdienlicher Steuerung von PV-Batteriespeichern*. Bachelorarbeit, 2024.
- [S9] Angellina Ebenezer Anitha, *Optimierung einer Photovoltaik Prognose mit Ensemble Learning Methoden*. Bachelorarbeit, 2024.
- [S10] Katrin Handel, *Linearer Optimierungsalgorithmus zur Minimierung der Leistungsrückspeisung eines Niederspannungsnetzes durch Nutzung von Haushaltssolarenergiespeichern*. Bachelorarbeit, 2022.
- [S11] Alen Murtovi, *Konzeptionierung der Sekundärtechnik einer digitalen Ortsnetzstation zur Automatisierung des Mittel- und Niederspannungsnetzes*. Bachelorarbeit, 2022.
- [S12] Majkel Straus, *Ladeplanoptimierung für Elektrofahrzeuge mittels nichtlinearer Optimierung*. Bachelorarbeit, 2022.

Betreute Projektarbeiten

- [S13] Angellina Ebenezer Anitha, *Untersuchung verschiedener maschineller Lernverfahren zum Clustering für die Optimierung eines PV-Prognosemodells*. Projektarbeit, 2024.
- [S14] Majkel Straus, *Lademanagementsystem für netzintegrierte Batteriespeicher mittels nichtlinearer Optimierung*. Projektarbeit, 2024.
- [S15] Katrin Handel, *Control of distributed energy storage systems for minimum reverse flow in a distribution grid with high share of photovoltaic*. Projektarbeit, 2023.
- [S16] Lars Engel, *Optimierung eines Prognosemodells zur Vorhersage der Photovoltaikleistung unter Verwendung eines durch Clustering bestimmten Korrekturfaktors*. Projektarbeit, 2023.
- [S17] Teresa Stratmann, *Prognosemodell zur Vorhersage des kurzfristigen Lastprofils eines Prosumer-Haushaltes unter Verwendung von künstlich neuronalen Netzen und Long Short-Term Memory*. Projektarbeit, 2022.
- [S18] Dennis Streich, *Modellierung eines PV-Modells in Python*. Projektarbeit, 2022.
- [S19] Hermann Tetzlaff, *Analyse elektrischer Netze in Deutschland und Frankreich*. Projektarbeit, 2022.
- [S20] Christoph Alexander Lübke und Katrin Handel, *Weiterentwicklung einer Web-App zur Kommunikation zwischen Mensch und Ladesteuerung einer intelligenten Mehrfachsteckdose*. Projektarbeit, 2021.
- [S21] Jonas Giebeler, *Weiterentwicklung eines Photovoltaik-Modells zur Berechnung der erzeugten Leistung*. Projektarbeit, 2021.
- [S22] Doreen Heuking, Sophia Linnemann und Katrin Handel, *Entwicklung einer Web-App für das Projekt Power2Load*. Projektarbeit, 2020.
- [S23] Tobias Kusch, *Entwicklung eines künstlichen neuronalen Netzes zur Zustandsabschätzung im Verteilnetz*. Projektarbeit, 2020.

Wissenschaftliche Publikationen

- [B1] K. Schulte, A. Ebenezer Anitha and J. Haubrock, "Reducing forecast errors of a day-ahead photovoltaic power forecast by ensemble learning," *IEEE PowerTech 2025*, Kiel, Germany, 2025, unpublished.
- [B2] A. Ebenezer Anitha, K. Handel, K. Schulte and J. Haubrock, "Improved Post Processing Model for Photovoltaic Power Forecasting Based on Clustering," in *Advances in Computational Intelligence: 18th International Work-Conference on*

- Artificial Neural Networks, IWANN 2025, I. Rojas, G. Joya, and A. Catala, Eds., Coruña, Spain: Springer, 2025, unpublished.*
- [B3] M. Gurcke, M. Cziomer, K. Schulte and J. Haubrock, "Grid-Assistive Use of Private Battery Energy Storage Systems based on Fuzzy Logic," *4th Mediterranean Conference on Power Generation Transmission, Distribution and Energy Conversion – MED POWER 2024*, Athens, Greece, 2024.
 - [B4] K. Schulte, L. Engel, L. Quakernack, F. Liegmann and J. Haubrock, "A comparison of machine learning algorithms for the optimization of a day-ahead photovoltaic power forecast," *2024 IEEE PES Innovative Smart Grid Technologies Europe (ISGT EUROPE)*, Dubrovnik, Croatia, 2024, pp. 1-5, doi: 10.1109/ISGTEUROPE62998.2024.10863315.
 - [B5] F. Liegmann, K. Schulte, F. Annen, J. Haubrock, M. Kelker, S. Joseph, S. Raczka and C. Rehtanz, "Concept of a test bench for research into automatic resupply to improve the resilience of critical infrastructure," *2024 IEEE PES Innovative Smart Grid Technologies Europe (ISGT EUROPE)*, Dubrovnik, Croatia, 2024, pp. 1-5, doi: 10.1109/ISGTEUROPE62998.2024.10863253.
 - [B6] L. Quakernack, M. Gurcke, K. Schulte and J. Haubrock, "Grid-Oriented Control of Vehicle Batteries in a Cellular Grid Setup Based on Fuzzy Logic," *2024 IEEE PES Innovative Smart Grid Technologies Europe (ISGT EUROPE)*, Dubrovnik, Croatia, 2024, pp. 1-5, doi: 10.1109/ISGTEUROPE62998.2024.10863623.
 - [B7] J. Dresselhaus, K. Schulte, K. Handel, J. Haubrock and J. Arens, "Potential of battery energy storage systems to avoid grid expansion in the low-voltage grid," *IEEE Power and Energy Student Summit PESS*, Dresden, Germany, 2024.
 - [B8] K. Handel, A. Ebenezer Anitha, R. Rigo-Mariani, K. Schulte and J. Haubrock, "Optimal decision tree design for near real-time control of battery energy storage systems," *26th International Symposium on Symbolic and Numeric Algorithms for Scientific Computing (SYNASC 2024)*, Timișoara, Romania, 2024.
 - [B9] K. Schulte, K. Handel und J. Haubrock, "Netzdienliche Steuerung von verteilten Batteriespeichern in einem Verteilnetz mit einem hohen Anteil an Photovoltaikanlagen für eine Reduzierung der Rückspeisung," *Poster beim ETG CIRED Workshop 2023 (D-A-CH)*, München, Deutschland, 2023.
 - [B10] K. Schulte, L. Engel, L. Quakernack and J. Haubrock, "Optimized photovoltaic power forecast using k-means clustering based error reduction," *2023 IEEE PES Innovative Smart Grid Technologies Europe (ISGT EUROPE)*, Grenoble, France, 2023, pp. 1-5, doi: 10.1109/ISGTEUROPE56780.2023.10407521.

- [B11] K. Handel, K. Schulte, R. Rigo-Mariani, J. Haubrock and J. Arens, "Control of distributed energy storage systems for minimum reverse flow in a distribution grid with high share of photovoltaic," *2023 IEEE PES Innovative Smart Grid Technologies Europe (ISGT EUROPE)*, Grenoble, France, 2023, pp. 1-5, doi: 10.1109/ISGTEUROPE56780.2023.10408443.
- [B12] T. Jungh, B. Steinhagen, M. Hesse and K. Schulte, "Comparison of Different Machine Learning Models for Short-Term Load Forecasting at Transformer Level with High Amounts of Photovoltaic Generation," *2023 IEEE PES Innovative Smart Grid Technologies Europe (ISGT EUROPE)*, Grenoble, France, 2023, pp. 1-5, doi: 10.1109/ISGTEUROPE56780.2023.10407216.
- [B13] K. Schulte, O. Runde, M. Kelker and J. Haubrock, "Prediction of the local cloud cover to optimize photovoltaic system power forecast," *2021 IEEE PES Innovative Smart Grid Technologies Europe (ISGT Europe)*, Espoo, Finland, 2021, pp. 01-05, doi: 10.1109/ISGTEurope52324.2021.9640074.
- [B14] M. Schwienheer, K. Schulte, K. Kröger and J. Haubrock, "Realization of a power distributing electric vehicle charging system," *NEIS 2021; Conference on Sustainable Energy Supply and Energy Storage Systems*, Hamburg, Germany, 2021, pp. 1-6.
- [B15] K. Schulte and J. Haubrock, "Linear programming to increase the directly used photovoltaic power for charging several electric vehicles," *2021 IEEE Madrid PowerTech*, Madrid, Spain, 2021, pp. 1-6, doi: 10.1109/PowerTech46648.2021.9494914.
- [B16] M. Kelker, K. Schulte and J. Haubrock, "State estimation in low-voltage grids by using artificial neural networks in consideration of optimal micro phasor measurement unit placement," *NEIS 2020; Conference on Sustainable Energy Supply and Energy Storage Systems*, Hamburg, Germany, 2020, pp. 1-6.
- [B17] K. Schulte, M. Kelker and J. Haubrock, "Artificial neural networks to predict the node voltages in a low-voltage grid," *NEIS 2020; Conference on Sustainable Energy Supply and Energy Storage Systems*, Hamburg, Germany, 2020, pp. 1-6.
- [B18] M. Kelker, A. Berrada, K. Schulte and J. Haubrock, "Entwicklung und Validierung eines optimalen Platzierungsalgorithmus für μ PMUS im Niederspannungsnetz," *EnInnov 2020; 16. Symposium Energieinnovation*, Graz, Austria, 2020.
- [B19] P. Lohmann, M. Kelker, K. Schulte and J. Haubrock, "Auslegung eines Antriebstranges für einen Batterie-Elektrischen Zug," *EnInnov 2020; 16. Symposium Energieinnovation*, Graz, Austria, 2020.

- [B20] M. Kelker, K. Schulte, K. Kröger and J. Haubrock, "Development and validation of a neural network for state estimation in the distribution grid based on μ PMU data," *2019 Modern Electric Power Systems (MEPS)*, Wroclaw, Poland, 2019, pp. 1-6, doi: 10.1109/MEPS46793.2019.9394975.
- [B21] K. Schulte, M. Kelker, and J. Haubrock, "Predicting the local generated photovoltaic power by creating a forecast model using artificial neural networks and verifying the model with real data," *IEEE Power and Energy Student Summit PESS*, Magdeburg, Germany, 2019.
- [B22] M. Kelker, K. Schulte, D. Hansmeier, F. Annen, K. Kröger, P. Lohmann and J. Haubrock, "Development of a forecast model for the prediction of photovoltaic power using neural networks and validating the model based on real measurement data of a local photovoltaic system," *2019 IEEE Milan PowerTech*, Milan, Italy, 2019, pp. 1-6, doi: 10.1109/PTC.2019.8810719.

D. Hinweis zur Nutzung generativer KI

Für die sprachliche Überarbeitung dieser Arbeit wurden die generativen KI-Werkzeuge DeepL Write und ChatGPT verwendet. Diese wurden ausschließlich zur Rechtschreib- und Grammatikprüfung eingesetzt und nicht zur inhaltlichen Erstellung der Arbeit.